

MTA DOKTORI ÉRTEKEZÉS

A MODELLEZÉS SAJÁTOSSÁGAI IDŐSORI ANOMÁLIÁK ESETÉN

RAPPAI GÁBOR

PÉCS, 2016

A MODELLEZÉS SAJÁTOSSÁGAI IDŐSORI ANOMÁLIÁK ESETÉN

Tartalom

1	Témaválasztás – motivációk.....	4
2	A dolgozatban használt legalapvetőbb idősor-modellek, illetve tesztek.....	13
2.1	Stacionárius idősorok és alapvető modelljeik.....	13
2.2	Okság-fogalmak és a kauzalitás tesztjei.....	18
2.2.1	Okság a filozófiában.....	18
2.2.2	Okság a statisztikában.....	21
2.3	Trendek.....	26
2.3.1	Determinisztikus és sztochasztikus trend.....	27
2.3.2	Kointegráció és hibakorrekció.....	31
2.4	Néhány, a későbbi szimulációk során használt idősor.....	37
2.4.1	AR1 folyamatok szimulálása.....	38
2.4.2	VAR-modellel szimulált idősorok.....	40
2.4.3	Egységgyököt tartalmazó idősorok szimulációja.....	41
2.4.4	Kointegrált idősorok generálása.....	42
3	Rendszertelen idősorok kezelése.....	44
3.1	A probléma kezelése hiányzó adatok feltételezésével.....	46
3.2	Folytonos idejű modellek a rendszertelenül megfigyelt adatok elemzésében.....	53
3.3	Rendszertelen idősorok kiegészítésével (interpolálással) elkövethető hibák.....	56
3.3.1	A stacionaritás, illetve a sztochasztikus trend felismerése nemekvidisztáns idősorok esetén.....	58
3.3.2	Okság, illetve együttmozgás felismerése rendszertelen, illetve nem azonos frekvenciájú idősorok esetén.....	60
3.4	A GDP és az export növekedése közötti kapcsolat Magyarországon az elmúlt közel két évtizedben.....	67
4	Outlierek az idősorokban.....	72
4.1	Az outlierek típusai.....	75
4.2	Az outlier-szűrés leggyakoribb módszerei.....	78
4.2.1	Modell-független outlier-szűrés.....	79
4.2.2	Modell-alapú outlier-szűrés.....	83
4.3	Outlierrel terhelt idősorok kezelése.....	96
4.4	Outlierek, illetve outlier-szűrés hatása az idősorok tulajdonságainak megítélésében.....	98

4.4.1	Outlier megjelenésének hatása az adatgeneráló folyamat felismerésére	98
4.4.2	Okozhatja-e az outlier-szűrés a DGP félrespecifikálását?	103
4.5	Illusztratív példa az elmúlt 20 év Magyarországról	105
5	Strukturális törést tartalmazó idősorok modellezése	109
5.1	A strukturális törés kimutatására szolgáló tesztek.....	110
5.2	Strukturális törés egységgyököt tartalmazó, vagy trend-stacioner idősor esetén.....	116
5.2.1	Egységgyök-teszt ismert időpontban bekövetkezett strukturális törés esetén	119
5.2.2	Egységgyök-teszt, amennyiben a törés időpontja nem ismert	120
5.3	Kointegráció tesztje strukturális törést tartalmazó idősor esetén	121
5.4	Az adatgeneráló folyamatok félrespecifikálásának lehetősége strukturális törést tartalmazó idősorokban	123
5.4.1	A sztochasztikus trend létének felismerése strukturális törést tartalmazó idősorokban.....	124
5.4.2	A hagyományos próbák ereje több töréspont esetén.....	132
5.4.3	A strukturális törés szerepe az együttmozgás elfedésében.....	136
5.5	Magyarország külkereskedelmének legfontosabb mutatói az elmúlt két évtizedben.....	139
6	Aggregált idősorok modellezése	143
6.1	Aggregálás módozatai, alapvető jelölések.....	144
6.2	Stacionárius folyamatok időbeli aggregálásának következményei	146
6.3	Az időbeli aggregálás hatása trendet is tartalmazó idősorokban.....	150
6.4	Aggregálás hatása az idősorok kointegráltságára.....	152
6.5	Az időbeli aggregálás modellre gyakorolt hatásának összegzése.....	154
6.6	Trend, illetve okság az aggregált idősorokban.....	155
6.6.1	Az aggregálás hatása a trend megítélésében.....	156
6.6.2	Okság megítélésének torzulása aggregált idősorok esetén	158
6.7	Árfolyam-idősorok és összefüggéseik az elmúlt 2 évtizedben	161
7	Konklúziók helyett az etikus idősor-modellezésről	168
8	Irodalom	182

1 Témaválasztás – motivációk

A tudományos kutatás talán legfontosabb célja, hogy az ember(iség) feltérképezze az őt körülvevő világegyetem törvényszerűségeit, megismerje a társadalom működési elveit, szabályszerűségét, kísérletet tegyen a jövőbeli események előrejelzésére. „*A tények vizsgálatának tudományos módszere nem különbözik sem a vizsgálat tárgya, sem a kutató személye alapján, egyaránt használható társadalom- és természettudományokban*” állította Pearson, a modern statisztika megalapozója.¹ Mindezt úgy kell értenünk, hogy a tudományban vannak olyan alapvető összefüggések, illetve módszerek, melyek alkalmazása/alkalmazhatósága nem függ attól, hogy éppen milyen szakterületen járunk.

Az elmúlt időszakban sokszor fellángolt, majd elült a vita arról, hogy vajon (köz)gazdaságtan, vagy (köz)gazdaságtudomány az adekvát megnevezés, és pl. Móczár is arra jut, hogy a közgazdaságtan tudománnyá válását jelentős mértékben segítette elő, a természettudományokban (is) alkalmazott módszertan, illetve modellezési eszköztár rendszeres használata.² Kijelenthető, hogy a jelenségek alakulásában kimutatható tendenciaszerűség, és a vizsgált változók közötti ok-okozati összefüggések feltárását célzó modellek ilyen metadisziplináris eszközök, sőt megkockáztathatjuk, hogy egy-egy vizsgálódási terület (ha úgy tetszik tanszak) éppen attól válik tudományággá, hogy képes megragadni az általános trendeket, vagy a vizsgált jelenségek közötti kauzális kapcsolatokat.

A gazdasági jelenségek alakulásának *trendjei*, vagy az események közötti *ok-okozati összefüggések* statisztikai-ökonometriai modelljei³ több szempontból is *megkerülhetetlenek* a közgazdaságtudományi kutatásokban. Valószínűleg már ezeknek a motívumoknak az összegyűjtése is önálló kutatás tárgyát képezhetné, most azonban csak három olyan szempontot említek, melyek témaválasztásomat meghatározták.

1. Köztudomású, hogy a világ eddig élt egyik legnagyobb gondolkodója, Albert Einstein úgy vélte, hogy a világegyetemben uralkodó állapotok determinisztikusan következnek egymásból, ugyanakkor a jövőbeni állapotok nem determinálhatók. Mindez tudományos érvényű megalapozást nyert a káosz-elmélet megjelenésével, pontosabban Lorenz 1963-as cikkével, melyben kifejti, hogy egymással függvényeszerű kapcsolatban álló elemekből álló rendszerek is lehetnek olyan bonyolultak, hogy a legkisebb hatások következményének

¹ Lásd Pearson (1892) 6. old.

² Lásd Móczár (2008, 2009).

³ A következőkben – nagyvonalúan – közös jelzőként kezelem a statisztikai-ökonometriai kifejezést a gazdasági modellek azonosítása során, elfogadva azt a definíciót, miszerint az ökonometria „matematikai és statisztikai módszerek alkalmazása a gazdasági adatok elemzésében, amelynek célja az, hogy a közgazdasági elméleteknek empirikus tartalmat adjunk, és megerősítsük, vagy megcáfoljuk őket”. (Maddala, 2004, 33. old.)

megjósolása is lehetetlen.⁴ Könnyen belátható, hogy a társadalom ilyen nagy bonyolultságú rendszer, így elemzése során a függvényszerű összefüggések csak „igazodási pontnak” tekinthetők. A káoszelmélet roppant szerteágazó módszertani apparátust (pl. fraktálgeometria) mozgósítva terjed szét a különböző tudományok által vizsgált jelenségek leírásában, mindezzel – hely-, és főleg kompetencia-hiány miatt – nem foglalkozom. Ugyanakkor elgondolkodtató, hogy a kaotikus rendszerek leírásának módszertana végül is *visszatér a statisztikához*.

2. Az elmúlt évtizedben egyre-másra jelentek meg azok a tudományos művek, tanulmánygyűjtemények, összefoglaló monográfiák, melyeknek tárgya a *statisztika filozófiája*.⁵ Közben valószínűleg minden statisztikával elmélyültebben foglalkozó kutató érzi, hogy mit értünk ezen a szóösszetételén, a kifejezést pontosan definiálni lehetetlen. Úgy gondolom, hogy a statisztika módszertana általánosságban két rendkívül fontos, részben megkülönböztető jellemzővel bír: egyrészt *induktív*, vagyis a vizsgálódási módszere az adatokból, vagyis az empiriából indul ki, és ebből próbálja feltárni az adatgeneráló folyamatot; másrészt *sztochasztikus*, vagyis vizsgálódásának tárgyában meghatározó részként tekint a véletlenre, véletlenszerűségekre, azt modelljeiben immanens részként szerepelteti. Az elmúlt években többször megjelent már a statisztika (mint tudomány) definiálása során, hogy „*a statisztika a társadalomtudomány matematikája*”. Szimpatikus ez a definíció, hiszen jól mutatja, hogy a statisztikai módszertan éppen úgy a tudományos megismerés során használt *nyelv*, mint a matematika, csak éppen olyan problémák esetén alkalmazandó, amikor a jelenségek, illetve szereplők nagy száma miatt nem tételezhetünk fel tisztán függvényszerű kapcsolatokat.
3. Felmerülhet az a kérdés is, hogy miért kellene a gazdasági jelenségek modelljeit a többi tudománytól elkülönítve, mintegy önálló kutatási problémaként kezelnünk. Nem állítom, hogy a disszertációmban bemutatandó módszertan csak gazdaságtudományi problémák esetén használható, de biztos, hogy találhatunk olyan specialitásokat, melyek pl. a fizika, az orvostudomány, vagy akár a jogtudomány területén nem jelentkeznének. Hicks könyvében (lásd Hicks, 1979) az egész első fejezetet ennek a kérdésnek szenteli, és három szempontot sorol fel, amelyek alapján a gazdasági események közötti okságot különlegesen tekinti, és feltárását a filozófia eszközrendszerén túl gondolja megoldani. Elsőként a közgazdasági törvényszerűségek rendkívül hézagos („*extremely imperfect*”) voltát említi, majd hosszan értekezik a közgazdaságtan és az idő sajátos viszonyáról, eközben kifejti, hogy a közgazdász(kutató), ellentétben a történésszel, a múlt helyett a jelen összefüggéseit kutatja („... *the past is to the historian, the present is to the economist*”). Végezetül azt fejtegeti, hogy a közgazdászt leginkább a döntéshozatal, illetve a döntések következményei izgatják, és a

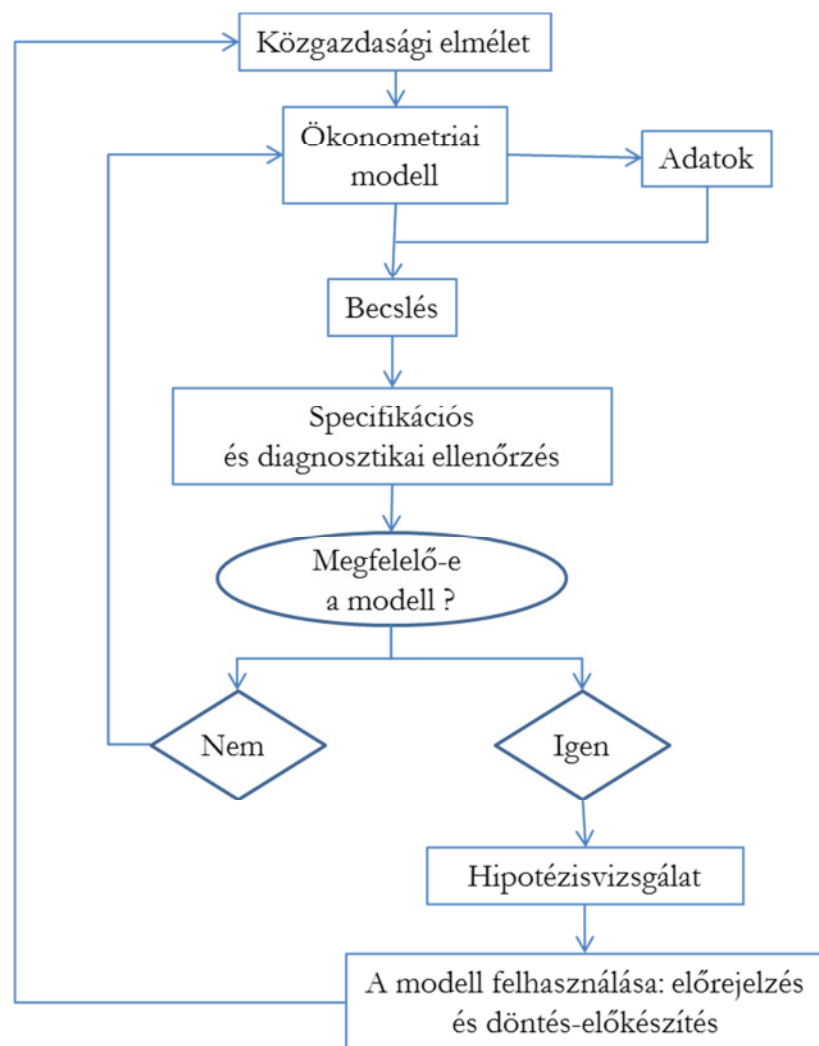
⁴ Lásd Lorenz (1963). A káoszelmélet és témánk kapcsolódásáról lásd pl. Nováky (1995), illetve Nováky (2015) tanulmányait.

⁵ Ezen munkák felsorolására nem vállalkozhatok. Talán kiindulásként érdemes megemlíteni Bandyopadhyay-Forster (2011) szerkesztésében megjelent könyvet, illetve ennek bevezető tanulmányát. A témával kapcsolatos akkori véleményem Rappai (2010) tanulmányban olvashatók.

köztük levő mélyebb összefüggések („*economics is specially concerned with the making of decisions, and with the consequences that follow from the decisions; the implications of this are profound*”).

Az előbbi gondolatok is azt valószínűsítik, hogy abban nincs vita a kutatók között, hogy a társadalmi-gazdasági jelenségek megismeréséhez a valóság lényeges elemeit jól megragadó, a lényegtelen részekről nagyvonalúan eltekintő *kanzális modellekre* van szükség. De hogyan „keletkezik” egy statisztikai-ökonometriai modell, milyen lépések szükségesek ahhoz, hogy eljussunk a probléma felmerülésétől a végkövetkeztetésig? Nem állíthatjuk, hogy az ökonometriai modellek folyamatábrája minden idevonatkozó szak-, illetve tankönyvben pontosan azonos, de a hangsúlyos lépések köre, illetve ezek sorrendje, szinte minden esetben nagyon hasonló.

Maddala (2004) könyvében javasolt felépítés többé-kevésbé általánosnak tekinthető, eszerint az ökonometriai modellezés vázlatos sémája az alábbi:



1-1. ábra: A statisztikai-ökonometriai modellezés menete (Maddala, 2004, 38. old. alapján)

Hasonló ábrát találunk Mills-Patterson (2006) 12. oldalán, de számos további ökonometria alapmű (pl. Greene, 2008) ugyanezen lépéseket sorolja fel, amikor a statisztikai-ökonometria modellezés fázisait bemutatja. A hazai szakirodalomban érdemes fellapozni Hajdu és mtsai (1987) könyvét, mely az első magyar nyelvű ökonometria tankönyvnek tekinthető. Figyelemre méltó, a modellépítés, elsősorban az előrejelzések gyakorlatát támogató megállapításokat tartalmaz Besenyei és mtsai (1977) könyve, vagy Mundruczó (1998) posztumusz műve, melyben a fenti ábra egyes elemei további, részletező felbontásban szerepelnek. Kevésbé meglepő módon hasonló ábra található Rappai (2013) 11. oldalán is.

A modellezés eszköztárát bemutató irodalom szinte végtelen, mind összefoglaló jellegű szakkönyvekkel, mind egészen speciális részleteket tartalmazó cikkekkel tele vannak a könyvtárak, folyóiratok. Ugyanakkor a témával foglalkozó kutatók mindegyike megerősítené azt a kijelentést, miszerint az ökonometria modellezés irodalma túlnyomó részt a paraméterbecslés módszertanával, illetve a specifikációanalízissel, a modell megfelelőségére vonatkozó hipotézisek ellenőrzésének eszköztárával foglalkozik.⁶ Éppen úgy, ahogy a modellspecifikáció előtti közgazdasági elméletek kidolgozását sem tekintik a modellezők a saját, szűken vett feladatuknak, *a megfelelően tisztított, a vizsgálat céljaira előkészített adatállomány összeállítását is „segédmunkának” tekintik*. Pedig, ha valaki foglalkozott már empirikus gazdaságmodellezéssel, akkor tudja, hogy – kifejezetten az informatikai támogató-eszközök rohamos fejlődését követően – az *adat-előkészítés* rendkívül időigényes, ugyanakkor kulcsfontosságú feladat. Éppen ezért disszertációmban a modellezés ezen fáziséval kívánok kiemelten foglalkozni, vagyis azokkal a feladatokkal, amely a *statisztikus* által végzendő, ha ökonometria modellt épít, vagy annak építésében részt vesz. (Nem gondolom, hogy a valóságban ez a két szereplő, a statisztikus és az ökonometria modellező ilyen szigorúan szétválasztható lenne, de úgy érzem a téma-lehatárolás szempontjából mindez egy jól használható megkülönböztetés!)

Érthető módon nem térhetek ki valamennyi gazdaságmodellezési problémára, ugyanakkor szeretnék olyan kérdéseket feldolgozni, amelyek gyakoriak, általános jellegűek. A dolgozatban három meglehetősen széles körben ismert modellel, vagy – talán így helyesebb – gazdaságmodellezési eszközzel foglalkozom:

1. a gazdasági jelenségek közötti ok-okozati összefüggések feltárásával (kauzalitás-modellekkel),
2. a változók alakulását meghatározó tartós tendenciák (trendek) kimutatásával,
3. a dinamikus egyensúlyi állapotok leírásával (közös trendek, kointegráló regresszió felírásával).

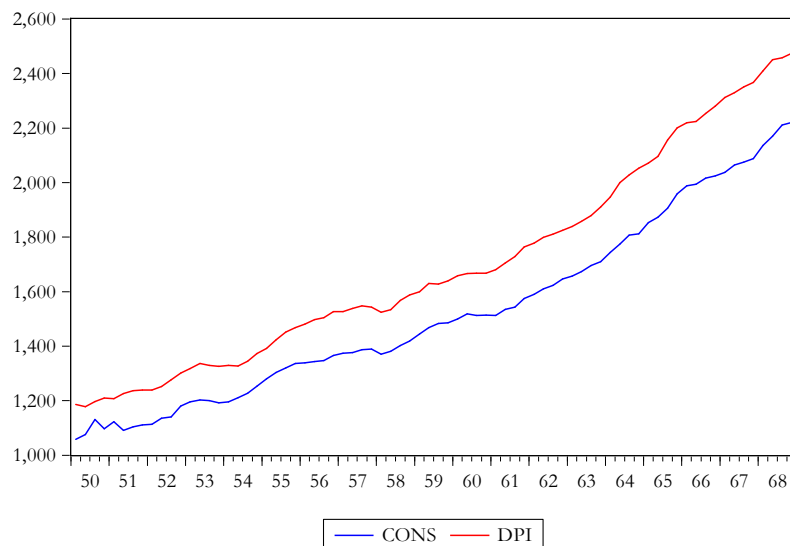
⁶ Elég csak végigtekintünk a Közgazdasági Nobel-Emlékdíjasok névsorán. A hivatalos méltatások alapján az ökonometria modellezésben elért eredményért kapta a díjat Frisch és Tinbergen (1969, dinamikus ökonometria modellek); Klein (1980, ökonometria modellek a gazdaságpolitika megalapozásában); Haavelmo (1989, az ökonometria valószínűségelméleti megalapozása); Heckman és McFadden (2000, mikroökonometria); Engle és Granger (2003, idősor-modellek az ökonometriában); valamint Sargent és Sims (2011, vektorautoregresszív modellek és okság). Valamennyi említett eredmény a modellspecifikáció, illetve a paraméterbecslés (esetleg specifikációanalízis) területén született.

A 2. fejezetben látni fogjuk, hogy az idősorokban kimutatható trendek, illetve a két, vagy több idősor között meglevő kointegráció tulajdonképpen szintén az okság megnyilvánulása (az első esetben az idő múlása és az időszaki érték közötti ok-okozati összefüggés, a második esetben a hibakorrekción keresztül megjelenő okság), vagyis leegyszerűsítve mindhárom esetben a kauzális relációkat vizsgálom. Ezen a ponton a téma körülhatárolásának egy roppant fontos pontjához értünk, ugyanis dolgozatomban *nem* az előbbieken felsorolt modellek, összefüggések definiálásával, a szükséges paraméterbecslési és hipotézisellenőrzési eljárások részletes bemutatásával kívánok foglalkozni, hanem azzal a kérdéssel, hogy *milyen* – a klasszikus értelemben vett modellezés előtti – *feladatokat kell elvégezniünk, hogy eredményeink relevánsak legyenek*. Az elmúlt mintegy három évtizedes modellezői tapasztalatom ugyanis azt mutatja, hogy a tényleges eredmények sokszor „elrejtőzhetnek”, vagy kifejezetten az elméletileg elvártnak ellentmondó megállapítások születhetnek azért, mert az *adattisztítás*, vagy, ahogy a dolgozatban néhányszor használom, az „adatmasszírozás” nem volt elég alapos és körültekintő.

A disszertáció motivációjának bemutatására álljon itt egy példa! Amennyiben Granger a kauzalitás irodalmát megalapozó cikkének⁷ születésekor megvizsgálta volna, hogy kimutatható-e ok-okozati összefüggés az USA makrofogyasztása és a lakosság rendelkezésére álló jövedelme között⁸ a második világháborút követő időszakban, akkor az 1950 és 1968 közötti időszak negyedéves adatai alapján (76 megfigyelés) azt tapasztalta volna, hogy a vizsgált időszakban a változók szinte párhuzamosan növekedtek, mint azt az 1-2. ábra mutatja:

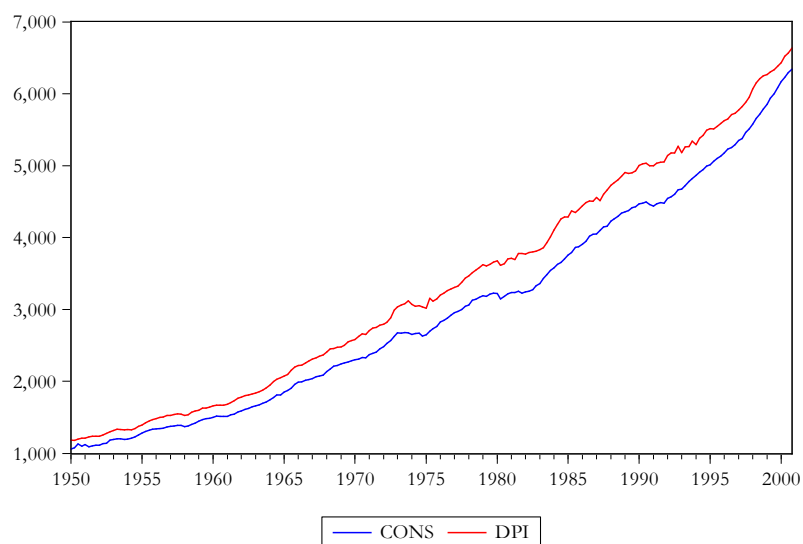
⁷ Granger (1969).

⁸ Mindez praktikusán azt jelenti, hogy érvényes a keynes-i fogyasztási modell (lásd pl. Mellár, 2003). Ebben a hangsúlyozottan illusztratív példában egyelőre nem foglalkozom a hamis regresszió lehetőségével, mindössze azt szeretném bemutatni, hogy azonos modellfeltevések teljesen eltérő eredményre vezethetnek, ha különböző időperiódusokat modellezünk.



1-2. ábra: Az USA makrofogyasztása (CONS) és a lakosság rendelkezésére álló jövedelem (DPI),
1950-1968 között negyedéves bontásban (változatlan áron, milliárd \$)

A Granger-okságot tesztelő klasszikus Wald-próba empirikus tesztstatisztikája $F = 12,375$ ($p < 0,001$), vagyis a jövedelem a fogyasztás előidejű okának tekinthető. Ugyanakkor, ha Granger a Közgazdasági Nobel-emlékdíj átvételére készülve ugyanezt a vizsgálatot, ugyanúgy a saját maga által javasolt próbával, de – mivel ekkor már erre lehetősége nyílt volna – egy jóval hosszabb időszakra 1950 és 2000 közötti adatokkal⁹ (szintén negyedéves bontású idősorra, azaz 204 megfigyelésre) végezné el, akkor egészen más eredményt kapott volna. A 1-3. ábra mutatja a két változó együttmozgását, ami – első szemrevételezés alapján – nem is nagyon tér el a korábbtól:



1-3. ábra: Az USA makrofogyasztása (CONS) és a lakosság rendelkezésére álló jövedelem (DPI),
1950-2000 között negyedéves bontásban (változatlan áron, milliárd \$)

⁹ Az adatok forrása mindkét esetben Greene (2008), 5.1 példa.

Az ábrák hasonlósága ellenére most azt tapasztaljuk, hogy mintegy 50 éves időhorizonton vizsgálva az oksági összefüggést az empirikus próbafüggvény-érték $F = 2,294$, amely minden ésszerű szignifikancia-szinten a nullhipotézis (azaz az okság hiánya) elfogadása mellett szól ($p = 0,104$).

Vajon mi lehet a magyarázat az előbbi jelenségre? Elméleti közgazdászok nyilván azzal érvelnének, hogy a neoklasszikus összefüggések idejétmúltak, és a jövedelem és a fogyasztás kapcsolatát más, lényegesen összetettebb modellek írják le (Friedman-féle permanens jövedelem hipotézis, Hall-féle életpálya jövedelem elmélet, racionális várakozások, stb.). A gazdaságtörténészek valószínűleg bizonyítani tudnák, hogy időközben olyan változások zajlottak le a társadalmi-gazdasági környezetben, amelyek minden bizonnyal alátámasztják az ok-okozati összefüggés megszűnését. (... és a további, általam végtelenül tisztelt gazdaságtudományi részterületekről még nem is ejtünk szót!)

Ebben a disszertációban mindezek helyett statisztikus magyarázatokat keresek, nyilván nem kikerülve az elméleti összefüggéseket, ám fő célom a gazdaság modellezésére szolgáló statisztikai-ökonometriai eszközök érzékenységtesztje. Ezért dolgozatomban azt kívánom megvizsgálni, hogy a gazdasági jelenségek között meglevő kauzalitási viszonyok, a változók egyedi, vagy közös trendjei milyen mértékben mutathatók ki, ha a jelenségeket leíró idősorok *zajosak*. Vizsgálni fogom

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

idősorokat, alapvetően azt tartva az elemzés fókuszában, hogy egy esetlegesen meglevő ok-okozati összefüggés, vagy sztochasztikus (közös) trend milyen módon, és milyen gyakorisággal képes „elbújni” annak következtében, hogy az idősoraink (vagy valamelyikük) a fenti problémák valamelyikével terhelt.

Gyakran halljuk, olvassuk, hogy korunk társadalma az ún. *információs társadalom*. Hihetetlen mennyiségű adat, információ keletkezik, és ezek egy része szinte azonnal hasznosul, míg más része szinte soha nem lesz feldolgozva. Új tudományágak születnek részben a statisztika eszköztárára használva, sőt – magát már a statisztika elemzési eszköztára nélkül definiálva – megjelent az *adatbányászat* módszertana. Az ún. *big data problémát*, mint korunk egyik legnagyobb jelentőségű módszertani kihívását emlegetik, és egyre inkább terjed az a nézet, hogy ez nem egyszerűen informatikai probléma, hanem az adatelemzés eszköztárának is meg kell újulnia. A big data megjelenésének egyik legszembeötlőbb jele, hogy szinte egy szempillantás alatt extrém hosszúságú idősorokat figyelhetünk meg, olyan új modellezési irányokra lehetünk figyelmesek, mint az extrém

nagy megfigyelési sűrűségű folyamatok analízise.¹⁰ Umberto Eco a *Velíná-k és a csönd* című írásában¹¹ azt fejtegeti, hogy korunk cenzúrája már nem az eltagadandó hírek letiltásával, hanem a hírcsatornák szinte kezelhetetlen *zajosságával* éri el hatását, az átadandó üzenet ugyanis már szinte alig hámozható ki a rengeteg szinte felesleges adat, információ közül. Dolgozatomban arra teszek kísérletet, hogy néhány, korábban említett, idősort „zajosító” hatás, a disszertáció szóhasználatával *anomália*, hatékony kiszűrési módozatait keressem, ezáltal néhány olyan makrogazdasági összefüggést is megtaláljak, melyek az elmúlt két évtized Magyarországon a homályban maradtak.

Hangsúlyozni kívánom, hogy a disszertáció statisztikai-módszertani indíttatású, vagyis a legfőbb célom annak bemutatása, hogy a hiányzó, vagy rendszertelen adatok, az outlierok, a strukturális törések, vagy az időbeli aggregáció miként torzítja az okság kimutatására szolgáló statisztikai próbákat. A fejezetekben (tipikusan ezek végén) található, a magyar nemzetgazdaság idősorai között kimutatandó ok-okozati relációk, vagy trendek, noha nem érdektelenek, de csak részben eredeztetik a disszertáció téziseit.¹² Alapvető célom a mára már standardnak számító *adatelőkészítési, szűrési eljárások* módszertanának árnyalása.

A következő (második) fejezetben vázlatosan összefoglalom a későbbiek megértéséhez szükséges idősor-elemzési terminológiát, bemutatom a kauzalitás legáltalánosabban használatos filozófiai megközelítéseit (definícióit), illetve bemutatom 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 többször említett anomáliák milyen mértékben képesek megváltoztatni a gazdasági idősorok alakulására vonatkozó hipotéziseink megítélését. Ezeknek a részeknek azonos a szerkezete, előbb bemutatom a statisztikai problémát, majd összefoglalom a (modellezés szempontjából) káros jelenség kimutatásának, illetve kiküszöbölésének módszereit. Ezt követi néhány illusztratív, szimulált példa, 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 tartom, hogy a dolgozat módszertani eredményeit felhasználva néhány gondolatot szánjunk a felvetődő modellezés-etikai kérdéseknek is. Úgy vélem, hogy a már említett adatmennyiség robbanás, valamint a hihetetlenül gyorsan szaporodó elemzési eljárások nagy kihívások elé állítják a kutatókat, hiszen azt sejtetik, hogy nincs álta-

¹⁰ Az extra nagy frekvenciájú idősorok modellezésének szép példája Engle (1996) tanulmánya.

¹¹ Megjelent Umberto Eco: *Ellenségét alkotni* (Európa Könyvkiadó, Budapest 2011) című esszégyűjteményében.

¹² A felhasznált magyar makroidősorok forrása egyrészt a Központi Statisztikai Hivatal a „KSH jelentő” című kiadványai (gördülő módon felhasználva), valamint a www.ksh.hu és a www.mnb.hu honlapok nyilvános adatállományai.

lános érvennyel 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.

Meggyőződésem, hogy a (statisztikai) modellezés kompromisszum-keresés, ahol a valósághűség versus egyszerűség (egyszerűsítés) dilemmájában keressük az optimális megoldást. Egy ilyen diszsertáció esetében is bizonyára fel fog merülni még számos módszer és technika, amelyet értelmetlenül kihagytam, illetve sok olyan összefüggés, amely nagyobb teret igényelt volna. Ahogy egy statisztikai modell, úgy egy akadémiai doktori értekezés is kompromisszumok eredményeként jön létre, kérem a Tisztelt Olvasót ennek szem előtt tartásával értékelje írásomat!

2 A dolgozatban használt legalapvetőbb idősor-modellek, illetve tesztek

Az akadémiai doktori értekezés műfaja nyilvánvalóan nem teszi lehetővé, hogy a felhasználandó idősor-elemzési módszereket tankönyvi alapossággal tárgyaljuk.¹³ Ugyanakkor szükségét érzem annak, hogy a dolgozatban használt fogalmakat, alapmodelleket, azok jelölésrendszerét röviden összefoglaljam annak érdekében, hogy a továbbiakban egyértelműen hivatkozhatassak az irodalomból ismert eredményekre. Hasonlóképpen, noha nem törekszem részletes filozófia-történeti eszmefuttatásokra, de ebben a fejezetben tömören ismertetem a későbbiekben alkalmazott okság-fogalmakat, illetve azok kialakulását. A fejezet harmadik részében bemutatom a kauzalitás meglétének, illetve hiányának tesztelésére kidolgozott statisztikai-ökonometria próbákat, ezek alkalmazási feltételeit.

2.1 Stacionárius idősorok és alapvető modelljeik

A sztochasztikus idősori modellezés kiemelt fogalma a stacionaritás. Egy sztochasztikus folyamatot (*adatgeneráló folyamatot*, data generating process, *DGP*) szigorúan *stacionáriusnak*, vagy szigorúan *stacionárnak* nevezünk, ha bármely $y_1, y_2, \dots, y_T$ megfigyelés-sorozat együttes eloszlása megegyezik az $y_{1+s}, y_{2+s}, \dots, y_{T+s}$ megfigyelés-sorozat együttes eloszlásával, bármely T és s esetén. A stacionaritás jelentősége a modellezés során úgy interpretálható, hogy amennyiben fennáll, úgy az empirikus idősor az elméleti sztochasztikus folyamat homogén mintájának tekinthető, így teljes hosszában alkalmazható a modell becslésére (ezáltal tulajdonképpen teljesül az idősormodellezésben elvárt *ergodicitási* feltétel is).

Az együttes eloszlások meghatározása meglehetősen körülményes eljárás, így a gyakorlatban általában megelégszünk a megfigyelés-sorozat első két momentumának vizsgálatával. Az általánosan használt, ún. *gyenge értelemben vett stacionaritás* definíciója szerint az idősor (kovariancia)stacioner, ha érvényes rá az alábbi három összefüggés

$$\begin{aligned} E(y_t) &= \mu \\ \text{Var}(y_t) &= E[(y_t - \mu)(y_t - \mu)] = \sigma^2 < \infty \\ \text{Cov}(y_t, y_{t-s}) &= E[(y_t - \mu)(y_{t-s} - \mu)] = \gamma_s \end{aligned}$$

Mindez szavakkal kifejezve annyit jelent, hogy egy folyamat stacionárius, ha várható értéke és varianciája időben állandó, valamint autokovariancia-struktúrája stabil.

¹³ Ha esetleg valaki ezt, mármint a tankönyvi alaposságot, igényli, akkor – többek között – javaslom Box-Jenkins (1970), Hamilton (1994), illetve Michelberger és mtsai (2001) munkáit, melyeken még azt is végigkövethetjük, hogy miként változik a „megkövetelt” matematikai apparátus az idősor-elemzés módszertanában.

A „leghíresebb” stacionárius folyamat a *fehér zaj* folyamat (*white noise process*). Definíciója szerint, ε_t fehér zaj, amennyiben

$$\begin{aligned} E(\varepsilon_t) &= \mu \\ \text{Var}(\varepsilon_t) &= \sigma^2 \\ \gamma_s &= \begin{cases} \sigma^2, & \text{ha } s = 0 \\ 0, & \text{ha } s > 0 \end{cases} \end{aligned}$$

vagyis praktikusán, ha az autokovarianciák értékei minden tényleges késleltetés esetén 0-k. Éppen ezért szokták a fehér zaj folyamatokat *korrelálatlan folyamatoknak* nevezni. Abban az esetben, ha a várható érték nem egyszerűen konstans, hanem 0, a folyamatot *nullaátlagú fehér zaj* folyamatnak (*zero mean white noise process*) nevezzük.¹⁴

A stacioner folyamatokra felírt sztochasztikus idősor-modellek legegyszerűbb válfaja az ún. *mozgóátlag modell* (moving average, *MA*). Induljunk ki egy roppant egyszerű stacionárius folyamatból

$$y_t = \mu + \varepsilon_t \quad (2.1)$$

ahol ε_t ($t = 1, 2, 3, \dots$) *iid véletlen változók* sorozata.¹⁵ A (2.1) folyamat várható értéke konstans, hiszen $E(y_t) = \mu + E(\varepsilon_t) = \mu$, varianciája egyenlő a véletlen változó (konstans) varianciájával, autokovariancia függvénye pedig csak ε_t -től függ, vagyis definíció szerint stacionárius.

Abban az esetben, ha feltételezzük, hogy a korábbi időszakokban megjelent véletlen értékek a mai időszori értéket már nem befolyásolják, akkor egy stacionárius modell legjobb előrejelzése a várható érték, ugyanakkor ez az „azonnali felejtés” általában nem reális feltételezés. Ha feltételezzük, hogy a korábbi időszakokban történtek befolyásolják még a mai megfigyeléseket, akkor van létjogosultsága a *mozgóátlag-modelleknek*, melyek

$$y_t = \mu + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} + \dots + \theta_q \varepsilon_{t-q} \quad (2.2)$$

formájúak, ahol a mai értékre vonatkozó modellben a mai és a korábbi hibatagok¹⁶ lineáris kombinációja (alkalmasan választott paraméterek mellett súlyozott átlaga) jelenik meg.

¹⁴ Általában (és ha külön nem jelzem a dolgozatban is) ez utóbbit hívjuk – kicsit pongyolán – fehér zajnak.

¹⁵ Az *iid* rövidítés – a szokásoknak megfelelően – a *független, azonos eloszlást* (independent identically distributed) jelöli. A későbbiekben többször használok a *Nid* jelölést, ami a *független normális eloszlást* (independent normally distributed) jelöli.

¹⁶ A véletlen változó egy korábbi realizációját majdnem tökéletes szinonimaként fogjuk *sokknak*, *innovációnak* vagy *sztochasztikus zajnak* (angolul *disturbance*-nek) nevezni.

Az $MA(q)$ folyamat várható értéke és a szórása konstans, autokovariancia értékei pedig csak a megfigyelések távolságától függnének, mindennek bizonyítását lásd pl. Brooks (2002) 236-238. pp. A mozgóátlag folyamatok tehát olyan stacionárius folyamatok, melyekben az idősor aktuális értéke a várható értékétől és a korábban bekövetkezett sokkoktól függ.

A sztochasztikus folyamatok modellezésének másik alapeszköze az *autoregresszív* (autoregressive, AR) *modell*, melyben feltételezzük, hogy az idősor mai értéke függ saját korábbi értékeitől. Egy p -ed rendű autoregresszív folyamat ($AR(p)$ -folyamat) leírható az alábbi formulával

$$y_t = \mu + \varphi_1 y_{t-1} + \varphi_2 y_{t-2} + \dots + \varphi_p y_{t-p} + \varepsilon_t \quad (2.3)$$

Bebizonyítható¹⁷, hogy egy $AR(p)$ folyamat stacionárius, ha az alábbi

$$1 - \varphi_1 z - \varphi_2 z^2 - \dots - \varphi_p z^p = 0$$

ún. *karakterisztikus egyenlet* z -re vonatkozó valamennyi megoldása abszolút értékében nagyobb 1-nél. (Gyakran használjuk az előbbi leírás szinonimájaként, hogy a karakterisztikus egyenlet valamennyi gyöke az egységsugarú körön kívül helyezkedik el.)

Az előző két sztochasztikus folyamat kézenfekvő kombinációjából alakul ki a stacionárius folyamatok esetén általánosan alkalmazott modell, az *autoregresszív mozgóátlag* ($ARMA$) *modell*. Alapgon dolata szerint a mai időszori érték nemcsak saját korábbi értékeitől, hanem a mai és az idősort korábban ért sokkok nagyságától is függ. Az általános $ARMA(p,q)$ folyamatot leíró egyenlet:

$$y_t = \mu + \varphi_1 y_{t-1} + \varphi_2 y_{t-2} + \dots + \varphi_p y_{t-p} + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} + \dots + \theta_q \varepsilon_{t-q} \quad (2.4)$$

ahol ε_t nulla várható értékű fehér zaj.

Érdemes ismét hangsúlyoznunk, hogy mivel a modell mozgóátlag tagja mindig stacionárius, így egy stacionárius folyamat modellezésben az $ARMA$ -specifikáció valódi alkalmazhatósága azon múlik, vajon az autoregresszív részben becsült paraméterek alapján felírható karakterisztikus egyenlet gyökei az egységkörön kívül vannak-e.

Az előbbi $ARMA$ -modellek egy alternatív felírásához vezessük be az ún. *késleltetési operátor*(oka)t, ami(k) az

¹⁷ A bizonyítást lásd Schlittgen-Streitberg (1991) 93-94. pp.

$$\begin{aligned}
Ly_t &= y_{t-1} \\
L^2 y_t &= y_{t-2} \\
&\vdots \\
L^r y_t &= y_{t-r}
\end{aligned}$$

átalakításokat generálja(ák). Ekkor egy (2.4) folyamat például felírható így is

$$y_t = \mu + \sum_{i=1}^p \varphi_i L^i y_t + \varepsilon_t + \sum_{j=1}^q \theta_j L^j \varepsilon_t \quad (2.5)$$

Egy további jelölésbeli egyszerűsítéssel még könnyebben használható formulát kaphatunk, hiszen amennyiben

$$\begin{aligned}
\varphi(L) &= 1 - \varphi_1 L - \varphi_2 L^2 - \dots - \varphi_p L^p \\
\theta(L) &= 1 + \theta_1 L + \theta_2 L^2 + \dots + \theta_q L^q
\end{aligned}$$

akkor az ARMA-folyamat

$$\varphi(L) y_t = \mu + \theta(L) \varepsilon_t \quad (2.6)$$

formában, vagy az ezzel ekvivalens

$$y_t - \frac{\mu}{\varphi(L)} = \frac{\theta(L)}{\varphi(L)} \varepsilon_t \sim MA(\infty) \quad (2.6a)$$

is felírható.

Az egyváltozós autoregresszív modellek többváltozós általánosításai a *vektor autoregresszív modellek* (VAR-modellek). A VAR-modellek többegyenletes rendszerek, praktikusán ez magával vonja, hogy egynél több eredményváltozó (jelenség) együttes vizsgálatára alkalmasak. A legegyszerűbb (kétváltozós, csak elsőrendű késleltetéseket tartalmazó) esetben a modell a következő

$$\begin{aligned}
y_{1t} &= \beta_{10} + \beta_{11} y_{1,t-1} + \beta_{12} y_{2,t-1} + \varepsilon_{1t} \\
y_{2t} &= \beta_{20} + \beta_{21} y_{1,t-1} + \beta_{22} y_{2,t-1} + \varepsilon_{2t}
\end{aligned} \quad (2.7)$$

ahol a két véletlen változó egymástól független fehér zaj (ez utóbbi feltételezés egyébiránt sokszor sérül, ami identifikációs problémákhoz vezethet).

Az előbbi VAR(1) modell mátrix alakja

$$\mathbf{y}_t = \begin{bmatrix} y_{1t} \\ y_{2t} \end{bmatrix} \quad \boldsymbol{\beta}_0 = \begin{bmatrix} \beta_{10} \\ \beta_{20} \end{bmatrix} \quad \mathbf{B}_1 = \begin{bmatrix} \beta_{11} & \beta_{12} \\ \beta_{21} & \beta_{22} \end{bmatrix} \quad \boldsymbol{\varepsilon}_t = \begin{bmatrix} \varepsilon_{1t} \\ \varepsilon_{2t} \end{bmatrix} \quad (2.8)$$

$$\mathbf{y}_t = \boldsymbol{\beta}_0 + \mathbf{B}_1 \mathbf{y}_{t-1} + \boldsymbol{\varepsilon}_t$$

A (2.8) modell több irányban is könnyen általánosítható:

- egyrészt figyelembe vehetünk elsőrendűnél magasabb rendben késleltetett eredményváltozó értékeket is,
- másrészt nyilván elképzelhető, hogy az egyenletek nem csak autoregresszív, hanem mozgóátlag tagot is tartalmaznak, ekkor jutunk a VARMA-modellekhez,
- továbbá kettőnél több változós VAR-modellek is felírhatók, ami komplexebb összefüggések vizsgálatát is lehetővé teszik.

A többváltozós, k -ad rendben késleltetett VAR-modell ($VAR(k)$ -modell) általános alakja:

$$\mathbf{y}_t = \boldsymbol{\beta}_0 + \mathbf{B}_1 \mathbf{y}_{t-1} + \mathbf{B}_2 \mathbf{y}_{t-2} + \dots + \mathbf{B}_k \mathbf{y}_{t-k} + \boldsymbol{\varepsilon}_t \quad (2.9)$$

ahol $\mathbf{y}_t$ tetszőleges számú eredményváltozó t -edik időpontban megfigyelt értékéből képzett vektor.

További kézenfekvő kiterjesztése a VAR-modelleknek, ha a jobb oldalon olyan változókat is szerepeltetünk, melyek a modell szempontjából teljes mértékben exogének, vagyis empirikus értékeik a modellen kívül határozódnak meg. Vegyünk például egy exogén változókat is tartalmazó VAR(2)-modellt:

$$\mathbf{y}_t = \boldsymbol{\beta}_0 + \mathbf{B}_1 \mathbf{y}_{t-1} + \mathbf{B}_2 \mathbf{y}_{t-2} + \mathbf{A} \mathbf{X}_t + \boldsymbol{\varepsilon}_t \quad (2.10)$$

ahol $\mathbf{X}_t$ az exogén változó(k), $\mathbf{A}$ pedig az ezekhez tartozó regressziós együtthatók mátrixa. Az ilyen (bizonyos irodalomban VARX-szel jelölt) modellek esetén is tesztelhető (a $\mathbf{A}$ mátrix elemeire megfogalmazott megszorítások alapján) akár az exogén változó szükségessége, akár az exogén és az endogén változók közötti okság megléte, mindez már átvezet az okság fogalmának és tesztjeinek ismertetéséhez.

2.2 Okság-fogalmak és a kauzalitás tesztjei

2.2.1 Okság a filozófiában

A következőkben rövid és meglehetősen szubjektív áttekintést adok arról, hogy az okság (kauzalitás) meglétét milyen mértékben fogadják el a klasszikus, illetve mai filozófusok, illetve hogyan definiálják azt.¹⁸ Hangsúlyozom, hogy a filozófusok, illetve irányzatok összeválogatása a dolgozat témájának megfelelően történt, így cseppet sem tekinthető átfogó filozófiatörténetnek. Biztos, hogy kevés olyan alapvető gondolat található a filozófiában, amelynek ennyire „végletes” megítélései élnek egymás mellett, mint az *okság* fogalomnak. A teljes elvetéstől a „mindennek ez az alapja” megközelítésig számos, külön-külön önmagában akár el is fogadható felfogást ismerünk, a választás közülük nyilvánvalóan nem témája a disszertációnak.

A *szeptikusok* egyenes tagadják, hogy egyáltalán létezhet elméleti okság, vagyis teljesen feleslegesnek tartják a jelenségek ok-okozati rendszerként történő megjelenítését. Sextus Empiricus (Szeptosz Empeirikosz) a Kr. u. II. században élő görög orvos, az ókori szeptikus filozófia „alapművének” számító *Adversus mathematicos* című művében tagadja a szillogisztikus bizonyítás lehetőségét és az okság meglétét (lásd pl. Kendeffy, 1998). Megítélése szerint az okság reláció, így ami valaminek az oka, az az okozata relatívumaként áll fönn. A relatívumnak azonban nincs önálló egzisztenciája, tehát nem rendelkezik az októl elvárható „előidejűség” vagy „önálló létezés” attribútumokkal. Hasonlóan szeptikus álláspontot képvisel Wittgenstein, aki egyenesen úgy fogalmaz „az oksági kapcsolat babona” (Wittgenstein, 2004. 5.13. szakasz). Sok más tudomány, például a fizika, kerüli az okság fogalmának használatát, Russel egyenesen azt írja: „Minden filozófus, bármelyik iskolához tartozzék is, úgy képzei, hogy az okozatiság a tudomány egyik alapvető axiómája, vagy posztulátuma; viszont az olyan fejlett tudományokban, mint például az égi mechanika, az „ok” szó furcsa módon soha nem fordul elő. ... Úgy hiszem, az okság törvénye, ... egy elmúlt kor emléke, s a monarchiához hasonlóan csak azért él tovább, mert – tévesen – ártalmatlannak vélik.” (Russel, 1976, 291. old.)

Az okságról alkotott filozófiai vélekedés másik „végpontja” talán Aquinói Szent Tamás, akinek értekezéseiben egyértelműen az ok-okozati láncolatok képezik minden bizonyítás alapját. Talán elég csak a „*Summa theologiae*” című művére utalni, melyben öt ésszerű okot sorol fel Isten létének bizonyítására. Ezek közül az egyik a létesítő okság premisszája: minden létezőnek van valamilyen létesítő oka, vagyis minden okozat egy okságot tételez fel. Aquinói Szent Tamás vélekedése szerint ez nem mehet a végtelenségig így, tehát létezik egy kezdeti ok, „akit mindenki Istennek nevez”.

David Hume szerint, a *kauzalitás* nem más, mint két dolog egymásra következésének az eszméje. Szerinte az ok és okozat „különböző létezők”, közöttük nincs szükségszerű kapcsolat. Az, hogy az ok és okozat között szükségszerűséget vélünk felfedezni, csak azért van, mert oly sokszor tapasztaltuk,

¹⁸ A filozófia-történet okság-felfogásainak vázlatos, jelenlegitől csak kismértékben eltérő ismertetése egy korábbi tanulmányomban (Rappai, 2011) már napvilágot látott.

hogy az egyik jelenség a másik után következik, hogy ez által a „*statisztikai bizonyítás*” véljük az okságnak. Hume szerint tehát az okság időbeli egymásra következés és állandó kapcsolat, de nem feltétlenül szükségszerűség.

Hume okságelméletének három sarkalatos pontja (lásd Hume, 2006, Absztrakt):

- a jelenségek térbeli vagy időbeli érintkezésének (szomszédosság) szükségessége,
- az ok időbeli elsőbbsége, valamint
- az okság szabályszerűsége, vagyis a jelenségek állandó együtt járása.

Mindez azt jelenti, hogy az ok olyan dolog, amit egy másik dolog, az okozat követ, ráadásul úgy, hogy az okhoz hasonló összes dolgot az okozathoz hasonló dolgok követik. Hume szerint tehát az állandó kapcsolat az okság elégséges feltétele, és viszont: ha a két jelenség oksági viszonyban van, akkor közöttük állandó a kapcsolat (szükséges feltétel).

Nyilvánvaló, hogy az elméletnek számos gyenge pontja van, melyre többen rámutattak, így például

- a természetben megannyi szabályosság található, melyek között nincs oksági viszony,
- ugyanakkor arra sincs érvünk, hogy nem létezhet egyszeri okság.

E. Szabó szerint (lásd E. Szabó, 2008) „*igaza van Hume-nak, hogy a kauzális kapcsolatot a regularitások észlelésén keresztül ismerjük fel, a regularitás azonban a kauzális kapcsolatok felismerésének a szükséges feltétele, és nem magának a kauzális kapcsolatnak.*” Mindez a dolgozatunk szempontjából azért roppant nagy jelentőségű, mert ha ezen állítást elfogadjuk, akkor a jelenségeket leíró idősorok közötti együttmozgás nem az ok-okozati viszonynak, hanem mindössze annak felismerhetőségének feltétele. Könnyen belátható, hogy a hume-i okságfelfogásnak egyik nemkívánatos következménye, hogy a gyakran ismétlődő, véletlenszerű állítások a törvényszerűségek kategóriájába esnek. Mindezen feloldására a filozófiatudományban kétféle válasz ismeretes, az egyik a *hume-i*, a másik a *nem hume-i* elmélet. Az előbbi elméletek szerint az oksági kapcsolatok állandó és nem feltétlenül szükségszerű kapcsolatok, a második álláspontot képviselők, azt állítják, hogy törvények és véletlenek között az a különbség, hogy a törvények szükségszerű kapcsolatokat jellemeznek, a véletlenek pedig nem.

A korábban említett, hume-i értelemben nehezen feltételezhető egyszeri okság értelmezésére fejlesztette ki a *tényellentétes* (kontrafaktuális) okság elméletét David Lewis (lásd Lewis, 1973). A tényellentétes okságfelfogás szerint *A* oka *B* eseménynek, ha igaz az állítás, miszerint „*ha A nem következett volna be, akkor B sem következett volna be*”. Lewis tényellentétes okságról szóló könyvének híres példája, egyben kiinduló gondolata, hogy „*ha a kengurunak nem lenne farka, felbukfencezne*”. Ezt persze úgy kell(ene) értelmeznünk, hogy egy olyan világban, ahol minden pont ugyanolyan, mint eb-

ben, csak az a különbség, hogy a kengurunak nincs farka, akkor a kenguru felbukfencezne. Jól látható, hogy ez a kiinduló tézis magában hordozza az elmélet leglényegesebb problémáját, vagyis azt a kérdést, hogy milyen mértékben kell hasonlítani egymásra a két világnak, amit vizsgálunk? A statisztikai modellezés nyelvén tehát azt kérdezzük, hogy az okság vizsgálata során alkalmazott modellben mely változóktól tekinthetünk el anélkül, hogy az „a világ megváltozását” okozná? Lewis bevezeti az ún. „*újra kombinálhatóság elvét*”, amivel itt nem foglalkozunk, ennek viszonylag részletes taglálása megtalálható E. Szabó korábban említett kitűnő tanulmányában. Érdekes megjegyezni, hogy ez a típusú okságfelfogás hívta életre az ún. *eseményanalízis (event study)* módszertanát. Ennek kialakulását, általános módszertanát, illetve a módszeren alapuló elemzéseket is közölünk Bedő-Rappai (2004), illetve Bedő-Rappai (2006) tanulmányokban.

Könnyen belátható, hogy a kontrafaktuális ok-definíció nem alkalmas az ok definiálására, vagyis szükséges, de nem elégséges feltétel, úgyis fogalmazhatunk, hogy a lewis-i értelemben definiált parciális faktorok csak összességükben tekinthetők olyan oknak, amely már kiváltja az okozatot. Ezen probléma kiküszöbölése érdekében alakult ki az ún. *elégséges feltétel* elmélet. Ennek értelmében, ha egy jelenség *szükséges* feltétele egy másiknak, azt jelenti, hogy ha az első jelenség nem következik be, akkor a másik sem. Az pedig, hogy egy jelenség *elégséges* feltétele a másiknak azt jelenti, hogy ha az első jelenség fennáll, akkor a második is. Mindezt úgy is megfogalmazhatjuk, hogy X pontosan akkor okozza Y -t, ha X szükséges és elégséges Y -hoz. Az elmélet cáfolatára könnyen található ellenpélda, hiszen egy jelenség több okból is előállhat.

Megoldást az elégséges feltétel elmélet problémájára elsőként Mackie (1965) javasolt, az ún. *INUS-elméletben*. Ezen elv értelmében az ok szükséges, de nem elégséges feltétel a feltételek egy olyan rendszerében, amely elégséges, de nem szükséges az okozathoz. (Az elmélet neve angol akronym: *Insufficient but Non-redundant part of an Unnecessary but Sufficient condition*.) Hasonló felfogást képvisel Pascal is, amikor úgy fogalmaz: „... *ugyanaz az esemény bizonyos esetekben véletlen eseménynek tekintendő, más esetekben pedig okozatilag teljes mértékben meghatározottnak, attól függően, hogy milyen körülmények között vizsgáljuk.*”¹⁹

Számos olyan példát lehetne találni, melyek felhívják a figyelmet arra, hogy az ok elégségesége még INUS-értelemben sem jelent determináltságot, mindez átvezet a *valószínűségi kauzalitás* (sztochasztikus kauzalitás) gondolatmenetébe, vagyis az okozat bekövetkezhet az ok nélkül is, és fordítva; előfordulhat, hogy az ok bekövetkezése ellenére sem lép fel az okozat. Mindezt úgy foglalkozhatunk össze, hogy a sztochasztikus okság értelmében az ok bekövetkezése megnöveli az okozat bekövetkezésének valószínűségét. Formalizálva a szokásos jelölésekkel A esemény B oka, ha fennáll, hogy

¹⁹ Pascal és Fermat levelezését Rényi (2004) idézi. Ezúton is köszönetet mondok a 2011-ben tragikus hirtelenséggel elhunyt egykori gazdaságtörténész professzoromnak, Tóth Tibornak, aki nemcsak erre a levelezésre hívta fel a figyelmet, hanem számos alkalommal foglalkozott idevonatkozó filozófiai ismereteim bővítésével.

$$\Pr\{B|A\} > \Pr\{B|\bar{A}\} \quad (2.11)$$

A valószínűségi okság megközelítéssel kapcsolatban számos problémát vetnek fel, ezek közül (ismét E. Szabó, 2008 művére hivatkozunk) a három legfontosabb

- megtévesztő, mert azt sugallja, hogy az események között aszimmetrikus viszony áll fenn, miközben könnyen belátható, hogy (2.11) csak annyit jelent, hogy A és B esemény között pozitív korreláció van,
- ráadásul a korreláció ténye nem feltétlenül jelent okságot, nagyon sokszor az ún. *közös ok* magyarázza a korreláció meglétét,
- mindemellett számos irodalmi példát találhatunk arra, hogy korrelációs kapcsolat mutat-ható ki két jelenség között, miközben közöttük sem direkt, sem közös ok típusú kauzali-tás sincs.

Nem kívánok részletesen foglalkozni a *Reichenbach-féle közös ok*-elvvel (első leírása Reichenbach, 1956, viszonylag részletes ismertetése Szabó, 2001), mindössze annyit érdemes rögzítenünk, hogy ennek értelmében nincs korreláció okság nélkül. A sztochasztikus kauzalitás elve értelmében tehát *a statisztikus korreláció tökéletes indikátora az oksági mechanizmusoknak*.

Az ok-okozati összefüggések filozófiai megközelítései könyvtárnyi irodalommal rendelkeznek, mindebbe csak belekontárokodtam, ugyanakkor az értekezés szempontjából az előzőekből három lényeges következtetést le kell vonnunk:

1. A gazdasági jelenségek közötti okság feltárása *idősor-modellezési feladat*.
2. A kauzalitás létét *korrelációs kapcsolatokon* (vagy mint azt a 2.1 alfejezetben áttekintettük VAR-modelleken) *alapuló próbák*kal tudjuk tesztelni.
3. A fennálló ok-okozati kapcsolatok feltárhatósága nagymértékben *függ* a modellezési kör-nyezettől, a statisztika nyelvén *a modellspecifikációtól*.

A következőkben az ökonometria által az okság feltárására kidolgozott, leggyakrabban alkalmazott próbákat, illetve modelleket tárgyaljuk.

2.2.2 Okság a statisztikában

Az előbbi alfejezet filozófiai fejtegetései legalábbis azt megmutatták, hogy már magának az okság fogalmának definiálása is meglehetősen komplex feladat, így nem lepődhetünk meg azon, hogy a kauzalitás tesztelésére alkalmas statisztikai próbák is csak mintegy négy és fél évtizede állnak rendelkezésre.

Elsőként bemutatom az okság klasszikus (granger-i) ökonometriai felfogását, majd röviden áttekintem az elmúlt 40 év legfontosabb „mérőföldköveit”, melyek a klasszikus felfogás kereteit feszegetik.

Az ökonometriában bevett gyakorlat értelmében, valószínűleg a könnyű operacionalizálhatóság okán, x változót y okának tekintjük, ha segítségével y -ra jobb becslést tudunk adni, mint nélküle.²⁰ Az elgondolás elsőként Wiener (1956) tanulmányában olvasható, ám általánossá válása Granger (1969) megjelenését követő időszakra tehető. Granger egy későbbi tanulmányában²¹ azt fejtegeti, hogy a különböző folyóiratokat tallózó citáció-keresőkben (Science Citation Index, vagy Social Science Citation Index) néhány év alatt több ezer olyan cikk bukkant fel, amelyek az oksággal foglalkoznak, mindez önmagában indokolta, hogy egy általánosan használható, tesztelhető definíció szülessen. Granger szerint az okság definiálását megelőzően mindössze két axiómát kell felállítani:

- a múlt és a jelen okozhatja a jövőt, de a jövő nem lehet oka múltbéli eseményeknek,
- az információhalmaz (információs állapot), amelyben az okságot vizsgáljuk, nem tartalmazhat redundáns információkat, vagyis, ha egy változó determinisztikus kapcsolatban van egy másik változóval, akkor valamelyiket el kell hagyni.

Mindezek alapján az okság legáltalánosabb definíciója x_t oka y_{t+1} -nek, ha

$$\Pr\{y_{t+1} \in A | \Omega_t\} \neq \Pr\{y_{t+1} \in A | (\Omega_t - x_t)\}$$

ahol Ω_t az adott környezeti állapotot leíró, a t -edik időpillanatban rendelkezésre álló információk összessége, A a lehetséges kimenetek halmaza.

Az előbbi definíciót általában a statisztikai hipotézisellenőrzés számára könnyebben kezelhető formában ismerjük. Ennek értelmében nullhipotézisünk szerint x *nem oka* y -nak, mivel segítségével nem adható jobb előrejelzés y -ra, mint akkor, amikor csak y múltbéli értékeit vizsgáljuk. Vagyis

$$\begin{aligned} H_0 : MSE(\hat{y}_{t+1} | y_t, y_{t-1}, \dots) &= MSE(\hat{y}_{t+1} | y_t, y_{t-1}, \dots, x_t, x_{t-1}, x_{t-2}, \dots) \\ H_1 : MSE(\hat{y}_{t+1} | y_t, y_{t-1}, \dots) &> MSE(\hat{y}_{t+1} | y_t, y_{t-1}, \dots, x_t, x_{t-1}, x_{t-2}, \dots) \end{aligned} \quad (2.12)$$

ahol MSE , az előrejelzés átlagos négyzetes hibáját jelöli.

²⁰ Az ún. Wiener-Granger okság fenti megfogalmazása akár nyolc különböző oksági viszonyt is előidézhethet (a nyolc esetet szemléletesen tárgyalja Kőrösi és mtsai, 1990), melyek közül az alábbiakban csak a legkézenfekvőbbet ismertettem.

²¹ Lásd Granger (1980).

A próba az alábbi regresszió becslését és paramétereinek együttes tesztelését igényli:

$$y_t = \alpha_0 + \alpha_1 y_{t-1} + \dots + \alpha_k y_{t-k} + \beta_1 x_{t-1} + \beta_2 x_{t-2} + \dots + \beta_m x_{t-m} + \varepsilon_t \quad (2.13)$$

ahol $\varepsilon_t \sim iid(0, \sigma^2)$.

Ekkor a hipotézisrendszer felírható:

$$\begin{aligned} H_0 : \beta_1 = \beta_2 = \dots = \beta_m = 0 \\ H_1 : \exists j \in 1, 2, \dots, k \rightarrow \beta_j \neq 0 \end{aligned} \quad (2.14)$$

aminek tesztelése Wald-típusú próbával viszonylag egyszerűen megoldható. (Könnyen belátható, hogy itt egymásba ágyazott modellek közötti választásról van szó, hiszen amennyiben a nullhipotézis igaz a (2.13) modell „mindössze” egy AR(k) modell.) A nullhipotézis *elvetése* számunkra azt jelenti, hogy vélelmezhető olyan ok-okozati viszony, melyben x magyarázza y értékét.

Jelölje $RSS_{\hat{\varepsilon}}$ a korlátozott modell (nullhipotézis alatti), $USS_{\hat{\varepsilon}}$ a korlátozás nélküli (alternatív hipotézisnek megfelelő) modell esetén keletkező reziduális eltérés-négyzetösszeget. Maddala (2004) bebizonyítja, hogy

$$F = \frac{RSS_{\hat{\varepsilon}} - USS_{\hat{\varepsilon}}}{USS_{\hat{\varepsilon}}} \times \frac{T - k - m}{m} \quad (2.15)$$

F-eloszlást követ, $(m, T - k - m)$ szabadságfok-párral, és így alkalmas a restriktió szükségességének tesztelésére. (A dolgozat Bevezető fejezetének példájában a (2.13) modell és a (2.15) képlet alapján végeztem el az okság tesztelését.)

Az elmúlt néhány évtizedben a Granger-okság számos módosítása látott napvilágot. Ezeket részben modellezés-filozófiai megfontolások, részben az eredeti felíráshoz képest korlátozott, vagy kiterjesztett modellek alkalmazása indokolta.²² Most csak a legfontosabb új gondolatokat tekintem át, nem időrendben haladva, hanem a módosítás jelentőségét szem előtt tartva.

Sims (1972) az előbbieken bemutatott próba egy alternatív változatát fejlesztette ki, ennek értelmében x nem *oka* y -nak, ha y_t -nek a késleltetett, egyidejű és jövőre vonatkozó x -ekre vonatkozó regressziójában az utóbbi együtthatói 0-k. Praktikusan, teszteljük az

²² Az okok és eredmények viszonylag részletes leírása olvasható Heckman (2008) tanulmányában.

$$y_t = \alpha_k x_{t+k} + \alpha_{k-1} x_{t+k-1} + \dots + \alpha_1 x_{t+1} + \alpha_0 + \beta_0 x_t + \beta_1 x_{t-1} + \beta_2 x_{t-2} + \dots + \beta_k x_{t-m} + \varepsilon_t \quad (2.16)$$

regresszióban a

$$\begin{aligned} H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_k &= 0 \\ H_1 : \exists j \in 1, 2, \dots, k \rightarrow \alpha_j &\neq 0 \end{aligned} \quad (2.17)$$

hipotézisrendszert. Amennyiben a nullhipotézist elfogadjuk, úgy kijelenthetjük, hogy y -nak az előrejelzése akkor sem javulna, ha x jövőbeli értékeit használnánk a prognózishoz, vagyis x nem oka y -nak. Chamberlain (1982) bebizonyítja, hogy – noha van némi eltérés a két próba között – a Granger, illetve a Sims által proponált esetben lényegében ugyanazt a hipotézist teszteljük.

Hsiao (1979) két, korábban nem tárgyalt problémára is kitér:

- egyrészt áttekinti az okság, a *feedback* (kölsönös okság), illetve a *pillanatnyi okság* fogalmát,
- másrészt kitér arra a kérdésre, hogy miként határozható meg a klasszikus Granger-regresszióban (lásd (2.13) képlet) az egyik, illetve másik változó késleltetésének rendje.

Definíciója szerint x oka y -nak, ha

$$\sigma^2(y_t | \Omega_t) < \sigma^2(y_t | (\Omega_t - X_t)) \quad (2.18)$$

ahol a már megismert jelölések mellett $X_t = \{x_s : s < t\}$.

Viszonylag kézenfekvő módon Hsiao szerint x és y között kölcsönös okság (feedback mechanizmus) érvényesül, ha

$$\sigma^2(y_t | \Omega_t) < \sigma^2(y_t | (\Omega_t - X_t)) \text{ és } \sigma^2(x_t | \Omega_t) < \sigma^2(x_t | (\Omega_t - Y_t)) \quad (2.19)$$

A pillanatnyi okság definíciójához vezessük be a $\tilde{X}_t = \{x_s : s \leq t\}$ jelölést. Ennek segítségével felírható, hogy x pillanatnyi oka y -nak, ha

$$\sigma^2(y_t | \Omega_t, \tilde{X}_t) < \sigma^2(y_t | \Omega_t) \quad (2.20)$$

Mivel a (2.18) definíció következményeként y akkor és csak akkor *pillanatnyi oka* x -nek, ha x *pillanatnyi oka* y -nak, így a pillanatnyi okság szintén egyfajta kölcsönös összefüggést jelez, vagyis nem gyarapítja jelentősen a kauzalitás-fogalmunkat.

A tanulmány sokat idézett eredménye, hogy – éppen az előbb említett feedback mechanizmus általános feltételezése miatt – az okságot egy VAR-modell keretrendszerében vizsgálja. Legyen

$$\begin{aligned} y_t &= \alpha_{10} + \alpha_{11}y_{t-1} + \dots + \alpha_{1k_1}y_{t-k_1} + \beta_{11}x_{t-1} + \beta_{12}x_{t-2} + \dots + \beta_{1m_1}x_{t-m_1} + \varepsilon_{1t} \\ x_t &= \alpha_{20} + \alpha_{21}y_{t-1} + \dots + \alpha_{2k_2}y_{t-k_2} + \beta_{21}x_{t-1} + \beta_{22}x_{t-2} + \dots + \beta_{2m_2}x_{t-m_2} + \varepsilon_{2t} \end{aligned} \quad (2.21)$$

az oksági reláció feltárásához használatos modell, ahol az egyirányú okságokat kézenfekvően a

$$H_0 : \beta_{11} = \beta_{12} = \dots = \beta_{1m_1} = 0$$

illetve

$$H_0 : \alpha_{21} = \alpha_{22} = \dots = \alpha_{2k_2} = 0$$

restrikciók külön-külön tesztelésével, a feedback kapcsolatot a

$$H_0 : \beta_{11} = \beta_{12} = \dots = \beta_{1m_1} = \alpha_{21} = \alpha_{22} = \dots = \alpha_{2k_2} = 0$$

nullhipotézis tesztelésével végezhetjük el.

Arra a kérdésre, hogy miként határozzuk meg az egyes változók késleltetésének rendjét, Hsiao az FPE-mutatót (lásd Akaike, 1969) javasolja, ami – például az y eredményváltozóra vonatkozó egyenlet esetén – az alábbi képlettel írható le:

$$FPE_y(k_1, m_1) = \frac{T + k_1 + m_1 + 1}{T - k_1 - m_1 - 1} \times \frac{\sum_{t=1}^T (y_t - \hat{y}_t)^2}{T} \quad (2.22)$$

Az okság-vizsgálat szempontjából optimális modellben az FPE-mutató értéke minimális.

A klasszikus okság teszt kiterjesztései között végezetül a Mosconi-Seri (2006) tanulmányban tárgyalt újítást tekintem át. A cikkben olyan bináris, kétváltozós Markov-modellt vizsgálnak a szerzők, melyben az okság (pontosabban az okság hiányának) tesztelése logisztikus regressziós modellek paramétereire vonatkozó megszorítások alapján elvégezhető. A gondolat újszerűsége abban áll, hogy amennyiben olyan idősorok között kell ok-okozati kapcsolatot feltárnunk, ahol az ered-

ményváltozó mindössze két állapotban (0,1) lehet, alkalmasan választott átmenet-valószínűségek becslésének segítségével a kauzalitás-teszt elvégezhető.

A tanulmány illusztratív példája is azt sugallja, hogy az ilyen típusú vizsgálatoknak ott van nagy jelentősége, ahol az eredeti idősor ún. *hard clipping* segítségével két-állapotúvá konvertálható (gondoljunk például az olyan tőzsdei árfolyam-idősorokra, melyekben a megfigyeléseink mindössze annyit tartalmaznak, hogy nőtt, vagy nem nőtt az árfolyam!). A cikkben tárgyalt eljárások messze túlmutatnak a disszertációban alkalmazott eszközrendszeren, ugyanakkor a későbbiekben – egy saját modellfejlesztés kapcsán – még visszatérek a kétállapotú (bináris) idősorok között kimutatható ok-okozati viszonyokra.

Az okság modellezése, illetve tesztelése nyilvánvalóan még számos további, itt részleteiben nem tárgyalandó statisztikai módszerfejlesztést indukált, ezek vázlatos felsorolása megtalálható pl. Lin (2008) műhelytanulmányában. A tanulmány azért is érdekes lehet, mert meglehetősen világos útmutatást ad abban a tekintetben, hogy mit szabad tenni, és mit érdemes elkerülni a modellezés során („*Do and Don't Do list*”).

Végképp nem képezi a disszertáció tárgyát, de rengeteg olyan tudományos eredmény született a kérdéskörben, amely még a statisztika- ökonometria tárgykörébe sem illeszkedett. Erre lehet példa Sheffrin-Triest (1995) gráfelméleti megközelítésen nyugvó, a gazdasági növekedés okait vizsgáló dolgozata. Dolgozatom további részében nem lépek ki az előbbiekben részletesebben tárgyalt, VAR-modelleken alapuló okság tesztek világából.

2.3 Trendek

A gazdasági életben a fejlődés, előrehaladás, vagy éppen a visszaesés egyértelmű jele, ha valamely jelenséget leíró idősorban trendet mutatunk ki. Valószínűleg, már amikor először hallottunk trendről azt gondoltuk, hogy az elmúlt több mint 200 évben, amióta 1798-ban Malthus először értekezett a népesség-növekedés exponenciális trendjéről, legalább a trend definíciója egyértelművé vált. Ehhez képest Phillips explicit kijelenti, „*senki sem érti a trendeket, de mindenki látja azokat az adatokban*”.²³ White és Granger néhány évvel ezelőtt a *Journal of Time Series Econometrics* hasábjain közel 40 oldalt szenteltek olyan kérdéseknek, hogy miként lehet a trendet definiálni, mit jelent a növekvő, vagy csökkenő trend (lásd White-Granger, 2011). Mindezzel mindössze azt szerettem volna érzékeltetni, hogy a dolgozatomban feszegetett kérdéseknél még sokkal alapvetőbb viták sincsenek végérvényesen lezárva, ennek ellenére a modellezők többé-kevésbé konszenzussal használják a következőkben bemutatandó trend-fogalmakat és trend-modelleket.

²³ Phillips (2005) 402. old.

2.3.1 *Determinisztikus és sztochasztikus trend*

A gazdasági idősorok modellezése során gyakran találkozhatunk olyan idősorokkal, amelyekről már egy egyszerű ábrázolást követően is kijelenthetjük, hogy adatgeneráló folyamatuk esetében a stacionaritás feltételezése nem helytálló. Ilyen pl. a *trend-stacionárius folyamat*, amely egy lineáris trend körüli véletlen ingadozást mutat:

$$y_t = \alpha + \beta t + \varepsilon_t \quad (2.23)$$

ahol t az idő múlását szemléltető trendváltozó, ε_t a szokásos fehér zaj.

Könnyen belátható, hogy amennyiben $\beta \neq 0$, akkor az előbbi idősor várható értéke időben nem állandó, vagyis a folyamat nem stacioner. Ezt a jelenséget a továbbiakban *determinisztikus nemstacionaritás*nak nevezzük.

A másik kitüntetett figyelmet érdemlő nemstacionárius folyamat, a *véletlen bolyongás*, melynek általános formája²⁴

$$y_t = \mu + y_{t-1} + \varepsilon_t \quad (2.24)$$

vagyis az idősor jelenlegi értéke egy konstanssal, valamint egy sztochasztikus taggal (fehér zajjal) tér el az eggyel korábbi értéktől.

Az (2.24) egyenlet még általánosabban is felírható, az alábbi módon

$$y_t = \mu + \varphi y_{t-1} + \varepsilon_t \quad (2.25)$$

Az (2.25) egyenletben szereplő összefüggés tulajdonképpen egy *elsőrendű autoregresszív modell*. Könnyen beláthatóan, hogy az összefüggés megengedi a stacionárius esetet ($|\varphi| < 1$), az egységgyök meglétét ($\varphi = 1$), ami a tulajdonképpeni véletlen bolyongás, és az *explozív* (szétrobbanó) esetet ($\varphi > 1$).

Fontos észrevétel a (2.24) egyenlettel felírt véletlen bolyongásról, hogy amennyiben elvégezzük a triviális behelyettesítéseket, akkor a *sztochasztikus trend* modelljéhez jutunk

$$y_t = y_0 + \mu t + \sum_{i=1}^t \varepsilon_i \quad (2.26)$$

²⁴ A bemutatott általános eset véletlen bolyongás eltolással (random walk with drift).

ami nagyban emlékeztet a trend-stacionárius folyamatra.

A trend-stacioner folyamat és véletlen bolyongás, illetve a kettő kombinációja, lehetővé teszi a stacionáriusagra (vagy a nemstacionaritást jelző egységgyökre) vonatkozó hipotézisek tesztelését. A disszertációban ezen a helyen csak a legalapvetőbb próbákat vezetem be, az egységgyök, illetve stacionaritás-tesztek alapos leírása pl. Hunyadi (1994) cikkében olvasható.

Az idősorok stacionaritásának vizsgálata során elsőként azt kívánjuk eldönteni, hogy egy adott idősor véletlen bolyongási folyamatból származik-e. Mindez egyszerűnek *látszik*, hiszen mindössze a (2.25) modellre kell paraméterbecslést végeznünk, és a becsült regressziós együttható alapján le kell tesztelnünk a

$$H_0 : \varphi = 1$$

$$H_1 : \varphi \neq 1$$

hipotézisrendszert, és azonnal választ kapunk a kérdésre.²⁵ Ugyanakkor a (2.25) modell direkt becslése számos technikai nehézséget vet fel, melyeket a későbbiekben tárgyaljuk. Mindezek miatt Dickey-Fuller (1979) megoldásként azt javasolta, hogy ne a (2.25) egyenletet, hanem annak módosított változatát az ún. *Dickey-Fuller regresszió*

$$\Delta y_t = \mu + (\varphi - 1)y_{t-1} + \varepsilon_t \quad (2.27)$$

paramétereit becsüljük meg, mellyel számos problémától kíméljük meg magunkat. Az említett tanulmány három különböző hipotézisrendszer mellett vizsgálja az egységgyök meglétét (a véletlen bolyongás, tehát a nemstacionaritás fennállását). A nullhipotézis mindhárom esetben eltolás nélküli véletlen bolyongást tételez fel

$$H_0 : y_t = y_{t-1} + \varepsilon_t$$

Az ezzel szemben álló alternatívák

- a) stacionárius, elsőrendű autoregresszív, tehát AR(1) folyamat

$$H_1 : y_t = \varphi y_{t-1} + \varepsilon_t \quad \varphi < 1$$

- b) eltolást (konstans tagot) is tartalmazó AR(1) folyamat

$$H_1 : y_t = \mu + \varphi y_{t-1} + \varepsilon_t \quad \varphi < 1$$

- c) eltolást és determinisztikus trendet is tartalmazó AR(1) folyamat

$$H_1 : y_t = \mu + \varphi y_{t-1} + \beta t + \varepsilon_t \quad \varphi < 1$$

²⁵ Érdemes megjegyezni, hogy a próbát egyoldálú ($\varphi < 1$) alternatív hipotézis mellett szoktuk elvégezni, hiszen a másik irány (*explozív alternatíva*) általában érdektelen, legalábbis a stacionaritás vizsgálata szempontjából.

Vegyük észre, hogy a próba egyszerűbben is felírható, ha a tesztelendő modellt átírjuk a csak sztochasztikus trendet tartalmazó (azaz determinisztikus trendet nem tartalmazó) alakra. Ekkor az $y_t = \mu + \varphi y_{t-1} + \beta t + \varepsilon_t$ egyenletben $\varphi = 1$ és $\beta = \mu = 0$, vagyis a hipotézisrendszer a

$$H_0 : \Delta y_t = \varepsilon_t$$

$$H_1 : \Delta y_t = \mu + (\varphi - 1) y_{t-1} + \beta t + \varepsilon_t$$

formájú. Ilyenkor csak a következő paraméter-megszorításokat kell tesztelnünk:

- a) stacionárius, elsőrendű autoregresszív (AR(1)) folyamat
 $\mu = \beta = 0, \varphi - 1 < 0$
- b) eltolást (konstans tagok) is tartalmazó AR(1) folyamat
 $\beta = 0, \varphi - 1 < 0$
- c) eltolást és determinisztikus trendet is tartalmazó AR(1) folyamat
 $\varphi - 1 < 0$

A próbák végrehajtása során nehézséget okoz, hogy – a modell előfeltevései következtében – a regressziós paraméterek 0-val való egyezőségének tesztelésére vonatkozó, általánosan használt t-próbák itt nem alkalmazhatók, a szerzőpáros cikkében a próbafüggvényeknek saját kritikus érték táblát közölt. További problémát jelenthet, hogy a Dickey-Fuller próba első változata csak akkor érvényes, ha a modellben szereplő ε_t véletlen változó valóban fehér zaj. A gyakorlatban azonban gyakran megfigyelhető, hogy a modellek véletlen tagjai autokorreláltak, ami a próbát gyengíti, sokkal többször fordul elő, hogy korrekt nullhipotézist hibás döntéssel elvetünk (vagyis első fajú hibát követünk el).

A megoldásra a DF-próba minimális módosítása kínálkozik. Az ún. *kiterjesztett Dickey-Fuller próbában* (*ADF-próba, augmented DF-test*) a becslendő regresszió

$$\Delta y_t = \mu + (\varphi - 1) y_{t-1} + \sum_{i=1}^p \gamma_i \Delta y_{t-i} + \beta t + \varepsilon_t \quad (2.28)$$

formájú, és ebben vizsgáljuk a korábban felírt paraméter-restrikciókat.²⁶

A Dickey-Fuller próbával (és az ezen alapuló többi egységgyök-próbával) szemben a legfontosabb kritikai észrevétel, hogy nem elég erősek a stacionárius folyamatokkal szemben, különösképpen akkor, amikor az autoregresszív taghoz tartozó regressziós paraméter közel van 1-hez. Éppen ezért érdemes fordított szemléletű eljárást is alkalmazni, vagyis olyan próbát, melyben a

²⁶ A kiterjesztett differenciák szerepeltetésének logikájáról lásd pl. Said – Dickey (1984).

stacionaritás a nullhipotézis, és az egységgyök meglelte az alternatív feltételezés. A leggyakrabban alkalmazott stacionaritás-teszt a *Kwiatkowski-Phillips-Schmidt-Shin (KPSS)-próba*.²⁷

A KPSS teszt lényege: legyen y_t az empirikus idősor, amelyről el kívánjuk dönteni, hogy stacionárius-e. Feltételezzük, hogy az idősor komponensekre bontható az alábbi módon:

$$y_t = x_t + \beta t + \varepsilon_t \quad (2.29)$$

ahol x_t véletlen bolyongást követ (vagyis $x_t = x_{t-1} + u_t$, ahol $u_t \sim Nid(0; \sigma_u^2)$, azaz *fehér zaj*). A (2.29) modell alapján tesztelhető a stacionaritás, ugyanis y_t várható értékét tekintve stacioner, ha $\sigma_u^2 = 0$ és $\beta = 0$, hiszen ekkor az első tag konstans tengelymetszetté válik ($x_t = x_0 = \mu$), és determinisztikus trend sem játszik szerepet értékének meghatározásában.

Ilyenkor tehát a

$$H_0 : \sigma_u^2 = \beta = 0$$

kiinduló hipotézis tesztelésére felírható a KPSS-próbafüggvény

$$KPSS = \frac{\sum_{t=1}^T S_t^2}{T^2 \hat{\sigma}_y^2} \quad (2.30)$$

ahol $S_t = \sum_{i=1}^t (y_i - \bar{y})$ és $\hat{\sigma}_y^2$ pedig y_t alkalmasan becsült varianciája. A próba nullhipotézisének elfogadása azt jelenti, hogy a vizsgált folyamat stacionárius. Érdeemes megemlíteni, hogy a KPSS-próbának is van determinisztikus trendet is feltételező változata (vagyis amikor $\beta \neq 0$), ilyenkor értelemszerűen más kritikus értékek mellett zajlik a döntés. (A próba alkalmazását ellenzők gyakran felvetik, hogy a szerzők által megadott aszimptotikus kritikus értékek csak nagy minta esetén érvényesek, mindez igaz, ám a gyakorlati felhasználók az intést gyakran figyelmen kívül hagyják...)

Láthatjuk, hogy az ADF és a KPSS próbák hipotézisrendszere ellentétes szemléletű, mindez azt jelenti, hogy lehetőségünk nyílik a hipotézisellenőrzés döntési szituációit tartalmazó szokásos táblázat valamennyi celláját kitölteni:

²⁷ Lásd Kwiatkowski et al. (1992).

2-1. táblázat: A stacionaritás-vizsgálat döntési szituációi

<i>Az ismeretlen valóságban az idősor</i>	<i>Döntésünk során az idősort</i>	
	<i>stacionernek tekintjük</i>	<i>nemstacionernek tekintjük</i>
<i>Stacionárius folyamatból származik</i>	helyes döntés ($1 - p_{KPSS}$)	hibás döntés (p_{KPSS})
<i>Egységgyökeket tartalmaz</i>	hibás döntés (p_{ADF})	helyes döntés ($1 - p_{ADF}$)

A következtetési statisztika alap gondolata, hogy a – mégoly korrektül végzett – hipotézisellenőrzés során is elkövethetünk a mintavételből eredő hibát. Így az egységgyök/stacionaritás-teszteknel is előfordulhat, hogy a két próba ellentétes következtetésre jut. Ilyen eset, ha mindkét próbánál azonosan ítélik meg (elfogadjuk, vagy elvetjük) a nullhipotézist. Carrion-i-Silvestre et al. (2001), illetve Kęblowski–Welfe (2004) kísérletet tettek az ADF és a KPSS próba „házasítására”, vagyis együttes alkalmazásukra. A tanulmányokban közölt kritikus érték táblázatok nem terjedtek el igazán az ökonometria modellezők körében, valószínűleg azért, mert a Monte Carlo eljárásokkal meghatározott értékek nem épültek be a szokásos programcsomagokba, ráadásul a próbákhoz tartozó szignifikancia-értékek kiszámítása nem megoldott.

A dolgozat további részében azokban az esetekben, ha a két próba egymásnak ellentmondó következtetést eredményezne, az óvatosság elve alapján a stacionaritás ellen foglalkozunk állást, vagyis feltételezem, hogy az idősor nemstacionárius.²⁸

2.3.2 Kointegráció és hibakorrekción

Ismeretes, hogy amennyiben egy idősor nem stacionárius, akkor differencia-képzéssel (nem feltétlenül elsőrendű differencia-képzéssel!) általában stacionáriussá tehető.²⁹ Egy sztochasztikus folyamat abban a rendben integrált, ahányadik rendű differenciája már stacionárius (amikor először stacionerré válik). Az *integráltsági rendet* szokásosan az $I(d)$ szimbólummal jelöljük. A definícióból következően a stacionárius folyamat nulla-rendben integrált, azaz ha y_t stacioner, akkor ez felírható $y_t \sim I(0)$ formában is. Ha egy idősor önmagában nem stacionárius, de az első rendű differenciája már az, akkor az idősor elsőrendű integrált, azaz $I(1)$ folyamatból származik.

Granger-Newbold (1974) megmutatta, hogy nemstacionárius idősorokból épített ökonometriai modellekben gyakori az ún. *hamis regresszió* (*spurious regression*), vagyis ilyen esetekben sokszor for-

²⁸ Annyiban tekinthető ez óvatos megközelítésnek, hiszen az esetlegesen stacionárius idősor első differenciája is $I(0)$, míg ha a nemstacionárius idősorból nem képezünk differenciát, akkor sérülnek a modellfeltevéseink.

²⁹ Nyilvánvaló, és gyakran használatos transzformáció a logaritmált változó differenciázása (ami még könnyen értelmezhető is, hiszen nem más, mint az eredeti változó relatív változásának idősor), így a lineáristól eltérő típusú (pl. exponenciális) trend is kiszűrhető.

dul elő, hogy nincs tényleges összefüggés a változók között, ám a modell magyarázó ereje mégis nagy (szignifikáns), a globális F-próba téves eredményre vezet. A nemstacionárius változók között felírt modellek sok kényelmetlen tulajdonsággal bírnak: a szokásos próbák gyengék, nem használhatók. (Ennek oka, hogy a standard hiba a szokásos normálással 0-hoz tart, ugyanis a konvergencia sebessége lényegesen magasabb, pl. elsőrendű integrált folyamat esetén T , a szokásos $\sqrt{T}$ helyett.)

Általában elmondható, ha két változó elsőrendű integrált, akkor a lineáris kombinációjuk is $I(1)$ folyamat lesz. Ennél is általánosabban fogalmazva, ha idősorok egy halmaza (a továbbiakban k idősor) esetén bevezetjük a következő jelölést $x_i \sim I(d_i)$, vagyis az i -edik idősor d_i rendű integrált, és képezzük ezen változók lineáris kombinációját, azaz a

$$z_t = \sum_{i=1}^k \alpha_i x_{it}$$

összeget, akkor általában igaz, hogy $z_t \sim I(\max d_i)$.

Engle-Granger (1987) megmutatja, hogy vannak olyan esetek, amikor a fenti, általánosságban megfogalmazott megállapítás nem igaz. Definíciójuk szerint, amennyiben mind x_t , mind y_t d -ed rendű integrált, és létezik olyan lineáris kombinációjuk, amely $(d-b)$ -ed rendű integrált, ahol $b > 0$, akkor x_t és y_t (d,b) rendű kointegráltak. Vagyis mindezt formalizálva: ha $x_t, y_t \sim I(d)$ és létezik olyan β_0, β_1 paramétervektor, melyre $y_t - \beta_0 - \beta_1 x_t \sim I(d-b)$, ahol $b > 0$, akkor x_t és y_t (d,b) rendű kointegráltak.

Természetesen ez a tulajdonság kettőnél több idősor lineáris kombinációja esetén is hasonlóan értelmezhető. Többváltozós esetben egy m darab $I(d)$ folyamatból álló $\mathbf{x}_t$ vektor komponenseire akkor mondjuk, hogy (d,b) rendű kointegráltak, azaz $\mathbf{x}_t \sim CI(d,b)$, ha létezik olyan $\beta \neq 0$ vektor, hogy $z_t = \beta' \mathbf{x}_t \sim I(d-b)$, ahol $b > 0$. Ebben az esetben a β vektort kointegráló vektornak nevezzük. Abban az esetben, ha $d = b$, a rendszert *tökéletesen kointegráltak* nevezzük. Könnyen belátható, hogy két változó esetén maximum egy független kointegráló vektor létezhet, egy k változós rendszerben pedig maximum (de nem feltétlenül!) $k-1$ kointegráló vektor képezhető.

Dolgozatomban a $d = b = 1$ esetnek lesz kiemelkedő jelentősége, vagyis amikor az $\mathbf{x}_t$ vektor komponensei $I(1)$ folyamatot követnek, viszont létezik olyan lineáris kombinációjuk, amely stacionárius, azaz $I(0)$ folyamat. Az ilyen *tökéletes kointegráltságot* nevezzük a *gazdaság dinamikus egyen-*

súlyi állapotának, vagy a kointegrált rendszert alkotó idősorok hosszú távú együttmozgásának, közös trendjének.

A nemstacionárius, első rendben integrált idősorok modellezésére korábban kínált megoldás az volt, hogy az eredeti idősorok helyett használjuk azok differencia-sorait. Ez egyváltozós idősori modellekben korrekt megoldás, ám ha a változók közötti kapcsolatot tartjuk fontosnak (elemzendőnek), akkor az eljárás nem megfelelő, hiszen amennyiben az eredeti idősorok helyett az első differenciáikat használjuk, akkor a hosszú távú együttmozgás helyett a rövid távú kapcsolatot vizsgáljuk.

Legyen x_t és y_t első rendben integrált! Ekkor a közöttük felírt, a rövid távú összefüggéseket vizsgáló regresszió

$$\Delta y_t = \beta_0 + \beta_1 \Delta x_t + \varepsilon_t$$

formájú. A modell paraméterei ugyan OLS-sel torzítatlanul becsülhetők, ám a hosszú távú együttmozgásról nem mondanak semmit, hiszen például ha egyik változóban sincs hosszabb ideig változás, vagyis amikor $y_t \approx y_{t-1} \rightarrow \Delta y_t = 0$ és $x_t \approx x_{t-1} \rightarrow \Delta x_t = 0$ a modell semmitmondó.

A problémára a megoldást – az Engle-Granger reprezentációs tétel alapján – az első differenciák és a kointegrált változók késleltetett értékeinek együttes használata adja. Tekintsük az alábbi, ún. *hibakorrekciós modellt* (*error correction model, ECM*):

$$\Delta y_t = \beta_0 + \beta_1 \Delta x_t + \beta_2 (y_{t-1} - \alpha_0 - \alpha_1 x_{t-1}) + \varepsilon_t \quad (2.31)$$

ahol $y_{t-1} - \alpha_0 - \alpha_1 x_{t-1}$ az ún. hibakorrekciós tag. Ha x_t és y_t kointegráltak $[1, -\alpha_0, -\alpha_1]$ kointegráló vektorral (a vektor most egyértelmű), akkor a hibakorrekciós tag is stacionárius, vagyis a (2.31) modell β paraméterei OLS-sel szintén torzítatlanul becsülhetők és a szokásos próbák használhatók.

A változók közötti kointegráció tesztelésére tehát kézenfekvő megoldás a kointegráló regresszió hibakorrekciós tagjának (véletlen változójának) vizsgálata. Amennyiben a

$$y_t = \beta_0 + \beta_1 x_t + \varepsilon_t \quad (2.32)$$

kointegráló regresszióban $y_t, x_t \sim I(1)$, de ε_t stacionárius, akkor a modellben szereplő változók (az előbbieken elmondottak alapján) kointegráltak, sőt tökéletesen kointegráltak.

A kointegráltság feltárása érdekében ε_t stacionaritását valamely ismert próbával tesztelhetjük. Leggyakrabban a kiterjesztett Dickey-Fuller próbát alkalmazzák, de meg kell jegyeznünk, hogy ebben az esetben a teszt kritikus értékei eltérnek az „egyszerű” egységgyök-tesztekben alkalmazottaktól. Az új kritikus értékeket Engle-Granger (1987) határozta meg, ezért a próbát *EG-teszt*nek szokás nevezni.

Könnyen átlátható, hogy amennyiben $k+1$ változó kointegráltságát vizsgáljuk, nem csak egy, hanem akár k különböző kointegráló regressziót is becsülhetünk (ennyi különböző kointegráló vektor létezhet). Mindez azt jelenti, hogy annak függvényében, hogy melyik változót állítjuk a kointegráló regresszió eredményváltozójának helyébe más és más paraméterbecslési eredményekre jutunk. Ezekben az esetekben választ kell kapnunk arra a kérdésre is, hogy melyik ezek közül a legjobb kointegrációs összefüggés, illetve pontosan hány független lineáris kombináció írható fel?³⁰ A megoldást az ún. *Johansen-próba* szolgáltatja, amely valamennyi ($r \leq k$) kointegráló regresszió becslésén alapul. A teszt első leírása megtalálható Johansen (1991) cikkében.

Tegyük fel, hogy $k+1$ változó mindegyike³¹ elsőrendű integrált, vagyis $I(1)$ folyamatot követ. Ezen változóhalmazra felírható egy vektor autoregresszív (VAR) modell:

$$\mathbf{x}_t = \mathbf{B}_1 \mathbf{x}_{t-1} + \mathbf{B}_2 \mathbf{x}_{t-2} + \dots + \mathbf{B}_r \mathbf{x}_{t-r} + \varepsilon_t \quad (2.33)$$

melyben $\mathbf{x}_t$ (és ebből adódóan, értelemszerűen valamennyi $\mathbf{x}_{t-j}$) $(k+1) \times 1$ rendű vektor, valamint valamennyi $\mathbf{B}_j$ paramétermátrix $(k+1) \times (k+1)$ -ed rendű. Az ún. Johansen-approximáció értelmében a fenti VAR-modell átírható egy *vektor-hibakorrekciós modellé* (VECM), vagyis

$$\Delta \mathbf{x}_t = \Pi \mathbf{x}_{t-r} + \Gamma_1 \Delta \mathbf{x}_{t-1} + \Gamma_2 \Delta \mathbf{x}_{t-2} + \dots + \Gamma_{r-1} \Delta \mathbf{x}_{t-(r-1)} + \varepsilon_t \quad (2.34)$$

ahol $\varepsilon_t \sim IN(\mathbf{0}, \Sigma)$

$$\Pi = \left(\sum_{i=1}^r \mathbf{B}_i \right) - \mathbf{I} ,$$

$$\Gamma_i = \left(\sum_{j=1}^i \mathbf{B}_j \right) - \mathbf{I} \text{ és}$$

$\mathbf{I}$ a $(k+1) \times (k+1)$ -ed rendű egységmátrix.³²

³⁰ Engle-Granger javaslata szerint az „exogénebb” változót kell magyarázó változóként használni, ennek eldöntése már visszavezet az okság vizsgálatához.

³¹ Az egyszerűbb jelölés kedvéért a $k+1$ változó alkotta vektort $\mathbf{x}_t$ -vel jelöljük, vagyis nem különböztetünk meg egy y_t -t, vagyis a módszer szempontjából most már egyáltalán nincs kitüntetett változó.

A kointegráció teszteléséhez, illetve annak megállapításához, hogy hány független kointegráló vektor írható fel a rendszerben, az előbbiekben definiált Π mátrix sajátértékeire van szükség. A mátrix rangja (vagyis a képezhető független kointegráló vektorok száma) megegyezik a zérótól különböző sajátértékek számával. Rendezzük a sajátértékeket (λ_i) csökkenő sorrendbe, vagyis legyen $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{k+1}$!³³ Ha a változók nem kointegráltak, akkor a Π mátrix rangja nem különbözik szignifikánsan nullától, vagyis $\lambda_i \approx 0 \quad \forall i \in 1, 2, \dots, k+1$. A Johansen által javasolt tesztstatisztika $\ln(1-\lambda_i)$ -re épít, kihasználva, hogy amennyiben $\lambda_i = 0$, akkor $\ln(1-\lambda_i) = 0$. Ha a mátrix rangja 1, akkor $\ln(1-\lambda_1)$ nem 0, és $\ln(1-\lambda_i) \approx 0 \quad \forall i > 1$.

Ezen az elven végighaladva Johansen két próbafüggvényt javasolt. Amennyiben a kointegráló vektorok feltételezett számát r -rel jelöljük, úgy a

$$\lambda_{trace}(g) = -T \sum_{i=g+1}^k \ln(1-\lambda_i) \quad (2.35)$$

próbafüggvénnyel tesztelhetjük a $H_0: r \leq g$ nullhipotézist, a $H_1: r > g$ alternatívával szemben. Amennyiben g értékét „végigfuttatjuk” a $0, 1, 2, \dots, k$ számsoron, dönteni tudunk arról a hipotézispárról, amelyben a nullhipotézis szerint legfeljebb g kointegráló vektort találunk a rendszerben, az alternatív szerint ennél több kointegráló vektor becsülhető. Láthatjuk, hogy az „alapkérdés” (kointegrált-e a rendszer, azaz van-e legalább egy kointegráló vektor) a $g=0$ helyettesítéssel válaszolhatjuk meg, amennyiben itt (is) elfogadjuk a nullhipotézist, akkor a rendszer nem kointegrált.

Emellett használatos a

$$\lambda_{max}(g, g+1) = \max\{-T \ln(1-\lambda_g); -T \ln(1-\lambda_{g+1})\} = -T \ln(1-\lambda_g) \quad (2.36)$$

statisztika, amivel tesztelhető a kointegráló vektorok száma egyenlő g -vel nullhipotézis ($H_0: r = g$), a vektorok száma egyenlő $g+1$ alternatívával ($H_1: r = g+1$) szemben. Johansen mindkét próba vonatkozásában közölte az aszimptotikus kritikus értékek táblázatát, melyek a standard programcsomagokban beépített beállításként szerepelnek.

³² A paraméterbecslés ML-elven történik (ezért is fontos a véletlen változók vektorára tett eloszlásbeli megkötés!).

³³ Komplex saját-értékek esetén azok abszolút értékét használjuk.

Abban az esetben, ha idősoraink (illetve az őket generáló folyamatok) nemstacionerek, ugyanakkor feltételezhetjük, hogy kointegráltak, (legalább) három paraméterbecslési eljárás közül választhatunk, ezek

- az Engle és Granger által javasolt kétlépcsős eljárás,
- az ún. Engle-Yoo módszer, illetve
- a Johansen-módszer.

Mindhárom módszer leírása megtalálható Hamilton (1994) könyvében.

Érdemes megjegyezni, hogy a vektor-autoregresszió és a kointegráció között egyértelmű kapcsolat áll fenn: ha egy kétváltozós VAR(1)-modell együtthatómátrixának³⁴ mindkét sajátértéke 1, akkor mindkét idősor I(1) folyamatot követő, de nem kointegrált, ha sajátértékek közül pontosan egy nem egyenlő 0-val, akkor az idősorok kointegráltak. Mindezt a későbbiekben többször kihasználom.

Granger több alkalommal is visszatért munkáiban saját okság-fogalmára, többször pontosította, illetve javította azt. 1988-as cikkében – többek között – azt vizsgálta, hogy milyen kapcsolat van kointegráltság és okság között.³⁵ Azt írja, hogy az többé-kevésbé evidencia, hogy a makrogazdasági idősorok nagy része I(1) folyamatot követ. Az okság általa megfogalmazott definíciója nem tesz különbséget I(0) vagy I(1) idősorok közötti kauzalitásban, bár az utóbbi esetben nem árt óvatosságnak lenni.

Korábban már láthattuk, hogy amennyiben x_t és y_t egyaránt első rendben integráltak, de létezik olyan lineáris kombinációjuk, amelyik stacionárius (vagyis kointegráltak), akkor közöttük hibakorrekciós összefüggés írható fel. Legyen az illusztratív modell

$$\begin{aligned}x_t &= \alpha q_t + \varepsilon_{1t} \\ y_t &= q_t + \varepsilon_{2t}\end{aligned}$$

formájú, ahol q_t , és ebből következően x_t, y_t első rendben integrált, $\varepsilon_{1t}, \varepsilon_{2t}$ stacionárius. Könnyen belátható, hogy x_t, y_t kointegrált rendszert alkotnak, hiszen lineáris kombinációjuk

$$z_t = x_t - \alpha y_t = \varepsilon_{1t} - \alpha \varepsilon_{2t}$$

³⁴ A (2.8) képlet $\mathbf{B}_1$ mátrixa.

³⁵ Lásd Granger (1988). A tanulmány emellett részletesen foglalkozik azzal a kérdéssel, hogy szükséges-e a pillanatnyi okság vizsgálata a közgazdaságtudományban, illetve hogy mennyiben hasznosítható az okság tesztje a gazdaságpolitikai döntéshozatalban.

stacionárius. Ebben az esetben a hibakorrekciós modell

$$\begin{aligned}\Delta x_t &= \beta_{10} + \beta_{11} \Delta y_t + \gamma_1 z_{t-1} + \varepsilon_{1t} \\ \Delta y_t &= \beta_{20} + \beta_{21} \Delta x_t + \gamma_2 z_{t-1} + \varepsilon_{2t}\end{aligned}$$

alakú, ahol $\gamma_1, \gamma_2 \neq 0$. Más szavakkal z_{t-1} oka vagy Δx_t -nek, vagy Δy_t -nek (vagy mindkettőnek), ugyanakkor z_{t-1} szükségszerűen x_{t-1} és y_{t-1} függvénye. Ebből következően vagy y_{t+1} okozata x_t -nek, vagy fordítva. Vagyis, ha két változó kointegrált, akkor szükségszerűen oksági kapcsolat is van köztük.

Mindemellett arra a némileg meglepő eredményre jutottunk, hogy a kointegráció a hosszú távú egyensúllyal, ugyanakkor az okság a rövid távú előrejelezhetőséggel foglalkozik. Ugyanezt a kérdést több cikkben is tárgyalja Dufour különböző szerzőtársaival.³⁶ A szerzők ugyan nem foglalkoznak a jelenség okaival, de kitűnő példákat mutatnak be arra, hogy releváns makrogazdasági változók (GDP, megtakarítás, hozamok, vagy infláció) között elvégzett okság-vizsgálatok eltérő következtetésre vezethetnek annak függvényében, hogy milyen időhorizontra vonatkozik az elemzés. (Mindez összecseng a bevezető fejezet illusztratív példájában tapasztaltakkal!)

2.4 Néhány, a későbbi szimulációk során használt idősor

A disszertáció szerkezetét taglaló első részben már kitértem arra, hogy a későbbiekben a különböző anomáliák hatásának szemléltetése érdekében, több helyen szimulációt alkalmazok. Annak érdekében, hogy a káros jelenségek hatását jól érzékelhessük, viszonylag egyszerű, ám az előbb bemutatott alapeseteknek megfelelő folyamatok generálására lesz szükség. Törekedtem arra, hogy az alkalmazott modellek (adatgeneráló folyamatok) összehasonlíthatóak legyenek, ennek érdekében a szimuláció során felhasznált konstansok (paraméterek) a különböző jellegű folyamatoknál azonosak, ahol ez nem lehetséges, hasonlóak legyenek. A dolgozatban leggyakrabban négy adatgeneráló folyamatot, illetve az ezek alapján szimulált idősorokat fogom elemezni. A vizsgált DGP-k

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

idősorokat eredményeznek. Az idősorok szimulációját az EViews 8.0 programcsomaggal végeztem. Valamennyi szimulált idősorra érvényes, hogy

³⁶ Lásd Dufour-Renault (1998), Dufour et al. (2006) és Dufour-Taamouti (2010).

- 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 (ennek jelölése egy folyamat esetén ε_t , két folyamat esetén $\varepsilon_{1t}, \varepsilon_{2t}$).

2.4.1 AR1 folyamatok szimulálása

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

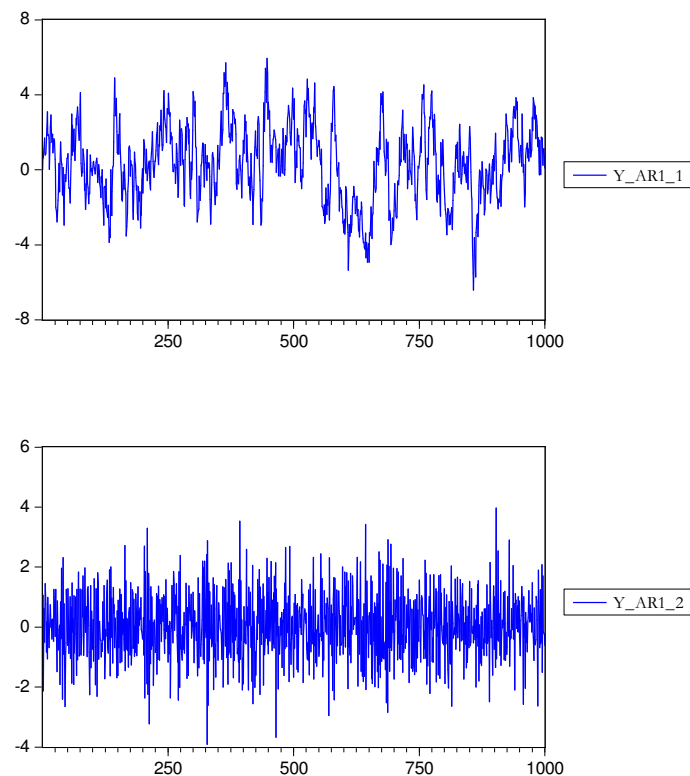
- pozitív, 1-hez közeli autokorrelációs együtthatóval

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

- negatív autokorrelációs együtthatóval

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

Az így keletkező két idősor jellemző képét a 2-1. ábra mutatja:³⁷

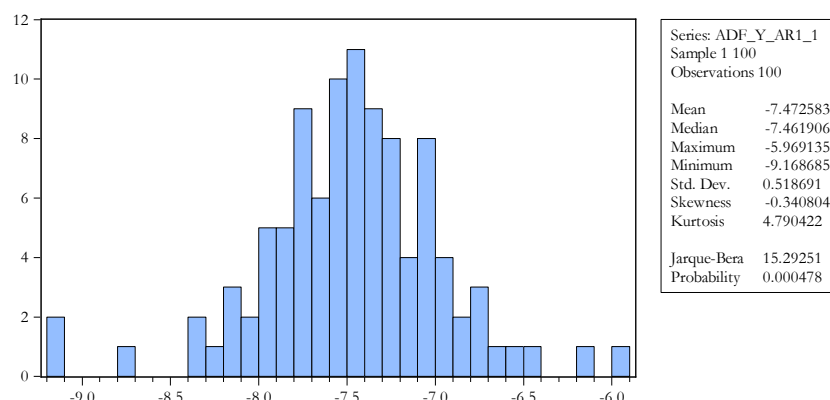


2-1. ábra: Példa elsőrendű autoregresszív folyamatokra

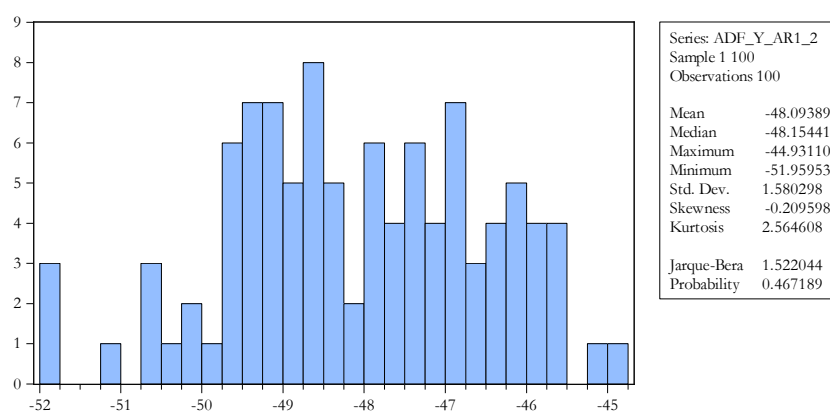
³⁷ A fejezet következő ábráin az idősorok kódja utal a generálás módjára, illetve – amennyiben teszt-eredményről van szó – a próba elnevezésére.

Az ábrákon látszik, hogy a várható érték és a szórás állandó, korábbi fejtegetéseink alapján pedig egyértelmű, hogy mindkét folyamat stacionárius. (A karakterisztikus egyenlet gyöke mindkét esetben az egységkörön kívül van.)

A későbbi szimulációk működését illusztrálandó 100-szor azonos feltételekkel elvégeztem az adatsorok generálását, és valamennyi esetben teszteltem az egységgyök létére vonatkozó hipotézist. A konstans tagot és determinisztikus trendet is feltételező Dickey-Fuller próbák teststatistikáiból képzett hisztogramokat mutatja a 2-2a és 2-2b ábra:



2-2a ábra: A 0,9 paraméterű AR1 folyamatra vonatkozó ADF-teszt próbafüggvény-értékeinek alapstatisztikái



2-2b ábra: A -0,4 paraméterű AR1 folyamatra vonatkozó ADF-teszt próbafüggvény-értékeinek alapstatisztikái

A próba kritikus értéke (500 elemet meghaladó hosszúságú idősorok esetén) 1%-os szignifikancia-szinten -3,96, vagyis látható, hogy mindkét modell esetén valamennyi esetben el kell vetni a nullhipotézist, azaz döntésünk értelmében a folyamatok nem tartalmaznak egységgyököt. (Az alacsonyabb abszolút értékű autokorreláció feltételezésével szimulált második esetben a próbafüggvény empirikus értékei távolabb esnek nullától, vagyis a szignifikancia-értékek kisebbek.)

2.4.2 VAR-modellel szimulált idősorok

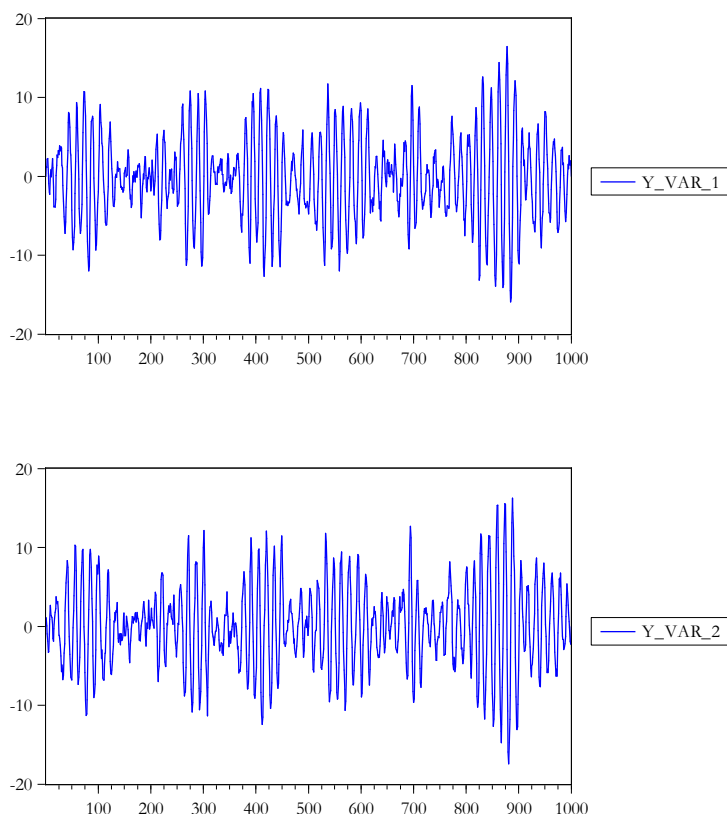
Az idősorok közötti ok-okozati összefüggések kimutatására általánosan használt Granger-teszt bemutatása során láthattuk, hogy amennyiben a két empirikus idősor vektor-autoregresszív modellel leírható, és a VAR- modell paraméterei szignifikánsan különböznek 0-tól, úgy az idősorok között kauzalitás mutatható ki. Illusztrációképpen szimuláltam az alábbi VAR(1) modellt

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

Mindez mátrix alakban is felírható:

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

Belátható, hogy a VAR-modellel felírható folyamatok akkor stacionáriusok, ha az együtttható-mátrixának valamennyi sajátértéke az egységkörön belül van. Mivel ez esetünkben teljesül az előbbi modell alkalmas lesz a Granger-okság meglétének behatóbb vizsgálatára. Példa a fenti VAR(1) modellel szimulált idősorokra:



2-3. ábra: VAR(1) modellel generált példa-idősorok

Ebben az esetben is megismétltem 100-szor az adatok generálását, annak érdekében, hogy megvizsgálhassuk, milyen eredményekkel jár a Granger-okság tesztje. A Wald-típusú F-próba empirikus értékei (akár azt tételeztem fel, hogy y_{1t} nem oka y_{2t} -nek, akár fordítva) minden esetben többes nagyságrendben találhatók, így a nullhipotézis valamennyi esetben elvethető, a szimulált idősorok kölcsönösen okai egymásnak.

2.4.3 Egységgyököt tartalmazó idősorok szimulációja

Az idősoros alapvetésben kiemelt figyelmet fordítottunk a véletlen bolyongás folyamatra, melynek két – a továbbiak szempontjából 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 szimuláltunk:

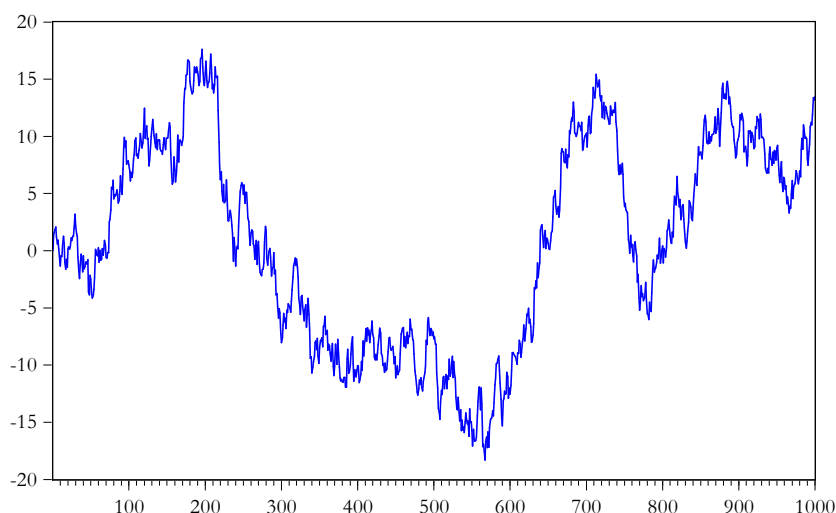
- véletlen bolyongás eltolás nélkül

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

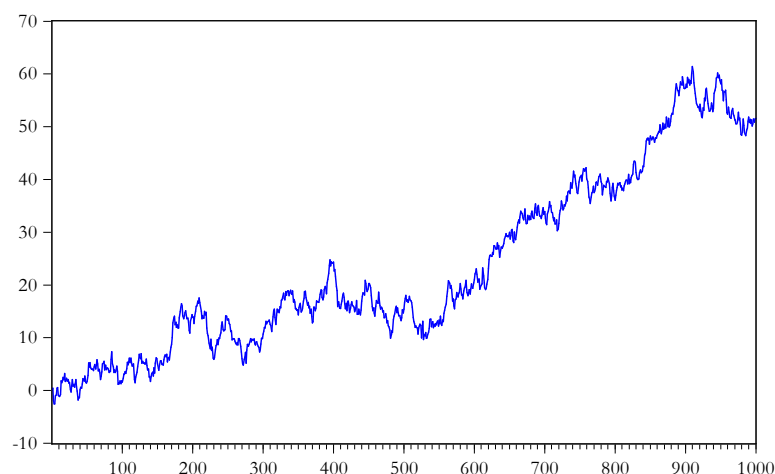
- véletlen bolyongás eltolással

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

A 2-4a, valamint 2-4.b ábra az előbbi idősorokra mutat egy-egy illusztratív példát:

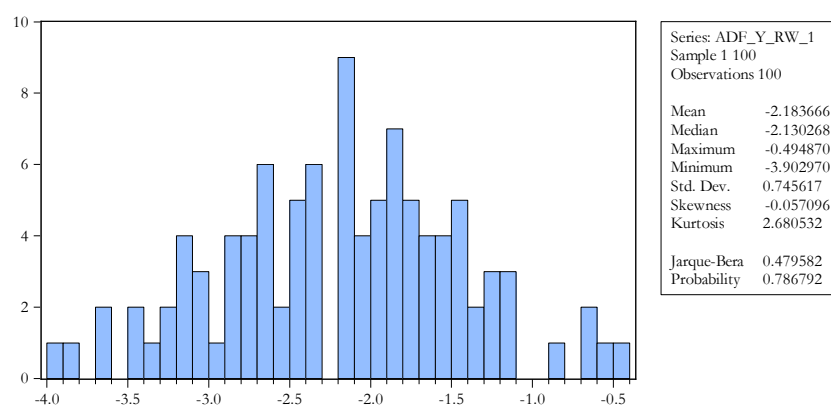


2-4a ábra: Véletlen bolyongás eltolás nélkül



2-4b ábra: Véletlen bolyongás eltolással

Az utóbbi ábrán egyértelműen felismerhető egy pozitív meredekségű trend (tendencia) az idősorban, mindez a 2-4a ábrán nem rajzolódik ki. Vizsgáljuk meg a véletlen bolyongásból származó szimulált idősorok ADF-tesztstatisztikáit (helytakarékoság okán csak az eltolás nélküli esetre):



2-5. ábra: Az eltolás nélküli véletlen bolyongás folyamatra vonatkozó ADF-teszt próbafüggvény-értékeinek alapstatisztikái

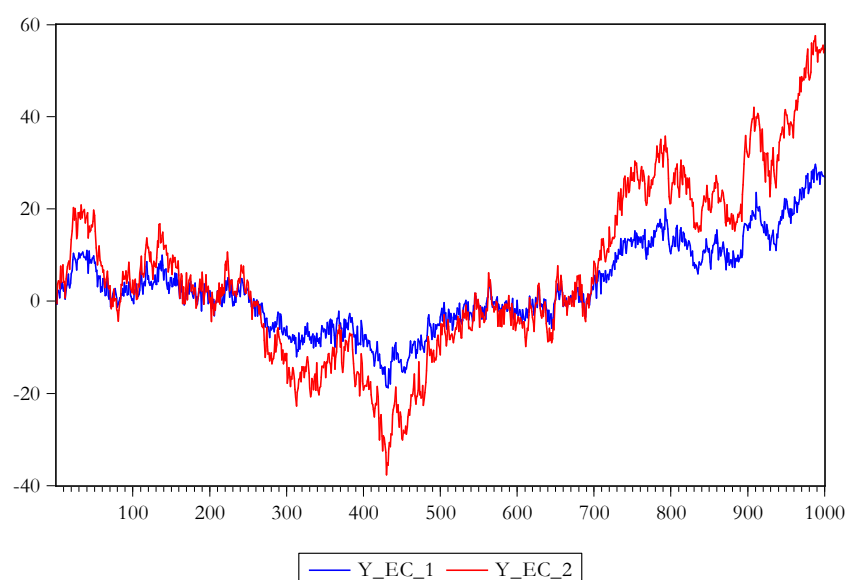
Láthatjuk, hogy a kiterjesztett Dickey-Fuller próba empirikus értékei túlnyomó részt (esetünkben 96-szor a 100 szimulációból) az elfogadási tartományban vannak, vagyis nem vehetjük el az egységgyök léteire vonatkozó nullhipotézist. Ebből is következtethetünk az ismert eredményre, miszerint a véletlen bolyongás nemstacionárius folyamat.

2.4.4 Kointegrált idősorok generálása

Granger ismertetett példájából kiindulva, közös trendet tartalmazó folyamatokat is generáltam, az alábbi paraméter-értékekkel:

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

A szimuláció egy kimenetele látható a 2-6. ábrán:



2-5. ábra: Példa kointegrált (közös trendet tartalmazó) idősorok

A kointegráltságra vonatkozó Engle-Granger egyenletes modellen alapuló teszt empirikus értékei valamennyi szimulált esetben kisebbek, mint -20, vagyis minden ésszerű szignifikanciaszinten el kell fogadnunk a közös trend meglétét.

A fenti szimulált idősorok a sztochasztikus folyamatok alapeseteinek számítanak, ugyanakkor a legfontosabb gazdaságtudományi kérdésfeltevések általában arra irányulnak, hogy a tapasztalati idősorok milyen mértékben feleltethetők meg ezeknek az alapeseteknek. A sztochasztikus idősorelemzés módszertanának vázlatos áttekintése után következzen a disszertáció érdemi része, vagyis annak vizsgálata, hogy az előbb bemutatott trendmodellek, kauzalitás-tesztek, illetve kointegrációs összefüggések mennyire érzékenyek az empirikus idősorokban kimutatható korábban említett anomáliákra.

3 Rendszertelen idősorok kezelése

Az idősor-elemzéssel foglalkozó munkák túlnyomó többsége olyan idősorokkal foglalkozik, amelyekben a megfigyelések időpontjai egymástól *azonos* távolságra vannak, vagyis két megfigyelt időpont között egyenlő idő telik el. Az ilyen idősorokat tartalmazó adatállományok esetében tulajdonképpen nincs szükség a megfigyelés dátumának feljegyzésére, teljes körű információt kapunk akkor is, ha csak a kezdő, illetve végső időpontot, valamint a megfigyelések *gyakoriságát* (az idősor *frekvenciáját*) tüntetjük fel.

Ezeket az idősorokat *ekvidisztánsnak*³⁸ nevezzük, és modellezésük során nagyon gyakran élünk azzal az egyszerűsítéssel (beláthatóan információvesztés nélkül), hogy a tényleges dátum helyett az időpontokat fiktív, számtani sorozatot alkotó egész számokkal jelöljük. (A leggyakrabban – dolgozatomban eddig is – alkalmazott időmegjelölés a szokásos $t = 1, 2, \dots, T$ vagy $t = 0, 1, 2, \dots, T$.)

Az idősor-elemzési technikák ugyanakkor nem szűkíthetők le az egyenletes (rendszeres) idősorok modellezésére, ugyanis a gyakorlatban számos olyan esettel találjuk szembe magunkat, amikor a megfigyelések nem egyenlő időközönként követik egymást. A gazdaságtudományok esetén talán a legeklektánsabb példa a pénzügyi piacokon mért árfolyam-idősorok esete, ahol – elsősorban a hétvégi, ritkábban az ünnepnapok tőzsde-szünete következtében – az idősorokban „lyukak” vannak, még az amúgy egymástól rendszeresen 24 óra távolságra levő záró árfolyam-adatok sem ismétlődnek szabályosan. Az egy napnál nagyobb gyakoriságú pénzügyi idősorok esetén pedig a rendszertelenség tekinthető általánosnak, hiszen például a valamennyi üzletkötést tartalmazó árfolyam-idősorok esetén az azonos távolság kritériuma úgysem teljesül, mivel semmi sem garantálja, hogy a brókerek pl. 10, vagy 20 másodpercenként kötnek üzletet.)³⁹

A nem egyenletes időközökben keletkező (angolul *nonequidistant*, vagy *unequally-spaced*, *unevenly-spaced*, a német irodalomban *unregelmäßige*) megfigyelések modellezésének széles tárházát használják a tengerbiológiában, az asztrofizikában, a meteorológiában. Önmagában megérne egy hosszabb fejtegetést, hogy mi okból keletkeznek ilyen rendszertelenül megfigyelt idősorok, mennyiben lenne javítható a helyzet a mérési módszer tökéletesítésével, ám ezzel dolgozatomban nem foglalkozom. A *rendszertelenséget* (*időbeli egyenetlenséget*) adottságként fogom fel, ugyanakkor előrebocsájtom, hogy a modellezést nehezítő anomáliának, kellemetlen tulajdonságnak tekintem. A *nemekvidisztáns* idősoroknak egyik legkényelmetlenebb tulajdonsága az, hogy modellezésükre ez idáig nem találtak

³⁸ Az equidistant kifejezés helyett az angol nyelvű szakirodalom gyakran használja az *evenly-spaced*, vagy *equally-spaced* kifejezést is. Bár tény, hogy ezek a megjelölések viszonylag ritkán kerülnek elő, ugyanis „hírértéke” az egyenetlenség nem teljesülésének van.

³⁹ Rendszertelenül keletkező árfolyam-adatokra vonatkozó autoregresszív modellt mutat be Engle-Russell (1998). A nem egyenletes időközönként keletkező tőzsdei árfolyamok tőkepiaci árfolyamok modelljével (CAPM) történő elemzésének sajátosságait fejtegettem magam is egy korábbi tanulmányomban (lásd Rappai, 2004).

általánosan érvényes metódust, ami azt is maga után vonja, hogy az idősorok modellezését támogató szoftverek sem mindig vannak felkészítve a probléma azonnali és hatásos kezelésére.⁴⁰

Az idők során különböző megoldások fejlődtek ki a változó frekvenciájú idősorok modellezésére:

1. Elsősorban pénzügyi idősorok modellezése esetén (de a csapadék-modellezésében) is használatos, hogy az idősorban meglevő lyukakat 0-kal (bizonyos esetekben az utolsó megfigyelt tényleges értékkel⁴¹) töltjük fel, és az így *kiegészített* idősorokon végezzük el a modellezést. Intuitív módon is könnyen belátható, hogy ez a megoldás nagy mennyiségű kiegészítés esetén teljesen tévútra vihet, ezért nem ajánlható.
2. Amennyiben el akarjuk kerülni a fiktív adatok használatát, nyilvánvalóan legkézenfekvőbb megoldásnak tűnhet, hogy keresünk az idősornak egy olyan frekvenciát, amelyre minden megfigyelés ráilleszthető, majd *interpolálunk* a hiányzó időpillanatokhoz tartozó, kvázi megfigyeléseket. A megoldás mindenképpen gyors és kényelmes (főleg, ha valamilyen egyszerű, pl. lineáris interpolációt alkalmazunk), ugyanakkor a nemlinearitásra vonatkozó tesztek ilyenkor kevésbé lesznek hatásosak, gyakran előfordul, hogy olyankor is nemlinearitást mutatnak, amikor egyébként az eredeti adatok között ez nem volt tapasztalható (Schmitz, 2000).
3. Bizonyos mértékig szofisztikáltabb megoldás, amennyiben az idősor kovariancia-struktúrájából (autokovariancia-függvényéből) kiindulva képezünk *becslőfüggvényt*, amivel a hiányzó helyeket ki tudjuk egészíteni. (Ilyen esztimátorként szokták használni például a Lomb-becslőfüggvényt, eredeti leírását lásd Lomb, 1976, jó áttekintést ad róla Schmitz, 2000.)

Napjainkra a gazdasági idősorok modellezésével foglalkozó szakirodalom két megoldást preferál:

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

⁴⁰ Valószínűleg a legtöbb idősor-modellezéssel foglalkozó szakember szembe találkozott már azzal a problémával, amikor egy modellben szeretne volna szerepeltetni a napi árfolyam adatot a havi bontású infláció idősorával és a negyedéves GDP-vel...

⁴¹ Az árfolyam-modellezésben a leggyakoribb feladat a hozam előrejelzése, ebben az esetben a hiányzó hozamadatok 0-val történő feltöltése ekvivalens a hiányzó árfolyam-adat legutolsó tényleges adattal való helyettesítésével. Az utolsó még ténylegesen létező adattal történő helyettesítés angol elnevezése *forward-flat interpolation*.

A továbbiakban ezekről a módszerekről adok áttekintést, előrebocsátva, hogy noha a rendszertelenül megfigyelt idősorok modellezési eszköztára sokkal szűkebb, mint a hagyományos módszerek tárháza, azért itt is meglehetősen széles a szakirodalmi bázis, mindennek kimerítő összefoglalása meghaladná ezen dolgozat kereteit.

3.1 A probléma kezelése hiányzó adatok feltételezésével

A nemekvidisztáns idősorok kezelésének egyik útja, ha azzal a feltételezéssel élünk, hogy létezik egy „eredeti” idősor, ami tulajdonképpen ekvidisztáns, csak nem ismerünk belőle néhány megfigyelt értéket. A hiányzó (missing) adatok kezelésének leggyakrabban alkalmazott módszere az időszori *interpoláció*.

Az interpoláció egyébként egy általánosabban használatos eljárás, vagyis nem csak akkor használatos, ha hiányzó, vagy vélt hiányzó adatot akarunk pótolni. Minden olyan becslést így nevezünk, amelyben egy az idősor „közepén” (értsd nem a megfigyelési időszakon túl) található időponthoz rendelünk hozzá egy *ex post* becslést. Dolgozatomban az interpoláció két gyakran használatos típusát mutatom be a

- lineáris (illetve az ezzel gondolatvilágában azonos log-lineáris), valamint a
- spline

közelítést.⁴²

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 adott egy természetes gyakoriság, csak valamely okból nem keletkezik minden elvárt időpillanatban adat: gondoljunk a már említett napi záró-árfolyam idősorban szombatonként és vasárnaponként keletkező lyukakra. (Ebben az esetben a természetes megfigyelési gyakoriság a naponkénti, ugyanakkor minden 6. és 7. érték hiányzik.) Más esetekben nincs ilyen kézenfekvő megoldás, hiszen 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:

⁴² A spline interpolációra vonatkozó fejtegetés megjelent Rappai (2014), illetve minimálisan módosítva Rappai (2015b) cikkekben.

- egyrészt kívánatos, hogy minél több eredeti megfigyelés ráessen a feltételezett idősorra, ami – könnyen átláthatóan – a minél sűrűbb megfigyelési gyakoriság mellett érv;
- másrészt el kell(ene) kerülni, hogy az interpolált értékek száma meghaladja (egyres 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 minden 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.

A szakirodalom az elmúlt két évtizedben sokat foglalkozott azzal a problémával, hogy az extrém nagy sűrűségű idősorok (ultra-high frequency data) modellezése esetén a fenti kritériumok nehezen teljesíthetők. Az ilyen, tipikusan tőzsdei ügyletkötéseket tartalmazó tőzsdei idősorok modellezési lehetőségeinek kitűnő összefoglalása olvasható Engle (2000) tanulmányában.

Vezessük be a következő jelöléseket: legyen a megfigyelt (empirikus) idősorunk

$$\mathcal{Y}_{t_1}, \mathcal{Y}_{t_2}, \dots, \mathcal{Y}_{t_i}, \dots, \mathcal{Y}_{t_T}$$

ahol $t_2 - t_1$ nem feltétlenül egyezik meg $t_3 - t_2$ távolsággal. Legyen Δ az a legnagyobb távolság, amelyre igaz, hogy valamennyi $t_i - t_{i-1}$ megegyezik Δ -val, vagy annak egész számú többszörösével.⁴³

Ekkor képezhető az alábbi hiányos idősor:

$$\mathcal{Y}_{t_1}, \mathcal{Y}_{t_1+\Delta}, \mathcal{Y}_{t_1+2\Delta}, \dots, \mathcal{Y}_{t_1+j\Delta}, \dots, \mathcal{Y}_{t_T}$$

ahol $\mathcal{Y}_{t_1+j\Delta}$ eredetileg nem megfigyelt, vagyis interpolációval előállítandó adat, ha $t_1 + j\Delta$ nem esik egybe egyetlen t_i -vel sem. Az interpoláció során a feladatunk tehát az, hogy valamilyen eljárással becsüljük azokat az értékeket, melyek olyan időpontokhoz tartoznak, amelyből eredetileg nem származik tapasztalati adatunk.

⁴³ Ez gyakran, de nem feltétlenül megegyezik a két tényleges megfigyelés közötti minimális távolsággal.

Triviálisnak tűnő, igaz a korábban felállított kritériumrendszernek nem feltétlenül megfelelő eljárás a *lineáris interpoláció*. Ekkor ha

$$y_{t_{i-1}} = y_{t_1+(j-1)\times\Delta} < y_{t_1+j\times\Delta} < y_{t_i} = y_{t_1+(j+1)\times\Delta}$$

vagyis az $y_{t_1+j\times\Delta}$ érték hiányzik, ám az előtte és utána következő értékek rendelkezésünkre állnak, akkor az interpolációval pótoltt érték

$$\hat{y}_{t_1+j\times\Delta} = y_{t_{i-1}} + \frac{y_{t_i} - y_{t_{i-1}}}{t_i - t_{i-1}} = y_{t_{i-1}} + \frac{y_{t_i} - y_{t_{i-1}}}{2\Delta} \quad (3.1)$$

Az utóbbi felírásból jól nyomon követhető, hogy amennyiben 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

$$\hat{y} = (1-\lambda) y_{t_{i-1}} + \lambda y_{t_i} \quad (3.2)$$

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. (Látható, hogy amennyiben egy érték hiányzik, akkor felezni kell az ismert különbséget, ha kettő hiányzó adat van, akkor harmadolni és így tovább.)

Könnyen belátható, hogy az így képzett egyszerű lineáris interpoláció meglehetősen „töredezett” folyamatot szolgáltat, az eljárással keletkező becslt függvény meredeksége gyakran és ugrásszerűen változik. Ennek a töredezettségnek a tompítására szokták alkalmazni a *log-lineáris interpolációt*, ahol a becslt érték a

$$\hat{y} = e^{(1-\lambda)\ln y_{t_{i-1}} + \lambda\ln y_{t_i}} \quad (3.3)$$

képlettel keletkezik. Miközben a logaritmálás variancia-stabilizáló jellegénél fogva az eljárás némiképpen 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.⁴⁴

⁴⁴ Az eljárások értelemszerűen továbbfejleszthetők: amennyiben nem csak a hiányzó adatot közvetlenül megelőző, illetve követő ismert értéket használjuk fel, akkor az interpoláció simasága javítható (ilyen pl. az EViews programban használatos Cardinal spline módszer).

A sima függvények megtalálására fejlesztették ki az ún. *spline interpolációt*.⁴⁵ 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:

$$(1) \quad S(t) = S_{t_i}(t) \quad t \in [t_i, t_{i+1}]$$

$$(2) \quad S(t_i) = y_{t_i}$$

$$(3) \quad S_{t_i}(t_{i+1}) = S_{t_{i+1}}(t_{i+1})$$

Az (1)-(3) feltételek tulajdonképpen a következőket jelentik:

- az interpoláció *szakaszokból áll*, és akár minden szakaszra különböző függvényt definiálhatunk;
- az interpoláló függvény a tényleges megfigyeléseket *képes reprodukálni*; valamint
- az interpoláció eredményképpen *kapott görbe folytonos* (hiszen a közbülső megfigyelések két szakaszhoz is tartoznak, de ott mindkét szakasz egyenlő értékkel bír).

Az előbbi három feltétel a spline-interpoláció általános definíciója. Annak függvényében, hogy milyen típusú $S(t)$ függvényeket használunk, más-más spline-eljárásokról beszélhetünk. (Noha az általános definíció megengedi, hogy akár minden szakaszon más-más függvénytípust használjunk, általában azonos függvényosztályból származtatjuk az interpoláló függvényeket.) A leggyakoribb megoldás, hogy $S(t)$ függvényeket a polinomok közül választjuk, mégpedig úgy, hogy magasabb fokszámú polinomok esetében a közbülső pontokban csatlakozó szakaszoknál a deriváltak egyezőségét is megköveteljük. Általánosságban egy spline k -ad fokú és m -ed rendű, ha szakaszonként legfeljebb k -ad fokú polinomokból áll és a közbülső pontokban a találkozó szakaszok deriváltjai m -ed rendig megegyeznek.

Ebben a részben két – a gyakorlatban viszonylag elterjedt – spline-interpolációt mutatok be:

- a lineáris spline-t és
- harmadfokú, másodrendű spline-t.

⁴⁵ A *spline* kifejezésnek ezidáig nem honosodott meg magyar megfelelője. A szó eredetileg a hajógyártásból származik, a hosszú, rugalmasan hajlítható, a hajótest formáját jól követő deszkákra (lécekre) használatos. A spline-okra vonatkozó első matematikai hivatkozás Schoenberg (1946) cikkében olvasható.

A *lineáris spline* tárgyalása során elsőként fókuszáljunk mindössze egy szakaszra: legyen a vizsgált intervallum $[t_i, t_{i+1}]$.⁴⁶ Definíció szerint a spline-ra igaz, hogy

$$\begin{aligned} S_{t_i}(t_i) &= \alpha_{t_i} t_i + \beta_{t_i} = y_{t_i} \\ S_{t_i}(t_{i+1}) &= \alpha_{t_i} t_{i+1} + \beta_{t_i} = y_{t_{i+1}} \end{aligned}$$

A fenti két ismeretlenes kétegyenletes rendszerből az ismeretlen paraméterek $(\alpha_{t_i}, \beta_{t_i})$ rendre meghatározhatóak, vagyis a spline felírható. A megoldás egyébiránt azonos az ismert, szokásos lineáris közelítéssel, vagyis:

$$S(t) = y_t = y_{t_i} + \frac{y_{t_{i+1}} - y_{t_i}}{t_{i+1} - t_i} (t - t_i) \quad t \in [t_i, t_{i+1}] \quad (3.4)$$

A felírásból jól látható, hogy a lineáris spline paraméterei minden $[t_i, t_{i+1}]$ intervallumban változnak, illetve változhatnak.

A gyakorlatban – mivel ésszerű számolásigénnyel megfelelő rugalmasságot biztosít – általában harmadfokú, másodrendű spline interpolációt⁴⁷ alkalmazunk. *Harmadfokú spline* esetén a korábban tárgyalt (1)-(3) feltételek újabb hárommal⁴⁸ egészülnek ki, melyek a megfigyelt értéknél találkozó függvények (polinomok) simaságát hivatottak garantálni:

$$(4) \quad S'_{t_i}(t_{i+1}) = S'_{t_{i+1}}(t_{i+1})$$

$$(5) \quad S''_{t_i}(t_{i+1}) = S''_{t_{i+1}}(t_{i+1})$$

$$(6) \quad S''(t_1) = 0 \quad S''(t_T) = 0$$

Az interpolációhoz szükséges görbék definiálása során tehát keressük a

$$S(t) = S_{t_i}(t_i) = y_{t_i} = \alpha_{t_i} + \beta_{t_i} (t - t_i) + \gamma_{t_i} (t - t_i)^2 + \delta_{t_i} (t - t_i)^3 \quad t \in [t_i, t_{i+1}] \quad (3.5)$$

⁴⁶ Mivel a vizsgált intervallum két végpontján ismert érték helyezkedik el, így ha meg tudjuk határozni a két empirikus érték között lefutó görbét (a sztochasztikus folyamat alakulását), akkor a szakaszon található hiányzó adatokat csak erről a függvényről kell leolvasnunk.

⁴⁷ Amikor a harmadfokú, másodrendű spline interpolációról esik szó, általában egyszerűen harmadfokú spline-ról (*cubic spline*) beszélünk. A továbbiakban a dolgozatban én is élek ezzel az egyszerűsítéssel.

⁴⁸ Az alábbi felírás az ún. *természetes spline*-ra vonatkozik, elvben nem kizárt, hogy a második derivált a kezdő, illetve a végső megfigyelésnél nem 0.

kifejezésekhez tartozó paramétereket. Belátható, hogy a (3.5)-ben szereplő összesen $4(T-1)$ darab ismeretlen paraméterhez a (2)-(6) feltételek pontosan ugyanennyi egyenletet határoznak meg, így a feladat megoldható.⁴⁹

Az előbbiekben bemutatott, általánosan definiált harmadfokú, másodrendű spline-interpoláció helyett gyakran alkalmazzák az ún. *Catmull-Rom-spline*-okat (első leírását lásd Catmull-Rom, 1974, továbbiakban *CRS*). Az eljárás akkor alkalmazható, ha feltételezhetjük, hogy a rendszertelen idő-sor tulajdonképpen nem más, mint egy ekvidisztáns idősor, melyből hiányoznak megfigyelések. Vezessük be a

$$\mathcal{Y}_{t_1+j\Delta} = \mathcal{Y}_j$$

jelölést, így az idősor első, biztosan megfigyelt értéke y_0 , a második értéke y_1 , és így tovább. Az eljárás lényege, hogy feltesszük minden (megfigyelt, vagy éppen hiányzó) érték egy harmadfokú polinomra fekszik, melynek az adott helyen nem csak az értékét, de a deriváltját is ismerjük, vagyis

$$\begin{aligned} y(t) &= \alpha_0 + \alpha_1 t + \alpha_2 t^2 + \alpha_3 t^3 \\ y'(t) &= \frac{\partial y(t)}{\partial t} = \alpha_1 + 2\alpha_2 t + 3\alpha_3 t^2 \end{aligned} \quad (3.6)$$

Nézzük az első két pont esetén mindez mit jelent:

$$\begin{aligned} y(0) &= \alpha_0 \\ y(1) &= \alpha_0 + \alpha_1 + \alpha_2 + \alpha_3 \\ y'(0) &= \alpha_1 \\ y'(1) &= \alpha_1 + 2\alpha_2 + 3\alpha_3 \end{aligned}$$

Oldjuk meg az egyenletrendszert az ismeretlen paraméterekre:

$$\begin{aligned} \alpha_0 &= y(0) \\ \alpha_1 &= y'(0) \\ \alpha_2 &= 3[y(1) - y(0)] - 2y'(0) - y'(1) \\ \alpha_3 &= 2[y(0) - y(1)] + y'(0) + y'(1) \end{aligned}$$

⁴⁹ A bizonyítást lásd pl. Mészárosné (2011).

Mindezt visszahelyettesítve az eredeti polinomba, és elvégezve a szükséges egyszerűsítéseket kapjuk az alábbi – könnyen átláthatóan – harmadfokú polinomot:

$$y(t) = (1 - 3t^2 + 2t^3)y(0) + (3t^2 - 2t^3)y(1) + (t - 2t^2 + t^3)y'(0) + (-t^2 + t^3)y'(1)$$

Láthatóan a nehézséget az okozza, hogy a különböző megfigyelt értékeknél nehezen adható meg az illesztett (illesztendő) görbe deriváltja (meredeksége). A Catmull-Rom eljárás lényege az, hogy felteszi, az előbbi deriváltak a megfigyelt értékekből egyszerűen meghatározhatóak. Keressük a spline-t az $[y_j, y_{j+1}]$ szakaszon! Legyenek a keresett meredekségek az alábbiak:

$$y'(j) = \frac{y_{j+1} - y_{j-1}}{2}$$

$$y'(j+1) = \frac{y_{j+2} - y_j}{2}$$

Így a korábban bemutatott harmadfokú polinom felírható mátrix alakban a következőképpen:

$$y(t) = \begin{bmatrix} 1 & t & t^2 & t^3 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ -3 & 3 & -2 & -1 \\ 2 & -2 & 1 & 1 \end{bmatrix} \begin{bmatrix} y_j \\ y_{j+1} \\ \frac{y_{j+1} - y_{j-1}}{2} \\ \frac{y_{j+2} - y_j}{2} \end{bmatrix} \quad (3.7)$$

Míndez minimálisan átalakítva:

$$y(t) = \begin{bmatrix} 1 & t & t^2 & t^3 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ -3 & 3 & -2 & -1 \\ 2 & -2 & 1 & 1 \end{bmatrix} \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ -\frac{1}{2} & 0 & \frac{1}{2} & 0 \\ 0 & -\frac{1}{2} & 0 & \frac{1}{2} \end{bmatrix} \begin{bmatrix} y_{j-1} \\ y_j \\ y_{j+1} \\ y_{j+2} \end{bmatrix} \quad (3.7a)$$

majd a két belső mátrixot összeszorozva

$$y(t) = \frac{1}{2} \begin{bmatrix} 1 & t & t^2 & t^3 \end{bmatrix} \begin{bmatrix} 0 & 2 & 0 & 0 \\ -1 & 0 & 1 & 0 \\ 2 & -5 & 4 & -1 \\ -1 & 3 & -3 & 1 \end{bmatrix} \begin{bmatrix} y_{j-1} \\ y_j \\ y_{j+1} \\ y_{j+2} \end{bmatrix} \quad (3.8)$$

A (3.8) egyenlettel meghatározott, viszonylag könnyen felírható görbe reprezentálja az idősor alakulását két kijelölt pont között. (Minden különösebb magyarázat nélkül látható, hogy a görbék minden szakaszon változnak, pontosabban változhatnak.)

A spline interpoláció alkalmas eszköznek látszik szinte minden rendszertelen idősor hiányzó adatainak pótlására: elegáns és viszonylag egyszerű a használandó matematikai algoritmus, könnyen interpretálható, szinte bármilyen frekvencián értelmezhető idősor az outputja. Ugyanakkor valószínűleg minden felhasználóban megfogalmazódik a gondolat, hogy szinte bármilyen fejlődési pályák kirajzolhatók, ha egy elég ritka idősort viszonylag nagy frekvenciájú idősorra próbálunk konvertálni. Erre a dolgotatomban szempontjából meghatározó fontosságú kérdésre hamarosan visszatérek.

3.2 Folytonos idejű modellek a rendszertelenül megfigyelt adatok elemzésében

Amennyiben az adatsorunk ténylegesen rendszertelenül keletkezik, vagyis nem egy ekvidisztáns idősorból hiányzik egy-egy megfigyelés, célszerű lehet bevezetni a *folytonos időt feltételező modelleket* (*continuous-time model*). A folytonos idő problémája tulajdonképpen az ekvidisztáns idősoroknál is felmerül, ám kezelése ebben az esetben nem oly szükségzerű, mint a rendszertelen megfigyeléseket tartalmazó adatsoroknál. A probléma már viszonylag korán megjelent a modellezési szakirodalomban, érdemben foglalkozott a kérdéssel Jones (1985), Bergstrom (1985), vagy Belcher et al. (1994). Mivel a dolgozatnak nem célja a folytonos időt feltételező modellekre vonatkozó valamennyi eredmény részletes tárgyalása, így az érdeklődőknek ajánljuk Brockwell (2001) kitűnő összefoglalóját. A következőkben csak a témánk szempontjából feltétlenül szükséges ismeretek szerepelnek, melynek összeállításakor alapvetően Cochrane (2012), illetve Wang (2013) tanulmányaira támaszkodtam.

Definiáljuk a folytonos időt feltételező autoregresszív modellt (*continuous-time autoregressive model*, a továbbiakban *CAR*) a következő módon: legyen $y(t_i)$ a megfigyelt idősori 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

$$Y^{(p)}(t) + \varphi_1 Y^{(p-1)}(t) + \dots + \varphi_{p-1} Y^{(1)}(t) + \varphi_p Y(t) = \varepsilon(t) \quad (3.9)$$

ahol $Y^{(i)}(t)$ az i -edik deriváltja $Y(t)$ -nek, $\varepsilon(t)$ pedig az első deriváltja a $B(t)$ Brown-mozgásnak, melynek varianciája $\text{var}_B[B(t+1) - B(t)]$.⁵⁰

⁵⁰ Wang (2013) cikkében a megfigyelt változó további additív zajt (véletlen változót) is tartalmaz, ám ettől most az egyszerűsítés kedvéért eltekintünk.

Vezessük be a deriváló operátort!

$$Dy_t = \frac{dy_t}{dt} \quad (3.10)$$

Kézenfekvő, hogy az ismert késleltetési operátor és a deriváló operátor között egyszerű összefüggés írható fel:

$$L = e^{-D}$$

$$D = -\ln L$$

Amennyiben felírjuk, hogy

$$\varphi(D) = D^p + \varphi_1 D^{p-1} + \dots + \varphi_{p-1} D + \varphi_p$$

úgy a kiinduló CAR(p) modellünk felírható a szokásos egyszerűbb formában:

$$\varphi(D)Y(t) = \varepsilon(t) \quad (3.11)$$

Jones (1985) bebizonyította, hogy a CAR-modell stabilitás vizsgálata egy alkalmasan választott karakterisztikus egyenlet felírásával, illetve megoldásával elvégezhető. Belcher et al. (1994) tanulmányukban részletesen megindokolták, hogy mennyiben szerencsésebb (könnyebben kezelhető) formulához jutunk, ha az előbbi egyenletet minimálisan módosítjuk

$$\varphi(D)Y(t) = (1 + D/\kappa)^{p-1} \varepsilon(t) \quad (3.11a)$$

ahol $\kappa > 0$ egy skálázó paraméter. (Az előbbi modell tulajdonképpen egy mozgóátlag szűrővel bővített változat, amivel a korábban említett additív zaj kiküszöbölhető.)⁵¹

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 már többször hivatkozott irodalom becslési eljárásaként Kalman-filter használatát javasolja (lásd Kalman-Bucy, 1961). A Kalman-filter alkalmazásához először át kell írni a becslendő modellt állapot-tér formába. Legyen az ismeretlen állapot-vektor

⁵¹ Cochrane (2012) kitűnő munkanyaga részletesen és következetesen felsorolja a folytonos időt feltételező modellek explicit formáját a CAR(1) folyamattól egészen a hibakorrektációs modellig. Ebben a dolgozatban csak CAR(p) folyamatokkal foglalkozom, így a témában elmélyedni kívánó Olvasónak javaslom az említett műhelytanulmányt.

$$\boldsymbol{\theta}(t) = [\zeta(t), \zeta^1(t), \dots, \zeta^{p-1}(t)]^T$$

és jelölje $\boldsymbol{\theta}'(t)$ az előbbi vektor első deriváltját. Ekkor állapot-egyenlet felírható a következő alakban

$$\boldsymbol{\theta}' = \mathbf{A}\boldsymbol{\theta} + \mathbf{R}\epsilon \quad (3.12)$$

ahol

$$\mathbf{A} = \begin{bmatrix} 0 & 1 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \\ -\varphi_p & -\varphi_{p-1} & \dots & -\varphi_1 \end{bmatrix}$$

és

$$\mathbf{R}^T = [0 \quad 0 \quad \dots \quad 1]$$

A megfigyelt értékekre vonatkozó egyenlet

$$y(t_k) = \mathbf{h}\boldsymbol{\theta}(t_k) \quad (3.13)$$

ahol az $(1 \times p)$ típusú $\mathbf{h}$ vektor i -edik eleme ($i = 1, 2, \dots, p$)

$$h_i = \binom{p-1}{i-1} / k^{i-1}$$

Innentől kezdve a Kalman-filterrel történő paraméterbecslés Wang (2013) alapján lépésről lépésre követhető. A becslés számolásigénye már nem túl hosszú empirikus idősorok esetén is jelentős, ennek kezelése érdekében a különböző algoritmusok különböző paraméterrestrikciókkal élnek.

A folytonos időt feltételező modellek, annak ellenére, hogy sok tanulmány foglalkozik a témával, és gyakorlatilag valamennyi a dolgozatomban tárgyalt folyamatnak és jelenségnek (pl. impulzus-válasz függvény, vagy kointegráció) megtalálható a folytonos idejű megfelelője, *nem terjedtek el igazán az empirikus adatelemzésekben*. A leggyakoribb alkalmazásuk pénzügyi idősorok esetén a következő

tranzakció legvalószínűbb helyének (időpontjának) előrejelzése, illetve az árfolyamok volatilitás-változásának prognosztizálása.

3.3 Rendszertelen idősorok kiegészítésével (interpolálással) elkövethető hibák

Ebben az alfejezetben arra mutatok szimulációval nyert eredményeket, hogy milyen következményei lehetnek annak, ha a rendszertelen (vagy nem azonos időpontban megfigyelt) idősorokat úgy próbáljuk meg modellezni, hogy első lépésben valamilyen interpolációs technikával kiegészítjük azokat, majd az új idősorokban keresünk trendeket.⁵²

A 2.4 pontban már áttekintettük azt a 8 idősori alapmodellt, melyek a dolgozat szimulációs fejezeteiben újfent megjelennek. A nemekvidisztáns idősorok vizsgálata során tehát az eredeti szimulációk egy lépéssel bővülnek, eszerint a generált 1000 elemű idősorainkból elhagyunk bizonyos megfigyeléseket, és az így kapott rendszertelen idősorokat elemezzük tovább. A megfigyelések elhagyása két módon történik majd

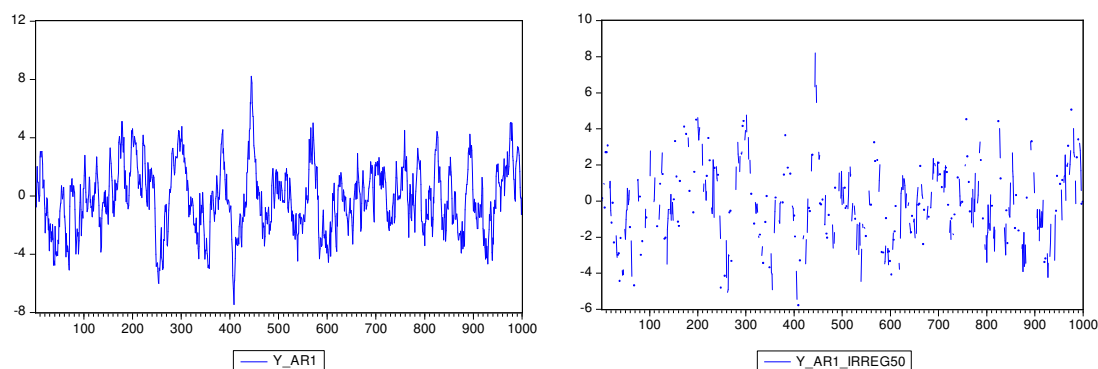
- *véletlenszerűen*, azaz egy véletlenszám-generálás eredményeképpen keletkeznek azok a sorszámok (megfigyelési egységek, vagyis tulajdonképpen időpontok), melyeknél az idősor elemeit elhagyom, méghozzá a vizsgálat céljának megfelelően az idősor 5, 10, 25, 50, 75, 90 százalékát fogom elhagyni, vagyis az eredeti 1000 megfigyelést tartalmazó szimulált idősor helyett rendre 950, 900, 750, 500, 250, illetve 100 elemű idősorokkal dolgozok;
- *szisztematikusan*, ekkor azt az esetet próbálom szimulálni, amikor a rendelkezésre álló idősorok nem azonos frekvenciájúak⁵³; ebben az esetben két kísérletet végeztem, az elsőben a szimulált idősorban csak minden harmadik megfigyelést hagytam meg (mintha egy havi bontású idősor egy negyedévessel „találkozna”), valamint amikor csak minden negyedik megfigyelés maradt az eredeti (negyedéves és éves felbontású idősorok közötti kapcsolat elemzésének szimulálása).⁵⁴

Remélhetőleg egyszerűbb átlátni az előbbieket, amennyiben ábrázoljuk az eredeti, illetve a kihagyott értékeket tartalmazó idősorokat.

⁵² A szimulációs eredmények egy része megjelent Rappai (2014), illetve Rappai (2015b) tanulmányokban.

⁵³ Itt nyilván csak az okság-vizsgálatnak, illetve a közös trend keresésének van értelme, hiszen csak több idősor együttes elemzése esetén reális az a felvetés, miszerint az egyik idősor „hozzá kell igazítani” a másik idősor frekvenciájához.

⁵⁴ Emlékeztetek arra, hogy szimulációról van szó, a valóságban nem túl valószínű, hogy pl. 1000 hónap és 333 negyedév, azaz közel 85 év idősorát módunk van vizsgálni. Az empirikus modellek a következő alfejezetben következnek.



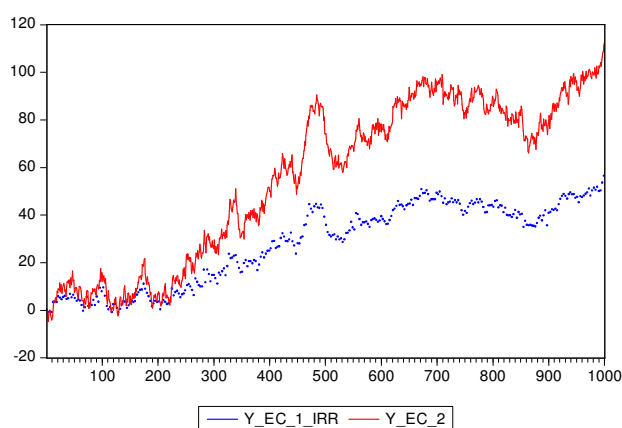
(a) eredeti

(b) véletlenszerű 50% elhagyásával

3-1. ábra: $AR1$ ($\varphi=0,9$) folyamatból származó idősor

A 3-1. ábra jól mutatja, hogy a megfigyelések véletlenszerű elhagyása következtében az idősor „nyomokban emlékeztet” az eredeti idősorra, valószínűleg egymás mellett látva a két ábrát úgy érezzük, hogy az eredeti stacioner folyamat is felismerhető, de érdemes belegondolni, hogy vajon akkor is triviálisnak tűnne-e a folyamat stabilitása, ha csak a (b) ábrát látnánk. Az ábra (b) része megfelelően illusztrálja, hogy a rendszertelenné tétel véletlenszerűen történt, vagyis nem úgy maradt meg 500 érték az eredeti idősorból, hogy pl. minden másodikat hagytuk el.

A 3-2. ábrán arra láthatunk példát, amikor két (egyébiránt kointegrált) idősorból az egyiknél csak minden harmadik elemet hagyjuk meg, vagyis azt szimulálom, hogy az egyik idősor havi, a másik negyedéves bontású.



3-2. ábra: Kointegrált idősorok szisztematikus kihagyással

Ebben az esetben az ekvidisztáns, de nem azonos frekvenciájú idősor viszonylag jól együttmozog a nagyobb gyakorisággal megfigyelt folyamattal, a szimuláció során azt fogom vizsgálni, hogy az interpolációval történő kiegészítés mennyiben rontja a kointegrációs tesztek erejét.

3.3.1 *A stacionaritás, illetve a sztochasztikus trend felismerése nemekvidisztáns idősorok esetén*

A második fejezetben már bemutattam a dolgozatban mindenütt újra visszatérő alapmodelleket, melyekből itt számunkra

- a $\varphi=0,9$ paraméterrel rendelkező Y_{AR1} , illetve a $\varphi=-0,4$ paraméterű Y_{AR2} stacionárius folyamatból származó, valamint
- az eltolás nélküli Y_{RW1} , illetve a $\mu=0,01$ eltolást (driftet) tartalmazó Y_{RW2} egységgyököt, tehát sztochasztikus trendet tartalmazó

idősorok a fontosak.

A szimulációk során elsőként azt vizsgálom, hogy véletlenszerűen elhagyva 5, 10, 25, 50, 75 és 90%-ot az 1000 elemű idősorokból, majd a hiányzó adatokat alkalmasan választott interpolációs technikával pótolva, 1000 szimulációból hányszor ítéljük meg rosszul a trend meglétét, illetve hiányát az ADF próba alapján. A korábban bemutatott interpolációs technikák közül a lineáris (*LIN*), a Catmull-Rom-spline (*CRS*) és a harmadfokú spline (*CUB*) módszereket alkalmaztam.

Az alábbi táblázatok címében és a következőkben is többször – némiképp utalva az orvos-kollégákkal folytatott nagyszámú közös kutatásra – *fals-positív*nak nevezem azokat az eseteket, amikor az ADF-teszt egységgyököt talált stacionárius folyamat esetén, és *fals-negatív*nak azokat az eseteket, amikor az ADF-próba alapján elvetjük a nullhipotézist, holott az idősor sztochasztikus trendet tartalmazó folyamatból származik, mindkét esetben 5%-os szignifikancia-szintet alkalmazva. A szimulációk eredményeit a 3-1. a-b táblázatok tartalmazzák.

3-1.a táblázat: *A fals pozitív esetek száma rendszertelenül megfigyelt elsőrendű autoregresszív folyamatok interpolációval történő pótlása esetén*

Véletlenszerűen kihagyott megfigye- lések aránya	A szimulált folyamat					
	AR1 ($\varphi=0,9$)			AR1 ($\varphi=-0,4$)		
	Interpoláció módja					
	LIN	CRS	CUB	LIN	CRS	CUB
5%	0	0	0	0	0	0
10%	0	0	0	0	0	0
25%	0	0	0	0	0	0
50%	0	0	0	0	0	0
75%	0	0	0	0	0	0
90%	4	0	2	0	0	0

3-1.b táblázat: *A fals negatív esetek száma rendszertelenül megfigyelt véletlen bolyongás interpolációval történő pótlása esetén*

Véletlenszerűen kihagyott megfigye- lések aránya	A szimulált folyamat					
	RW ($\mu=0$)			RW ($\mu=0,01$)		
	Interpoláció módja					
	LIN	CRS	CUB	LIN	CRS	CUB
5%	36	36	39	39	41	43
10%	46	52	60	39	44	55
25%	50	59	53	54	61	45
50%	40	55	58	50	58	54
75%	36	48	76	36	62	84
90%	46	61	113	38	58	104

Láthatjuk, hogy amennyiben a stacionárius adatgeneráló folyamatból származó idősorokból csak rendszertelen megfigyelések állnak rendelkezésünkre, és a nemekvidisztáns idősorokat interpolációval kiegészítjük, gyakorlatilag csak minimális esély van arra, hogy a kiegészített idősor sztochasztikus trendet tartalmazzonak tűnjön, vagyis attól nem kell igazán tartanunk, hogy interpolációval trendet generálunk.

Kis mértékben rosszabb a helyzet az egységgyök megítélése során: ugyanis, ha egy véletlen bolyongásból származó idősorból sok elemet (75-90%-ot) elveszítünk, vagy nem tudunk felmérni,

akkor előfordulhat, hogy az interpolációkkal kiegészített idősor helytelen módon stacionáriusnak tűnik. Hangsúlyozandó, hogy mindehhez a megfigyelések nagyobb részének hiánya szükséges (ne tévesszenek meg a viszonylag magas számok az alacsonyabb régiókban, hiszen az ADF-próba 5%-o szignifikancia-szint mellett 1000 esetből átlagosan 50-szer téved akkor is, ha az idősorból nem hagyunk ki elemeket!).

Érdeemes arra is rámutatni, hogy a lineáris folyamatból származó nemekvidisztáns idősorok viszapótlása esetén az összetettebb interpolációs technikák (Catmull-Rom, illetve harmadfokú spline) rosszabbul teljesítettek.

Összefoglalva kijelenthető, hogy *a rendszertelenül megfigyelt idősorok interpolációval történő kiegészítése, az így már ekvidisztánssá vált idősorok stacionaritásának megítélése szempontjából nem aggályos.*

3.3.2 Okság, illetve együttmozgás felismerése rendszertelen, illetve nem azonos frekvenciájú idősorok esetén

Ebben a részben azt vizsgálom, az immár talán szokásosnak nevezhető szimulációs technikával, hogy mennyiben okoz problémát az okság (együttmozgás) megítélésében, ha

- a 2.4.2 alpontban bemutatott, VAR-modellből származó, egymással kölcsönösen ok-okozati viszonyban álló idősorok, illetve
- a 2.4.4 alpontban bemutatott kointegrált (közös trendet tartalmazó) idősorok esetében

hagyunk el előbb véletlenszerűen, majd szisztematikusan megfigyeléseket, majd ezeket interpolációs technikákkal pótolva, a kiegészített idősorokra vizsgáljuk az okság, illetve kointegráció meglétét.

A 3-2. táblázatok azt mutatják, hogy milyen következményei lesznek a Granger-teszt szempontjából, ha az eredetileg VAR-modellből származó idősorokból véletlenszerűen, különböző mennyiségű elemet elhagyunk⁵⁵, majd ezeket interpolációval pótoljuk és így vizsgáljuk az okságot (az alábbi táblázatok egyirányú Granger-okságot vizsgálnak, ahol y_{1t} az ok és y_{2t} az okozat):

⁵⁵ Hangsúlyozandó, hogy nem feltétlenül ugyanaz a megfigyelés származik a két idősorból, hiszen a véletlenszerű megfigyelés elhagyás ezt a legkevésbé sem támogatja.

3-2.a táblázat: Okság hiányát mutató próbák száma, különböző arányban hiányos idősorok
lineáris interpolációval történő pótlása esetén

Véletlenszerűen kihagyott megfigyelések aránya az y_{1t} idősorból	Véletlenszerűen kihagyott megfigyelések aránya az y_{2t} idősorból					
	5%	10%	25%	50%	75%	90%
5%	0	0	0	0	0	0
10%	0	0	0	0	0	0
25%	0	0	0	0	0	1
50%	0	0	0	0	0	2
75%	450	443	383	138	6	29
90%	921	905	882	726	393	428

A 3-2a. táblázatból leolvasható eredmények roppant érdekesek: láthatjuk, hogy nem fogadjuk el helytelenül az okság hiányát jelentő nullhipotézist, csak olyankor, ha az ok szerepét betöltő változóból hiányzik viszonylag nagyszámú megfigyelés. Az ilyen esetekben szembeötlő, hogy a Granger-okságot tesztelő Wald-próba leginkább arra az esetre érzékeny, amikor az ok szerepét játszó változóból kevés megfigyelésünk van, azaz sokat kell interpolációval pótolnunk, mindeközben a másik stacionárius változó szinte hiánytalan.

Hasonló eredményekre jutunk akkor is, ha nem lineáris, hanem harmadfokú spline interpolációval pótoljuk az időszori értékeket:

3-2.b táblázat: Okság hiányát mutató próbák száma, különböző arányban hiányos idősorok
harmadfokú spline interpolációval történő pótlása esetén

Véletlenszerűen kihagyott megfigyelések aránya az y_{1t} idősorból	Véletlenszerűen kihagyott megfigyelések aránya az y_{2t} idősorból					
	5%	10%	25%	50%	75%	90%
5%	0	0	0	0	0	31
10%	0	0	0	0	0	43
25%	0	0	0	0	0	50
50%	0	0	0	0	0	50
75%	145	148	183	190	11	196
90%	783	789	807	810	687	690

Az előbbieket alapján kijelenthető, ahhoz, hogy egymással ok-okozati viszonyban álló idősorok esetén a Wald-próbával *elkövessük a másodfajú hibát*, az idősorokból *nagyszámú* (a szimulációk alap-

ján az eredeti elemszám felét meghaladó mértékű) *megfigyelésnek kell hiányozni*. Szintén felfigyelhetünk az *ok és az okozat közötti aszimmetriára*, vagyis a hibás döntés nagyobb valószínűséggel fordul elő, ha az ok szerepét játszó idősorból hiányoznak az adatok, mintha az okozatból.

Ezt követően megismételtem a vizsgálatot nemstacioner (sztochasztikus trendet tartalmazó), kointegrált változók esetében. Azt vizsgáltam, hogy amennyiben a két közös trendet tartalmazó változót eredményező DGP-ből (ennek specifikációját a 2.4.4 alpontban bemutattam) csak rendszertelenül keletkeznek adatok, ezért az így megfigyelt nemekvidisztáns idősorokat különböző (lineáris, illetve harmadfokú spline) interpolációkkal ki kell egészítenünk, majd a kiegészített idősorokra EG-tesztet futtattam, és megszámláltam, hogy milyen arányban „tűnik el” az együttmozgás a kointegrált idősorokból. Az eredményeket 3-3.a és 3-3.b táblázatok tartalmazzák:

3-3.a táblázat: Kointegráltság hiányát mutató próbák száma, különböző arányban hiányos idősorok lineáris interpolációval történő pótlása esetén

Véletlenszerűen kihagyott megfigyelések aránya az y_{1t} idősorból	Véletlenszerűen kihagyott megfigyelések aránya az y_{2t} idősorból					
	5%	10%	25%	50%	75%	90%
5%	4	2	3	4	9	55
10%	3	2	4	2	11	33
25%	4	6	5	5	11	38
50%	9	5	10	6	8	45
75%	9	14	15	20	31	52
90%	50	48	52	46	66	133

3-3.b táblázat: Kointegráltság hiányát mutató próbák száma, különböző arányban hiányos idősorok harmadfokú spline interpolációval történő pótlása esetén

Véletlenszerűen kihagyott megfigyelések aránya az y_{1t} idősorból	Véletlenszerűen kihagyott megfigyelések aránya az y_{2t} idősorból					
	5%	10%	25%	50%	75%	90%
5%	7	2	5	8	16	57
10%	2	6	7	3	7	65
25%	3	10	3	8	15	70
50%	8	6	10	12	10	48
75%	15	15	12	18	26	58
90%	63	66	84	66	60	109

Az eredmények viszonylag egyértelműek: *kointegrált idősorok esetén még meglehetősen nagyarányú hiányzó megfigyelés esetén sem növekszik kezelhetetlen méretűre a fals negatív esetek aránya*, mindössze abban az esetben téved a próba 100-nál többször (10%-ot meghaladó mértékben), ha mindkét idősorból hiányzik az eredeti értékek kilenc-tizede. Láthatjuk, hogy – egyáltalán nem meglepő módon – a korábban tapasztalt *aszimmetria itt nem jelentkezik*, a kointegráltság kölcsönös jellegéből adódóan gyakorlatilag érzéketlenek az eredményeink arra, hogy melyik idősor hiányos. Szintén nem tudok határozott állítást megfogalmazni az interpolációs módszer választása tekintetében, szimulációs eredményeim alapján a lineáris, illetve a harmadfokú spline interpoláció közel azonos erősségű eredményekre vezetett.

Mivel az előbbieken azt tapasztaltuk, hogy az ok-okozati viszonyok megállapíthatósága szempontjából, igazából csak attól kell tartanunk, ha nagy arányban hiányoznak megfigyelt értékek a szimulált idősorokból, ezért igazán érdekes az a kérdés, hogy mennyiben kapunk hasonló eredményeket abban az esetben, ha az egyik idősor frekvenciája szisztematikusan eltér a másiktól. Az alábbiakban azt vizsgálom, hogy amennyiben az egyik idősorom csak harmad-, illetve negyedéves megfigyelést tartalmaz, mint a másik, de a kimaradó értékek szisztematikusan helyezkednek el az időtengelyen, mennyiben csökken az okság, illetve együttmozgás megítélésére kidolgozott próbák ereje.

Elsőként tekintsük a korábban már használt VAR-modellt, és vizsgáljuk meg, hogy 1000 esetből hányszor mutatja helytelenül az okság hiányát a Wald-próba, vélhetően annak következtében, hogy az eredeti, ritkábban megfigyelt idősort interpolációs technikákkal kiegészítem.⁵⁶

⁵⁶ Teljesen világos, hogy a nem azonos frekvenciájú idősorok közötti kapcsolatok vizsgálata megoldható lenne úgy is, hogy a sűrűbben megfigyelt idősor esetén valamilyen alkalmas aggregálást végzünk (összegzés, átlagolás, stb.), ezzel a problémával nem itt, hanem a 6. fejezetben foglalkozom. Értelemszerűen nem vizsgálom itt azt az esetet sem, amikor az egyik idősorból minden harmadik, a másiktól minden negyedik megfigyelés áll rendelkezésre, hiszen ennek empirikus megfelelője nehezen lenne megideologizálható.

3-4. táblázat: Az okság hiányát mutató próbák száma, különböző frekvenciájú idősorok interpolációval történő pótlása esetén

Ok szerepét betöltő y_{1t} idősor megfigyelési gyakorisága		Okozat szerepét betöltő y_{2t} idősor megfigyelési gyakorisága						
		Eredeti	Minden harmadik, melynek pótlása			Minden negyedik, melynek pótlása		
			LIN	CRS	CUB	LIN	CRS	CUB
Eredeti		0	0	0	0	0	0	0
Minden harmadik, melynek pótlása	LIN	0	0	1	0			
	CRS	0	1	0	0			
	CUB	0	2	1	0			
Minden negyedik, melynek pótlása	LIN	0				1	0	0
	CRS	0				0	0	0
	CUB	0				1	1	1

A táblázat alapján kimondható, hogy ettől az esettől *a valóságban nem nagyon kell tartanunk*. Amennyiben a két vizsgált idősor a korábban definiált VAR-folyamatból származik, *az okság kimutatása nem érzékeny* arra, ha az ok, illetve az okozat szerepét betöltő idősor más-más frekvencián áll rendelkezésre. Hangsúlyozni kívánom, hogy igen hosszú megfigyelés-sorozatot szimuláltam, az empirikus idősor-elemzésben viszonylag ritka a 250 év/1000 negyedév, vagy a 333 negyedév/1000 hónap (~83 év) hosszúságú idősor, ezért a szimulációkat elvégeztem eredetileg 100 elemű, tehát a kihagyás után 25, illetve 33 megfigyelést tartalmazó idősorokon is. Az eredmények lényegében nem változtak, vagyis az eredetileg ok-okozati viszonyban levő idősorok esetében a Granger-okságot tesztelő Wald-próba nem produkált érdemi számban fals negatív esetet, vagyis nem „tűnt el” az okság, az interpoláció hatására.

Ezt követően a kointegrált idősorok esetét vizsgáltam, szinte azonos módon a stacioner idősoroknál bemutatottakkal. A 3-5. táblázatban itt is a fals negatív esetek számát látjuk 1000 iteráció esetén, vagyis ennyiszer „fedi el” az interpolációs kiegészítés az amúgy létező együttmozgást.

3-5. táblázat: *A kointegráltság hiányát mutató próbák száma, különböző frekvenciájú idősorok interpolációval történő pótlása esetén*

<i>A DGP alapján kointegrált y_{1t} idősor megfigyelési gyakorisága</i>		<i>A DGP alapján kointegrált y_{2t} idősor megfigyelési gyakorisága</i>						
		<i>Eredeti</i>	<i>Minden harmadik, melynek pótlása</i>			<i>Minden negyedik, melynek pótlása</i>		
			<i>LIN</i>	<i>CRS</i>	<i>CUB</i>	<i>LIN</i>	<i>CRS</i>	<i>CUB</i>
<i>Eredeti</i>		28	32	34	31	50	54	50
<i>Minden harmadik, melynek pótlása</i>	<i>LIN</i>	34	34	83	49			
	<i>CRS</i>	62	64	189	140			
	<i>CUB</i>	50	85	161	181			
<i>Minden negyedik, melynek pótlása</i>	<i>LIN</i>	48				173	74	213
	<i>CRS</i>	54				118	178	112
	<i>CUB</i>	52				179	115	156

Láthatjuk, hogy közös trendet tartalmazó idősorok esetén már óvatosabban kell kezelnünk a kointegrációs tesztek (Johansen-próba) eredményét. *Eltérő frekvenciájú, viszonylag nagy arányban hiányos idősorok esetén gyakran kaphatunk tévesen a kointegráltság hiányára utaló teszt-eredményt*, vagyis a kointegráltság hiányát jelentő nullhipotézis elfogadását sokszor okozhatja a hiányos idősorok interpolálása. Érdemes azt is hangsúlyozni, hogy az *interpolációs technikák közül az egyszerűbb (lineáris) jobban teljesít*, ha viszonylag kevés érték hiányzik tényleges megfigyelés között, de ez a módszertani előny nagyobb lyukak esetén már eltűnik.

Az ok-okozati viszony téves elvetése (fals negatív eset) mellett ugyancsak fontos az a kérdés, hogy vajon nem okoz-e a hiányos idősorok interpolációval történő kiegészítése tévesen okságnak látszó állapotot (fals pozitív eset). Ennek megvizsgálása érdekében elvégeztem a szimulációt „fordított” esetre is, vagyis, két véletlenszerűen generált idősor (a korábbi jelöléssel ε_{1t} és ε_{2t}) elemeiből hagytam ki véletlenszerűen, illetve szisztematikusan megfigyeléseket, majd a rendszertelen idősorokat interpolációval pótoltam és az így létrejött új idősorok között ellenőriztem az okság meglétét. A szimuláció érdekesebb eredményeit tartalmazza a 3-6.a és 3-6.b táblázat.

3-6.a táblázat: A tévesen okságot mutató (fals pozitív) próbák száma rendszertelen idősorok lineáris interpolációval történő pótlása esetén

ϵ_{1t} idősor megfigyelési gyakorisága		ϵ_{2t} idősor megfigyelési gyakorisága					
		Eredeti	Véletlenszerűen kihagyott			Szisztematikusan kihagyott	
			10%	50%	90%	minden harmadik maradt	minden negyedik maradt
Eredeti		56	55	50	37	58	51
Véletlenszerűen kihagyott	10%	56	45	39	37	62	41
	50%	53	48	46	47	52	62
	90%	57	52	45	40	71	51
Szisztematikusan kihagyott	3.	41	47	46	53	157	49
	4.	57	43	51	51	54	183

3-6.b táblázat: A tévesen okságot mutató (fals pozitív) próbák száma rendszertelen idősorok harmadfokú spline interpolációval történő pótlása esetén

ϵ_{1t} idősor megfigyelési gyakorisága		ϵ_{2t} idősor megfigyelési gyakorisága					
		Eredeti	Véletlenszerűen kihagyott			Szisztematikusan kihagyott	
			10%	50%	90%	minden harmadik maradt	minden negyedik maradt
Eredeti		51	44	54	45	62	60
Véletlenszerűen kihagyott	10%	52	47	47	45	58	55
	50%	37	47	39	49	56	64
	90%	49	44	37	54	60	73
Szisztematikusan kihagyott	3.	69	52	55	51	193	70
	4.	51	54	55	52	45	207

Az előbbi táblázatok értékelésénél ismét ne feledkezzünk meg arról, hogy 5%-os szignifikancia-szinten hozzuk a döntést, vagyis mintegy 50 fals pozitív eset keletkezése megfelel a várakozásainknak. Ezt is figyelembe véve, mindkét táblázat, vagyis mindkét interpolációs technika hasonló eredményre vezetett: igazából *nem kell attól tartanunk, hogy az egymástól független idősorok között csak*

azért keletkezik látszat-okság, mert a rendszertelenül megfigyelt idősorokat interpolációval pótoljuk. Kivétel ez alól, ha a szisztematikusan kihagyott értékeket tartalmazó idősorokat pótolunk, majd az interpolált – egyébként eredetileg egymástól független – idősorokra végezzük a Granger-okság tesztjét, ilyenkor sokszor fordulhat elő, hogy tévesen ok-okozati viszonyt állapítunk meg ott is, ahol ez az adatgeneráló folyamatok szerint nem létezik. A szimulációk alapján megállapítható, hogy az interpolációs technikák eredményei között jelentős eltérések nincsenek, tendenciózus különbségek nem mutathatók ki, egyik, vagy másik eljárás javára.

Összefoglalva a nemekvidisztáns idősorokra vonatkozó szimulációs eredményeket a következő lényegi megállapításokat (téziseket) fogalmazhatjuk meg:

1. **A 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-okság 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 okság „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-különbözőséget próbáljuk orvosolni az interpolációs technikával.**

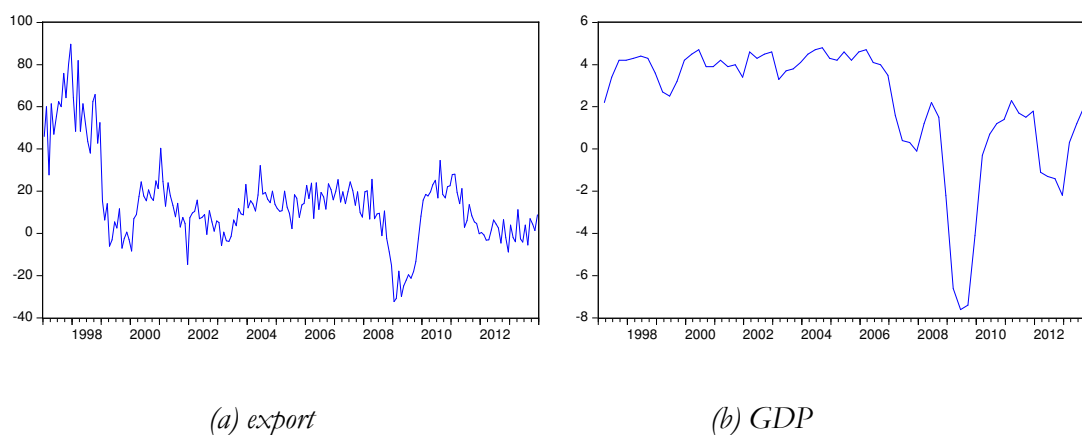
A következő alfejezetben vizsgálom, hogy mennyire relevánsak az előbbieken megfogalmazott állítások valós empirikus idősorok és egy konkrét közgazdasági probléma esetén.

3.4 A GDP és az export növekedése közötti kapcsolat Magyarországon az elmúlt közel két évtizedben

Dolgozatom elején már hangsúlyoztam, hogy az egyes módszereket, illetve anomália-kezeléseket bemutató példák illusztratív jellegűek, ezekben csak kevésbé törekszem makroökonómiai, illetve pénzügytani novumok felismerésére. Nyilvánvalóan igyekeztem olyan jelenségeket, összefüggéseket vizsgálni, melyek nem irrelevánsak, vagyis amelyek modellezésére lehet empirikus kutatásokat találni a szakirodalomban. Téziseimet ugyanakkor nem ezen gazdaságtudományi kérdések megválaszolásaként kívántam megfogalmazni.

A gazdasági növekedés és a külkereskedelmi teljesítmény, elsősorban az export nagysága közötti kapcsolatot, szigorúbb értelemben véve ok-okozati viszonyt számos fórum, és tudományos közlemény vizsgálja. Az ismeretterjesztő gazdasági újságírás esetében elég, ha csak az „exportvezérelt növekedés” kifejezésre utalok, amely valószínűleg napi gyakorisággal jelenik meg a sajtóban. Ennél nyilván meggyőzőbb és módszertanilag korrektebb tanulmányok sora jelent meg a kérdésről a szakirodalomban (alaposabb kutatás nélkül elég, ha csak Xu, 1996, vagy Narayan-Smyth, 2009 cikkeire utalok).

Magyarországon az elmúlt évtizedekben egymást váltó kormányok és nemzetgazdasági elemzők egyet értettek abban, hogy az export dinamikus növekedése jótékony hatással van a gazdaság teljesítményére, amit dolgozatomban szóhasználatával, illetve operacionalizálva úgy is fogalmazhatunk, hogy *az export-változás Granger-oka a GDP-változásnak*. A két jelenség 1997 és 2013 közötti változását az alábbi ábra-pár⁵⁷ mutatja:



3-3. ábra: A vizsgált változók alakulása 1997 – 2013 (előző év azonos időszakának %-ban)

A statisztikai adatgyűjtés, illetve adatszolgáltatás korlátai azt eredményezik, hogy míg az export nagyságát havi rendszerességgel meg tudják figyelni, vagyis módunk van hó/előző év azonos hó típusú idősorokat képezni, addig a GDP közelítése csak negyedéves gyakorisággal valósítható meg, így a bruttó hazai termék volumenváltozásának idősorai ritkábbak. A bemutatandó illusztratív példában 17 év adatai állnak rendelkezésre, ami több mint 200 havi export adatot, illetve 68 negyedéves GDP volumenindexet jelent. Amennyiben a két jelenség közötti oksági viszonyra vagyunk kíváncsiak, két módon járhatunk el

⁵⁷ Ügyeljünk arra, hogy a két ábra eltérő léptékű, míg a GDP-változás esetében a terjedelem mindössze 12 százalékpont, addig az export-változásnál a range közel tízszeres! Mindezzel semmiképpen sem megtévesztetni kívántam az Olvasót, inkább azért választottam ezt a megoldást, mert így viszonylag jól „sejtetik” az ábrák a hasonló lefutású tendenciákat.

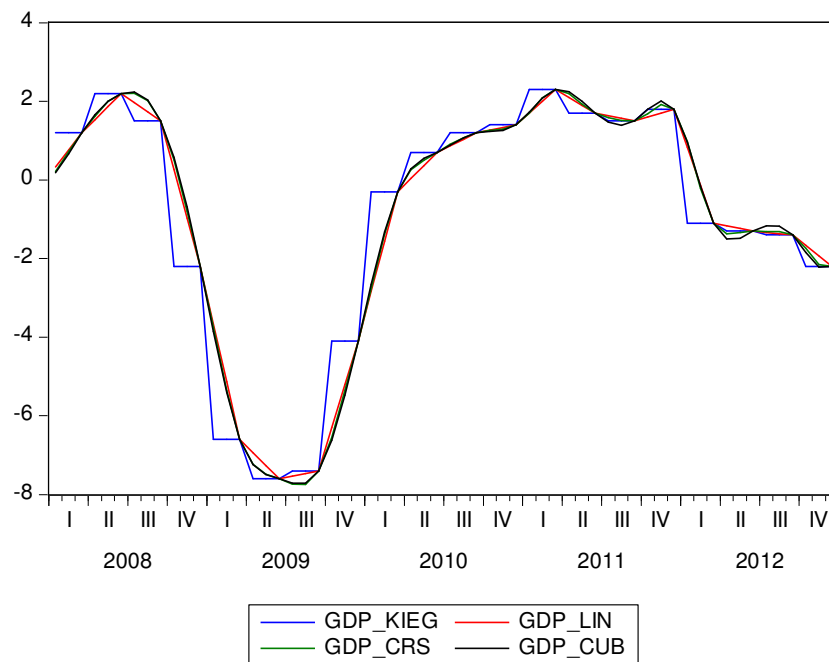
- *három hónaponként aggregálva a havi exportvolument* létrehozuk a negyedéves export-nagyságokat, majd ezekből dinamikus viszonszámot számolva a 72 negyedévre vonatkozó két idősort elemezzük tovább;
- *a GDP negyedéves idősorát felhasználva, interpolációval* képezzük a kiegészített, immár valameny-nyi hónapra értékkel bíró GDP-volumenváltozás idősort és a kéthavi bontású idősor között teszteljük az okságot.

Mindkét megoldás mellett lehet érveket és ellenérveket felhozni. Nyilvánvalóan az első megoldás ellen szól, hogy a sztochasztikus idősor-modellezés eszközszerrendszere, a kedvezőtlen kismintás tulajdonságok kivédése érdekében, általában minimálisan 100 elemű idősorokat kíván meg, ami itt nem teljesül. Ugyanakkor ebben az esetben (aggregált) tényadatokra támaszkodunk, szemben a második esettel, amikor a GDP-változás idősora becsült értékekkel lesz feltöltve.

Dolgozatom logikájából következően a második megoldást választottam, és azt vizsgáltam, hogy okoz-e problémát ebben a konkrét esetben, ha az eltérő frekvenciájú idősorok között úgy keresek ok-okozati összefüggést, ha az egyik idősort interpolációval kiegészítem. Négy megoldást alkalmaztam a negyedéves GDP-volumenindex „sűrítésére”:

- egyszerű kiegészítést, vagyis a negyedévre vonatkozó adatot hozzárendeltem a negyedév mindhárom hónapjához,
- lineáris interpolációt,
- Catmul-Rom spline-t, és
- harmadfokú, másodrendű spline interpolációt.

A négy eljárás eredményeképpen létrejött idősorok kis mértékben, de eltérnek egymástól, amit a 3-4. ábra illusztrál (az ábrán mindössze a válság 2008-2012-es időszakát ábrázoltam, annak érdekében, hogy a rövidebb időszakon plauzibilisebben elváljanak a különböző módszerekkel kiegészített idősorok egymástól).



3-4. ábra: A GDP-volumenindex kiegészített idősorai

A 3-4. ábra jól mutatja az interpolációk és az egyszerű kiegészítés közötti szemléletbeli különbséget, mindemellett az is jól felismerhető, hogy a különböző interpolációs technikákkal kiegészített idősorok sem esnek teljes mértékben egybe. Ezt követően elvégeztem az ok szerepét betöltő export-változás és az okozatnak vélelmezett GDP-változás idősorai között a Granger-okságot tesztelő Wald-próbát, mind a négy kiegészített GDP-időssorral (értelmszerűen a teljes, tehát a 17 évre vonatkozó, havi bontású idősorokon, a VAR-modellekben a havi bontást figyelembe véve 12-ed rendű késleltetésekkel számoltam).

Az eredményeket a 3-7. táblázat tartalmazza:

3-7. táblázat: Granger-okságot tesztelő Wald-próbák eredményei

Az okozatnak feltételezett GDP-változás kiegészítésének módja	Wald-próba	
	F-érték	p-érték
kiegészítés	1,091	0,371
lineáris interpoláció	0,890	0,558
Catmul-Rom spline	0,865	0,584
harmadfokú, másodrendű spline	1,896	0,038

A fenti eredmények önmagukért beszélnek: az export-változás és a GDP-volumenváltozás közötti ok-okozati összefüggés megtalálása akár attól is függhet, hogy milyen interpolációs technikával „sűrítettük” a bruttó hazai termék változására vonatkozó adatokat. Láthatjuk, hogy az egyszerű kiegészítés, a lineáris vagy a CRS-interpoláció esetén a nullhipotézist el kell fogadni, vagyis ilyen-

kor nem sikerült elvetnünk az okság hiányára vonatkozó feltevést, ezzel ellenkezően ok-okozati összefüggést találunk, ha a havi bontású export-változáshoz harmadfokú, másodrendű spline-nal igazítjuk a GDP-változást. Mindez azért is rendkívül elgondolkodtató, mert az egyébiránt mégoly korrekt szakirodalomban is viszonylag gyakran fordul elő, hogy a különböző interpolációs technikával előállított ekvidisztáns idősorok vizsgálata esetén a tanulmány nem tér ki az interpoláció mikéntjére.

Mivel empirikus példám is alátámasztotta, ezért a rendszertelen idősorok kezeléséről írt fejezet zárómondataként megismétlem: elsősorban az idősorok közötti oksági viszony feltárása során kell körültekintően használni a rendszertelen, vagy ritkábban megfigyelt idősor kezelhetősége érdekében alkalmazott interpolációs technikákat.

4 Outlierek az idősorokban

Az outlierek, azaz a kiugró értékek⁵⁸ vizsgálata nem új keletű, ugyanakkor az egyik legnehezebben kezelhető statisztikai probléma a gazdasági idősorok modellezésében. Durbin egyenesen úgy vélte, hogy „*bárki, aki gazdasági idősorok modellezésével foglalkozik, aggódhat az outlier-probléma miatt*”.⁵⁹ A kiugró értékek kezelése a korai statisztikai elemzésekben is gyakran vizsgált probléma volt, az irodalomban hemzsegnek az ezzel kapcsolatos írások, ezek hasznos összefoglalása megtalálható Barnett és Lewis (1994) monográfiájában. A hosszú múlt ugyanakkor nem eredményezte, hogy konszenzus szülessen akár a definícióban, akár a jelenség káros voltának és kezelésének megítélésben.

Abban többé-kevésbé egyetértés mutatkozik a szakirodalomban, hogy a kiugró értékek két alapvető okból is keletkezhetnek, és ezeket eltérően kellene megítélni. Létrejönnek anomáliák mintavételi, vagy adatrögzítési hibákból, sőt néha szándékos csalás következtében. Ezeket, az angol nyelvű forrásokban *gross error*-ként nevesített kiugró értékeket, ebben a dolgozatban nem soroljuk az outlierek közé, a mérési, rögzítési hibákat ki kell (vagy ki kellene) küszöbölni, így az outlier-kérdés megoldhatóvá válik.

Az „igazi” (*true*) outlierek a folyamatokban meglevő természetes eltérések, esetleg (adatfelvételi, vagy gazdasági) rezsim-váltás következményeiként jönnek létre, a megfigyelt érték korrekt, csak valamiért meglepő, a várakozásokból kilógó. Nyilvánvalóan még ezt követően is szükséges valamilyen pontosabb meghatározás arra vonatkozóan, hogy mit értünk természetes eltérésen. Az idősorokban fellelhető outlierek vonatkozásában a legáltalánosabban használatos definíció Hawkinstól származik (Hawkins, 1980), eszerint az *outlier egy olyan megfigyelés, amely annyira különbözik az összes többitől, hogy vélelmezhetően egy másik adatgeneráló folyamatból származik*. Disszertáciomban az előbbi definícióra alapozva vizsgálom a kiugró, kevésbé hihető értékek kiszűrésének, illetve kezelésének módszereit. Ugyanakkor érdemes megjegyeznünk, hogy Hawkins felteszi, hogy

- az outlierek egyváltozós modellekben értelmezettek,
- a vizsgált sztochasztikus folyamat DGP-je „normális”, jól kezelhető, és
- kiugró értékek ritkán keletkeznek.

⁵⁸ Az outlier kifejezésnek a magyarban nincs teljeskörűen elfogadott megfelelője. Darvas (2005) „rendellenes” értékekről ír, én inkább a kiugró érték kifejezést használom, el akarom ugyanis kerülni, hogy úgy tűnjön, az outlier megjelenése az idősorban valamilyen a modellező által elkövetett hiba következménye. Mundruczó (1998) az outlier szinonímájaként használja a „kilógó”, illetve a „befolyásos” adat kifejezést, de ezek sem váltak általánosan elfogadottá.

⁵⁹ Durbin (1979) 339. pp. Amúgy a probléma a keresztmetszeti adatállományokban éppen ilyen nagy, vagy néha nagyobb gondot okoz (pl. egy-egy kiugróan nagy vállalat, vagy extrém nagy ország); a keresztmetszeti adatállományokból kilógó adatok kezelését ugyancsak régóta kutatja az irodalom.

Természetesen ennél jóval komplexebb megközelítések is léteznek, melyek az outlier (vagy outlierek) megjelenését sokkal több szempontból vizsgálják. Kítűnő gondolatokat olvashatunk Kriegel et al. (2010) tanulmányában, ahol a szerzők az outliereket sokkal szofisztikáltabban osztályozzák, így különbséget tesznek azon esetek között, ahol

- *globális*, vagy *lokális* kiugró érték szűrése a feladat,
- *megjelölni* (megtalálni), vagy „csak” *valószínűsíteni* akarjuk az outliert,
- *minta*-, vagy *modell*-szempontból tekintünk kiugrónak egy értéket.

Az első alternatíva azért fontos, hiszen ha globális kiugró értéket keresünk, akkor feltételezzük, hogy csak egy DGP létezik, és az esetleges többi outlier csupán torzítja ezt a feltételezett adatgeneráló folyamatot. Ehhez a gondolatsorhoz kívánczok az ún. elfedési-probléma (*masking problem*), amelyben egy kiugró érték csak akkor tűnik anomáliának, ha egy másik (pl. a globális megközelítésben outliernek tartott) megfigyelést kiveszünk az idősorból.⁶⁰ Abban az esetben, ha nem a globálisan értelmezett outlier megtalálása a cél, hanem engedélyezünk több, lokális outliert, akkor a probléma úgy kezelendő, mintha az idősorunk több darabból állna, akár mindegyik szakaszra más DGP is lehet jellemző, és ezen részekben keresünk kiugró értéket. Nem könnyű megválaszolni a kérdést, hogy egy adott tapasztalati idősort hány részre oszthatunk (praktikusan ez ekvivalens azzal, hogy hány outliert keresünk), a szerzők érzékenységvizsgálatokat javasolnak, de nem oldják meg megnyugtatóan a problémát.

Érdekes kérdés az is, hogy mi a célunk az outlier-kereséssel: egyértelműen meg kívánjuk jelölni az általunk kiugrónak tartott megfigyelést (kiugró értéket produkáló időpontot), vagy egy folytonos valószínűséget akarunk hozzárendelni minden időponthoz, amely megmutatja, hogy milyen eséllyel tekinthetünk egy értéket outliernek. Az előbbi céllal működő eljárások végül egy bináris változót eredményeznek, melyek két kimenetele a normális érték (0), illetve az outlier (1), ez a változó a probléma kezelése során jól használható. Amennyiben *outlierség-valószínűséget* rendelünk az egyes megfigyelésekhez általában akkor is végül egy eldöntendő kérdést akarunk megválaszolni, vagyis valamilyen küszöbérték (cut-point) felhasználásával ilyenkor is bináris változó lesz az eredmény. A valószínűség hozzárendelése inkább a modellezési lehetőségek számát növeli, hiszen a küszöbérték változtatásával több, illetve kevesebb érték sorolódik az outlierek közé.

Az outlierek megtalálása, szükség esetén kiszűrése szempontjából rendkívül fontos azt a kérdést eldönteni, hogy milyen megközelítés alapján tekintünk egy kiugró értéket outliernek. Három, gondolatvilágban teljesen különböző megfontolás is elképzelhető:

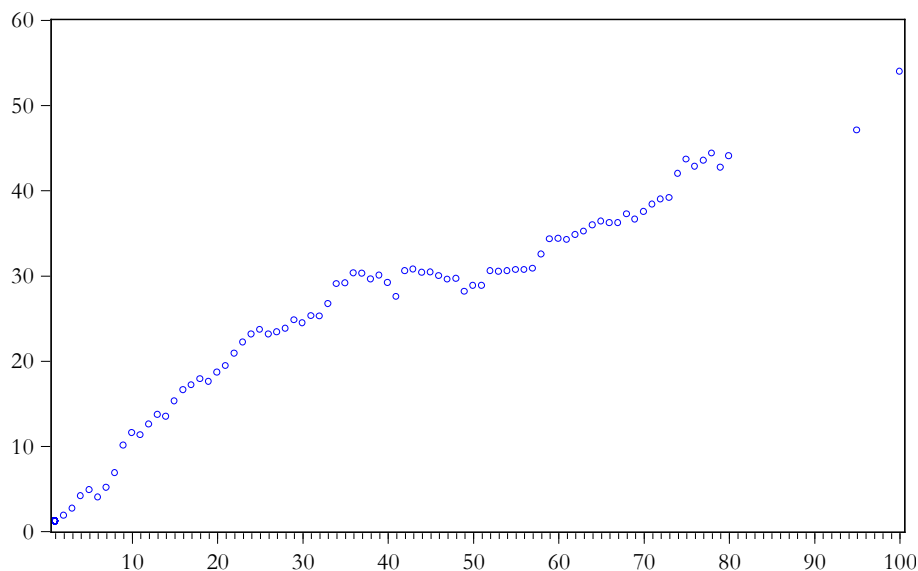
- *modellre alapozó*, vagyis illeszteni kell egy modellt, majd a nem megfelelően illeszkedő empirikus értéket outliernek tekintjük (könnyen átlátható, hogy az eljárás túlságosan kitett an-

⁶⁰ A probléma „ellentettjét” is szokás vizsgálni: az ún. mocsár-effektus (*swamping-problem*) akkor lép fel, amikor egy elem csak akkor nem látszik outliernek, ha egy másik elem megjelenik az adatok között.

nak, hogy milyen korrekt modellt specifikálunk, mennyire hatásos becslési módszert választunk, stb., vagyis meglehetősen érzékeny a modellező döntéseire);

- *távolság alapú*, amikor egy pont (érték) és a szomszédjának (szomszédjainak) a távolságát vizsgálva jutunk a végkövetkeztetéshez;
- különösen idősorokban keresett outlierok esetén kedvelt a *homogenitásra (simaságra) alapozó* megközelítés, melynek geometriai interpretációja szerint össze kell kötni minden pontpárt, és megkeresni azt a megfigyelést, amelynél a leginkább ingadozik a szakaszok által bezárt szög.

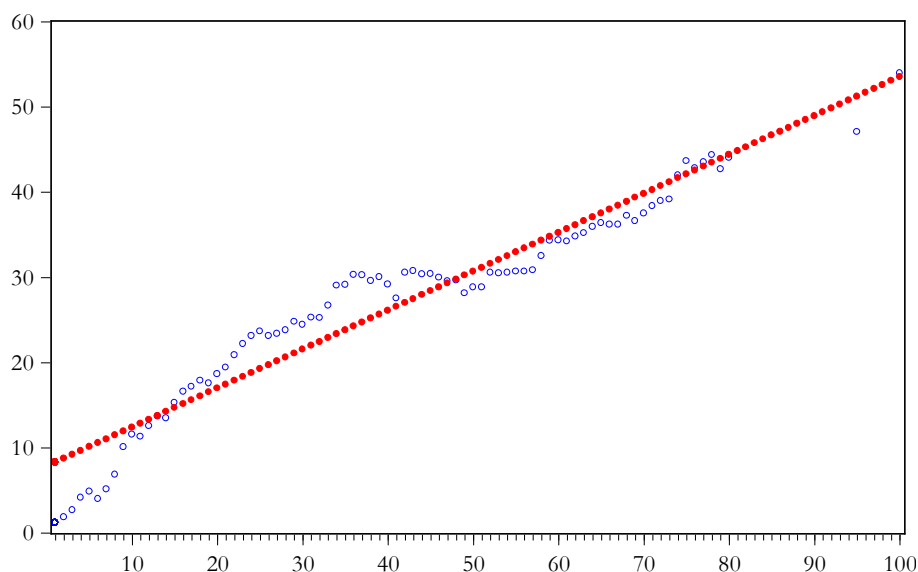
Fontos megjegyezni, hogy az első esetben a modellválasztás alapvetően befolyásolhatja az outlier megítélését, míg a másik két eset modell-független módon keresi a kiugró értéket. Mindennek illusztrálására tekintsük a következő, szimulált adatokat tartalmazó ábrát!



4-1. ábra: Outlierek az idősorban

A 4-1. ábrán látható utolsó két megfigyelés megítélése nagymértékben függ attól, hogy az előbbiekben taglalt megfontolások melyikét alkalmazzuk. Ha ugyanis a pontok többi ponttól mért távolságát, vagy a környezetükben levő pontok sűrűségét vizsgáljuk (minta alapú megközelítés), akkor feltétlenül outliernek tűnnek.

Ugyanakkor, ha az idősorra egy roppant egyszerű modellt, pl. egy lineáris (determinisztikus) trendet illesztünk, akkor az alábbi ábrát nyerjük:



4-1a. ábra: Megfigyelt időszori értékek és lineáris trendjük

A 4-1.a ábra alapján, ha egy megfigyelést annak alapján sorolunk az outlierok közé, hogy az empirikus idősorra illesztett modellhez milyen közel van (mekkora az adott időponthoz tartozó reziduum), akkor az utolsó megfigyelésünk semmiképpen sem tűnik kiugrónak, ugyanakkor az időszak közepe táján több adat is „gyanús”.

Tekintsük át tehát azokat a statisztikai eljárásokat, amelyekkel a kiugró értékek, anomáliák viszonylag egzakt módon megtalálhatók, majd ezt követően vegyük sorra az outlierok kezelésének módszertanát!

4.1 Az outlierok típusai

Outlierek számos módon megjelenhetnek a vizsgált idősorokban. A kiugró értékek *általában* 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 eredeti cikkben Fox *I.* és *II. típusú* outliereket definiál, és vizsgálatát az autoregresszív modellek esetére korlátozza. Dolgozatomban nem szűkítem le a modelleket erre az osztályra, ugyanakkor az outlierok tipizálása során semmilyen mértékben sem foglalkozom (nem használom ki), hogy a vizsgált idősor eredeti adatgeneráló folyamata milyen.

Legyen általános modellünk⁶¹ a következő:

⁶¹ A (4.1) modell annyiban nem általános, hogy mindössze egyetlen (a t időpontban bekövetkező) outliert tételez fel, de mindjárt látni fogjuk, hogy ez az egyszerűsítés könnyen feloldható.

$$y_t = z_t + f(t) \quad (4.1)$$

ahol z_t bármilyen szokásos ARMA-modell⁶², vagyis

$$\varphi(L)z_t = \theta(L)\varepsilon_t$$

Az outlier az $f(t)$ tagon keresztül kerül a modellbe, hatásmechanizmusa a (4.2) modell-felíráson jobban látszik:

$$y_t = \frac{\theta(L)}{\varphi(L)}\varepsilon_t + \omega_0 \frac{\omega(L)}{\delta(L)}I_t(d) \quad (4.2)$$

ahol $\omega(L)$ és $\delta(L)$ szintén késleltetési polinomok, ω_0 az outlier nagysága (*magnitude*), $I_t(d)$ pedig egy dummy változó, ami $I_t(d) = 1$ értéket vesz fel, ha $t = d$ és $I_t(d) = 0$ egyébként. Látható, hogy a modell implicit módon feltételezi, hogy az outlier helyét (időpontját) ismerjük (ezt jelzi d értéke). A késleltetési polinomok típusa határozza meg az outlier korábban említett két típusát.

Az első outlier típus (Fox ezt nevezte I. típusúnak) gyakrabban szerepel a témával foglalkozó irodalomban. Feltételezésünk szerint ebben az esetben 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. Additív outlier esetén

$$\frac{\omega(L)}{\delta(L)} = 1,$$

tehát modell praktikusán a

$$y_t = z_t + \omega_{0t}$$

$$\omega_{0t} = \begin{cases} 0 & t \neq d \\ \omega_0 & t = d \end{cases} \quad (4.3)$$

⁶² Alapmodellünkben sem konstant, sem determinisztikus trendet nem szerepeltetünk, mindez a mondanivalót nem befolyásolja. Értelemszerűen, amennyiben az idősor nemstacionárius, akkor a megfelelő differenciája, illetve stationerré tett transzformáltja szerepel itt.

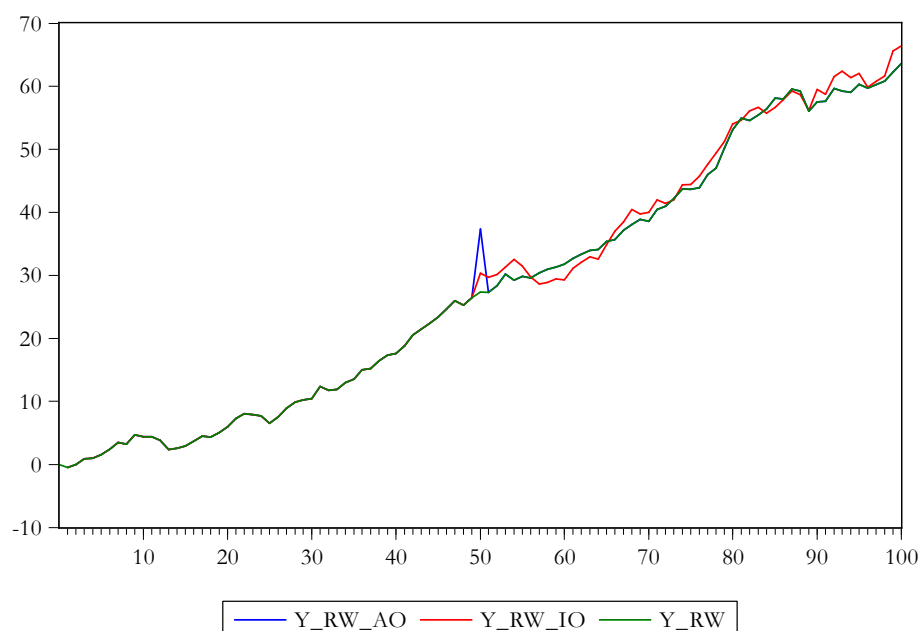
formára egyszerűsödik.

Ezzel éles ellentétben áll Fox II. típusú outliere. Az ún. *tovagyűrűző outlier* (innovational outlier, továbbiakban *IO*) több megfigyelés értékét is befolyásolja.⁶³ Ha az előbbi modell-keretbe kívánjuk mindezt beépíteni, akkor a kiugró (meglepő) értéknél jelentkező, nem várt eltérés nem az ARMA-taghoz adódik hozzá, hanem a modell véletlen (ϵ_t) változójába épül be, és így hatással van valamennyi $t \geq d$ megfigyelésre. Ebben az esetben a késleltetési polinomok hányadosára – a szokásokon túl – tehát nem teszünk megkötést, így az IO-t tartalmazó modell felírható az alábbi formában:

$$y_t = z_t + \omega_{0t}$$

$$\omega_{0t} = \begin{cases} 0 & t < d \\ \omega_0 \frac{\omega(L)}{\delta(L)} & t \geq d \end{cases} \quad (4.4)$$

Egy véletlen bolyongás folyamatban (Y_RW) megjelenő AO, illetve IO típusokat illusztrálja a 4-2. ábra:



4-2. ábra: Különböző típusú (AO, IO) outlierek hatása az eredeti idősor alakulására

Láthatjuk, hogy szimulált – egyébiránt sztochasztikus trendet tartalmazó – idősorunkat a középső megfigyelésnél éri valamilyen extrém hatás, az additív outliert tartalmazó modell mindössze abban

⁶³ Mák (2011) alapvetően strukturális törésekkel foglalkozó cikkében (melyre a következő fejezetben még visszatérek) az innovatív outlierek hatását csillapodónak nevezte, ami a hatás vonatkozásában plauzibilis elnevezés, viszont az outlier maga nyilván nem lehet csillapodó, ezért használom a tovagűrűző elnevezést.

az időpontban tér el az eredeti értékektől, a tovagyűrűző hatású IO-t tartalmazó modellben ettől kezdve mindvégig érződik az egyszeri kiugró érték hatása.

Fox úttörő cikkében négy modell-feltevés mellett vizsgálja az outliereket:

1. csak additív outlier tartalmazó,
2. csak tovagyűrűző outliereket beépítő
3. hol ilyen, hol olyan, de egy megfigyelésnél csak egy típusú outlierrel terhelt, valamint
4. kevert

modelleket eleméz. A kiugró értékekkel foglalkozó alapirodalom⁶⁴ empirikus vizsgálatait azt mutatják, hogy amennyiben egy idősorban több outlier is kimutatható, akkor általában mindkét típusú zavaró hatással számolnunk kell, vagyis a reális feltételezés a Fox által említett negyedik típus, még akkor is, ha ettől a modellek akár nagyon bonyolultakká is válhatnak.

Tsay a téma szempontjából mérföldkőnek számító cikke (Tsay, 1988) további három outlier típust definiál, ezen típusok az előbb bemutatott esetek olyan kiterjesztései, melyben az outlier hatása modell-szinten nem csak egy időpillanatban jelentkezik. A szerző szerint elképzelhető *szinteltolódás* (level shift), *időszakosan ható* (temporary change), illetve a *varianciára ható* (variance change) outliereket. Wu és társai tanulmányukban (Wu et al., 1993) azt az esetet is taglalják, amikor több outlier terheli az idősort, ám ezek együttes hatása zéró. Számos, elsősorban az utóbbi évtizedben született műhelytanulmányban az outlierok GARCH-folyamatokra gyakorolt hatását vizsgálják.⁶⁵ Sajnos a magyar szakirodalom viszonylag szegényes, a témával módszertani szempontból foglalkozó cikkek közül említést érdemel Csereháti (2004) tanulmánya.

Dolgozatomban, ahol az kívánom megvizsgálni, hogy milyen mértékben képes elfedni egy, vagy több outlier az idősorokban meglevő trendet, vagy az idősorok között kimutatható ok-okozati összefüggést (esetleg együttmozgást) az előbbieken részletesebben tárgyalt AO, illetve IO típusokon túlmutató modellekre nem lesz szükség.

4.2 Az outlier-szűrés leggyakoribb módszerei

Az előbbi részben, az outlierok tipizálása során már kitértem arra, hogy egy megfigyelést kiugró értéknek minősíthetünk azért, mert az empirikus idősorból „kilóg”, de azért is, mert a vélelmezett adatgeneráló folyamatból kisebb valószínűséggel származik, más szavakkal: az idősorra illesztett modell által becsült értéktől távol esik. Mindkét eset feltételezi, hogy valamilyen módon megfigyelt értékek távolságát számszerűsítjük, és ebből következtetünk az outlier megjelenésére.

⁶⁴ Lásd pl. Tsay (1986), Ljung (1993), Balke-Fomby (1994).

⁶⁵ Lásd pl. Doornik-Ooms (2005), vagy Ardelean (2012).

A következőkben az outlierok megtalálására irányuló módszerek két nagy csoportját tekintem át, vizsgálom a

- modell-független, illetve a
- reziduális változón alapuló

outlier-felfedő módszereket (*outlier labeling methods*). A bemutatott eljárások kiválasztása nyilvánvalóan szubjektív, a szakirodalom nagyszámú további módszert ismertet, ezek bemutatása megtalálható a korábban már hivatkozott alapművekben.⁶⁶ Az alfejezetet néhány, az outlierok feltételezett hatásaira vonatkozó megállapítás zárja, itt térek ki röviden a szűrési eljárások szubjektív jellegére is.

4.2.1 Modell-független outlier-szűrés

A kiugró értékek megtalálására vonatkozó *modell-független* (értsd sztochasztikus időszori modellt nem használó) *eljárások* sem tekinthetők minden feltétel nélkül működőknek. Alapgondolatuk ugyanis feltételez legalább egy eloszlást (leggyakrabban nyilvánvalóan normális eloszlást), és az ennek alapján várt értékektől mért eltérések alapján szűrik (detektálják) az outlierokat.

A következőkben bemutatandó módszerek leírása, érzékenység-vizsgálata, továbbfejlesztése számos tanulmányban megtalálható. Az eljárások alapgondolata először Tukey (1977) könyvében olvasható, majd a folyamatosan burjánzó továbbfejlesztések korrekt összefoglalása megtalálható Seo (2006), illetve Tolvi (1998) tanulmányában.

A legelső outlier-szűrési technikák az ismert 2, illetve 3 szigma-szabályra építenek (elnevezésük *SD-módszer*). A Csebisev-egyenlőtlenség alapján köztudott, hogy egy μ várható értékkel és σ^2 varianciával jellemezhető véletlen változóra (y), valamennyi $k > 0$ szám esetén igaz, hogy

$$\begin{aligned} \Pr\{|y - \mu| \geq k\sigma\} &\leq \frac{1}{k^2} \\ \Pr\{|y - \mu| < k\sigma\} &\geq 1 - \frac{1}{k^2} \end{aligned} \tag{4.5}$$

Alkalmasan választott k esetén a $\mu \pm k\sigma$ intervallumon kívül eső értékek (megfigyelések) outliernek tekinthetők. Minden különösebb magyarázat nélkül belátható, hogy a módszer meglehetősen robusztus, ugyanakkor az outlier-szűrés előtt egy viszonylag kényes feladat a modellezőre vár, meg kell határoznia a kiugró értékek „kívánatos” arányát.

⁶⁶ Aguinis et al. (2013) nem kevesebb, mint 39 módszert ismertet, mellyel az outlierok megtalálhatók. A tanulmány másképp csoportosítja a kiugró értékek felismerésének technikáit, megkülönbözteti a változó momentumain, a távolságon, illetve a hatáson alapuló outlierokat, de a felsorolt módszerek nagyjából megfeleltethetők az itt tárgyaltakkal.

Az előbbiekkal teljes mértékben ekvivalens, ugyanakkor normális eloszlású véletlen változót feltételező eljárás a *z-score módszer*. Ennek végrehajtása során képezzük az eredeti – vélelmezhetően normális eloszlásból származó - értékek transzformáltjait

$$z_t = \frac{y_t - \bar{y}}{\sigma} \quad (4.6)$$

és az így keletkezett, standardizált értékekre kell kijelölni a 2-es, vagy 3-as határt.⁶⁷ Általánosságban érdemes megjegyezni, hogy

$$z_{\max} = \frac{T-1}{\sqrt{T}}$$

ami kis mintánál nem eredményezne outliert, ugyanakkor az általunk vizsgált több 100 elemű idősoroknál mindez már nem okoz problémát. Annál inkább nehézséget okoz, hogy a *z-score módszer* az empirikus időszori átlag, illetve szórás használatából következően *erősen érzékeny* a kiugró értékekre, aminek következtében az outlier gyakran elfedődik.

Az extrém értékek esetében is használható eljárásként javasolja Iglewicz-Hoaglin (1993) a mediánon alapuló szűrést. Képezzük a

$$M_t = \frac{0,6745(y_t - Me)}{\text{median}(|y_t - Me|)} \quad (4.7)$$

transzformált értékeket, ahol Me az idősor mediánja⁶⁸ és $\text{median}(\cdot)$ az adott argumentum mediánját jelöli. A szerzők küszöbértékként – némiképp szubjektív módon – a $\pm 3,5$ értékpárt jelölik meg.

Hasonló elven képezhető⁶⁹ az ún. MAD_E -*módszer* transzformált változója is

$$MAD_E = 1,483 \times \text{median}(|y_t - Me|) \quad (4.8)$$

ebben az esetben a küszöbértékek a Csebisev-egyenlőtlenségre emlékeztetnek

$$Me \pm k \times MAD_E \quad (4.9)$$

⁶⁷ Az eljárás végrehajtását nem befolyásolja, de mintaként továbbra is y_t idősort feltételezünk, ahol $t = 1, 2, \dots, T$.

⁶⁸ Ügyeljünk rá, hogy a medián értelemszerűen a nagyságát tekintve a középső érték, és nem a középső megfigyelés!

⁶⁹ A módszereket jól és tömören foglalja össze Olewuezi (2011).

A helyzeti középérték, illetve általánosságban a kvantilisok használata visszavezet Tukey eredetileg javasolt, box-plot-on alapuló módszeréhez (Tukey, 1977). Ebben az eljárásban képeznünk kell a megfigyelések alsó, illetve felső kvartiliséit (Q_1 , illetve Q_3), valamint ezek különbségeként az interkvartilis terjedelem mutatóját (IQR). Előbbiek felhasználásával meghatározható egy ún. *belső* korlát

$$\{Q_1 - 1,5 \times IQR; Q_3 + 1,5 \times IQR\} \quad (4.10a)$$

illetve egy ún. *külső* korlát:

$$\{Q_1 - 3 \times IQR; Q_3 + 3 \times IQR\} \quad (4.10b)$$

Tukey szerint egy megfigyelés *lehetséges* (*possible*) *outlier*, ha a belső és külső korlát között van. A külső korlátokon kívüli értékeket a szerző *valószínű* (*probable*) *outlier*nek tekinti. Megjegyzendő, hogy az alkalmazott 1,5, illetve 3 kritikus értékek magyarázataként nem feltétlenül találunk statisztikus megfontolásokat. Az eljárás érzékeny a ferdeségre, így *csak szimmetrikus eloszlások esetén* javasolt.

Az előbbi probléma kiküszöbölése érdekében fejlesztették ki Brys és szerzőtársai (lásd Brys et al., 2004) a módosított Tukey-eljárást (amelyet néhány irodalom *adjusted boxplot* néven ismeri). Rendezzük a mintaelemeket növekvő sorrendbe, jelölje az így képzett elemeket $y_{(i)}$, vagyis $y_{(1)} \leq y_{(2)} \leq \dots \leq y_{(T)}$. Ekkor az outlier detektálásához alkalmazandó transzformáció

$$MC = \text{median} \left(\frac{(y_{(j)} - Me) - (Me - y_{(i)})}{y_{(j)} - y_{(i)}} \right), \quad i < j \quad (4.11)$$

és az outlier-szűréshez használt intervallum

$$[L, U] = \begin{cases} Q_1 - 1,5e^{-3,5MC} IQR; Q_3 + 1,5e^{4MC} IQR, & \text{ha } MC \geq 0 \\ Q_1 - 1,5e^{-4MC} IQR; Q_3 + 1,5e^{3,5MC} IQR, & \text{ha } MC < 0 \end{cases} \quad (4.12)$$

Belátható, hogy $-1 \leq MC \leq 1$ és szimmetrikus eloszlások esetén $MC = 0$, ilyenkor egyébiránt az outlier-szűrés esetén használt határok megegyeznek a Tukey által korábban megadottal.

A bemutatott egyváltozós eljárások meglehetősen rapszodikusak abban a tekintetben, hogy az idősor mekkora hányadát tekintik outliernek. Ferde eloszlások esetén egyértelműen $2 \times MAD_E$ -

módszer szűri ki a legtöbb kiugró értéket, természetesen azt állítani, hogy ezáltal ez a legjobban teljesítő módszer, nem lenne megalapozott!

Dolgozatom szempontjából nem számít főcsapásnak, de érdemes megemlíteni, hogy többdimenziós eloszlások (többváltozós adatállományok) esetén a *megfigyelések távolságára alapozó módszereket* érdemes használni, az outlierok szűrésére. A leggyakrabban alkalmazott eljárás az általánosított Mahalanobis-távolságra épít (eredetiben lásd Mahalanobis, 1936).

Az outlier-szűrés szabálya szerint kiugró értéknek kell tekinteni azokat a megfigyeléseket (időpontokat), melyekre érvényes, hogy

$$MDist(\mathbf{y}, \boldsymbol{\mu}) = (\mathbf{y} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1} (\mathbf{y} - \boldsymbol{\mu}) > \chi^2(0,975)_m \quad (4.13)$$

ahol $\mathbf{y}$ az adott időponthoz tartozó időszori értékek vektora, $\boldsymbol{\mu}$ az idősorok átlagából képzett vektor, $\boldsymbol{\Sigma}^{-1}$ az idősorok kovarianca-mátrixának inverze, m az adatállományban levő idősorok száma. A módszer számos problémát vet fel:

- nagy változóság esetén alig talál (praktikusan nagyon sokszor nem talál) outliert,
- nem robusztus,
- végrehajtásához ismerni kell(ene) az együttes eloszlást,
- nem stacionárius idősorok esetén a kovariancia-mátrix torzított.

A problémák ellenére számos további távolságon alapuló eljárás ismeretes, melyek igyekeznek kiküszöbölni a fenti problémákat. Knorr-Ng (1997) alapgondolata nem a megfigyelés és az átlagos érték, hanem a *megfigyelés és a szomszédjainak a távolságán* alapul (*distance-based outlier*). Vételezzük, hogy a normális (nem kiugró) pontok körül nagy sűrűségben vannak értékek, az outlierok körül viszont nincsenek. A módszer értelmében, ha definiálunk egy ε sugarat és egy π százalékot (arányt), akkor a y_i megfigyelés akkor outlier, ha a többi megfigyelés kevesebb, mint π százaléka van közelebb hozzá, mint ε . A küszöbértékek meghatározásának szubjektív (és sokszor ad hoc jellegű) voltára ismét érdemes felhívni a figyelmet.

Összességében megállapítható, hogy a modell-független outlier-szűrési technikák rengeteg vitatható elemet tartalmaznak, ráadásul – ezt idáig nem hangsúlyoztam – kidolgozóik nem feltétlenül idősoros környezetben gondolkodtak, vagyis nem használták (használhatták) ki az adatok rendezettsége kínálta lehetőségeket. Ezután talán már nem meglepő, hogy a sztochasztikus folyamatok (idősor-modellezés) esetén hatékonyan alkalmazható detektálási módszerek általában valamilyen modellfeltevésre építenek.

4.2.2 Modell-alapú outlier-szűrés

Az outlierok modellre alapozó szűrésének is több módja ismeretes. Eleve más-más módszert alkalmazunk, ha egy, vagy több outliert tételezünk, sőt akkor is más eljárással dolgozunk, ha ismerjük a kiugró értékek helyét, illetve ha nem. Ne feledjük, hogy a detektálás szükséges feltétele, hogy jól specifikált, a tényleges adatgeneráló folyamatot leíró modellel dolgozzunk! A szakirodalomban tárgyalt modell alapú outlier-szűrés eljárások általában ARMA, illetve dinamikus regressziós (ARMAX) modellekkel operálnak. Noha számos módszer ismert, ezek nem, illetve nem nagyon térnek el egymástól, valamint bizonyos kapcsolatban állnak az előző fejezetben tárgyalt hiányzó adat, vagy nem ekvidisztáns idősorok problémájával is. A modell alapú (regresszióra építő) outlier szűrés eljárások hasznos összefoglalója olvasható Ledolter (1990), illetve Ljung (1993) cikkeiben.

A következőkben előbb két eljárást mutatok be, melyek széles körben használtak a kiugró értékek keresése során, ezek

- a *likelihood-arány teszt*, illetve
- az *érzékenység-vizsgálatra alapozó eljárás*.

Harmadikként aztán bemutatok egy, fiatal kollégáimmal (Kehl Dániellel és Abaligeti Gallusszal) közösen kifejlesztett⁷⁰ saját módszert, amely dummy változók modellbe építését követően a *reziduális szórás minimalizálására épít*.

A likelihood-arány tesztre építő outlier-szűrés alapgondolata Fox (1972) már hivatkozott cikkéből származik. Az eljárás lényege, hogy minden megfigyeléshez rendelünk hozzá egy dummy változót, határozzuk meg az összes így képezhető likelihood-ot, és keressük meg ezek maximumát. Amelyik modellenél 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.

Tekintsük a korábban már vizsgált modelleket! Legyen az eredeti idősorból ARMA-szűrővel képzett idősor

$$Y_t = \frac{\varphi(L)}{\theta(L)} y_t \quad (4.14)$$

és a $t = d$ helyen ω_0 nagyságú outliert feltételező idősor

⁷⁰ A módszert bemutattuk a Gazdaságmodellezési Társaság 2012-es éves Szakértői Konferenciáján (lásd <http://www.gazdasagmodellezes.hu/images/stories/konferenciak/rappaigkehldabaligetig2012.pdf>), illetve kollégáim prezentálták a Bernoulli Society 2013-as éves gyűlésén (http://ems2013.eu/conf/upload/BEK086_006.pdf).

$$X_t = \frac{\varphi(L)}{\theta(L)} \frac{\omega(L)}{\delta(L)} x_t \quad (4.15)$$

ahol feltesszük, hogy

$$Y_t = \omega_0 X_t + \varepsilon_t \quad (4.16)$$

(4.16) egy egyszerű lineáris regresszió, melynek paraméterei akár OLS-sel becsülhetők, a $\hat{\omega}_0$ -ra vonatkozó becslőfüggvény a szokásos módon alakul:

$$\begin{aligned} \hat{\omega}_0 &= \frac{\sum Y_t X_t}{\sum X_t^2} \\ \text{Var}(\hat{\omega}_0) &= \frac{\sigma_\varepsilon^2}{\sum X_t^2} \end{aligned} \quad (4.17)$$

Mindezek után az outlier hatása már becsülhető. Az alapgondolat, hogy

- elsőként illesszük a vélelmezetten helyes ARMA-modellt az idősorra, majd az eredeti idősor és a modell különbségeként keletkeztessük a becsült reziduumokat,
- ezt követően az eltérés-négyzetösszegek segítségével képezhetők az LR-tesztek.

Fontos megjegyezni, hogy a két outlier típus (AO, illetve IO) esetén más-más LR-teszt adódik. Az egyszerűbb eset a tovagyrúzó hatású outlier (IO). A korábbi modellfeltevéseinkből következően itt X_t értéke megegyezik az indikátor (dummy) változóval, vagyis egyedül a $t = d$ helyen 1, minden más esetben 0. Ekkor a (4.17)-ben bemutatott becslés roppant egyszerű lesz

$$\begin{aligned} \hat{\omega}_{I,d} &= Y_d \\ \text{Var}(\hat{\omega}_{I,d}) &= \sigma_\varepsilon^2 \end{aligned} \quad (4.18)$$

Ezek után teszteljük a nullhipotézist, miszerint a $t = d$ helyen outlier van, az itt nincs outlier ellenhipotézissel szemben. A likelihood-arány teszt empirikus próbafüggvény értéke

$$\lambda_{I,d} = \frac{\hat{\omega}_{I,d}}{\sigma_\varepsilon} \quad (4.19)$$

lesz, ami a nullhipotézis teljesülése esetén aszimptotikusan standard normális eloszlást⁷¹ követ, feltéve, hogy az ARMA-modell helyes.

Némileg bonyolultabb az additív outlier esete, mivel itt az „inverz-szűrés” nem eredményez olyan szép eredményeket. Az outlier nagyságára, illetve annak varianciájára vonatkozó becslőfüggvények:

$$\begin{aligned}\hat{\omega}_{A,d} &= \rho_{A,d}^2 \left(Y_d - \sum_{t=1}^{T-d} \pi_t Y_{t+d} \right) \\ \text{Var}(\hat{\omega}_{A,d}) &= \rho_{A,d}^2 \sigma_\varepsilon^2\end{aligned}\tag{4.20}$$

alakúak, ahol a π_t értékek az ARMA-szűrőből, tehát a $\frac{\varphi(L)}{\theta(L)}$ kifejezésből eredő súlyok és

$$\rho_{A,d}^2 = (1 + \pi_1^2 + \dots + \pi_{T-d}^2)^{-1}$$

Az LR-próba empirikus értéke

$$\lambda_{A,d} = \frac{\hat{\omega}_{A,d}}{\rho_{A,d} \sigma_\varepsilon}\tag{4.21}$$

melynek eloszlása szintén aszimptotikusan standard normális.

Abban az esetben, ha a kiugró érték helyét nem tudjuk előre (értsd: nincs róla preconcepciónk), akkor a (4.19), illetve (4.21) próbafüggvény-értékeket ki kell számolni valamennyi időpont esetében, és ahol az empirikus értékek maximuma található, ott van legnagyobb eséllyel az outlier. (Több outlier feltételezése, illetve megengedése esetén értelemeszerűen kell eljárunk.)

Az outlier-szűrés előbb bemutatott LR-próbára alapozó módozata a legelterjedtebb, mondhatni ez a standard módszer, a legelterjedtebb szoftverek is ezt használják. Mindez valószínűleg annak köszönhető, hogy alapgondolata egyszerűen átlátható, és akár iteratív úton is használható (erre mutat példát Tsay, 1988). 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!

⁷¹ Szokatlan lehet, hogy likelihood-arány tesztnél határeloszlásként nem χ^2 -eloszlás adódik, de vegyük észre, hogy itt egy 1 szabadságfokú χ^2 -eloszlású valószínűségi változó gyökéről, vagyis tényleg standard normális eloszlásról van szó!

Egészen más intuícióból keletkezett az outlier-szűrés érzékenység-vizsgálaton alapuló módozata. Könnyen belátható, hogy egy outlier (kiugró érték) jelentősen meg tudja változtatni az adatgeneráló folyamat feltárására irányuló becslésünket.⁷² 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

$$I(F, \beta(F), x) = \lim_{\varepsilon \rightarrow 0} \frac{\beta((1-\varepsilon)F + \varepsilon\delta_x) - \beta(F)}{\varepsilon} \quad (4.22)$$

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 (4.22) azt vizsgálja, hogy *ennyiben változnak az eloszlás becsült paraméterei, ha az x megfigyelés értéke kis mértékben megváltozik*. Anélkül, hogy különösen sokat kívánnék foglalkozni a hatás-függvény tulajdonságaival, könnyen belátható, hogy amennyiben a (4.22) képletben δ értékét 0-nak definiáljuk, akkor a hatás-függvény azt méri, hogy mennyiben változik a becslésünk attól, hogy az x megfigyelést elhagyjuk.

Az előző, rendszertelen idősorok modellezésére vonatkozó fejezet fejtegetései alapján nem meglepő, hogy két módon is vizsgálhatjuk mekkora hatást gyakorol a becslésünkre, ha feltesszük, hogy a $t = d$ helyen outlier van:

- y_d értékét kihagyva az idősorból, ezáltal nemekvidisztáns idősort előállítva, és azt modellezve, vagy
- y_d értékét alkalmas módon helyettesítve (például a két szomszédos érték átlagával).

Peña (1987) tanulmányában azt javasolja, hogy az egyes megfigyelések hatását legegyszerűbben a következő mutatóval mérhetjük:

$$P_d = \frac{(\hat{y} - \hat{y}_{(d)})^T (\hat{y} - \hat{y}_{(d)})}{k\hat{\sigma}_\varepsilon^2} \quad (4.23)$$

⁷² A fejezet második részében pont ezt fogom vizsgálni különböző szimulációk segítségével.

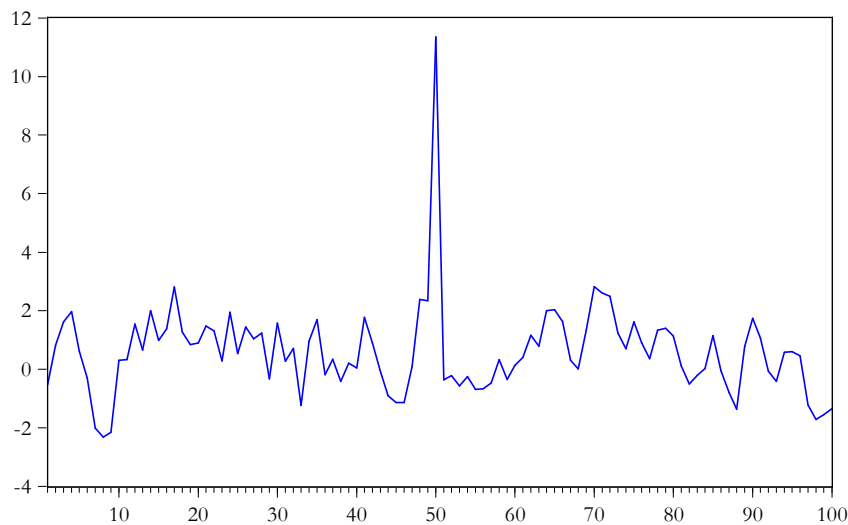
ahol P_d a $t = d$ megfigyelés hatása, $\hat{\mathbf{y}}$ a teljes idősor figyelembe vételével becsült időszori értékek, $\hat{\mathbf{y}}_{(d)}$ az időszorra vonatkozó becslés, amennyiben a $t = d$ megfigyelést elhagyjuk, k a modellben található becsülendő paraméterek száma, $\hat{\sigma}_\varepsilon^2$ a becsült reziduális szórásnégyzet. Nem kapunk ugyan pontos iránymutatást arra nézve, hogy mikor tekintünk egy hatást elég nagynak, de a (4.23) mutató feltétlenül alkalmas az egyes megfigyelések hatásainak összevetésére, esetleg a hatások rangsorolására.

A legelterjedtebb outlier-szűrő eljárás nem az idősor becsült értékeire, hanem a becsült paraméterekre gyakorolt hatás mérésére alapoz. Képezzük a következő egyszerű mértéket:

$$I_{j(d)} = \frac{\hat{\beta}_j - \hat{\beta}_{j(d)}}{\hat{\sigma}_{\varepsilon(d)}^2 \sqrt{\text{var}(\hat{\beta}_j)}} \quad (4.24)$$

ahol $\hat{\beta}_j$ a teljes idősor felhasználásával becsült paraméter, $\text{var}(\hat{\beta}_j)$ ennek varianciája, $\hat{\beta}_{j(d)}$ a $t = d$ megfigyelés elhagyásával becsült paraméter, $\hat{\sigma}_{\varepsilon(d)}^2$ a megfigyelés elhagyása után becsült modell reziduális varianciája, $I_{j(d)}$ pedig a $t = d$ megfigyelés elhagyása által okozott hatás. Jól látható, hogy amennyiben a modellünk egynél több becsülendő paramétert tartalmaz, akkor annyi ilyen hatás-mérték képezhető, amennyi a paraméterek száma.

Tekintsünk minderre egy illusztratív példát! A 4-3. ábra egy AR1 folyamatból ($\varphi = 0,7$, ε_t a szokásos *Nid* fehér zaj) származó 100 elemű idősort mutat, amelynek 50. megfigyelése $\omega_0 = 10$ nagyságú outlierrel terhelt:



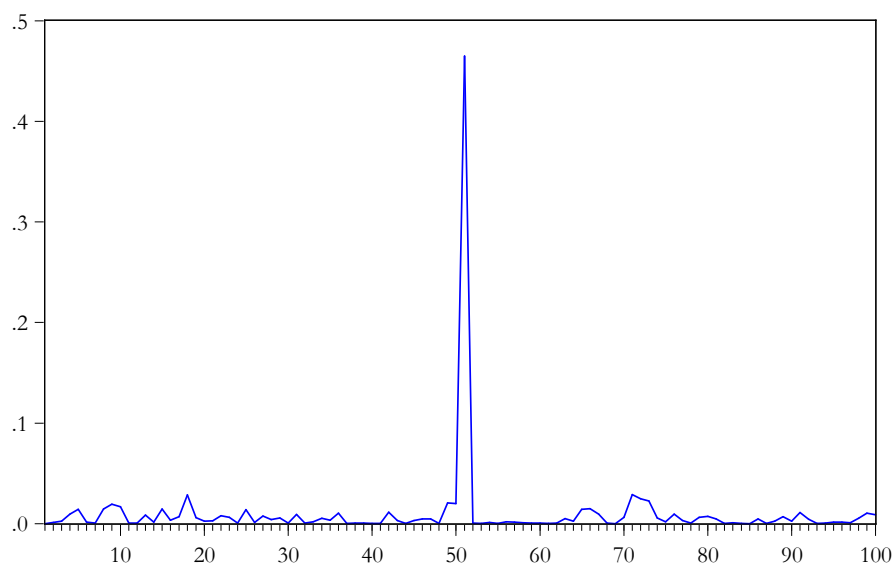
4-3. ábra: AR1 folyamat outlierrel

Illesztettem a fenti folyamatra egy AR1 modellt, a paraméterbecslés eredménye:

$$y_t = 0,433 y_{t-1} + \varepsilon_t$$

$$\sigma_\varepsilon = 1,514$$

Láthatjuk, hogy az outlier jelentősen eltérítette a becsült paramétereket az adatgeneráló folyamatban szereplő értékeket. Vizsgáljuk meg a (4.24) képletben bemutatott indikátor értékét a teljes időhorizonton:



4-4. ábra: A paraméter megváltozását mérő hatás-függvény

A 4-4. ábra jól illusztrálja, hogy az outlier azonosítása a bemutatott módszerrel viszonylag hatékonyan megoldható.

Az előzőekben bemutatott két modellalapú eljárás hiányosságaként említhetjük, hogy egy lépcsőben mindössze egy outlier megtalálását célozzák, amennyiben több kiugró értéket vélelmezünk, akkor több iterációra lehet szükség, azaz sor kerülhet a módszer többszöri ismételtesére. Részben ennek a problémának a megoldására fejlesztettünk ki egy eljárást, amely egy lépcsőben jelöli meg valamennyi outliernek tűnő időszori értéket, illetve időpontot. (Nem térek ki részletesen a kérdésre, de jelzem, hogy a korábban említett elfedési probléma így *in time* kezelhető.)

Induljunk ki a szokásos ARMA-modellből⁷³ és tételezzük fel, hogy megbecsültük a legkisebb eltérés-négyzetösszeget eredményező paramétereket. Formalizálva

⁷³ 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 idősoros regressziós modell alkalmazható, de ez az eljárás is érzékeny a félrespecifikált, azaz rosszul illeszkedő modellekre!

$$\begin{aligned} \frac{L(\varphi)}{L(\theta)} y_t &= \varepsilon_t \\ \sum_{t=1}^T \varepsilon_t^2 &\rightarrow \min \end{aligned} \quad (4.25)$$

ahol $\varepsilon_t \sim iid$ fehér zaj.

Keressük azt a D_t bináris dummy változót (vagyis $D_t \in 0;1$), melyre

$$\sum_{t=1}^T (\varepsilon_t - \alpha D_t)^2 \rightarrow \min \quad (4.26)$$

azaz, keressük azt a kétkimenetelű változót, amelyet additív módon, konstanssal szorozva az eredeti ARMA-modellbe illesztve a modell eltérés-négyzetösszege a legnagyobb mértékben csökkenthető. Intuíciónk szerint azok az időszori értékek illeszkednek a legkevésbé az eredeti idősor adatgeneráló folyamatába, ahol a bináris dummy változó 1-es értéket vesz fel.

Fejtsük ki a (4.26) kifejezésben szereplő négyzetes tagot!

$$\sum_{t=1}^T (\varepsilon_t - \alpha D_t)^2 = \sum_{t=1}^T \varepsilon_t^2 - 2\alpha \sum_{t=1}^T \varepsilon_t D_t + \alpha^2 \sum_{t=1}^T D_t^2 \quad (4.26a)$$

Vegyük észre, hogy mivel (4.25) meghatározta azt a modellt, melynél $\sum_{t=1}^T \varepsilon_t^2 \rightarrow \min$, ezért a (4.26a) kifejezés minimumhelye szempontjából ez indifferens, tehát a feladatunk nem más, mint

$$\alpha^2 \sum_{t=1}^T D_t^2 - 2\alpha \sum_{t=1}^T \varepsilon_t D_t \rightarrow \min \quad (4.27)$$

megkeresése, amely feladat az α paraméterre kényelmesen megoldható:

$$\frac{\partial \left[\alpha^2 \sum_{i=1}^T D_i^2 - 2\alpha \sum_{i=1}^T \varepsilon_i D_i \right]}{\partial \alpha} = 2\alpha \sum_{i=1}^T D_i^2 - 2 \sum_{i=1}^T \varepsilon_i D_i = 0$$

$$\alpha = \frac{\sum_{i=1}^T \varepsilon_i D_i}{\sum_{i=1}^T D_i^2} \quad (4.28)$$

ami – figyelembe véve, hogy D_i értéke csak 0 és 1 lehet – nem más, mint a $D_i = 1$ értékekhez tartozó reziduumok egyszerű számtani átlaga.

Rendezzük ε_i reziduumot növekvő sorrendbe, jelölje a rendezett reziduumot $\varepsilon_{(i)}$! Mivel az (4.27) szélsőérték-feladat ε_i -ben lineáris, így könnyen belátható, hogy az 1-es értéket felvevő dummy változó értékek vagy a rendezett reziduum legalacsonyabb (triviális, de érdemes megjegyezni, hogy ezek a legnagyobb abszolút értékű negatív reziduumok!), vagy legmagasabb értékeinél találhatók.

Bontsuk két részre a feladatot!

1. Legyen a megoldást eredményező D_i értéke 1, a legalacsonyabb (legnagyobb abszolút értékű negatív) k reziduumnál, ekkor a minimalizálandó (4.27) kifejezés

$$k\alpha^2 - 2\alpha \sum_{(i)=1}^k \varepsilon_{(i)},$$

ahova behelyettesítve az α -ra (4.28)-ban kapott megoldást kapjuk a minimum értékét:

$$k\alpha^2 - 2\alpha \sum_{(i)=1}^k \varepsilon_{(i)} = k \frac{\left(\sum_{(i)=1}^k \varepsilon_{(i)} \right)^2}{k^2} - \frac{2 \left(\sum_{(i)=1}^k \varepsilon_{(i)} \right)^2}{k} = - \frac{\left(\sum_{(i)=1}^k \varepsilon_{(i)} \right)^2}{k}$$

Látható, hogy az első k reziduum mindegyike negatív érték lesz, ugyanis amennyiben megtörténik az előjelváltás a fenti kifejezés értéke növekedni kezd.

Ugyanakkor tudjuk, hogy a $k+1$ -edik legkisebb reziduum szerepeltetése a kifejezésben már növeli annak értékét (távolodunk a minimumtól), vagyis

$$\begin{aligned}
-\frac{\left(\sum_{(t)=1}^k \varepsilon_{(t)}\right)^2}{k} &\leq -\frac{\left(\sum_{(t)=1}^{k+1} \varepsilon_{(t)}\right)^2}{(k+1)} \\
\frac{\left(\sum_{(t)=1}^k \varepsilon_{(t)}\right)^2}{k} &\geq \frac{\left(\sum_{(t)=1}^k \varepsilon_{(t)}\right)^2 + 2\varepsilon_{(k+1)} \sum_{(t)=1}^k \varepsilon_{(t)} + \varepsilon_{(k+1)}^2}{k+1} \\
\frac{\left(\sum_{(t)=1}^k \varepsilon_{(t)}\right)^2}{2\varepsilon_{(k+1)} \sum_{(t)=1}^k \varepsilon_{(t)} + \varepsilon_{(k+1)}^2} &\geq k
\end{aligned}$$

A számítás egyszerűen elvégezhető, így a keresett k érték könnyen megtalálható.

2. Hasonló a megoldás abban az esetben is, amikor a megoldást eredményező D_i értékek az m legnagyobb (pozitív) reziduumoknál keletkeznek. A minimalizálandó kifejezés

$$m\alpha^2 - 2\alpha \sum_{(t)=T-m+1}^T \varepsilon_{(t)},$$

ahova behelyettesítve az α -ra (4.28)-ban kapott megoldást, ismét kényelmesen adódik a minimum értéke

$$m\alpha^2 - 2\alpha \sum_{(t)=T-m+1}^T \varepsilon_{(t)} = -\frac{\left(\sum_{(t)=T-m+1}^T \varepsilon_{(t)}\right)^2}{m}$$

Hasonlóan az előbb alkalmazott gondolatmenethez m értékére az alábbi korlát adódik.

$$\frac{\left(\sum_{(t)=T-m+1}^T \varepsilon_{(t)}\right)^2}{2\varepsilon_{(T-m)} \sum_{(t)=T-m+1}^T \varepsilon_{(t)} + \varepsilon_{(T-m)}^2} \geq m$$

A (4.26) szélsőérték-feladatban keresett D_i dummy változó, tehát

- a legkisebb (legnagyobb abszolút értékű negatív) k reziduumnál 1, amennyiben

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

- a legnagyobb m reziduumnál 1, amennyiben

$$\frac{\left(\sum_{(t)=1}^k \varepsilon_{(t)}\right)^2}{k} < \frac{\left(\sum_{(t)=T-m+1}^m \varepsilon_{(t)}\right)^2}{m}$$

Az eljárás egyébként – a szélsőérték-feladat additív jellegéből adódóan – kézenfekvő felfogható 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 (4.26) szélsőérték-feladat az alábbiak szerint módosul:

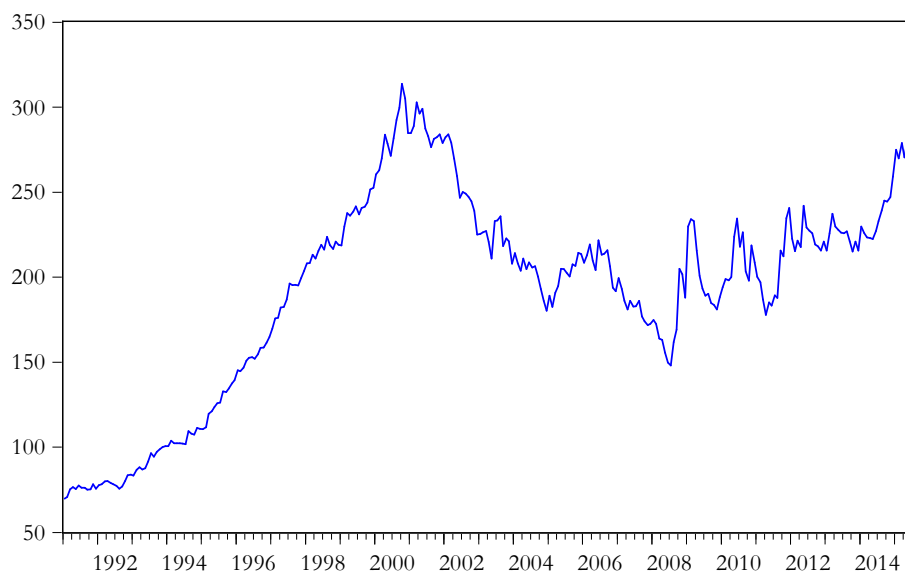
$$\sum_{t=1}^T (\varepsilon_t + \alpha D_{1t} - \beta D_{2t})^2 \rightarrow \min \quad (4.29)$$

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

$$O_t = D_{2t} - D_{1t} \quad (4.30)$$

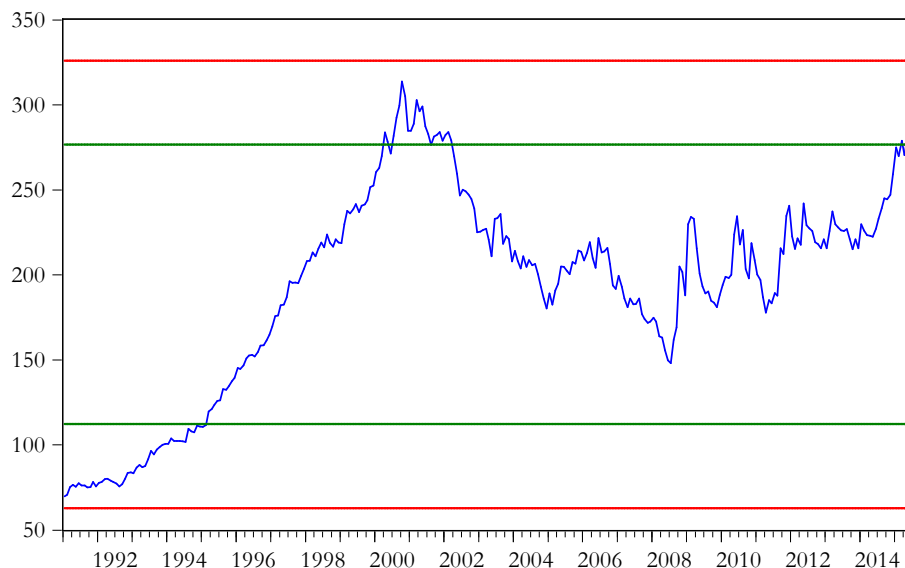
a továbbiakban az *idősr nyomának* nevezzük, és felhasználjuk az idősorban található outlier detektálás során. A (4.30)-ban meghatározott O_t idősr -1 értékénél lefelé kiugró értékről, +1 értékénél felfelé kiugró értékről beszélhetünk, ahol pedig az idősr 0 értéket vesz, ott outlier-mentes időszakot vélelmezünk.

Tekintsünk egy egyszerű példát a fent bemutatott eljárással! Az amerikai dollár (USD) forint-árfolyamának 1991. január és 2015. június közötti idősorát mutatja a 4-5. ábra:



4-5. ábra: Az USD/Ft árfolyam 1991. január – 2015. június (Ft)

Az ábrán is jól láthatjuk, hogy az árfolyam nem stacionárius⁷⁴, így a modell-független outlier szűrésektől nem várhatunk túl sokat. Csupán az illusztráció kedvéért bejelöltem az előbbi ábrán a széles körben alkalmazott Tukey-módszerrel (lásd (4.10a) és (4.10b) képlet) meghatározott belső, illetve külső korlátot:



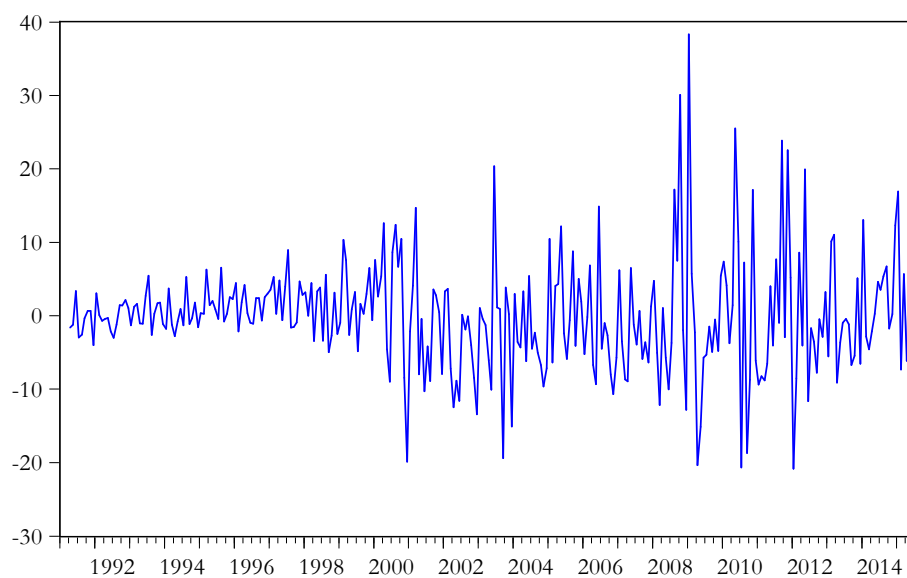
4-6. ábra: Az USD/Ft árfolyam és a vélt outlierok

Jól látható, hogy a könnyen felismerhető nemstacionaritás okán a lefelé kiugró értékek az időszak elején (alacsony árfolyamok), míg a felfelé kilógó értékek a 2000-es évek elejének és a vizsgált

⁷⁴ Az ADF-próba szignifikancia-értéke 0,778, vagyis az egységgyök léte nem vethető el.

időszak legvégének extrém magas árfolyamai között találhatók. Minden mélyebb fejtegetés nélkül is érzékelhető, hogy a kiugró értékek gondolatvilágával ellenkezik, ha egy-egy tömbben találjuk az outliereket!

Meghatároztam az árfolyamra legjobban illeszkedő Box-Jenkins modellt, az AIC kritérium alapján⁷⁵ az $ARIMA(2,1,2)$ tűnt a legtakarékosabb specifikációnak, így ezt használva keletkezett a reziduumok idősora:



4-7. ábra: Az USD árfolyamra illesztett $ARIMA(2,1,2)$ reziduumai (Ft)

A reziduumokra elvégezve az előbb bemutatott eljárást az USD árfolyam outlierai meghatározhatók. A módszer segítségével például az utolsó 24 hónapról (2013. július – 2015. június) az alábbi táblázat állítható össze:

⁷⁵ Az optimális ARIMA-identifikációról lásd pl. Rappai (2013).

4-1. táblázat: Outlierek az USD/Ft árfolyamban

Időszak	Árfolyam (Ft)		
	lefelé kiugró	normál	felfelé kiugró
2013. július		225,79	
2013. augusztus		226,92	
2013. szeptember	221,06		
2013. október	215,06		
2013. november		220,99	
2013. december	215,67		
2014. január			229,8
2014. február		226,03	
2014. március	223,38		
2014. április		223,11	
2014. május		222,40	
2014. június		227,13	
2014. július		232,95	
2014. augusztus		239,21	
2014. szeptember		245,13	
2014. október		244,50	
2014. november		247,21	
2014. december			259,13
2015. január			274,91
2015. február	269,94		
2015. március		278,94	
2015. április	270,37		
2015. május			282,35
2015. június		282,75	

Az eredmények önmagukért beszélnek: például a 2014-es év viszonylagos nyugalomát decemberben, majd 2015 januárjában drasztikus árfolyam-növekedés (az árfolyam-idősor szempontjából felfelé kiugró értékek) zavarta meg, majd a februári korrekciót követően a 2015-ös év elején hektikus ingadozásokat tapasztalhattunk. Érdeemes azt is észrevenni, hogy nem feltétlenül az árfolyam abszolút nagysága, hanem az árfolyam-érték modellbe (adatgeneráló folyamatba) illeszkedése határozza meg, hogy egy adott megfigyelés outlier, vagy sem. Így fordulhat elő, hogy pl. 2014 március-áprilisában két szinte azonos árfolyamérték közül a módszer az egyiket (a „túl gyorsan esőt”) outliernek, míg a másikat normál értéknek tekinti.

Az outlier-szűrésre kifejlesztett módszerek egyik nagy problémája, és sajnos ez alól az általunk javasolt módszer sem kivétel, hogy elsősorban a trendet tartalmazó idősorokban meglehetősen gyakran „jeleznek”, ezáltal a modellezőnek sokszor az az érése, hogy szinte nincs is a jelenséget generáló folyamatból származó megfigyelés. A kiugró értékek kezelésével a következő pontban foglalkozom.

4.3 Outlierrel terhelt idősorok kezelése

Az előző alfejezetekben bemutattam az outlierok definiálásának, illetve kiszűrésének leggyakoribb módszereit, ám nem került szóba, hogy mit lehet tenni olyankor, ha kiderül, hogy a modellezni kívánt idősor outlierrel, vagy akár több outlierrel terhelt.

A szakirodalom az outlierok kezelése tekintetében szinte végtelen megoldást kínál, sajnos viszonylag kevés olyan forrás található, amely felvállalja, hogy az ezek közötti választás dilemmáiba is elmélyedjen. A leggyakrabban használatos technikákat áttekintve nagyjából három⁷⁶ fő kezelési irány kristályosodik ki:

- az outliernek minősülő érték elhagyása, és az ezáltal nemekvidisztánssá vált idősor modellezése (v.ö. a 3. fejezetben tárgyalt módszerekkel, illetve megfontolásokkal),
- a kiszűrt kiugró érték helyettesítése 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 becslési eljárás alkalmazása.

Ennek a dolgozatnak a kereteit és a fókuszát messze meghaladná, ha az outlierok kezelésére kidolgozott technikákat részletesen bemutatnám, így csak egy rövid, és vélhetőleg szubjektív felsorolásra vállalkozom, hogy milyen megoldások képzelhetők el, ha felismerjük az outliert a modellezni kívánt idősorban. (Az outlier végleges elhagyásával nem foglalkozom, ezt ugyanis – amikor nem hibás adatról van szó – kifejezetten károsnak tartom. Megítélésem szerint ugyanis valamilyen empirikus idősori érték fontos információt hordoz, bármelyik – ráadásul szubjektív – elhagyása információvesztés, ami a modellezési bizonytalanságot növeli, a modellező korrektségébe vetett bizalmaz pedig csökkenti!)

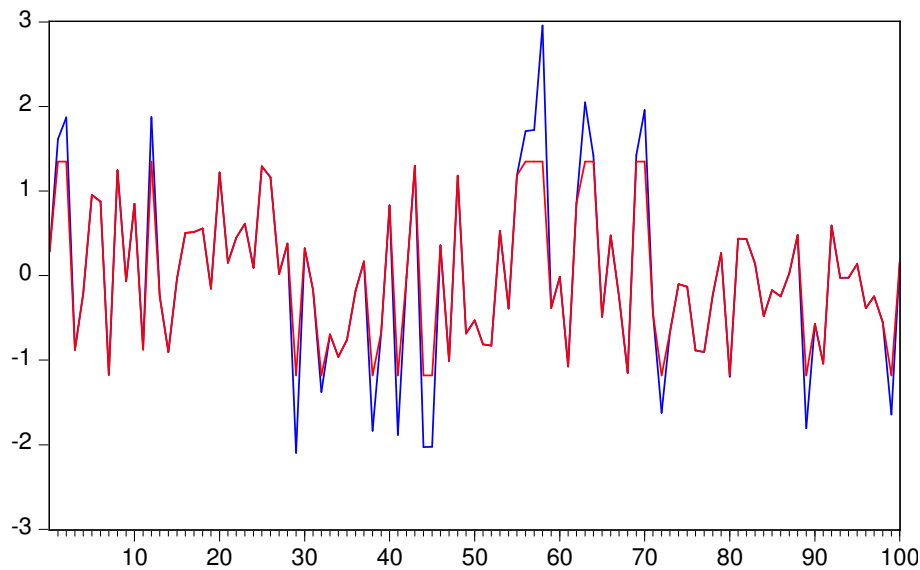
Az egyik leggyakrabban alkalmazott módszer az adatok *winsorizációja*⁷⁷, lényege, hogy a kiugró értékeket az idősor egy alkalmasan választott kvantilisével helyettesítjük, azaz

$$y_t^{W_\alpha} = \begin{cases} y_t^{(\alpha)}, & \text{ha } y_t < y_t^{(\alpha)} \\ y_t & \text{egyébként} \\ y_t^{(1-\alpha)}, & \text{ha } y_t > y_t^{(1-\alpha)} \end{cases} \quad (4.31)$$

⁷⁶ Természetesen most csak arról esik szó, amikor az outlier nem elütés, hibás adatfelvétel, stb. eredménye, hiszen ilyenkor a megoldás nyilvánvalóan a hibás adat hibátlanra cserélése...

⁷⁷ A winsorising (winsorisation) kifejezésnek nincs bevett magyar megfelelője. A módszer a nevét első proponálójától, a XX. század első felében élt és alkotott C. P. Winsorról kapta.

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. Egy 100 elemű fehér zaj idősorán illusztrálva mindezt:



4-8. ábra: Fehér zaj és a winszorizált idősora ($\alpha=0,1$)

Láthatjuk, hogy az outlier-szűrés eredményeképpen kisebb ingadozást mutató idősorhoz jutottunk, hiszen a kiugró értékeket a 10., illetve 90. percentilissel helyettesítettük. (Érdeemes megjegyeznünk, hogy az adatgeneráló folyamat szimmetrikus jellegéből adódóan a transzformált idősor átlaga – várható értékében – azonos az eredeti folyamat átlagával, ugyanakkor a szórása nyilvánvalóan csökken.)

Amennyiben a megoldást nem abban keressük, hogy a megtalált kiugró értéket elhagyjuk, vagy helyettesítjük, akkor számos modellbecslési eljárás ismert, melyekkel az outlier(ek) hatása csökkenthető. Korántsem a teljességre törekedve néhány példa:

- *legkisebb abszolút eltérés (least absolute deviation, LAD)* módszere, az OLS-hez hasonló módszer, de a négyzetre emelés elmaradása miatt az outlier nem „húzza” annyira magához az illesztett görbét (a módszer előnyeiről, illetve hátrányairól lásd pl. Chen et al., 2008),
- *csökkentett legkisebb négyzetek (least trimmed squares, LTS)* módszere, a paraméterbecslés során a legnagyobb, illetve legkisebb reziduumokat elhagyva keressük a négyzetösszeg minimumát (outlier szűrésben való alkalmazásáról ír Nguyen-Welsch, 2010),
- *M-becslések*, robusztus becslések egy ismert osztálya, a reziduumok négyzet-összegét a reziduumok valamilyen más függvényével helyettesítjük, és azt minimalizáljuk (idősor-modellezésben viszonylag gyakori, ha tovagyrúzó (IO) outliert detektáltunk),
- *kétfokozatú robusztus eljárás*, ahol a megfigyelt adatokat a modelltől mért Mahalanobis-távolságuk reciprokával súlyozzuk, vagyis az outliereket kisebb súllyal vesszük figyelembe

(a robusztus statisztikákat számos kézikönyv és tanulmánykötet tárgyalja, az egyik ilyenből származik az ezt az eljárást is bemutató tanulmány, lásd Yohai et al., 1991).

Az említetteken túl még számos nemparaméteres eljárás ismert, terjedőben vannak a minta másodlagos hasznosításán alapuló módszerek (pl. bootstrap), illetve kitűnő terepet lát az outlierrel terhelt adatok modellezésében a bayes-i statisztika is.⁷⁸

A következő alfejezetben, a már ismert szimulációs technikával mindebből csak azt fogom megvizsgálni, hogy milyen hatással van az idősor alapvető tulajdonságainak megítélésére, ha nem vesszük észre az outliert, illetve ha drasztikus szűréssel (winsorizációval) próbáljuk eliminálni a problémát.

4.4 Outlierek, illetve outlier-szűrés hatása az idősorok tulajdonságainak megítélésében

Az alfejezet első részében azt vizsgáljuk, hogy mi történik a próbákkal, ha az ismert tulajdonságú adatgeneráló folyamatból származó idősorokat outlierekkel „szennyezzük”, a második alpont arról szól, hogy tarthatunk-e félrespecifikálástól, ha túlzottan erőltetjük az outlier-szűrést.

4.4.1 Outlier megjelenésének hatása az adatgeneráló folyamat felismerésére

A következőkben bemutatandó szimulációk logikája viszonylag egyszerű: először a már korábban ismertetett módon létrehozuk az ismert adatgeneráló folyamatból származó idősort, illetve idősorokat; ezt követően az idősorokba véletlen helyen outliereket helyezünk el; majd a szokásos tesztekkel megvizsgáljuk, hogy az eredeti adatgeneráló folyamatok felismerhetők-e. A szimulációs eredmények így egyrésztől különbözni fognak abban a tekintetben, hogy

- elsőrendű autoregresszív folyamatot,
- véletlen bolyongást,
- VAR-modellből származó két idősort,
- kointegrált két idősort

tételezünk fel, másrésztől abban, hogy az idősort (idősorokat) milyen jellegű, gyakoriságú és méretű (méretű) outlierrel szennyezzük. Az utóbbi tulajdonságokat tekintve az 1000 elemű szimulált idősorokban feltételezünk

- IO, vagy AO típusú outliert, amely
- az 1000 megfigyelésből véletlenszerűen kiválasztott 25, 50, 100 adatnál jelenik meg és
- nagysága átlag ± 3 vagy 5, vagy 10 szórás nagyságú.

⁷⁸ Az alkalmazható módszerekről alapos felsorolást tartalmaz a már említett Aguinis et al. (2013) cikk.

Mindemellett megvizsgáltam, hogy milyen lesz az outlier hatása, ha a VAR-modellnél, illetve kointegrációnál csak az egyik, illetve mindkét idősor szennyezett.

Akárcsak a rendszertelen idősorok interpolációval történő kiegészítése esetén, az outlierekkel szennyezett folyamatoknál is elmondható, hogy az ADF-próba nagyon erős az AR1 folyamatoknál, még az idősor 10%-át érintő, nagymértékű (a szórás tízszeresét kitevő) outlierok esetén sem kaptam egyetlen esetben sem fals pozitív esetet, vagyis az elégségesen hosszú stacionárius idősorok esetén az outlierok megjelenése nem eredményezte az adatgeneráló folyamat félrediagnosticszálását. Ezért az AR1 folyamatokra vonatkozó szimulációs eredményeket nem rendezem a szokásos táblázatos formába.

Tekintsük ezek után az eltolásos véletlen bolyongásra vonatkozó szimulációs eredményeket, ahol – akárcsak korábban – valamennyi paraméterkombinációból 1000 esetet futtattam! A 4-2. táblázat az ADF-próba által 5%-os szignifikancia-szinten hibásan stacionáriusnak (egységgyököt nem tartalmazónak) minősített esetek számát mutatja.

4-2. táblázat: Az ADF-próba által hibásan stacionárius idősortnak minősített esetek száma eltolásos véletlen bolyongás esetén

Outlier típusa	Outlierek nagysága	Outlierek aránya (%)		
		2,5	5,0	10,0
AO	$\bar{y} \pm 3\sigma$	59	59	82
	$\bar{y} \pm 5\sigma$	64	82	89
	$\bar{y} \pm 10\sigma$	97	108	145
IO	$\bar{y} \pm 3\sigma$	80	111	180
	$\bar{y} \pm 5\sigma$	119	254	395
	$\bar{y} \pm 10\sigma$	414	618	833

A szimulációs eredményekből egyértelmű és teljesen logikus tendencia olvasható ki: *minél több és minél nagyobb mértékű az idősort szennyező outlier, annál gyakrabban fordul elő, hogy az ADF-próba fals negatív eredményt ad*, vagyis az outlierok következtében eltűnik a sztochasztikus trend. Szembeötlő, hogy a *tovagyűrűző (IO) outlierok sokkal nagyobb gondot okoznak*, ebben az esetben már a viszonylag kisszámú, nem túl nagy kiugró érték is képes jelentősen megzavarni a DGP feltárását.

Következő vizsgálatom a VAR-modellből származó idősorok között meglevő Granger-okság elfedhetőségére irányult, vagyis azt vizsgáltam, hogy milyen gyakran mutatja az okság hiányát az egyébként egymással oksági kapcsolatban álló változók esetén hibásan a Wald-próba. Feltevésem szerint mindkét változót szennyezhetik kiugró értékek, ezeknek azonos a nagyságrendje, ám nem azonos a helyük. Az eredményeket a 4-3.a – 4.3.c táblázatok tartalmazzák:

4-3.a táblázat: Fel nem tárt ok-okozati összefüggések száma – az ok nem szennyezett, az okozat igen

Outlier típusa az okozatban	Outlierek nagysága	Outlierek aránya (%)		
		2,5	5,0	10,0
AO	$\bar{y} \pm 3\sigma$	0	0	0
	$\bar{y} \pm 5\sigma$	0	0	0
	$\bar{y} \pm 10\sigma$	0	0	0
IO	$\bar{y} \pm 3\sigma$	0	0	0
	$\bar{y} \pm 5\sigma$	0	0	0
	$\bar{y} \pm 10\sigma$	0	10	167

4-3.b táblázat: Fel nem tárt ok-okozati összefüggések száma – az ok szennyezett, az okozat nem

Outlier típusa az okban	Outlierek nagysága	Outlierek aránya (%)		
		2,5	5,0	10,0
AO	$\bar{y} \pm 3\sigma$	0	0	4
	$\bar{y} \pm 5\sigma$	1	10	180
	$\bar{y} \pm 10\sigma$	194	523	764
IO	$\bar{y} \pm 3\sigma$	17	204	559
	$\bar{y} \pm 5\sigma$	353	679	825
	$\bar{y} \pm 10\sigma$	823	899	919

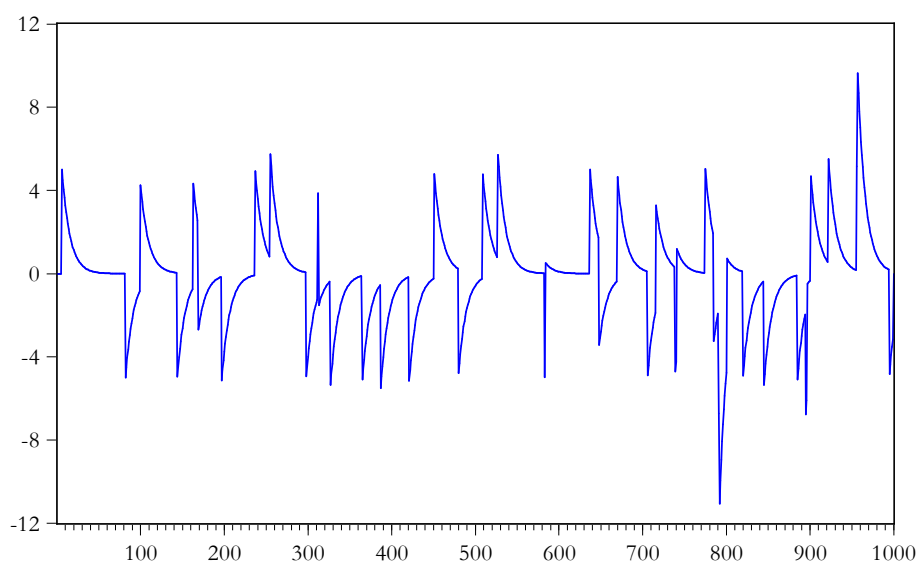
4-3.c táblázat: Fel nem tárt ok-okozati összefüggések száma – az ok is és az okozat is szennyezett

Outlier típusa az okban	Outlierek nagysága	Outlier típusa az okozatban					
		AO, melynek aránya (%)			IO, melynek aránya (%)		
		2,5	5,0	10,0	2,5	5,0	10,0
AO	$\bar{y} \pm 3\sigma$	0	0	0	0	0	0
	$\bar{y} \pm 5\sigma$	0	0	0	0	0	16
	$\bar{y} \pm 10\sigma$	0	8	156	21	232	622
IO	$\bar{y} \pm 3\sigma$	25	114	253	0	0	15
	$\bar{y} \pm 5\sigma$	161	305	528	4	72	385
	$\bar{y} \pm 10\sigma$	404	659	853	442	764	900

A Granger-okság outlierok általi elfedésével kapcsolatban három lényeges megállapítást tehetünk:

- triviálisnak tűnik, ennek ellenére hangsúlyozandó, hogy *minél nagyobb arányban* tartalmaz kiugró értékeket egyik, vagy másik idősor, tendenciájában annál nagyobb az esély arra, hogy egymással ok-okozati relációban levő idősor-pár esetén a Wald-próba *tévesen* a Granger-okság hiányát mutatja (elfogadjuk a nullhipotézist),
- amennyiben az outlierok, mint az *ok szerepét betöltő változó* jelennek meg, *nagyobb eséllyel jutunk fals negatív eredményhez*, és
- a *hibás döntés* valószínűsége az ok szerepét betöltő változóban megjelenő *tovagyűrűző (IO) outlier esetén magasabb*, ebben az esetben az F-próba már viszonylag kevés, és nem extrém nagy kiugró érték esetén is gyakorlatilag használhatatlan.

Mindez nyilván nem meglepő eredmény, ha belegondolunk, hogy tovagűrűző outlier esetén már viszonylag kis szennyezettségi arány (pl. 5%) mellett is az 1000 elemű idősor szinte valamennyi eleme, kisebb, vagy nagyobb mértékben eltér az eredeti adatgeneráló folyamatból származó értéktől. A 4-9. ábra illusztrálja az IO-outlier és az eredeti DGP-ből származó értékek eltérését 5%-os outlier-sűrűség és 5 szórásnyi eltérés esetén:



4-9. ábra: IO-outlier hatása egy 0 várható értékű stationer idősorra

Összegezve kijelenthető, hogy az ok-okozati összefüggés kimutathatósága szempontjából nagy jelentőséggel bír nem csak a kiugró érték megtalálása, hanem az annak hatását (utóéletét) leíró folyamat korrekt specifikálása is.

Ebben a gondolkörben utolsóként megvizsgáltam, hogy *miként torzítja az egységgyököket tartalmazó idősorok együttmozgását a kiugró értékekkel történő szennyezettség*. A feltevéseim megegyeztek a VAR-

modelleknél bemutatottakkal, természetesen azzal a triviális eltéréssel, hogy ebben az esetben kointegrált idősorokat eredményező adatgeneráló folyamatot tekintettem kiindulásként.

A szimulációk eredményét a 4-4 táblázat tartalmazza:

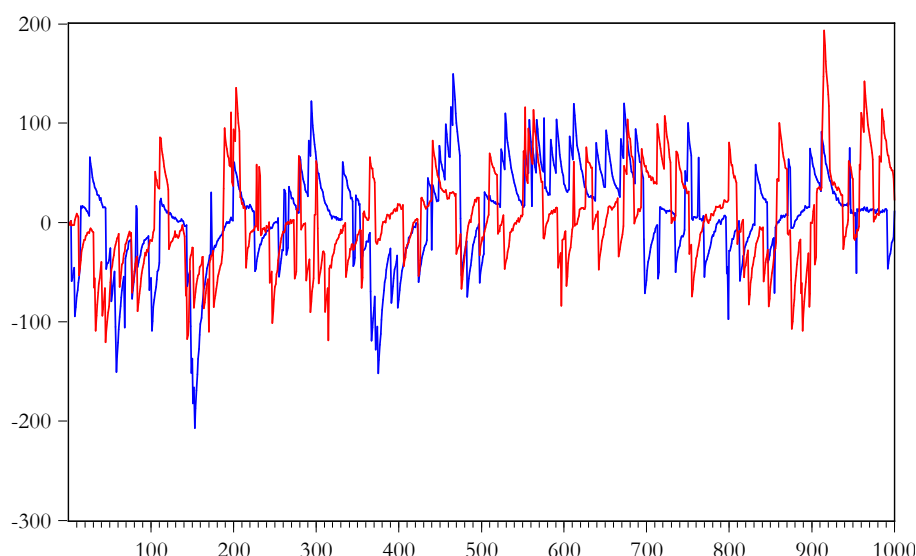
4-4. táblázat: 1000 hibakorrekciós modell esetén a Johansen-próba által fel nem ismert kointegrációs összefüggés, különböző típusú, nagyságú és gyakoriságú outlier esetén

Outlier típusa az y_{1t} változóban	Outlierek nagysága	Outlier típusa az y_{2t} változóban								
		nincs			AO			IO		
		Outlier gyakorisága (ha van valamelyik változóban)								
		2,5	5,0	10,0	2,5	5,0	10,0	2,5	5,0	10,0
nincs	$\bar{y} \pm 3\sigma$				5	5	5	61	72	70
	$\bar{y} \pm 5\sigma$				5	5	4	64	98	82
	$\bar{y} \pm 10\sigma$				8	6	6	67	60	64
AO	$\bar{y} \pm 3\sigma$	7	3	5	9	5	5	65	60	50
	$\bar{y} \pm 5\sigma$	6	1	5	5	5	2	52	57	47
	$\bar{y} \pm 10\sigma$	5	3	5	8	5	5	61	29	14
IO	$\bar{y} \pm 3\sigma$	78	61	73	59	56	61	65	34	14
	$\bar{y} \pm 5\sigma$	73	62	61	51	75	75	42	15	4
	$\bar{y} \pm 10\sigma$	54	66	68	62	67	55	4	0	0

A hibakorrekciós modellel generált idősorok együttmozgásának vizsgálata két figyelemre méltó, igaz nem túl váratlan eredménnyel járt:

- a fals negatív (hibásan nem együttmozgónak vélt) esetek előfordulása *szimmetrikus* a két változó szempontjából,
- a *tovagyűrűző hatás* (IO) ez esetben is *nagyobb mértékben csökkenti a próba erejét*, vagyis a modellezés során ennek kezelése fontosabbnak tűnhet.

Az egyetlen tulajdonképpen meglepő eredmény, hogy amennyiben mindkét idősor IO-típusú outlierrel szennyezett, akkor a fals negatív ráta magas outlier értékeknél sem látszik veszélyesnek. Illusztrációként tekintsük a 4-10. ábrát:



4-10. ábra: Két eredetileg kointegrált, IO outlierrel szennyezett idősor

Az ábra azt sejteti, hogy a hibakorrekciós összefüggés pontosan azért marad meg (mutatható ki), mert az egyes nagy abszolút értékű outlierok utáni, jellegében azonos pályájú tovagyrúzó periódusok dominálják az eredeti adatgeneráló folyamatok kisebb volatilitását.

4.4.2 Okozhatja-e az outlier-szűrés a DGP félrespecifikálását?

Megvizsgáltam, hogy az ellentétes irányú statisztikai eljárás, vagyis az outlier-szűrés (jelen esetben winsorizálás) eredményezheti-e az eredeti (elvben ismert) adatgeneráló folyamat elfedését, vagyis előfordulhat-e az az eset, amelyben például nem találjuk meg a sztochasztikus trendet egy véletlen bolyongásban. Ebben az alpontban mindössze a winsorizálás hatását vizsgálom, a következő alfejezet egy konkrét empirikus példában is ennek praktikus alkalmazhatóságát illusztrálom.

Az alkalmazott szimulációs eljárás viszonylag könnyen átlátható, ugyanis akárcsak korábban, most is a 2.4 alfejezetben bemutatott hat adatgeneráló folyamatból (két elsőrendű autoregresszív modell, két random walk, egy VAR- és egy hibakorrekciós modell) származó idősorokon vizsgáltam, hogy a kiterjesztett Dickey-Fuller-próba, a Granger-oktságot vizsgáló F-próba, illetve a kointegráció Johansen-tesztje a DGP természetének megfelelő eredményre jut-e.

Az 1000 esetből hibásan specifikált folyamatok számát, különböző winszorizálásnál használt α értékek mellett a 4-5. táblázat tartalmazza.⁷⁹

4-5. táblázat: Hibásan specifikált folyamatok száma

α	Adatgeneráló folyamat					
	AR1 ($\varphi = 0,9$)	AR1 ($\varphi = -0,4$)	RW ($\mu = 0$)	RW ($\mu = 0,01$)	VAR	EC
0,05	0	0	30	30	0	13
0,10	0	0	37	37	0	29
0,15	0	0	42	42	0	47
0,20	0	0	63	63	0	66
0,25	0	0	89	89	0	84

Az eredmények egyértelműek:

- a stacionárius folyamatok helyes felismerését nem zavarja a winszorizálás, nyilvánvaló, hogy az amúgy is konstans szórás további csökkentése nem okoz problémát;
- ezzel szemben, amennyiben a trendet tartalmazó idősről (idősorokból) kiszűrjük az elején, illetve végén található kis, illetve nagy értékeket, akkor az idősorok egyre jobban hasonlítanak a stacionáriusra, így a sztochasztikus (kointegráció esetén közös) trend elfedődik, a hibás specifikáció valószínűsége növekszik a winszorizálás szigorításával, ugyanakkor az outlier-szűrés a fals negatív rátát nem növeli kezelhetetlenül nagyra.

⁷⁹ Hibás specifikáció az AR-folyamatoknál az egységgyök hipotézis elfogadása, az RW folyamatoknál az egységgyök elvetése, VAR-modellben az okság hiányát jelentő nullhipotézis elfogadása, hibakorrekciós modellekben a kointegráció hiányának elfogadása. A winszorizáció során az időszori értékek közül a legkisebb, illetve a legnagyobb α százalékot helyettesítettem a megfelelő percentilissel.

Összefoglalva az outlierekkel kapcsolatos szimulációs eredményeket:

1. Pusztán az outlierek megjelenésétől 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 tovagyűrűző outlierok (IO) 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 winsorizálás elfedheti az egységgyököt, illetve az idősorok együtt bolyongását.

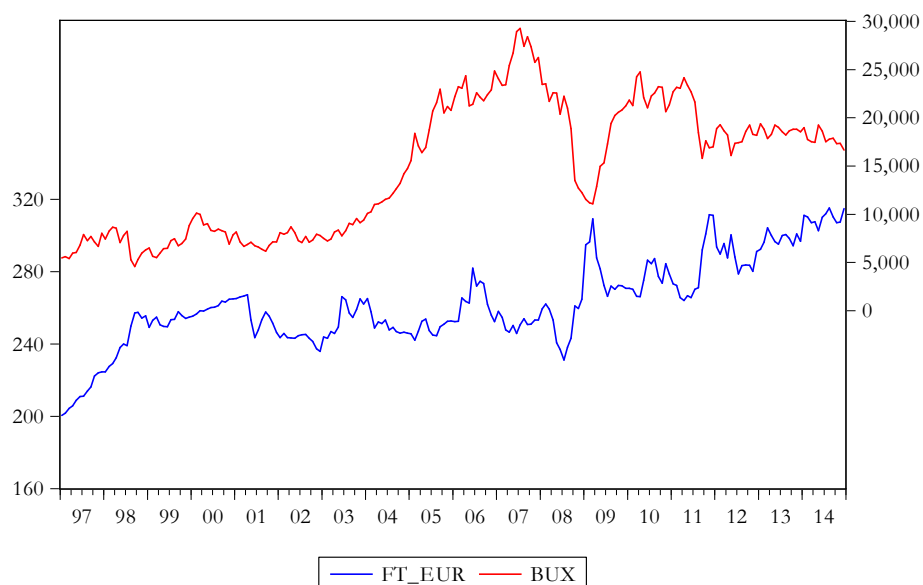
A következőkben egy empirikus példán illusztrálom az idősorok között csak outlier-szűrést követően felismerhető összefüggéseket.

4.5 Illusztratív példa az elmúlt 20 év Magyarországról

A befektetés-elemzők régóta kutatják, hogy vajon kimutatható-e összefüggés a feltörekvő tőkepiacok teljesítménye és az adott ország nemzeti valutájának értéke között.⁸⁰ A tanulmányok zöme arra a megállapításra jut, hogy – noha az elmélet emellett szólna – nem mutatható ki kointegrációs kapcsolat a tőzsdeindex és a devizaárfolyam között, vagy amit ezzel gyakran ekvivalensnek tekintenek, nincs ok-okozati összefüggés az indexhozam és az árfolyam-változás között.

Ha megvizsgáljuk az elmúlt közel két évtized magyarországi adatait, akkor a Budapesti Értéktőzsde meghatározó részvényindexének (BUX), illetve az euró-árfolyamnak hó végi záróértékei az alábbi módon alakultak:

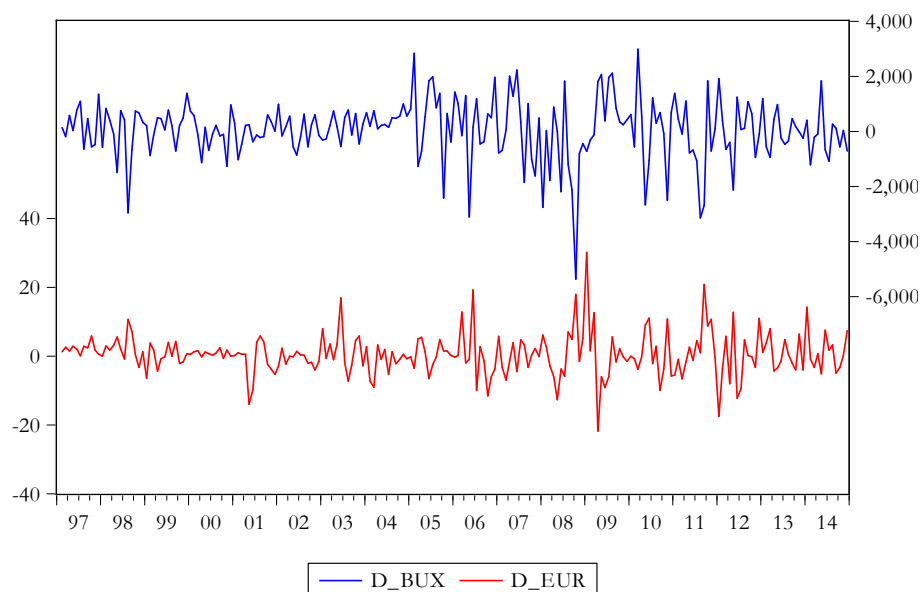
⁸⁰ Az következő illusztratív példához leginkább köthető tanulmány Srivastav et al. (2010), gyakran hivatkozzák Kim (2003) cikkét is.



4-11. ábra: A BUX és az euró-árfolyam hó végi záróértékei 1997. január 2014. december
(bal oldali tengely Ft, jobb oldali tengely pont)

A 4-11. ábrából leolvasható, hogy az euró-árfolyam a vizsgált periódusban 200 és 315 forint közötti sávban mozgott, összességében emelkedő tendenciát mutatott, mindeközben a tőzsdeindex (BUX) 4500 és 30 000 pont tartományban járt, az első évtized (1997-2007) masszív emelkedése után a válság idején harmadára esett vissza, azóta inkább az „oldalazás” jellemzi. Mindkét idősor egységgyököt tartalmaz (az ADF-próba szignifikancia értéke a BUX-nál 0,492; az euró-árfolyamnál 0,229), vagyis közöttük – elvben – nem kizárt a kointegráltság. A grafikonon ugyanakkor nem látunk szembeötlő együttmozgást, amit a Johansen-próba is alátámaszt a szignifikancia-érték meglehetősen magas ($p_{\lambda_{true}} = 0,442$, $p_{\lambda_{max}} = 0,692$), vagyis a közös trend nem volt felfedezhető. Megpróbálkoztam az eredeti idősorok winszorizálásával (különböző kvantilisek mellett), de az ilyen módon szűrt idősorok között sem lehetett kointegrációs kapcsolatot találni.

Megvizsgáltam, hogy kimutatható-e az összefüggés a havi árfolyam-változás és a tőzsdeindex azonos havi elmozdulása között.



4-12. ábra: A BUX és az euró-árfolyam havi változása 1997. január 2014. december között
(bal oldali tengely Ft, jobb oldali tengely pont)

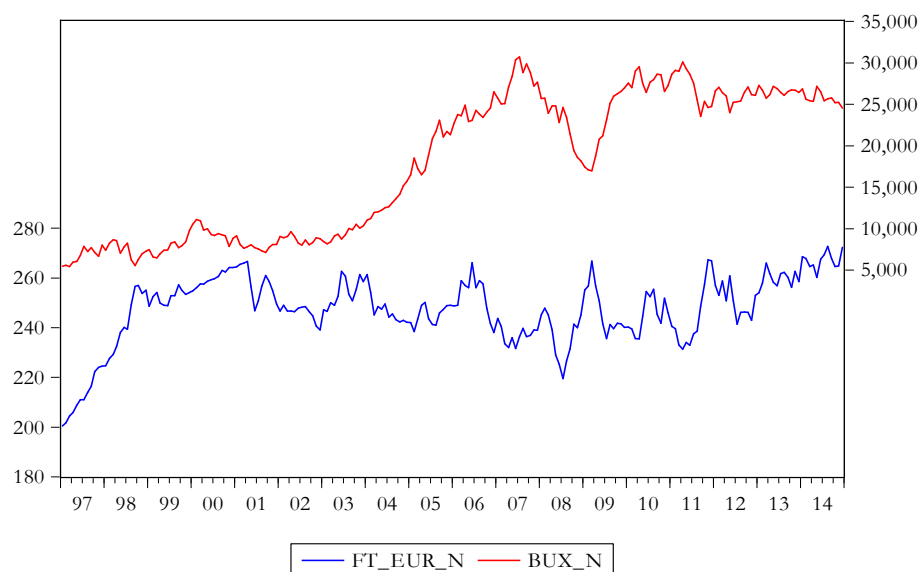
Mindkét differencia-idősor stacionáriusnak bizonyult (az ADF-teszt szignifikancia értékei minden ésszerű szintnél alacsonyabbak), így a Granger-okság a szokásos Wald-próbával tesztelhető. Amennyiben feltételezzük, hogy az árfolyamváltozás nem okozója a tőzsdei mozgásoknak a próbafüggvény értéke $F = 1,532$ ($p = 0,115$), vagyis a nullhipotézist nem vethetjük el, tehát a differencia-idősorok között sem találtunk oksági viszonyt.

Ennek ellenére sem tettem le arról, hogy az elméletileg indokolt, a grafikus ábra alapján sem kizárható összefüggést kimutassam, és a következő transzformációt alkalmaztam

$$\begin{aligned}\tilde{y}_1 &= y_1 \\ \tilde{y}_t &= \tilde{y}_{t-1} + (\Delta y_t)^{w_\alpha}\end{aligned}$$

vagyis az eredeti időszori értékből kiindulva a winszorizált differenciák kumulálásával nyertem az új, transzformált idősort.⁸¹ A 4-13. ábra mutatja az így nyert két módosított idősort.

⁸¹ A winzorálás során az $\alpha = 0,1$ paramétert alkalmaztam.



4-13. ábra: A BUX és az euró-árfolyam hó végi transzformált záróértékei
(1997. január 2014. december, bal oldali tengely Ft, jobb oldali tengely pont)

Az ábra – első ránézésre – csak kicsit tér el a korábban látottól, ám az eltérés hatása jelentős: amennyiben a transzformált idősorokra vonatkozóan végezzük el a Johansen tesztet, a kointegráció kimutatható ($p_{\lambda_{\text{traz}}} = 0,021$, $p_{\lambda_{\text{max}}} = 0,024$). Mindezt nyilván az okozza, hogy az eredeti idősorokban levő drasztikus változásokat az eljárás kisimította, vagyis a havi változás nagyságát limitálta. Ennek következtében a közös trend nem „rázódott” szét a magas volatilitású periódusokban, aminek természetesen ára is van: mivel mindkét differencia idősorunk jobbra elnyúló és vastagszélű, ezért több nagy visszaesést szűrtünk ki, mint nagy növekedést, így a 18. év végére a transzformált idősori értékek jelentősen elértek a ténylegesen megfigyelt adatoktól. (Mentségemre szóljon, hogy a cél nem valamifajta mozgóátlagolás, vagy simítás volt, hanem outlier-szűrés.)

Összefoglalva a hangsúlyozottan illusztratív példa tanulságait kijelenthető, hogy *körültekintő outlier-szűréssel kimutathatók lehetnek rejtőzködő trendek*, illetve ok-okozati összefüggések, de mindenképpen fontos, hogy tudjuk: *az idősorok „transzformálgatása” sok szempontból aggályos lehet!*

5 Strukturális törést tartalmazó idősorok modellezése

Ismeretes, hogy a standard lineáris modellben a paraméterek időben állandók. Ugyanakkor az empirikus vizsgálatok jelentős részében ez a feltevés nem tartható, így az alkalmazott idősorelemzés során gyakran élünk a strukturális törés feltevésével. Dolgozatomban *strukturális törés* alatt az idősor(ok)ra vonatkozó adatgeneráló folyamat(ok) paramétereinek a mintahorizonton történő megváltozását értem, és – noha nem teljes mértékben tekinthetők ekvivalensnek – szinonimaként használok a *törésmentes*, illetve a (paramétereiben) *stabil modell* kifejezéseket.

A téma az elmúlt 50 évben meglehetősen nagy figyelmet kapott, mind a statisztika, mind pedig az ökonometria irodalmában. A probléma kezelhetőségéről talán elég annyi, hogy az ökonometria elméletét összefoglaló kézikönyv (Mills-Patterson, 2006) e témával foglalkozó fejezetének címe azt sugallja, hogy a problémával együtt kell élnünk, a legfontosabb a szakszerű „bánásmód”, ugyanis az anomália nem feltétlenül szüntethető meg.⁸² Az említett fejezet kimerítően tárgyalja a strukturális törés vizsgálatának történetét, a törés időpontjának (időpontjainak) becslésére vonatkozó eljárásokat, a törés-teszteket, külön alpontot szentel a sztochasztikus, illetve determinisztikus trend esetén alkalmazandó eljárásoknak, valamint a kointegrált rendszerekben alkalmazandó próbáknak. A fejezet irodalomjegyzéke 140 tételt tartalmaz, amelyben minden érdeklődő megtalálja az esetleg még fennmaradó nyitott kérdéseire a válaszokat. A strukturális törés kimutatásának és kezelésének talán legérthetőbb empirikus alkalmazása Hansen (2001) tanulmányában olvasható. A gazdasági folyamatok közötti együttmozgások (kointegráltság) strukturális törések miatt bekövetkező megszűnését, illetve törést követő kialakulását vizsgálja PhD-disszertációjában, illetve ennek összefoglalását tartalmazó tanulmányában egykori doktoranduszom Ács Barnabás (lásd Ács, 2012).

Dolgozatomban – lévén a célom nem a strukturális törés kezelésének alapos bemutatása, hanem annak vizsgálata, hogy mennyiben lehet félrevezető, ha elmulasztjuk megvizsgálni és modellezni a törést – nem követem egyetlen korábban említett módszertani alpmű gondolatmenetét, de bizonyos helyeken értelemszerűen ezekkel azonos logika mentén tárgyalom a legfontosabb kérdéseket. Látni fogjuk, hogy a strukturális törést tartalmazó idősorok modellezése során bizonyos kérdésekben vissza kell utalnunk az előző fejezetben bemutatott IO-, illetve AO-modellekre, sőt a téma tárgyalása során egészen vissza kell nyúlnunk az egységgyök-, illetve trend-stacioner folyamatok megkülönböztetésére.

⁸² A fejezet címe „*Dealing with Structural Breaks*”, lásd Perron (2006).

5.1 A strukturális törés kimutatására szolgáló tesztek

A strukturális töréssel terhelt idősorok elemzésének kiinduló kérdése, hogy a törés időpontja *a priori* ismert-e, vagy sem. Amennyiben nincs előzetes információnk a törés helyéről, akkor egy további feladat ennek megbecsülése, az erre vonatkozó eljárásokat ismerteti pl. Bai (1997) és Bai-Perron (1998). Ebben az alfejezetben előbb áttekintem az *a priori* ismert helyen levő strukturális törés(ek) létének kimutatására szolgáló legfontosabb teszteket, majd bemutatom azokat az eljárásokat, melyekkel az ismeretlen időpontban levő törések helye megbecsülhető, végül röviden ismertetem azokat a módszereket, melyekkel egy strukturális törést tartalmazó idősorra vonatkozó modell paraméterei megfelelően becsülhetők.

A strukturális törés kimutatására szolgáló első tesztek alapgondolata az volt, hogy a törés a T megfigyelést tartalmazó idősort két részre osztja, melyek hossza rendre T_1 és $T - T_1$. A kezdeti próbák feltételezték, hogy a törés valamilyen ismert esemény (háború, természeti katasztrófa, kormányváltás, stb.) következtében keletkezett, így időpontjának *a priori* ismerete nem tűnt irreális feltevésnek.

Kézenfekvő volt tehát Chow (1960) kiinduló gondolata, miszerint végezzük el a modellbecslést mindkét részintervallumra és hasonlítsuk össze a becslési eredményeket, teszteljük, hogy közöttük szignifikáns eltérés mutatkozik-e. Amennyiben igen, vagyis az idősor elejére, illetve végére illesztett modell paraméterei jelentősen eltérnek egymástól, úgy feltételezhető, hogy az idősor strukturális törést tartalmaz. A Chow által javasolt teszt (tulajdonképpen egy ANOVA) próbafüggvénye az alábbi

$$F = \frac{\frac{RSS - (RSS_1 + RSS_2)}{k}}{\frac{RSS_1 + RSS_2}{T - 2k}} \quad (5.1)$$

ahol RSS, RSS_1, RSS_2 rendre a teljes, az első, illetve a második időszakra vonatkozó becslés eltérésnégyzet-összege, 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. Dolgozatomban nem foglalkozom vele, de könnyen belátható, hogy ugyanezen a gondolatmeneten felírható lenne egy LR-típusú próba a korlátozott (teljes időszakra vonatkozó), illetve korlátozás nélküli (két részmintán eltérő paramétert is megengedő) modellek likelihoodjainak maximumai alapján; hasonlóképpen egy Wald-típusú próba a paraméterkorlátozás tarthatóságára vonatkozóan.

Tekintsünk egy végletesen egyszerű példát: származzon idősorunk egy AR1 folyamatból, azaz az adatgeneráló folyamat

$$y_t = \mu + \phi y_{t-1} + \varepsilon_t \quad (5.2)$$

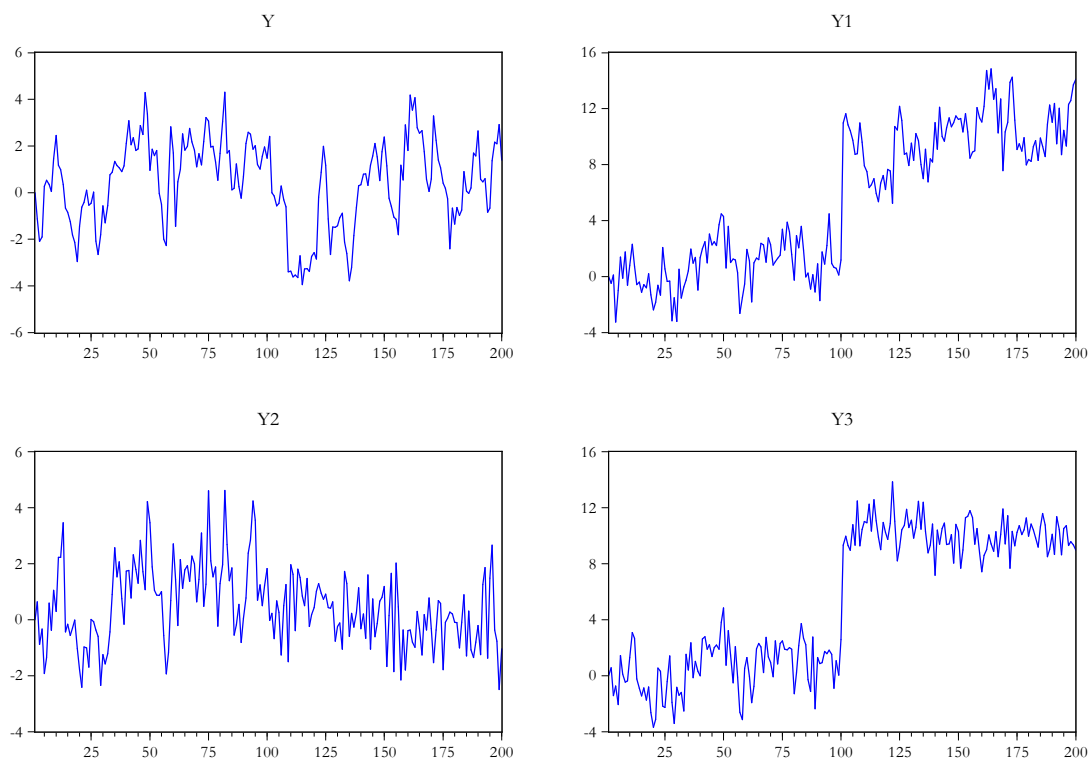
formájú, ahol az idősor szokásos jelölése mellett μ és ϕ a DGP paraméterei, ε_t pedig a szokásos *iid* véletlen. Amennyiben az idősor a T_1 időpontban strukturális törést tartalmaz, akkor az alábbi specifikációk képzelhetők el:

$$y_{1t} = \begin{cases} \mu + \phi y_{t-1} + \varepsilon_t & \text{ha } t < T_1 \\ \mu_1 + \phi y_{t-1} + \varepsilon_t & \text{ha } t \geq T_1 \end{cases} \quad (5.2a)$$

$$y_{2t} = \begin{cases} \mu + \phi y_{t-1} + \varepsilon_t & \text{ha } t < T_1 \\ \mu + \phi_1 y_{t-1} + \varepsilon_t & \text{ha } t \geq T_1 \end{cases} \quad (5.2b)$$

$$y_{3t} = \begin{cases} \mu + \phi y_{t-1} + \varepsilon_t & \text{ha } t < T_1 \\ \mu_1 + \phi_1 y_{t-1} + \varepsilon_t & \text{ha } t \geq T_1 \end{cases} \quad (5.2c)$$

Egy 200 elemű idősort, valamint annak felénél ($T_1 = 100$) bekövetkezett törést feltételezve a fenti esetek az alábbi ábrák illusztrálják:



5-1. ábra: Különböző paraméterekkel generált AR1-folyamatok ($\mu = 0, \mu_1 = 10, \phi = 0,9, \phi_1 = -0,2$)

Szimulált idősorainkra vonatkozóan a 100. megfigyelésnél feltételezett törést tesztelő Chow-próba eredményeit az alábbi táblázat tartalmazza:

5-1. táblázat: Chow-próba eredmények strukturális törést tartalmazó DGP-k esetén

Idősor	DGP	F-érték	p-érték
y_t	$0,9 y_{t-1} + \varepsilon_t$	0,853	0,427
y_{1t}	$0,9 y_{t-1} + \varepsilon_t$ ha $t < 100$ $10 + 0,9 y_{t-1} + \varepsilon_t$ ha $t \geq 100$	20,813	0,000
y_{2t}	$0,9 y_{t-1} + \varepsilon_t$ ha $t < 100$ $-0,2 y_{t-1} + \varepsilon_t$ ha $t \geq 100$	12,472	0,000
y_{3t}	$0,9 y_{t-1} + \varepsilon_t$ ha $t < 100$ $10 - 0,2 y_{t-1} + \varepsilon_t$ ha $t \geq 100$	36,234	0,000

Láthatjuk (noha az egyetlen szimulációból kiindulva semmilyen bizonyító ereje sincs!), hogy a Chow-próba empirikus eredményei *teljes mértékben igazolták* a várakozásainkat: az első (alap) esetben nem találtunk törést, míg valamennyi, az adatgeneráló folyamat paramétereinek változtatásával előállított idősor esetében a teszt szignifikáns törést jelez.

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 feltételezve⁸³),
2. majd alkalmasan 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.

⁸³ A strukturális törés lehetséges helyét behatároló T_1 és T_2 időpontok kijelölésének nincs egzakt szabálya, az irodalomban leggyakrabban 15-20%-os „trimmelést” javasolnak.

Az egyedi statisztikák összegzésére alkalmas lehet egy likelihood-arány próba (ebben az esetben összehasonlítjuk a korlátozott, illetve korlátozás nélküli eltérés-négyzetösszegeket), de ugyanígy használhatunk Wald-próbát (megvizsgálva, hogy valamennyi részmintában azonosak-e a DGP paraméterei). Quandt hivatkozott tanulmányában megmutatta, hogy lineáris adatgeneráló folyamat esetén az LR- és a Wald-próbák ekvivalensek.

A Quandt-Andrews próbáknak három fajtája⁸⁴ használatos: a maximumon, az átlagon, illetve a természetes alapú logaritmuson alapuló tesztek:

$$MaxF = \max_{T_1 \leq t \leq T_2} F(t) \quad (5.3)$$

$$AveF = \frac{1}{T_2 - T_1 + 1} \sum_{t=T_1}^{T_2} F(t) \quad (5.4)$$

$$ExpF = \ln \left(\frac{1}{T_2 - T_1 + 1} \sum_{t=T_1}^{T_2} \exp \left(\frac{1}{2} F(t) \right) \right) \quad (5.5)$$

ahol $F(t)$ a t -edik időpontot töréspontnak feltételező Chow-próba (F-próba) empirikus próba-függvény értéke. A fenti tesztstatisztikák nem rendelkeznek standard határeloszlásokkal, ugyanakkor Andrews (1993), illetve Hansen (1997) közöl aszimptotikus p-értékeket.

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 $m+1$ rezsim 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 rezsimre (ebben a $T_{j-1}+1, T_{j-1}+2, \dots, T_j-1, T_j$ időpontokhoz tartozó megfigyelések vannak) vonatkozó modell általános alakja

$$y_t = \mathbf{x}_t' \boldsymbol{\beta} + \mathbf{z}_t' \boldsymbol{\delta}_j + \varepsilon_t \quad (5.6)$$

ahol a megszokott jelöléseken túl $\mathbf{x}_t$ a nem rezsim-specifikus, $\mathbf{z}_t$ a rezsim-specifikus magyarázó változók értéke a t -edik időpontban, $\boldsymbol{\beta}$ és $\boldsymbol{\delta}_j$ paraméterek, melyek közül az előbbieket mindvégig,

⁸⁴ Lásd Andrews (1993), illetve Andrews-Ploberger (1994).

az utóbbiak csak a j -edik rezsimben befolyásolják az idősor értékét.⁸⁵ Az (5.6) alapmodellre építve háromféle töréspont-teszt ismeretes:

- a globális próbák,
- a szekvenciális próba, és
- a hibrid teszt.

Bai-Perron (1998) megmutatta, hogy az (5.6) 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, akkor meghatározva az

$$SS(\hat{\beta}, \hat{\delta} | \{T\}_m) = \sum_{j=0}^m \left(\sum_{t=T_j}^{T_{j+1}-1} (y_t - \mathbf{x}'_t \hat{\beta} - \mathbf{z}'_t \hat{\delta}_j)^2 \right) \quad (5.7)$$

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

Könnyen belátható, hogy mind m , mind T növekedése esetén a megvizsgálandó esetek száma drasztikusan növekszik, így a keresés végrehajtása még megfelelő számítástechnikai háttér mellett is nehézkes. Éppen ezért a hivatkozott tanulmány szerzői azt javasolják, hogy egyszerűsítsük a hipotézisrendszerünket arra, hogy mindössze azt tételezzük fel, hogy vagy l töréspont van, vagy egy sem. Ebben az esetben nullhipotézisünk:

$$H_0 : \delta_0 = \delta_1 = \dots = \delta_l$$

ami úgy interpretálható, hogy nincs töréspont. A próbafüggvény általános formája Bai-Perron (1998) tanulmányában olvasható. Természetesen ebben az esetben újabb problémával szembesülünk: meg kell állapítanunk az l töréspont optimális számát. A korábban hivatkozott tanulmányok alapján kijelenthető, hogy a Schwarz-féle (bayes-i) információs kritériumra alapozott mo-

⁸⁵ A rezsimek száma – értelemszerűen – eggyel több, mint a töréspontok száma, hiszen az első időszak y_0 és y_{T_1} között van.

dellválasztás, kellően általános feltételek esetén kielégítő megoldást ad, vagyis a töréspontok számát annyinak kell választani, hogy a korábbi általános modell esetén BIC értéke minimális legyen.

Bai (1997) bemutat egy jól használható intuitív megoldást az egynél több strukturális törés kimutatására. Az eljárást *szekvenciális töréspont-próbának* hívják, ugyanis a következő – rendkívül szellemes – lépéssorozatból áll:

1. teszteljük a modell paramétereinek stabilitását a teljes mintaidőszakra becsült modell alapján,
2. amennyiben a próba elveti a nullhipotézist, vagyis van töréspont, meghatározzuk annak legvalószínűbb helyét, és két részre osztjuk az idősort (töréspont előtti, illetve utáni rész)
3. elvégezzük a törésteztet mindkét részmintára, mindkét próbát úgy fogjuk fel, mintha a hipotézisrendszerünk az alábbi lenne

$$H_0 : l = 1$$

$$H_1 : l = 2$$

4. bármely részminta esetén elvetjük a nullhipotézist, ott ismét két részre osztjuk az idősort, és mindezt addig ismételtetjük (természetesen l feltételezett számát folyamatosan növelve), amíg nem lesz minden részmintánk (idősor-darabunk) törésmentes.

Nilvánvaló, hogy amennyiben előre ismert a töréspontok száma, akkor annival egyszerűsödik a fenti eljárás, hogy az l versus $l+1$ töréspont létezésére vonatkozó tesztet nem kell ennyiszer elvégezni.

Végezetül Bai-Perron (1998) annival árnyalta a megoldást, hogy az előzőekben megmutatott szekvenciális eljárás minden lépésében megkereste – a globális minimumra alapozva – a töréspont(ok) legvalószínűbb helyét, és ezek felhasználásával lépett tovább egy törésponttal. (Praktikusan ez annyit jelent, hogy az eljárás megengedi, hogy egy korábban választott töréspont-hely elmozduljon.) Mivel itt a globális, illetve szekvenciális eljárás előnyeinek kombinálása történik, ezért szokták ezt az eljárást *hibrid-módszernek* is nevezni.

Az eddig bemutatott eljárások feltételezték, hogy a töréspont(ok) helyét

- vagy a priori ismerjük,
- vagy valamilyen eljárással becsüljük.

Aue-Horváth (2012) kitűnő összefoglaló cikkükben részletesen tárgyalják azt az esetet, amikor a strukturális törés létének teszteléséhez nem használjuk fel a töréspont helyét (vagyis nem törekszünk az idősor törés előtti és törés utáni szakaszainak elkülönítéséhez). Nem megismételve a tanulmány megállapításait, belátható, hogy a Brown és szerzőtársai által már mintegy 40 éve javasolt ún. CUSUM-próba (lásd Brown et al., 1975) ebben az esetben is kedvező hatásfokkal alkalmazható.

A strukturális törésre vonatkozó próbák *általános* áttekintését követően vizsgáljuk meg, hogy milyen problémák merülhetnek fel strukturális törést tartalmazó folyamatok esetén a dolgozatom szempontjából leglényegesebb két esetben, vagyis

- sztochasztikus trendet tartalmazó idősor,
- illetve kointegrált idősori rendszer

esetén.

5.2 Strukturális törés egységgyököt tartalmazó, vagy trend-stacioner idősor esetén

Az idősor-modellezés alapjait bemutató 2. fejezetben részletesen foglalkoztam a determinisztikus, illetve sztochasztikus trendet tartalmazó modellekkel, vagyis a trendstacioner-, illetve egységgyök folyamatokkal. Láthattuk, hogy a két folyamat típus alapvetően úgy interpretálható, hogy az előbbi esetben elméletileg a trend *sobasem*, míg a második típus esetén *mindig* változik. Perron többször hivatkozott könyvfejezetében viszonylag hosszan értekezik a „mindig”, „soha”, illetve „néha” fogalmak tartalmáról, ezt dolgozatomban mellőzöm. Témám szempontjából jelentőséggel az bír, hogy *meg kell különböztetnünk* a determinisztikus trend paramétereinek megváltozását, ami egy (strukturális) törés következtében jön létre, attól a változástól, amelyet a véletlen bolyongás okoz.

Perron (1989, 1990) tanulmányaiban, a trendfüggvényben egyszer bekövetkező változásra négy esetet specifikál:

- a) szinteltolás (szint-eltolódás) az idősorban, amely nem tartalmaz trendet,
- b) szinteltolás a trendet tartalmazó idősorban,
- c) a trend-meredekség változása,
- d) mind a trend meredekségének, mind pedig a konstans tagjának megváltozása.

Az előbbi négy eset mindegyikét két különböző típusú modellbe építve vizsgálhatjuk: az előző fejezet terminológiáját használva additív, illetve tovagyűrűző outliert feltételezve, vagyis AO-, illetve IO-modellbe ágyazva.⁸⁶

A dolgozatomban bemutatandó modellezési problémák szempontjából nem bír különösebb jelentőséggel, hogy hány strukturális törést tartalmaz az idősor, ráadásul az előző alfejezet végén láthattuk, hogy a szekvenciális próbák logikájával a több törést tartalmazó idősorok modellezése mindig visszavezethető a két szakaszra bontott idősorok modellezésére, így a továbbiakban Perron modellvariánsai közül csak azokkal foglalkozom, melyek egyetlen töréspontot tételeznek fel.

⁸⁶ Mák (2011) a szinteltolás esetére ábrákkal is demonstrálja, hogy miként jelentkezik a törés AO, illetve IO-modell esetén.

Legyen a töréspont helye T_1 , amely időpillanat tekintetében egyelőre nem foglalkozunk azzal, hogy ez a hely a priori ismert-e, vagy sem! Additív outlierrel terhelt folyamatok esetén, a korábbi négy esetet feltételezve az alábbi négy modellt írhatjuk fel:

$$y_t = \mu_1 + (\mu_2 - \mu_1) D_{1t} + u_t \quad (5.8a)$$

$$y_t = \mu_1 + \beta t + (\mu_2 - \mu_1) D_{1t} + u_t \quad (5.8b)$$

$$y_t = \mu_1 + \beta_1 t + (\beta_2 - \beta_1) D_{2t} + u_t \quad (5.8c)$$

$$y_t = \mu_1 + \beta_1 t + (\mu_2 - \mu_1) D_{1t} + (\beta_2 - \beta_1) D_{2t} + u_t \quad (5.8d)$$

ahol

$$D_{1t} = \begin{cases} 0, & \text{ha } t \leq T_1 \\ 1, & \text{egyébként} \end{cases}$$

$$D_{2t} = \begin{cases} 0, & \text{ha } t \leq T_1 \\ t - T_1, & \text{egyébként} \end{cases}$$

és a hibatagra érvényes, hogy

$$\Delta u_t = C(L) \varepsilon_t$$

ahol $\varepsilon_t \sim Nid(0, \sigma_\varepsilon^2)$ és $C(L) = \sum_{j=0}^{\infty} c_j L^j$, feltételezve, hogy $c_0 = 1$ és $\sum_{j=1}^{\infty} j |c_j| < \infty$.

A tesztelendő hipotézisrendszer valamennyi esetben

$$H_0 : c_1 \neq 0$$

$$H_1 : c_1 = 0$$

formájú. Más szavakkal, ha specifikálunk egy autoregresszív polinomot melynek formája

$$A(L) = (1 - L) C(L)^{-1}$$

akkor a nullhipotézis azt jelenti, hogy ennek a gyöke 1 (egységgyököt tartalmaz), míg az alternatív hipotézis szerint a folyamat stacionáriusan fluktuál a trend körül, és az említett $A(L)$ valameny-nyi gyöke az egységkörön kívül van.

IO-modellek esetén – akárcsak korábban – a hipotézisrendszerek egyszerűbben átláthatók.⁸⁷ Ebben a modellosztályban csak az *a)*, *b)* és *d)* eseteket vizsgáljuk (ezért az alábbi képletek szokatlan számozása) ugyanis lineáris modellt feltételezve a *c)* esetre nehéz lenne empirikus alkalmazást találni. A nullhipotézis itt is mindig az egységgyök feltételezése, ilyenkor modelljeink az alábbi alakúak:

$$y_t = y_{t-1} + C(L)(\varepsilon_t + \delta D_{3t}) \quad (5.9a)$$

$$y_t = y_{t-1} + \tilde{\beta} + C(L)(\varepsilon_t + \delta D_{3t}) \quad (5.9b)$$

$$y_t = y_{t-1} + \tilde{\beta} + C(L)(\varepsilon_t + \delta D_{3t} + \eta D_{1t}) \quad (5.9d)$$

ahol a még nem ismert dummy változó képzése

$$D_{3t} = \begin{cases} 1, & \text{ha } t = T_1 + 1 \\ 0, & \text{egyébként} \end{cases}$$

szerint történik. Az ezzel szemben álló alternatív hipotézis (trendstacioner folyamat) szerint a modellek formája a következő:

$$y_t = \mu + C(L)^*(\varepsilon_t + \theta D_{1t}) \quad (5.10a)$$

$$y_t = \mu + \beta t + C(L)^*(\varepsilon_t + \theta D_{1t}) \quad (5.10b)$$

$$y_t = \mu + \beta t + C(L)^*(\varepsilon_t + \theta D_{1t} + \vartheta D_{2t}) \quad (5.10d)$$

ahol $C(L)^* = (1-L)^{-1}C(L)$. A strukturális törést tartalmazó trendstacioner modellek feltevése szerint a konstans tagban bekövetkező közvetlen hatás θ , melynek hosszútávon érvényesülő

⁸⁷ Lássuk be, hogy ez szerencsés dolog, hiszen a csillapodó hatás (tovagyűrűző outlier) realisabb modellfeltevés, mind az egyetlen időpontban megjelenő, majd azonnal eltűnő sokk!

következménye $C(1)^* \theta$, ugyanígy a meredekségben bekövetkező azonnali változás, illetve hosszú távú változás rendre ϑ és $C(1)^* \vartheta$.

Korábban már tárgyaltuk, hogy az egységgyök hipotézis ellenőrzésének standard tesztje a kiterjesztett Dickey-Fuller (ADF) próba. Perron legfontosabb megállapításai arra irányultak, hogy bebizonyítsa (szimulációval igazolja), hogy amennyiben a trendfüggvényben strukturális törés következik be az egységgyök tesztek ereje drasztikusan csökken. Későbbi tanulmányok (lásd pl. Lee et al., 1997) megmutatták, hogy amennyiben a problémát nem az egységgyök, hanem a stacionaritás irányából vizsgáljuk (pl. KPSS-tesztet alkalmazunk) a probléma akkor sem szűnik meg. Összességében tehát elmondható hogy amennyiben a tengelymetszetben (szintben), illetve a meredekségben elég nagy változás következik be a vizsgált időhorizonton, úgy a standard egységgyök, illetve stacionaritás tesztek csak nagy körültekintéssel használhatók. Nyilván további (és a dolgozat következő fejezetében vizsgálandó) kérdés, hogy mit értünk „elég nagy” változáson, és mit teszünk, ha több kisebb törés eredményez egy összességében nagy változást, de az kijelenthető hogy *az egységgyök tesztekben a törésre tekintettel kell lenni.*

5.2.1 Egységgyök-teszt ismert időpontban bekövetkezett strukturális törés esetén

Megfordítva a korábbi tárgyalási sorrendet, tekintsük előbb a reálisabb feltételezésnek tűnő IO-modelleket. Ebben az esetben az (5.10.a, b, d), illetve a kiterjesztett DF-próbát leíró (2.28) modellek felhasználásával felírhatjuk a legáltalánosabb esetet leíró egyenletet:

$$y_t = \mu + \theta D_{1t} + \beta t + \vartheta D_{2t} + \delta D_{3t} + \varphi y_{t-1} + \sum_{i=1}^k \gamma_i \Delta y_{t-i} + \varepsilon_t \quad (5.11)$$

Látható, hogy (5.11) modellbe beágyazottként megjelenik a korábban tárgyalt összes eset nullhipotézise, hiszen

- (5.9a) modell $\varphi = 1, \theta = \mu = 0, \delta \neq 0$
- (5.9b) modell $\varphi = 1, \theta = \beta = 0, \delta \neq 0$
- (5.9d) modell $\varphi = 1, \vartheta = \beta = 0, \delta \neq 0$

paraméter-restrikciók mellett előállíthatók. Hasonlóképpen felírhatók az (5.10a)-(5.10b)-(5.10d) alternatív hipotézisek is, ilyenkor az (5.11) modellben érvényre kell juttatni az $\varphi < 1, \delta = 0$ paraméter-megkötéseket.

Amennyiben a

$$H_0 : \varphi = 1$$

$$H_1 : |\varphi| < 1$$

hipotézisrendszert akarjunk tesztelni, akkor – feltéve, hogy a törés időpontját jól határoztuk meg – használható a Perron által javasolt $t_\varphi(\lambda_1)$ statisztika, ahol $\lambda_1 = T_1/T$.

Az additív outliert tartalmazó AO-modell esetén a teszt lebonyolítása logikájában is más, kétlépcsős eljárást igényel. Az első lépésben OLS-sel meg kell becsülni az idősorra illeszkedő trendfüggvényt, majd ennek értékeitől tisztítjuk az eredeti idősort. Amennyiben $\tilde{y}_t$ a trendtől tisztított idősort jelöli, akkor a korábbi (5.8a)-(5.8d) modellek helyébe felírható egyenletek

$$y_t = \tilde{\mu} + \tilde{\theta}D_{1t} + \tilde{y}_t \quad (5.12a)$$

$$y_t = \tilde{\mu} + \tilde{\beta}t + \tilde{\theta}D_{1t} + \tilde{y}_t \quad (5.12b)$$

$$y_t = \tilde{\mu} + \tilde{\beta}t + \tilde{\vartheta}D_{2t} + \tilde{y}_t \quad (5.12c)$$

$$y_t = \tilde{\mu} + \tilde{\beta}t + \tilde{\theta}D_{1t} + \tilde{\vartheta}D_{2t} + \tilde{y}_t \quad (5.12d)$$

A második lépésben elimináljuk az előbbi modellekben még szereplő D_{1t} változókat⁸⁸ és felírható, a tesztelésre alkalmas egyenlet

$$\tilde{y}_t = \varphi\tilde{y}_{t-1} + \sum_{j=0}^k \delta_j D_{3,t-j} + \sum_{i=1}^k \alpha_i \Delta\tilde{y}_{t-i} + \varepsilon_t \quad (5.13)$$

A próbafüggvény, akárcsak az IO-modell esetén, a $\varphi=1$ nullhipotézis tesztelésére szolgáló t -statisztika, amelynek határeloszlása megegyezik az előbb már bevezetett $t_\varphi(\lambda_1)$ -vel.

5.2.2 Egységgyök-teszt, amennyiben a törés időpontja nem ismert

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. Az utóbbi két évtizedben számos tanulmány látott napvilágot, melyek azt fejtették, hogy a globális világgazdaságban már nem lehet egyértelműen beazonosítani a strukturális törések időpontját, hiszen nincsenek olyan kirívó események, mint az 1929-es válság, vagy a második világ-

⁸⁸ Szép, érthető levezetést közöl Mák többször hivatkozott cikke (Mák, 2011).

háború, esetleg az 1973-as olajválság (a szerzők nyilván nem látták/láthatták előre a 2008-ban kirobbanó globális pénzügyi válság erős szakaszhatár jellegét). Ráadásul az is nehezíti a törés időpillanatának konkrét megadását, hogy még a makroidősorok is egyre gyakrabban évnél sűrűbb gyakoriságúak (negyedévesek, havi bontásúak), és ilyen nagy frekvencián beazonosítani a strukturális törés konkrét időpontját szinte lehetetlen.

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 $t_{\phi}(\lambda_1)$ statisztika maximumát. Eredetileg ezek a töréspont-kereső eljárások nyesett mintát használtak (az általánosan alkalmazott ökölszabály szerint az idősor 15%-át célszerű levágni), ám például Zivot-Andrews (1992) megmutatta, hogy a trimmelésre nincs feltétlenül szükség. Perron (2006) többszörösen hivatkozott összefoglaló fejezete részletesen taglalja a különböző kiterjesztéseket, illetve problémákat, ám ez dolgozatom alapgondolata szempontjából irreleváns.

5.3 Kointegráció tesztje strukturális törést tartalmazó idősor esetén

A 2.3.2 alponthan összefoglaltam a közös trenddel rendelkező idősorok modellezésének legfontosabb tudnivalóit, bemutattam az Engle-Granger, illetve Johansen által javasolt próbákat, melyek ebben az alfejezetben is lényeges szerephez jutnak. A strukturális törés kointegrált rendszerek esetén, több módon is megjelenhet:

1. Egyrészt elképzelhető, hogy *az idősorok trendjében változás áll be*, ám a köztük levő *együttmozgás nem változik*. Intuitív módon is könnyen belátható, hogy amennyiben csak az eredeti idősorok konstans tagjában következik be változás (szinteltolódás), az nem befolyásolja a köztük levő kapcsolatot, és a kointegráció tesztjeinek ereje nem csökken (bizonyítást lásd Campos et al., 1996). Ugyanakkor a trend-meredekségben bekövetkező változásokra a próbák érzékenyebbek, gyakrabban fordul elő, hogy szükségtelenül elvetjük a kointegráltság hiányára vonatkozó nullhipotézist, azaz pusztán az ilyen jellegű strukturális törés gyakran mutat ki látszat-kapcsolatot.
2. Másrészt előfordulhat, hogy a strukturális törés a hosszú távú összefüggést leíró modellben (*kointegráló regresszió*) jelenik meg, annak *paramétereiben következik be szignifikáns elmozdulás*. Gregory és Hansen megmutatták, hogy ebben az esetben a standard tesztek ereje csökken, végigvizsgálták az előforduló eseteket, és javasoltak egy dummy változó beépítésén alapuló eljárást, a próba hatékony végrehajtása érdekében (lásd Gregory-Hansen, 1996).

A probléma kezelésében legelterjedtebb próbák a Bartley és mtsai által javasolt modellre épülnek (Bartley et al., 2001). A próba során alkalmazandó modell (a korábban már bevezetett jelölésekkel):

$$y_{2t} = \alpha_1 + \alpha_2 D_{1t} + \gamma_2 D_{2t} + \beta_1 y_{1t} + \beta_2 y_{1t} D_{1t} + \varepsilon_t \quad (5.14)$$

Látható, hogy (5.14) két változó esetére lett felírva, természetesen amennyiben y_{2t} vektorváltozó, akkor az értelemszerű módosulások után ugyanezt az egyenletet használhatjuk a kettőnél több változó közötti kapcsolat vizsgálata során is.

Akárcsak korábban itt is problémát okozhat, hogy a modell ismert időpontban bekövetkezett törés esetén írható fel, amennyiben nem ismert a törés ideje, akkor – hasonlóan az előzőekhez – az eltérés-négyzetösszeg minimalizálásával választandó ki. A fenti regresszióból származó becsült reziduális változóra elvégezve az EG-tesztet képet kapunk a kointegráció meglétéről.

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

$$\Delta y_t = (\Pi, \Pi_j) \begin{pmatrix} y_{t-1} \\ t \end{pmatrix} + \mu_j + \sum_{i=1}^{k-1} \Gamma_i \Delta y_{t-i} + \varepsilon_t \quad (5.15)$$

Az (5.15) 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 próba továbbfejlesztésére tucatnyi kísérlet történt, ezeket Perron többször hivatkozott könyvfejezete ismerteti.

Az 5. fejezet rövid elméleti összefoglalójából láthattuk, hogy a strukturális törés kimutatására, időpontjának meghatározására számos teszteljárás, modell-megfontolás született. A probléma kezelése általában megfelelő dummy változók modellbe építésével viszonylag jó hatásfokkal 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 idősorainkban különböző időpillanatokban következett be a 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.

A következő alfejezetben ilyen jelenségek szimulálására teszünk kísérletet.

5.4 Az adatgeneráló folyamatok félrespecifikálásának lehetősége strukturális törést tartalmazó idősorokban

Ebben az alfejezetben először azt vizsgálom, hogy milyen mértékű strukturális törés kell ahhoz, hogy a (trend)stacionárius folyamat egységgyököt tartalmazónak tűnjön, illetve milyen rezsimváltás(ok)ra van szükség ahhoz, hogy a sztochasztikus trendet tartalmazó idősorunk stacionáriusnak tűnjön. A második alpontban azzal foglalkozom, hogy hány (milyen „rendszerességű”) töréspont-ra van ahhoz szükség, hogy a klasszikus Chow-próba, illetve az erre alapozó Quandt-Andrews teszt hibásan ne ismerje fel a törést. A harmadik alpontban kitérek arra is, hogy mennyire különböző helyen kell megjelennie a törésnek két, egymással együtt bolyongó idősorban ahhoz, hogy a kointegráltságot hibásan ne ismerjük fel.

Mindhárom kérdéskörben ugyanazzal a kellően általános (5.11) modellel dolgoztam, csak emlékeztetőül

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

ahol D_{1t} a szinteltolást, D_{2t} a determinisztikus trend meredekségének változását és D_{3t} törés egyszeri hatásának lecsengését meghatározó dummy változó. Az előbbi modellt, annak érdekében, hogy a szimulációk során vizsgált esetek száma még kezelhető legyen, még tovább egyszerűsítettem, és éltem a $\mu = \beta = \gamma_1 = \dots = \gamma_k = 0$ paraméterrestrikcióval. Ennek eredményeképpen a vizsgálandó modell az alábbi egyszerűbb alakot ölti

$$y_t = \begin{cases} \phi y_{t-1} + \varepsilon_t, & \text{ha } t < T_1 \\ \theta + \vartheta t + \delta + \phi y_{t-1} + \varepsilon_t, & \text{ha } t = T_1 \\ \theta + \vartheta t + \phi y_{t-1} + \varepsilon_t, & \text{ha } T_1 < t \leq T \end{cases} \quad (5.16)$$

ahol T_1 a strukturális törés szimulációimban változó (ismeretlen) helye.

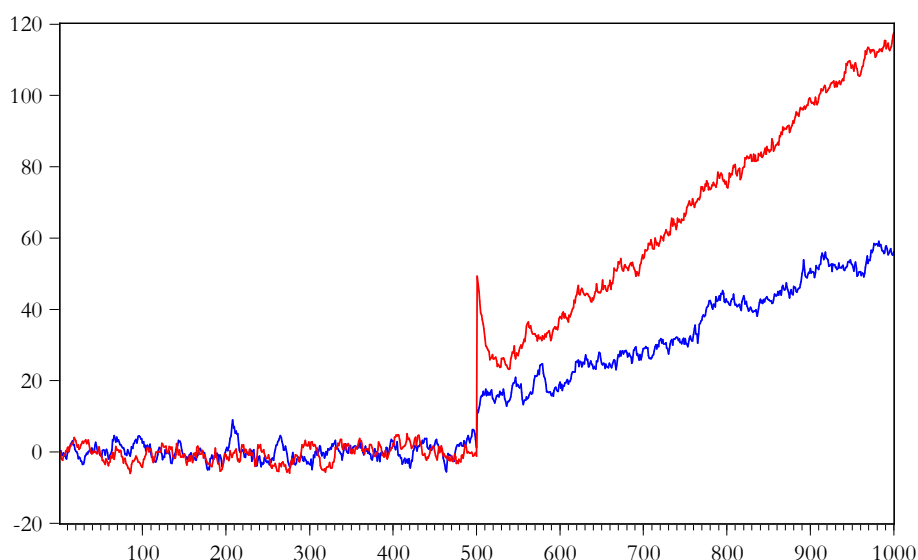
5.4.1 A sztochasztikus trend létének felismerése strukturális törést tartalmazó idősorokban

A következőkben előbb megvizsgálom, 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.⁸⁹

Illusztrációképpen tekintsünk két (5.16) által generált folyamatot, ahol a törés helye az idősor közepe ($T_1 = 500$), $\varphi = 0,9$ és $\varepsilon_t \sim N(0,1)$, miközben

- az elsőnél $\theta = 1$ $\vartheta = 0,01$ $\delta = 10$
- a másodikonál $\theta = 2$ $\vartheta = 0,02$ $\delta = 50$

értéket vesz fel:



5-2. ábra: Különböző paraméterekkel generált strukturális törést tartalmazó folyamatok

Az ábrát érdemes azzal a szemmel is nézni, hogy a szinteltolás nagysága az első (kék) folyamatnál megegyezik a véletlen varianciájával, a másodikonál ennek kétszerese, az egyszeri kiugró érték ugyanezen variancia 10-, illetve 50-szerese, míg a determinisztikus trendek meredeksége a véletlen varianciájának század, illetve ötvened része. Ha elvégezzük a kiterjesztett Dickey-Fuller-próbát konstans és determinisztikus trend feltételezésével, egyik esetben sem tudjuk elvetni az egységgyököt feltételező nullhipotézist, holott az hamis. (Mintegy önmagam megnyugtató megvizsgáltam ugyanezen idősorokat Phillips-Perron próbával is, az egységgyök léte itt sem volt elvethető.)

⁸⁹ Viszonylag sokat töprengtem azon, hogy ezeknél a szimulációknál Dickey-Fuller, vagy Phillips-Perron tesztet alkalmazzak, hiszen a modellspecifikáció az utóbbinak felel meg. Végül mégis az ADF-próbánál maradtam, hiszen a dolgozat teljes egészében azzal foglalkozik, hogy mennyiben téveszthető meg a standard eszközöket alkalmazó modellező az idősorokban levő anomáliák által.

Érdemes lenne tehát képet alkotni arról, hogy melyik paraméterre érzékeny strukturális törés esetén a stacionaritást vizsgáló próba, mitől függ a tévedések gyakorisága! Az alábbi szimulációkban – mindvégig a korábbi standard normális eloszlású véletlen változót feltételezve – végigfuttattam az (5.16) modell paramétereit az alábbi értékkészleteken

- $\varphi \in [0, 9; 0, 5; 0, 1; -0, 1; -0, 5; -0, 9]$
- $\theta \in [0, 5; 1; 2; 5; 10]$
- $\vartheta \in [0; 0, 002; 0, 005; 0, 01; 0, 02]$
- $\delta \in [10; 20; 50; 100]$

és megvizsgáltam, hogy 1000 esetből hány alkalommal fogadjuk el konstanst és determinisztikus trendet feltételező ADF-próba esetén az egységgyököt létét feltételező nullhipotézist.⁹⁰ A strukturális törés helye szimulációnként véletlenszerűen változott, de a realitások talaján maradva mindvégig az idősor középső 50%-ában volt ($250 < T_1 < 750$).

A szimulációk eredményeit az 5.2.a-d táblázatok tartalmazzák:

⁹⁰ Aki már végzett szimulációs elemzést (érzékenységi vizsgálatot), az tudja, hogy a paraméter-értékek nem véletlenül vették fel a fenti értékeket, hanem ezen kombinációk mellett tűnnek leginkább plauzibilisnek a későbbi megállapítások. Noha a szimulációk száma így sem túl kevés ($6 \times 5 \times 5 \times 4 = 600$ különböző paraméter-kombináció esetén 1000 generált idősor, azaz 600 ezer futás), ennél sokkal többet végeztem, amíg kialakultak a fenti jól interpretálható paraméter-értékek.

5-2.a táblázat: Strukturális törést tartalmazó AR1-idősorok esetén hibásan detektált egységgyök-folyamatok száma 1000 szimuláció esetén ($\varphi, \theta, \vartheta$ változik, $\delta=10$)

θ	ϑ	φ					
		0,9	0,5	0,1	-0,1	-0,5	-0,9
0,5	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	330	10	0	15	15	10
	0,010	1000	400	510	425	520	400
	0,020	1000	1000	1000	1000	1000	1000
1	0,000	0	0	0	0	0	0
	0,002	100	5	0	0	0	0
	0,005	410	60	45	40	30	40
	0,010	1000	395	430	425	405	420
	0,020	1000	1000	1000	1000	1000	1000
2	0,000	945	0	0	5	10	0
	0,002	610	40	20	5	5	35
	0,005	615	155	100	125	115	135
	0,010	955	400	405	430	460	415
	0,020	1000	1000	995	1000	995	1000
5	0,000	1000	680	705	665	705	715
	0,002	1000	645	655	635	685	685
	0,005	1000	535	615	645	585	565
	0,010	990	585	485	585	605	530
	0,020	1000	820	875	825	845	860
10	0,000	1000	1000	1000	1000	1000	1000
	0,002	1000	995	1000	1000	1000	995
	0,005	1000	1000	1000	995	1000	995
	0,010	1000	985	1000	1000	1000	985
	0,020	1000	835	860	810	855	815

5-2.b táblázat: Strukturális törést tartalmazó AR1-idősorok esetén hibásan detektált egységgyök-folyamatok száma 1000 szimuláció esetén ($\varphi, \theta, \vartheta$ változik, $\delta = 20$)

θ	ϑ	φ					
		0,9	0,5	0,1	-0,1	-0,5	-0,9
0,5	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	245	0	5	0	0	0
	0,010	1000	235	300	315	290	240
	0,020	1000	1000	1000	1000	1000	1000
1	0,000	0	0	0	0	0	0
	0,002	15	0	0	0	0	0
	0,005	380	5	5	10	10	0
	0,010	1000	240	185	250	300	270
	0,020	1000	1000	1000	1000	995	1000
2	0,000	505	0	0	5	0	0
	0,002	430	0	0	10	5	0
	0,005	575	65	60	85	75	85
	0,010	840	315	375	315	345	315
	0,020	1000	995	990	1000	995	990
5	0,000	1000	495	435	480	440	535
	0,002	1000	560	505	510	440	435
	0,005	1000	485	450	435	500	435
	0,010	960	455	545	540	525	505
	0,020	1000	720	760	765	770	750
10	0,000	1000	1000	995	1000	1000	1000
	0,002	1000	990	995	1000	1000	1000
	0,005	1000	995	995	985	995	995
	0,010	1000	965	980	975	945	970
	0,020	1000	760	770	740	825	750

5-2.c táblázat: Strukturális törést tartalmazó AR1-idősorok esetén hibásan detektált egységyök-folyamatok száma 1000 szimuláció esetén ($\varphi, \theta, \vartheta$ változik, $\delta = 50$)

θ	ϑ	φ					
		0,9	0,5	0,1	-0,1	-0,5	-0,9
0,5	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	0	0	0	0	0	0
	0,010	435	0	0	0	0	0
	0,020	1000	565	500	575	565	535
1	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	0	0	0	0	0	0
	0,010	330	5	0	15	0	5
	0,020	1000	530	500	500	500	465
2	0,000	0	0	0	0	0	0
	0,002	5	0	0	0	0	0
	0,005	225	0	0	0	0	0
	0,010	425	15	25	15	30	50
	0,020	1000	430	480	445	410	505
5	0,000	1000	0	5	15	10	0
	0,002	1000	10	10	5	5	10
	0,005	1000	50	80	35	50	65
	0,010	705	260	255	215	215	220
	0,020	855	480	390	510	435	410
10	0,000	1000	725	725	765	720	760
	0,002	1000	730	705	745	790	800
	0,005	1000	700	710	725	700	735
	0,010	1000	645	585	575	605	630
	0,020	1000	550	565	580	525	620

5-2.d táblázat: Strukturális törést tartalmazó AR1-idősorok esetén hibásan detektált egységgyök-folyamatok száma 1000 szimuláció esetén ($\varphi, \theta, \vartheta$ változik, $\delta=100$)

θ	ϑ	φ					
		0,9	0,5	0,1	-0,1	-0,5	-0,9
0,5	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	0	0	0	0	0	0
	0,010	0	0	0	0	0	0
	0,020	665	0	0	0	0	0
1	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	0	0	0	0	0	0
	0,010	0	0	0	0	0	0
	0,020	670	0	0	0	0	0
2	0,000	0	0	0	0	0	0
	0,002	0	0	0	0	0	0
	0,005	0	0	0	0	0	0
	0,010	65	0	0	0	0	0
	0,020	485	5	0	5	0	0
5	0,000	30	0	0	0	0	0
	0,002	135	0	0	0	0	0
	0,005	225	0	0	0	0	0
	0,010	385	0	0	0	0	0
	0,020	590	90	60	45	80	115
10	0,000	1000	0	0	0	0	0
	0,002	1000	0	0	5	0	0
	0,005	1000	15	5	5	15	5
	0,010	1000	75	80	75	80	85
	0,020	790	245	200	260	225	240

Az eredmények meglehetősen érdekesek: láthatjuk, hogy az egységgyök-teszt két paraméterre érzékeny

- egyrészt – logikusan – *gyakrabban fordul elő másodfajú hiba*, amennyiben az *autoregresszív együtttható abszolút értéke közel van az 1-hez* (érthetően a pozitív előjel esetén a tévedés való-

színűsége nagyobb, mint a negatív együtttható esetén, az aszimmetria valószínűleg abból fakad, hogy negatív esetben az oszcilláció valamelyest stabilizáló tényező⁹¹),

- másrészt az ADF-teszt *jóval többször jelez hibásan egységgyököt*, amikor a *szinteltolás nagyobb*, ám ezt némiképpen kompenzálja, ha jelentős mértékű tovagyrúzó outlier is megjelenik a strukturális törés helyénél.

Noha a törés helyével az 5-2.a-d. táblázatokban explicit módon nem foglalkoztam, ám a szimulációs eredmények alapján megállapítható, hogy egy törés esetén, amennyiben ez az idősor középső felében jelentkezik, az adatgeneráló folyamat felismerésének hatékonysága *nem függ attól, hogy mikor következett be a törés*.

A strukturális törést tartalmazó stacionárius folyamatokat vizsgálva kimondató, hogy amennyiben a törést követően jelentősen megváltoznak az adatgeneráló folyamat paraméterei, akkor kevésbé kell tartanunk félrespecifikálástól, mint olyankor, amikor a viszonylag kis szinteltolás, illetve nem túl jelentős outlier a két egyébként (trend) stacionárius szakaszt összemossa, és sztochasztikus trendet vélelmez.

Nilvánvalóan elvégezhető az előbbi szimuláció pandanja is, azaz megvizsgálhatjuk, hogy miként viselkedik az ADF-próba egységgyököt és strukturális törést tartalmazó adatgeneráló folyamat esetén. Mivel ez esetben az (5.16) modellben Φ nem változik (pontosabban valamennyi modellben $\Phi = 1$), ezért itt mindössze 150 paraméter-kombináció létezik, melyekre vonatkozó fals negatív esetek számát tartalmazza az 5-3. táblázat.

⁹¹ Erre az összefüggésre „örökös mentorom” Hunyadi professzor hívta fel a figyelmemet, amit ezúton is köszönök!

5-3. táblázat: Strukturális törést tartalmazó RW-idősorok esetén hibásan stationernek vélt folyamatok száma 1000 szimuláció esetén

θ	ϑ	δ			
		10	20	50	100
0,5	0,000	81	59	10	0
	0,002	0	0	0	0
	0,005	0	0	0	0
	0,010	0	0	0	0
	0,020	0	0	0	0
1	0,000	269	266	143	21
	0,002	0	0	0	0
	0,005	0	0	0	0
	0,010	0	0	0	0
	0,020	0	0	0	0
2	0,000	200	255	291	216
	0,002	142	178	143	3
	0,005	0	0	0	0
	0,010	0	0	0	0
	0,020	0	0	0	0
5	0,000	53	80	181	295
	0,002	67	75	162	274
	0,005	34	58	135	203
	0,010	0	0	1	0
	0,020	0	0	0	0
10	0,000	9	21	67	153
	0,002	23	25	54	154
	0,005	15	19	69	149
	0,010	1	2	40	130
	0,020	0	0	0	0

A szimulációk eredménye meglehetősen egyértelmű képet fest: *jelentősebb arányú első fajú hibára (fals negatív ráta) viszonylag nagy színteltolás és/vagy tovagűrűző outlier, valamint nem túl markáns (tehát 0-hoz közeli meredekségű) trendstationaritás esetén kell számítanunk.*

Összességében a kétirányú vizsgálat eredménye nem igazán kedvező, hiszen azt láthatjuk, hogy

- kis szinteltolás, nem túl jelentős outlier és markáns determinisztikus trend esetén attól kell tartanunk, hogy a kiterjesztett Dickey-Fuller próba az egyébként stacionárius folyamatot hibásan egységgyököt tartalmazónak véli, míg
- az ellenkező paraméterkombináció (nagy szinteltolás, jelentős outlier, alacsony meredekségű determinisztikus trend) pedig a fordított irányú tévedés valószínűségét növeli.

A kérdés az, hogy mit tehetünk? Amennyiben ismerjük a strukturális törés helyét (vagy legalább nagyjából ismerjük), ami nem irreális feltételezés, ha valamilyen fundamentális ok áll a törés mögött, akkor célszerű a még mindig nem túl bonyolult (5.16) modell paraméterbecslését elvégezni, és különböző esetekhez tartozó paraméterrestrikciókat alkalmazva Wald-próbákkal tesztelni. Természetesen a modellező helyzete tovább bonyolódik, ha nem ismert a strukturális törés helye, illetve vélelmezhetjük, hogy egynél több törést tartalmaz az idősor, ilyenkor – tapasztalatom szerint – *a kisebb hibát akkor követjük el, ha a folyamatot integrálnak tekintjük* és az eredeti idősor helyett annak valamilyen értelmes transzformáltjával dolgozunk tovább.

5.4.2 *A hagyományos próbák ereje több töréspont esetén*

Dolgozatom alapgondolata (szemlélete) szerint azt vizsgálom, hogy milyen anomáliák merülhetnek fel, amikor a modellező a hagyományos (standard) tesztekkel használja az előfeltételek teljesülésének vizsgálata során. A gazdasági idősorok modellezésének módszertanával foglalkozó irodalom szinte minden esetben felhívja a figyelmet arra, hogy a modellezendő adatsornak törésmentesnek kell lenni, ugyanakkor – mivel az esetleges strukturális törés(ek) helyét általában nem ismerjük – a stabilitás vizsgálatára Quandt-Andrews próbát javasol (lásd pl. Kőrösi és mtsai, 1990). Mára szerencsére a modellezők körében alapstratégiának számít a modellstabilitás vizsgálata, ennek során azonban általában megelégszünk azzal, hogy az említett QA-próba jelez-e valahol strukturális törést, és amennyiben egyik időpontnál sem találunk magas F-értéket (alacsony p-értéket), akkor kimondjuk, hogy a modell törésmentes, a továbbiakban használható.

Ebben az alponthoz azt kívánom megvizsgálni, hogy mi történik akkor, ha egy hosszabb idősorban esetleg több töréspont is jelentkezik (valljuk be, ez az általam alapesetként vizsgált 1000 elemű idősorok esetén azért nem irreális feltevés!), milyen mértékben kell számítanunk arra, hogy az egy (legvalószínűbb) pontra vonatkozó töréstesztjeink nem jeleznek, ha egynél több töréspontot tartalmaz a jelenség. Alapmodellem a már korábban alaposan körüljárt (5.16) modell módosítása lesz:

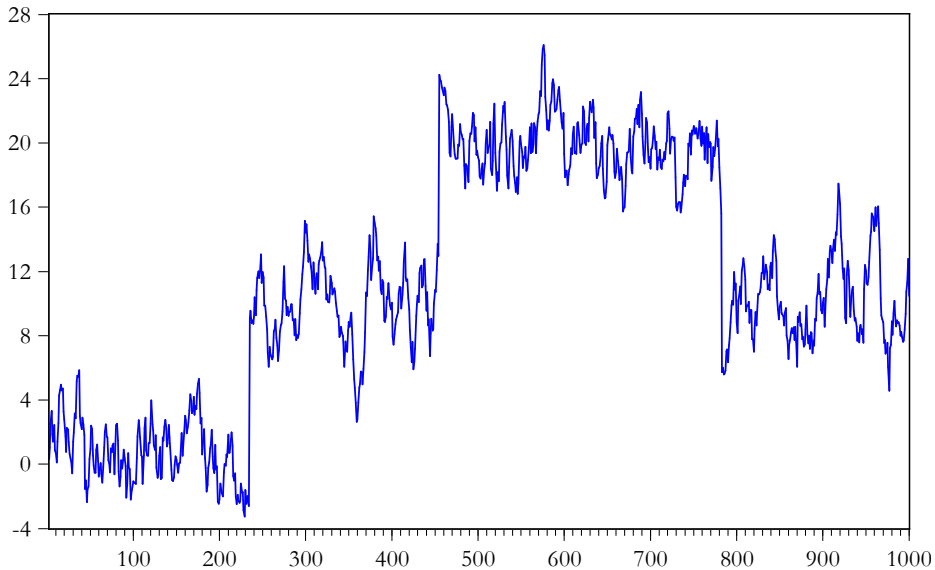
$$y_t = \begin{cases} \varphi y_{t-1} + \varepsilon_t, & \text{ha } t < T_1 \\ \theta_1 + \delta_1 + \varphi y_{t-1} + \varepsilon_t, & \text{ha } t = T_1 \\ \theta_1 + \varphi y_{t-1} + \varepsilon_t, & \text{ha } T_1 < t < T_2 \\ \theta_2 + \delta_2 + \varphi y_{t-1} + \varepsilon_t, & \text{ha } t = T_2 \\ \theta_2 + \varphi y_{t-1} + \varepsilon_t, & \text{ha } T_2 < t < T_3 \\ \theta_3 + \delta_3 + \varphi y_{t-1} + \varepsilon_t, & \text{ha } t = T_3 \\ \theta_3 + \varphi y_{t-1} + \varepsilon_t, & \text{ha } T_3 < t \leq T \end{cases} \quad (5.17)$$

Láthatjuk, hogy az előzőhöz képest, az egyszerűsítés érdekében kihagytam az adatgeneráló folyamatból a determinisztikus trendet, ugyanakkor az (5.17) modell megengedi a maximum három törés kialakulását is.⁹²

Illusztrációként tekintsünk egy három töréspontot tartalmazó adatgeneráló folyamatot, amelyben (ilyenekre vonatkozó szimulációkat fogok mutatni) a DGP paramétereire érvényesek az alábbi restriktciók:

$$\begin{aligned} \theta_1 &= \theta_2 = -\theta_3 \\ \delta_1 &= \delta_2 = -\delta_3 \end{aligned}$$

Az így keletkező idősort az 5-3. ábra mutatja:



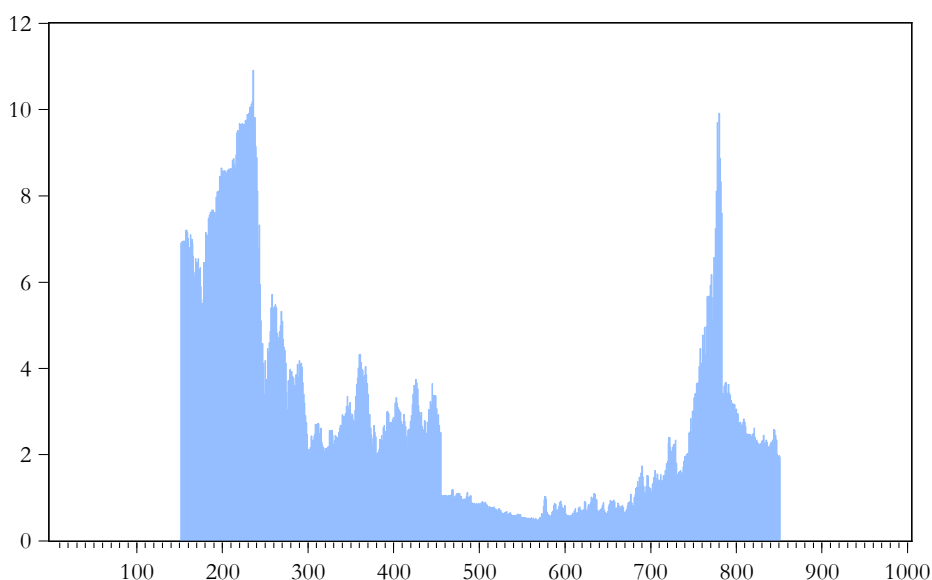
5-3. ábra: Három strukturális törést tartalmazó idősor

⁹² Vegyük észre, hogy a törések száma kevesebb is lehet a maximumnál, hiszen például amennyiben $T_2 = T$, akkor „visszajutottunk” az egy töréspont esetéhez!

Az ábrán jól kirajzolódik a vizsgált jelenség: az idősor négy stacionáriusnak tűnő szakaszból áll, melyek közül a második és a negyedik szakasz várható értéke azonosnak látszik. Vajon miként vélekedik a strukturális törésről a Quandt-Andrews próba, hol vél töréspontot felfedezni, illetve a későbbi szimulációk során azt vizsgálom, hogy mennyire függ a fals negatív ráta (fel nem ismert törés aránya) a töréspontok számától, illetve az adatgeneráló folyamat paramétereinek nagyságrendjétől?

Az 5-3. ábrán bemutatott konkrét esetben a paraméterek értéke $\theta = 1$, $\delta = 10$ volt, és a strukturális törést kereső QA-próba törést valószínűsített a 236-ik megfigyelésnél, az $LR \max F$ próba-függvény értéke $F = 10,9$, a szignifikancia-érték $p = 0,0005$.

A véletlenül generált töréspontok helye egyébiránt rendre 235, 455, 783 voltak, vagyis az ismeretlen helyen levő töréspontot kereső próba az első törést találta meg. Érdekes megvizsgálni, hogy milyen próbafüggvény-értékeket rendelt az egyes időpontokhoz a QA-próba:



5-4. ábra: A Quandt-Andrews LR maxF-próba F-értékei három törést tartalmazó idősorban

Láthatóan az első és harmadik töréspontot megtalálja a próba (kritikus F-érték mintegy 5,6), de a második töréspontnál rosszul teljesít, és teszi mindezt viszonylag drasztikus paraméterváltozások mellett, ami rendkívül elgondolkodtató.

Megvizsgáltam, hogy mennyire érzékeny a Quandt-Andrews próba a töréspontok számára, illetve az adatgeneráló folyamatban szereplő paraméterek nagyságára. A szimulációban az előzőekhez képest kevesebb paraméterkombinációval számoltam, ugyanis – megítélésem szerint – a próba

használhatóságára vonatkozó megállapítások így is jól kirajzolódnak.⁹³ Az 1000 elemű idősorokat ebben az esetben is 1000-szer generáltam, és megszámláltam, hogy egy adott paraméterkombináció mellett hányszor nem talál a hagyományos QA-teszt strukturális törést. Az eredményeket az 5-4. táblázat tartalmazza:

5-4. táblázat: Hibásan törésmentesnek vélt folyamatok száma 1000 szimuláció esetén

δ	θ	φ	Töréspontok száma		
			1	2	3
0,5	0,1	0,9	940	438	667
		0,5	930	432	690
	0,25	0,9	365	1	95
		0,5	276	1	46
	0,5	0,9	0	0	1
		0,5	0	0	0
1	0,1	0,9	944	406	655
		0,5	937	455	711
	0,25	0,9	330	1	88
		0,5	264	2	44
	0,5	0,9	0	0	0
		0,5	0	0	0
2	0,1	0,9	933	419	673
		0,5	938	474	681
	0,25	0,9	321	2	106
		0,5	271	0	58
	0,5	0,9	0	0	0
		0,5	0	0	0

Az eredmények plauzibilisek:

⁹³ A vizsgált paraméterspektrum is kisebb volt, ugyanis kifejezetten azt akartam megvizsgálni, hogy miként teljesít a próba kevésbé markáns törések esetén. Amennyiben egy töréspont volt az idősorban, annak helye a 250. és a 750. megfigyelés között volt, amennyiben két törés volt, akkor azok rendre a 201-500, illetve 501-800 tartományban voltak, és három törésnél ezek rendre a 201-400, 401-600, 601-800 tartományban fordultak elő, értelemszerűen véletlenszerűen sorsolva. A Quandt-Andrews próbát mindvégig 15%-os csonkolással hajtottam végre, de láthatóan még így is elégséges volt a távolság a töréspont és az első/utolsó vizsgált potenciális töréspont között.

- a várható értékben megjelenő törésre lényegesen *érzékenyebb* a próba, mint az egyszeri (igaz tovagyrúzó) sokkra,
- amennyiben több töréspont van, akkor még azokban az esetekben is *erősebb* a próba, amikor egyébként az adatgeneráló folyamatban csak viszonylag kis változások következnek be,
- rosszabbul teljesít a QA-próba olyankor, amikor egy korábbi törést követően az idősorban *visszarendeződés* látszik, vagyis amikor egy harmadik töréspontot követően az idősor egy korábbi állapotába tér vissza.

Úgy ítélem meg, hogy különösen az utolsó megállapítás nagy jelentőségű, mindegyre már a modell-specifikáció végiggondolása során érdemes odafigyelni.

5.4.3 A strukturális törés szerepe az együttmozgás elfedésében

A gazdasági elemzésekben, nemritkán még szakcikkekben is gyakran találkozunk azzal a kijelentéssel, hogy a közgazdasági/üzleti/pénzügyi elmélet értelmében együttmozgó jelenségeket leíró empirikus idősorok nem tartalmazznak közös trendet, és a kointegráltság hiánya nehezen magyarázható. Érdemes tehát megvizsgálni, hogy milyen mértékben vezethető vissza az együttmozgás hiánya esetlegesen különböző mértékű strukturális törésekre, illetve – és ez talán még izgalmasabb – nem azonos időpillanatokban bekövetkezett rezsimváltásokra.

Ebben az alponban a már többször vizsgált, a 2.4.4 fejezetben bemutatott kointegrált idősorokat elemzem, azt vizsgálva, hogy az egyik, illetve mindkét idősorban bekövetkezett strukturális törés milyen mértékben okozhatja a közös trend elfedését. Az alábbi négy esetet vizsgáltam:

- 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 rezsimváltás lép fel.

Az alapmodell tehát a korábban már bemutatott:

$$x_t = x_{t-1} + \varepsilon_t$$

$$y_{1t} = x_t + \varepsilon_{1t}$$

$$y_{2t} = 2x_t + \varepsilon_{2t}$$

és ez változik a fent bemutatott eseteknek megfelelően, vagyis

$$\begin{aligned}
 a) \quad & \tilde{y}_{1t} = y_{1t} \\
 & \tilde{y}_{2t} = y_{2t} + \theta D_t \\
 b) \quad & \tilde{y}_{1t} = y_{1t} + \theta D_t \\
 & \tilde{y}_{2t} = y_{2t} + \theta D_t \\
 c) \quad & \tilde{y}_{1t} = y_{1t} + \theta D_{1t} \\
 & \tilde{y}_{2t} = y_{2t} + \theta D_{2t} \\
 d) \quad & \tilde{y}_{1t} = y_{1t} + \theta_1 D_{1t} \\
 & \tilde{y}_{2t} = y_{2t} + \theta_2 D_{2t}
 \end{aligned}$$

ahol D_t , D_{1t} és D_{2t} a szokásos dummy-k, melyek a töréspontig 0, azt követően 1 értéket vesznek fel, és amelyeknél a töréspont helye különböző. A szimulációs részben azt vizsgálom, hogy mennyire érzékeny a Johansen-féle kointegrációs teszt a strukturális törésre, miként reagál az eltolás mértékének, illetve helyének megváltozására.

Tekintsük a szimulációs eredményeket! Az 5-5. táblázatban az első három modellre vonatkozó eredményeket gyűjtöttem össze, ezek mindegyikénél véletlenszerű a törés időpontja, a $c)$ modell esetében a második idősorhoz tartozó töréspont véletlen hosszúságú távolságra van az első idősor töréspontjától, az érzékenységvizsgálat során azt vizsgáltam, hogy milyen hatással van a próba erejére a szinteltolás nagysága. Mivel trendet tartalmazó idősorokról van szó, ezért a strukturális törés okozta változást (praktikusan θ értékét) nem a véletlen változó varianciájának, hanem az y_{1t} változó átlagának konstans-szorosaként definiáltam és azt vizsgáltam, hogy a változás növekedése miként hat a fals negatív rátára (elfedett együttmozgások arányára). Az eredményeket az 5-5. táblázatban láthatjuk:

5-5. táblázat: Helytelenül a közös trend hiányát vélelmező próbák száma 1000 szimuláció esetén

θ ($\bar{y}_{1t}$ százalékában)	Modell		
	a)	b)	c)
1	6	6	25
5	173	171	388
10	489	476	626
25	742	726	783
50	806	821	839

Az eredmények *egyértelműek és nem túl biztatóak*: már amikor a várható érték viszonylag kismértékű elmozdulását okozza a strukturális törés, akkor is drasztikusan lecsökken a próba ereje, gyakorlatilag az átlag 10%-ának megfelelő nagyságú θ paraméter esetén már az esetek közel felében hibásan a kointegráció hiányát mutatja a Johansen-teszt. Láthatjuk, hogy a különböző időpillanatokban bekövetkezett strukturális törések még tovább rontják a helyzetet, vagyis ilyenkor még többször döntünk helytelenül, ami nyilvánvalóan arra vezethető vissza, hogy itt három szakasza van az együttmozgásnak

- az első rezsimváltásig az eredeti modell az érvényes, majd
- a két strukturális törés közötti időszakban az *a)* modellre emlékeztet az adatgeneráló folyamat, végül
- a *b)* modell szerinti állapotba kerül az összefüggés.

Az 5-6. táblázatban a *d)* modellre vonatkozó szimulációs eredmények találhatók. Ezekben a futtásokban arra összpontosítottam, hogy mi történik akkor, amikor a két – eredetileg együttmozgó – idősort azonos időpillanatban, de különböző mértékű változást okozó strukturális törés módosítja. A törések okozta változást relatív módon (θ_2 -t θ_1 százalékában) definiáltam, és arra voltam kíváncsi, hogy az első idősorban bekövetkező változás nagysága, vagy a két változás különbségének mértékére érzékenyebb-e a Johansen-teszt.

5-6. táblázat: Fals negatív esetek száma 1000 szimuláció esetén

θ_1 ($\bar{y}_1$ százalékában)	θ_2 (θ_1 százalékában)			
	75	90	110	125
1	18	19	21	19
5	397	409	420	427
10	601	642	651	623
25	770	786	789	801

Az eredmények összecsengenek az előzőekkel: *a közös trendet elfedi a strukturális törés*, amennyiben elég nagy változást okoz akár csak az egyik idősorban. Ezen a helyzeten már nem sokat, nem lényegesen változtat az, hogy a két idősort eltérő mértékben befolyásolja a törés, mindez lényegi erőcsökkenést már nem okoz a Johansen-próbában.

Összefoglalva a strukturális törés(ek) okozta legjelentősebb félrespecifikálási veszélyeket kimondhatjuk:

1. **Rezsimváltás(ok) esetén a stacioner 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 a törésre, még akkor is, ha nem vagyunk biztosak abban, hogy ennek hatása szignifikáns.**
2. **Amennyiben az idősorban több strukturális törés van, ezt viszonylag jól kezelik a hagyományos eljárások, ám ügyelnünk kell 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 időközleltetést követően követi a változást.**

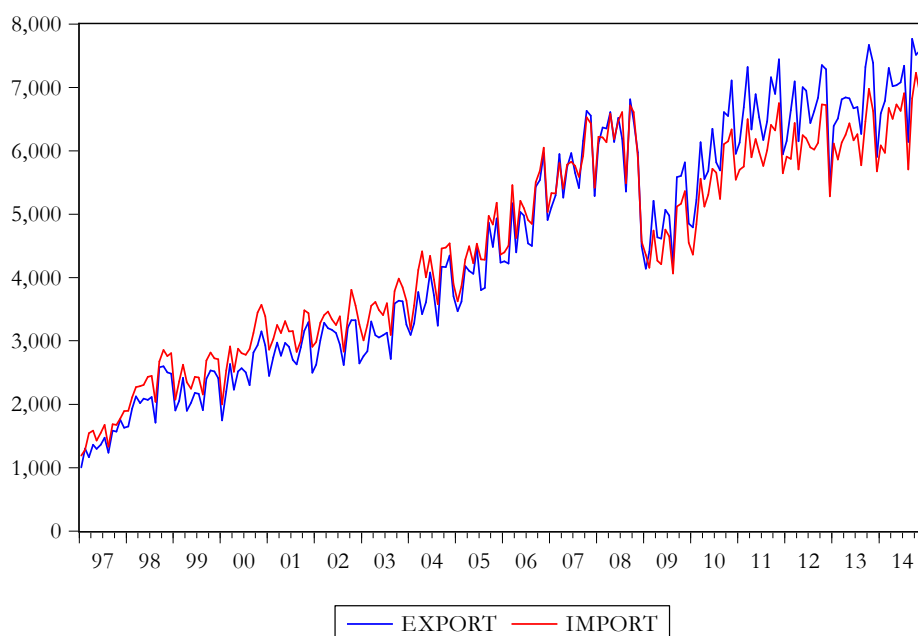
Összességében tehát az adatgeneráló folyamat feltárására irányuló *hagyományos eljárások nagyon érzékenyek a strukturális törés(ek)re*, a félrespecifikálás elkerülése ezekben az esetekben nagy körültekintést igényel. A következő alfejezetben egy empirikus példát látunk az említett veszélyek illusztrálására.

5.5 Magyarország külkereskedelmének legfontosabb mutatói az elmúlt két évtizedben

A strukturális törésekről szóló fejtegetéseket lezáró alfejezetben a magyar export, illetve import adatsoraira felírt modellt vizsgálom. Az adatok az 1997-2014-es időszakból származnak, havi bontásúak (összesen 18 év, azaz 216 hónap megfigyelés), mindkét idősor folyó áron, euróban mért.⁹⁴

Az 5-5. ábra szemlélteti az adatsorok viselkedését a vizsgált időszakban:

⁹⁴ Nyilván felmerülhet, hogy miért nem került sor az árváltozásoktól történő tisztításra, ám úgy gondoltam, hogy az euró-forint árfolyam-változás jelentős mértékben tisztítja az adatsorokat a hazai infláció hatásától (v.ö. fedezett kamatparitás elve). Másrészt ismét hangsúlyozom, hogy a példa illusztratív, célja nem valamilyen jelentős világgazdasági novum bemutatása, hanem a strukturális törés okozta anomáliák hatásának demonstrálása.



5-5. ábra: A folyó áras export és import alakulása Magyarországon
(1997. január -2014. december, millió €)

Kézenfekvő a kérdés: vajon kointegrált-e a két idősor? A közös trend megléte (amit első ránézésre az ábrán látható szinte párhuzamos lefutás nem cáfol) számos jól interpretálható makrogazdasági összefüggést eredményezne. Amennyiben az export és az import között hibakorrekciós mechanizmus lenne kimutatható, az kb. azt jelentené, hogy a külkereskedelmi egyensúly stacionárius, *a kivitelt és a behozatalt dinamikus egyensúlyi állapotban van.*

A kointegráció tesztelését elvégeztem az Engle-Granger-, valamint mindkét Johansen-teszttel, az eredményeket az 5-7. táblázat tartalmazza. (Mivel havi bontású adatokról van szó, ráadásul az ábrán is látható egy ciklikusra emlékeztető ingadozás, ezért a próbákat nemcsak az eredeti, hanem a szezonálisan tisztított idősorokra⁹⁵ is elvégeztem.)

⁹⁵ A szezonális tisztítást – kihasználva az Eviews nyújtotta lehetőségeket – egyszerű mozgóátlagolással, illetve TRAMO/SEATS eljárással is elvégeztem. Az eljárásokról lásd Sugár (1999).

5-7. táblázat: Az export és az import kointegráltságának tesztjei

Idősorok	Szignifikancia érték (p-érték) különböző próbáknál		
	EG	Johansen λ_{trace}	Johansen λ_{max}
eredeti	0,844	0,460	0,721
szezónálisan tisztított (M4)	0,814	0,567	0,820
tisztított (TRAMO/SEATS)	0,272	0,538	0,772

Látható, hogy minden előzetes várakozásunk ellenére *nem találtunk* közös trendet (nem tudtuk elvetni a kointegráció hiányát feltételező nullhipotézist), ami nemcsak az elméleti összefüggések tekintetében meglepő, de az előbbi ábra alapján sem volt sejthető. A dinamikus egyensúlyt megzavaró egyetlen anomália, hogy a globális pénzügyi válság kitörésekor (2008 őszén) mindkét idő-sorban jelentős visszaesés látszik, majd ezt követően egy alacsonyabb szintről újra indul a viszonylag dinamikus növekedés, a korábbihoz képest azzal a különbséggel, hogy innentől kezdve inkább az export-többlet a jellemző.

Megvizsgáltam, hogy mennyiben módosulnak az előbbi eredmények, ha strukturális törést feltételező vektor-hibakorrekciós modellt írok fel. Az alkalmazott specifikáció a Johansen által javasolt (5.15) modell konkretizálása, azaz

$$\begin{bmatrix} \Delta x_t \\ \Delta y_t \end{bmatrix} = \begin{bmatrix} \pi_{11} & \pi_{12} \\ \pi_{21} & \pi_{22} \end{bmatrix} \begin{bmatrix} x_{t-1} \\ y_{t-1} \end{bmatrix} + \begin{bmatrix} \mu_{11} & \mu_{12} & \mu_{13} & \mu_{14} \\ \mu_{21} & \mu_{22} & \mu_{23} & \mu_{24} \end{bmatrix} \begin{bmatrix} 1 \\ t \\ D_{1t} \\ D_{2t} \end{bmatrix} + \begin{bmatrix} \gamma_{11}^x & \gamma_{12}^x & \gamma_{13}^x \\ \gamma_{21}^x & \gamma_{22}^x & \gamma_{23}^x \end{bmatrix} \begin{bmatrix} \Delta x_{t-1} \\ \Delta x_{t-2} \\ \Delta x_{t-3} \end{bmatrix} + \begin{bmatrix} \gamma_{11}^y & \gamma_{12}^y & \gamma_{13}^y \\ \gamma_{21}^y & \gamma_{22}^y & \gamma_{23}^y \end{bmatrix} \begin{bmatrix} \Delta y_{t-1} \\ \Delta y_{t-2} \\ \Delta y_{t-3} \end{bmatrix} + \begin{bmatrix} \epsilon_{1t} \\ \epsilon_{2t} \end{bmatrix}$$

formájú, ahol x_t az export, y_t az import a t -edik hónapban, és a korábbihoz hasonlóan

$$D_{1t} = \begin{cases} 0, & \text{ha } t \leq T_1 \\ 1, & \text{egyébként} \end{cases}$$

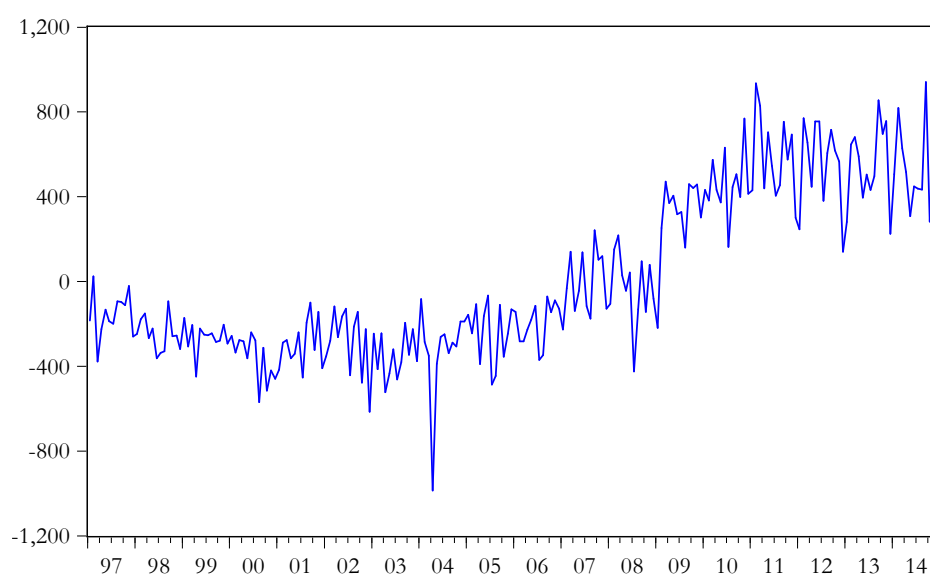
$$D_{2t} = \begin{cases} 0, & \text{ha } t \leq T_1 \\ t, & \text{egyébként} \end{cases}$$

a törés időpontját (T_1) 2008 októberére tettem. Elvégezve a paraméterbecslést a vizsgálat szempontjából leglényegesebb Π mátrixban szereplő paraméterek ML-becslése:

$$\Pi = \begin{bmatrix} 0,987 & -1,361 \\ 1,053 & -1,452 \end{bmatrix}$$

melynek sajátértékei 0,465 és gyakorlatilag 0. (Ez utóbbi nem meglepő, hiszen két független kointegráló vektor nyilván nem képzelhető el.) Johansen et al. (2000) tanulmány alapján viszonylag egyszerűen meghatározható a próbafüggvény empirikus értéke, illetve a kritikus érték (ez utóbbi kiszámítható a <http://web.uvic.ca/~dgiles/downloads/johansen/index.html> linken található programmal). A próbafüggvény értéke esetünkben 135,1, ami minden ésszerű szignifikanciaszint mellett a nullhipotézis elvetését eredményezi, vagyis *strukturális törést feltételezve már kimutatható a közös trend megléte az export és az import idősorában*.

Érdekes lehet egy másik oldalról is végiggondolni a fentieket! A kivitel és a behozatal kointegráltsága tulajdonképpen a külkereskedelmi egyenleg stacionaritását jelenti, tehát megvizsgálhatjuk ezt is:



5-6. ábra: A folyó áras külkereskedelmi egyenleg alakulása Magyarországon
(1997. január -2014. december, millió €)

Első ránézésre nem tűnik ugyan stacionáriusnak, ám alaposabban megfigyelve már észrevehető, hogy az idősor két, eltérő várható stacioner szakaszból tevődik össze, melyek között a rezsimváltás a 2009-es év elején valósult meg. Ezt támasztják alá a szokásos próbák eredményei is. A kiterjesztett Dickey-Fuller-próba (determinisztikus trend és konstans feltételezésével) az eredeti idősorra egységgyököt jelez ($p = 0,481$). Ugyanakkor, ha elvégezzük a Perron által javasolt próbát (lásd a korábbi 5.8a modellt), akkor a szignifikancia érték gyakorlatilag 0 lesz, vagyis a külkereskedelmi egyenleg szinteltolódásos stacioner folyamatból származik.

Újabb példát láthattunk tehát arra, hogy a körültekintő adatelőkészítés és modellspecifikáció árnyalhatja a gazdasági jelenségekről alkotott képünket, vagyis semmiképpen sem tekinthető statisztikus „szőrszál-hasogatásnak”!

6 Aggregált idősorok modellezése

Az idősor-elemzéssel foglalkozó művek, így dolgozatom is, szinte magától értetődő természetességgel úgy kezdődik, jelöljön y_t egy idősort, ahol t a megfigyelés időpontját jelenti. A további fejtegetések aztán az idősor megfigyelési gyakoriságát (*frekvenciáját*) a közgazdasági törvények, vagy a statisztikai gyakorlat által determinálnak tételezik fel, azaz objektív adottságként kezelik. Ugyanakkor sok esetben a megfigyelési gyakoriságot az elemző megválaszthatja, gondoljunk csak a pénzügyi idősorokra, ahol a hozamok lehetnek napi, heti, havi, vagy éves bontásban egyaránt. A kérdés, amit ebben a fejezetben vizsgálni fogok az, hogy vajon a megfigyelési gyakoriság (illetőleg dolgozatom fő mondanivalója szempontjából érdekesebb, hogy a megfigyelések a modellező által megválasztott gyakorisága), vajon befolyásolja-e a becslési eredményeinket, a vizsgált modelljeink paramétereit, illetőleg az ezekből levonható következtetéseket?

Intuitív módon viszonylag gyorsan megválaszolnánk a kérdést, és azt mondanánk, ahogy például az éves aggregáltságú adatok a negyedéves adatok függvényében alakulnak, úgy valószínűsíthető, hogy az éves bontású idősorokra vonatkozó modellek sem függetlenek a negyedéves adatokra illesztett modellektől. Ugyanakkor fontos kérdés, hogy vajon a negyedéves adatokra épülő modell információ-többletet hordoz-e, pusztán abból adódóan, hogy négyszer annyi megfigyelésre alapoztuk a becslést?

A pénzügyi ökonometria egyik leggyakoribb dilemmája, hogy az idősor megfigyelési gyakoriságának növelésével gyarapodnak a megfigyeléseink, ami *valószínűleg* javítja a becsléseink tulajdonságait, ugyanakkor a magasabb frekvencia általában nagyobb *volatilitást* (szóródást) eredményez, ami rontja a modellek illeszkedését. Sok példát és ellenpéldát lehetne hozni annak illusztrálására, hogy az időbeli aggregálás hatásának megítélése nem triviális, ám ezek helyett talán elég mindössze azt megemlíteni, hogy egy Kanadában dolgozó, új-zélandi kolléga, David E. Giles 2014 nyarán az elmúlt 60 évből 135 elemű bibliográfiát állított össze olyan, nagyrészt vezető statisztikai és ökonometriai lapokban megjelenő publikációkból, melyek az aggregálás hatásaival foglalkoztak.⁹⁶

Ebben a fejezetben először röviden áttekintem az aggregálás módozatait, illetve legalapvetőbb következményeit, majd megvizsgálom, hogy milyen hatással van az időbeli aggregálás az AR, ARMA, illetve ARIMA-modellekre, ezt követően felvázolom a hamis okság kialakulásának lehetőségét aggregált idősorok esetén. A fejezetet – szokás szerint – szimulációs eredmények, illetve hazai árfolyam adatokra vonatkozó illusztratív példa zárja.

⁹⁶ Lásd <http://web.uvic.ca/~dgiles/downloads/NZAE/biblio.pdf> (letöltve 2015. január 7.)

6.1 Aggregálás módozatai, alapvető jelölések

Közismert, hogy aggregált („felösszegzett”) idősoros változó két módon keletkezhet a statisztikában: 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. (Ezt az eljárást Marcellino (1999) *point-in-time sampling*nek nevezi, ezáltal azt jelezve, hogy aggregálás helyett tulajdonképpen csak egy hosszabb időperiódushoz rendeljük hozzá valamely tényleges „pontoszerű” megfigyelésünket, valójában ilyenkor összegzésről nincs szó.) Így járunk el, ha például a Budapesti értéktőzsde indexének napi záró-értékeiből kiválasztjuk minden ötödiket (vagyis a pénteki záró-értékeket) és az így keletkezett idősort heti bontásuként kezeljük.

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. (Marcellino korábban említett tanulmánya ezt az esetet *average sampling*ként tárgyalja, minden különösebb bizonyítás nélkül utalva arra a triviális tényre, miszerint az aggregátum, illetve az abból képzett átlag modellezése elméletileg ugyanaz a feladat.) Általánosításként elmondható, hogy a ráták, indexek (pl. munkanélküliségi ráta, vagy infláció) *stock* változók, míg a gazdasági teljesítmény mérésére szolgáló idősorok (GDP, hiány, stb.) *flow* idősorok.

Fontos előrebocsájtani, hogy a dolgozatomban tárgyalt időbeli aggregáláson túl a témának még legalább *három lényeges kapcsolódása* van:

- az aggregálás nyilvánvalóan sajátos kapcsolatot mutat a nem megfigyelt, vagy *hiányzó adatok* témakörével (illetőleg az ezek pótlására vonatkozó interpolációs eljárásokkal, erre vonatkozó eredményeket tartalmaz dolgozatom 3. fejezete);
- az időbeli összegzésnél talán gyakrabban (sőt, a mikroszimulációs technikák elterjedését követően biztos, hogy egyre gyakrabban) találkozhatunk *keresztmetszeti aggregálással*, az ilyen egyidejű (*contemporaneous*) aggregálással keletkező adatok elemzésére lehet példa Kőrösi és mtsai (1993), vagy Rappai (2016); ez utóbbiban az Európai Unió 2020-ra kitűzött céljainak teljesülését vizsgálom úgy, hogy figyelembe veszem az egyes tagországok (aggregálandó megfigyelési egységek) egyenkénti teljesítményét is;
- az aggregálás eddigi tárgyalása során mindvégig feltételeztük, hogy a sűrűbb megfigyelések, illetve az ezekre épülő modell létezik és korrektül specifikálható, ugyanakkor ez nem feltétlenül van így, ennek a *modell-bizonytalanságnak* a kezelése is érdekes, és akár etikai megfontolásokat felvető terület, ám ezzel itt nem foglalkozom.

A fejezetben alkalmazott jelölések⁹⁷ a következők: legyen szokás szerint y_t , ahol $t = 0, 1, 2, \dots$ egy nagy gyakorisággal (magas frekvencián) megfigyelt idősor t -edik megfigyelése. Képezzük ebből az aggregált (alacsonyabb frekvenciájú) idősort, vagyis meg kívánjuk határozni a minden k -adik megfigyeléshez tartozó aggregált idősori értékeket, praktikusán azon y_t értékeket, ahol $t = k, 2k, 3k, \dots$

Definiáljuk az aggregált változót a

$$y_t^* = \sum_{j=0}^A w_j y_{t-j} = W(L) y_t \quad (6.1)$$

formulával, vagyis a jelenlegi és korábbi értékek lineáris kombinációjával. A (6.1) képletből számunkra két roppant fontos eset adódik:

1. Amennyiben tartam-idősorok aggregálását végezzük, akkor $A = k - 1$, és $w_j = 1$ valamennyi j -re. Mindez ekvivalens a $W(L) = 1 + L + L^2 + \dots + L^{k-1}$ felírással, azaz

$$y_t^* = \sum_{j=0}^{k-1} L^j y_t.$$

Könnyen belátható, hogy k jelöli az aggregálás gyakoriságát, más néven az *aggregálás rendjét*, például a havi bontású tartam idősról a negyedéves gyakoriságra történő „átállás” harmadrendű aggregálással oldandó meg. Érdeemes megjegyeznünk, hogy az angolszász irodalom csak ezt az esetet nevezi időbeli aggregálásnak (*temporal aggregation*).

2. Állapot idősorok esetén a (6.1) képletben $A = k - 1$ és $w_0 = 1, w_1 = w_2 = \dots = w_{k-1} = 0$, ami praktikusán előállítja a minden k -adik megfigyelésből álló idősort, vagyis $y_t^* = y_{kt}$. Az állapot idősorok esetén alkalmazandó aggregálás – mint már említettem – megfeleltethető egy szisztematikus mintavételnek (*systematic sampling*).

A fejezetben tárgyalt kérdéskör szempontjából marginális probléma, de érdemes megjegyezni, hogy w_j alkalmas megválasztásával (pl. $w_j = \frac{1}{k}, j = 0, 1, \dots, k - 1$) a (6.1) képlet átlagolására is felhasználható, illetve amennyiben nem csak a k -adik megfigyelésekből állítjuk elő idősorunkat, úgy ezzel a képlettel nyerjük a *mozgóátlagolt* idősori értékeket is.

Végezetül vezessük be az alábbi jelöléseket is! Legyen a szokásos módon

⁹⁷ A fejezetben alkalmazott jelölések összhangban vannak a dolgozatban korábban használt képletekkel, ahol azok kiegészítésre szorultak, ott átvettem Silvestrini-Veredas (2005) összefoglaló munkájának jelöléseit.

$$\varphi(L) y_t = \theta(L) \varepsilon_t \quad (6.2)$$

az eredeti, vagyis magasabb frekvencián megfigyelt idősor, és jelölje

$$\beta(B) y_\tau^* = \eta(B) \varepsilon_\tau^* \quad (6.3)$$

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

A fejezet következő részeiben azt vizsgálom, hogy milyen változásokat, illetve anomáliákat okoz(hat) az aggregált idősor használata az eredeti (nagyobb gyakorisággal megfigyelt) idősorok helyett. Akárcsak a korábbi részekben itt is csupán azt elemzem, hogy módosulhat-e a trend létezéséről, illetve az ok-okozati összefüggésről alkotott képünk pusztán abból adódóan, hogy az idősor megfigyelési gyakorisága megváltozik.

6.2 Stacionárius folyamatok időbeli aggregálásának következményei

Az időbeli aggregálás időszori modellekre gyakorolt hatásának vizsgálata már Slutsky (1937) tanulmányában megjelent, ugyanakkor az írás lényeges megállapításai sokáig elfelejtődtek. Az aggregált adatokból képzett autoregresszív modellek vizsgálata Amemiya-Wu (1972) cikkében szerepelt először. A szerzők bebizonyították, hogy egy p -ed rendű autoregresszív folyamatból származó idősor aggregálásával keletkezett kisebb gyakoriságú idősor ARMA folyamatból származik, ahol az autoregresszivitás rendje p , a mozgóátlag tag rendje pedig az aggregálás rendjétől függ. Noha az elmúlt több mint 40 évben számos továbbfejlesztés látott napvilágot, a dolgozatomban szempontjából lényeges elemeket érdemes ezen tanulmány alapján áttekinteni!

Származzon y_t egy $AR(p)$ folyamatból, vagyis a korábbi jelölésekkel és a véletlenre vonatkozó szokásos megfontolásokkal

$$\varphi(L) y_t = \varepsilon_t \quad (6.4)$$

Legyen $\delta_j, j = 1, 2, \dots, p$ az előbbi $\varphi(L)$ polinom gyökeinek inverze (reciproka), melyek – ha a folyamat valóban stacionárius – az egységsugarú körön belül vannak. Alakítsuk szorzattá $\varphi(L)$ polinomot

$$\varphi(L) = \prod_{j=1}^p (1 - \delta_j L)$$

Ekkor belátható, hogy az y_t aggregálásával keletkező y_τ^* idősor (ahol $\tau = kt$) egy $ARMA(p, r)$ folyamatként adódik, ahol

$$r = \text{int} \left[\frac{(p+1)(k-1)}{k} \right] \quad (6.5)$$

ahol $\text{int}[\cdot]$ az argumentum egészrészét jelenti.⁹⁸

A bizonyítás legfontosabb elemét kiemelve, elmondható, hogy az eredeti és az aggregált idősor közötti kapcsolatot az alábbi polinom jelenti

$$T(L) = \prod_{j=1}^p \left(\frac{1 - \delta_j^k L^k}{1 - \delta_j L} \right) \left(\frac{1 - L^k}{1 - L} \right) \quad (6.6)$$

Ugyanis, ha megszorozzuk (6.4) mindkét oldalát a fenti polinommal, akkor

$$\prod_{j=1}^p (1 - \delta_j^k L^k) y_\tau^* = T(L) \varepsilon_\tau^*$$

adódik, amiből jól látható, hogy a p -ed rendű autoregresszív folyamat aggregálása ARMA-folyamatot eredményez, vagyis az aggregálás hatására az eredeti folyamat mozgóátlag tagokkal bővül.

A dolgozatom szempontjából legfontosabb stacionárius folyamat, az elsőrendű autoregresszív idősorok szempontjából mindez rendkívül kézenfekvően alkalmazandó. Ha y_t -re igaz, hogy $AR(1)$, vagyis

$$(1 - \phi L) y_t = \varepsilon_t$$

a szokásos megkötésekkel, akkor az aggregálásával nyert idősor

⁹⁸ A bizonyítás mind csak „szavakkal”, mind megfelelő formulákkal kitűnően nyomon követhető Silvestrini-Veredas idézett művében. Ugyanez a tanulmány tartalmazza a fejezet következő bizonyítását is.

$$(1 - \varphi^k L^k) y_\tau^* = (1 + \eta_1 B + \dots + \eta_r B^r) \varepsilon_\tau^*$$

azaz $ARMA(1, r)$ ahol

$$r = \text{int} \left[\frac{2(k-1)}{k} \right] = \text{int} \left[2 - \frac{2}{k} \right] = 1$$

Tehát az az elsőrendű autoregresszív idősorok aggregálásával, az aggregálás rendjétől függetlenül $ARMA(1, 1)$ folyamatokból származó idősorok keletkeznek, azaz például negyedrendű aggregálás esetén (negyedévesről éves bontásra történő átálláskor) az alábbi

$$(1 - \varphi^4 L^k) y_\tau^* = (1 + \eta_1 L^k) \varepsilon_\tau^*$$

adatgeneráló folyamattal kell számolnunk.

Tekintsük ezután a stacionárius folyamatok általános esetének számító ARMA-modellekkel jellemezhető idősorokat, vagyis legyen a vizsgálandó eredeti (nagy frekvenciájú) idősor (6.2) alakú! Ebben az esetben belátható, hogy amennyiben az eredeti idősor p -ed rendű autoregresszív és q -ad rendű mozgóátlag tagokból álló vegyes folyamatból származik, akkor az ennek aggregálásával keletkező idősor $ARMA(p, r)$ folyamatból származik, ahol

$$r = \text{int} \left[\frac{(p+1)(k-1) + q}{k} \right] \quad (6.7)$$

A bizonyítás a korábbiakkal ekvivalens módon arra épít, hogy (6.2) mindkét oldalát megszorozva (6.6) polinommal

$$\prod_{j=1}^p (1 - \delta_j^k L^k) y_\tau^* = T(L) \left(\sum_{j=0}^q \theta_j L^j \right) \varepsilon_\tau^*$$

kifejezést kapjuk, ahol $\theta_0 = 1$. Az előbbi kifejezés jobb oldalán álló polinom felírható

$$T(L) \left(\sum_{j=0}^q \theta_j L^j \right) = \left(\prod_{j=0}^p \sum_{i=0}^{k-1} \delta_j^i L^i \right) \left(\sum_{j=0}^q \theta_j L^j \right)$$

szorzat formában, ahol a szorzást elvégezve láthatjuk, hogy (6.7) tagszámú összeget kapunk.

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

Az alfejezet végén tekintsük át, hogy milyen következményekkel jár az aggregálás a stacionárius idősorok között esetlegesen kimutatható Granger-okság szempontjából! Granger (1990) viszonylag lakonikus rövidséggel kijelenti, hogy 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. Minde mellett 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.

Még szkeptikusabban fogalmaz Tiao professzor egy vele készült interjúban (lásd Chan, 1999), szerinte „az oksági viszonyok kuszák (*muddled*), ha az adatok aggregáltak”. A problémát abban látja, hogy amennyiben az adatainkat időintervallumon figyeljük meg, akkor a dinamikus modellek nem megfelelően működnek.

Breitung és Swanson (2002) tanulmányában részletesen végigköveti 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).

Az előbbieket illusztrálására álljon itt egyetlen példa!⁹⁹ Legyen x_t , y_t és z_t három eredeti (nem aggregált) stacionárius idősor, melyeket az alábbi elven képzünk:

$$\begin{aligned}x_t &= \alpha y_{t-1} + \beta z_{t-1} + \varepsilon_t^x \\y_t &= \varepsilon_t^y \\z_t &= \varepsilon_t^z\end{aligned}$$

ahol $\varepsilon_t = [\varepsilon_t^x, \varepsilon_t^y, \varepsilon_t^z]'$ fehér zajokból képzett vektor, melynek kovariancia-mátrixa diagonális. (Praktikusan tehát független véletlen változókról van szó.)

Ebben az esetben definíció-szerint nincs Granger-okság y_t és z_t között. Ugyanakkor bebizonyítható (lásd az előbb említett Breitung-Swanson, 2002), hogy

⁹⁹ A fejezet későbbi részében több szimulációs eredményt is megmutatok, hasonló problémára.

$$\text{Corr}(y_{\tau}^*, z_{\tau}^*) = \frac{-\alpha\beta}{\alpha^2 + \beta^2 + 1}$$

vagyis az aggregált idősorok között akár számottevő korreláció is elképzelhető.

Az alfejezetben leírtakat összefoglalva, tehát az alábbi elméleti összefüggésekre jutottunk:

- az eredetileg stacionárius idősorok az aggregálás következtében elméletileg nem válhatnak nemstacionáriussá, tehát *hamis trend nem keletkezik* az aggregálás következtében, de *ügyelniünk kell az idősört leíró vegyes modell identifikációjára*, ez ugyanis megváltozhat az összegzés hatására,
- az eredeti (nagyobb frekvenciájú) idősorok közötti *ok-okozati összefüggések* az aggregált idősorokban akár *meg is változhatnak*, mind az okság elfedődése, mind a látszat-okság kialakulása reális lehetőség az aggregátumok között.

A fejezet szimulációt tartalmazó részében az előbbi megállapításokban rejlő bizonytalanságokat próbálom csökkenteni.

6.3 Az időbeli aggregálás hatása trendet is tartalmazó idősorokban

A stacionárius folyamatok körüljárása után vizsgáljuk meg, milyen hatással van az időszakonkénti megfigyelések összegzése a véletlen bolyongást (sztochasztikus trendet) tartalmazó idősorokra! Dolgozatomban szempontjából nyilván az lenne az érdekes eset, ha az eredetileg trendet tartalmazó idősor(ok)ról kimutatható lenne, hogy az aggregálás következtében akár stacionáriussá is válhatnak, vagyis várható értékük konstanssá egyszerűsödik.

Amennyiben az időbeli aggregálás hatását nemstacioner idősorok esetében vizsgáljuk, akkor a fejezetben meghatározó szerepet betöltő (6.2) helyett a

$$\varphi(L)(1-L)^d y_t = \theta(L)\varepsilon_t \quad (6.8)$$

formájú $ARIMA(p, d, q)$ modellt használhatjuk, ahol a jelölések a szokásosak.

Amennyiben y_t egységgyököt tartalmaz, vagyis első differenciája stacionárius, akkor az előbbi (6.8) modell helyett az alábbi, minimálisan módosított $ARIMA(p, 1, q)$ modellt írhatjuk fel:

$$\varphi(L)\Delta y_t = \theta(L)\varepsilon_t \quad (6.8a)$$

Bebizonyítható¹⁰⁰, hogy a (6.8)-ban felírt ARIMA-modell aggregálást követő identifikációja egyszerűen következik, ha (6.6)-ot minimálisan módosítjuk

$$\bar{T}(L) = \prod_{j=1}^p \left(\frac{1 - \delta_j^k L^k}{1 - \delta_j L} \right) \left(\frac{1 - L^k}{1 - L} \right)^{d+1} \quad (6.9)$$

Könnyen belátható, hogy az integrált vegyes modellre felírt (6.9) stacionárius esetben, vagyis amikor $d=0$ visszaadja a korábban használt (6.6) polinomot, elsőrendű integrált, vagyis egy egységgyököt tartalmazó idősorok esetén pedig az alábbi formát veszi fel

$$\bar{T}(L) = \prod_{j=1}^p \left(\frac{1 - \delta_j^k L^k}{1 - \delta_j L} \right) \left(\frac{1 - L^k}{1 - L} \right)^2 \quad (6.9a)$$

Szorozzuk meg a (6.8) egyenlet mindkét oldalát (6.9) kifejezéssel, ekkor

$$\prod_{j=1}^p \left(\frac{1 - \delta_j^k L^k}{1 - \delta_j L} \right) \prod_{j=1}^p (1 - \delta_j L) \left(\frac{1 - L^k}{1 - L} \right)^{d+1} (1 - L)^d y_t = \prod_{j=1}^p \left(\frac{1 - \delta_j^k L^k}{1 - \delta_j L} \right) \left(\frac{1 - L^k}{1 - L} \right)^{d+1} \left(\sum_{j=0}^q \theta_j L^j \right) \varepsilon_t$$

Egyszerűsítsünk az egyenlet bal oldalát azáltal, hogy a korábbiaknak megfelelően áttérünk az aggregált idősorra, amit – éppúgy, mint eddig – jelöljön y_τ^* , miközben az aggregált modellben szereplő véletlen tagot továbbra is ε_τ reprezentálja.

Ekkor a

$$\prod_{j=1}^p (1 - \delta_j^k L^k) (1 - L^k)^d y_\tau^* = \prod_{j=1}^p \left(\frac{1 - \delta_j^k L^k}{1 - \delta_j L} \right) \left(\frac{1 - L^k}{1 - L} \right)^{d+1} \left(\sum_{j=0}^q \theta_j L^j \right) \left(\sum_{j=0}^{k-1} L^j \right) \varepsilon_\tau$$

egyenletet nyerjük, melynek bal oldala egy $ARI(p, d)$ folyamat, természetesen a ritkább megfigyelésekkel (vagyis az idő múlását reprezentáló változó τ).

Egyszerűsítve az egyenlet jobb oldalán található jelöléseket, vezessük be ξ_l -t, ahol $l = 0, 1, \dots, p(k-1) + (d+1)(k-1) + q$, amely kifejezés a korábban már használt δ_j ($j = 1, \dots, p$) és θ_j ($j = 1, \dots, q$) paraméterek függvénye! Ekkor korábbi egyenletünk

¹⁰⁰ Weiss (1984)-re is hivatkozva, lásd Silvestrini-Veredas (2005).

$$\prod_{j=1}^p (1 - \delta_j L^k) (1 - L^k)^d y_\tau^* = \left(\sum_{l=0}^{p(k-1) + (d+1)(k-1) + q} \xi_l L^l \right) \varepsilon_\tau$$

formájú, amelynek a jobb oldalán található polinom $\text{int} \left[p(k-1) + (d+1)(k-1) + q \right]$ rendű.

Mindebből következik, hogy amennyiben azonosan $ARIMA(p, d, q)$ rendű idősorok aggregálásával létrejött idősorokat kívánunk modellezni, a megfelelő modell $ARIMA(p, d, r)$ identifikációjú, ahol

$$r = \text{int} \left[\frac{p(k-1) + (d+1)(k-1) + q}{k} \right] \quad (6.10)$$

A sztochasztikus trendnek megfeleltetett *véletlen bolyongás*, vagyis az $ARIMA(0, 1, 0)$ modell esetén (6.10) az

$$r = \text{int} \left[\frac{(p+2)(k-1) + q}{k} \right] = \text{int} \left[\frac{2(k-1)}{k} \right] = 1 \quad (6.10a)$$

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

6.4 Aggregálás hatása az idősorok kointegráltságára

A sztochasztikus trendet tartalmazó idősorok vizsgálata után tekintsük át, hogy milyen következményei vannak az időbeli aggregálásnak a közös trendre, vagyis az idősorok kointegráltságára. A következőkben csak a legegyszerűbb esettel foglalkozom, amelyben két $I(1)$ idősor kointegrált, azaz lineáris kombinációjuk stacioner. A szakirodalom mindkét irányú kiterjesztést is tartalmaz, így Marcellino (1999) a magasabb rendben integrált idősorok közötti kointegrációt, Rajaguru-Abeyasinghe (2008) a többdimenziós esetet tárgyalja részletesen.

Legyen tehát a szokásos jelölésekkel y_t és x_t egységgyököt tartalmazó nemstacioner folyamat, miközben $z_t = y_t - \beta x_t$ lineáris kombinációjuk stacioner. Mindez – az ismert okokból – úgy következhet be, ha y_t és x_t tartalmaz közös, első rendben integrált komponenst, valamint tartalmaznak egymástól független, stacionárius komponenst is. Vagyis

$$\begin{aligned} y_t &= \beta w_t + u_t^y \\ x_t &= w_t + u_t^x \end{aligned} \quad (6.11)$$

ahol $w_t \sim I(1)$ és $u_t^y, u_t^x \sim I(0)$.

A korábbi definíciókból adódóan, ha az idősorok kointegráltak, akkor közöttük létezik hibakorrekciós mechanizmus, amit az alábbi módon írhatunk fel:

$$\begin{aligned} \Delta y_t &= -\gamma_1 z_{t-1} + \sum_{j=1}^{n_1} \alpha_j^y \Delta y_{t-j} + \sum_{j=1}^{m_1} \alpha_j^x \Delta x_{t-j} + \varepsilon_t^y \\ \Delta x_t &= -\gamma_2 z_{t-1} + \sum_{j=1}^{n_2} \beta_j^x \Delta x_{t-j} + \sum_{j=1}^{m_2} \beta_j^y \Delta y_{t-j} + \varepsilon_t^x \end{aligned} \quad (6.12)$$

ahol a paraméterekre vonatkozó minimális megkötés, hogy γ_1, γ_2 nem zéró.

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.

A bizonyítás roppant egyszerű, hiszen a fejezet elején bemutatott módon például y_t k -ad rendű aggregáltja

$$y_t^* = \sum_{j=0}^{k-1} L^j y_t$$

amiből következik, hogy (6.11) az aggregált idősorok esetén az alábbiak szerint módosul

$$\begin{aligned} y_t^* &= \beta \sum_{j=0}^{k-1} L^j w_t + \sum_{j=0}^{k-1} L^j u_t^y = \beta w_t^* + u_t^{y*} \\ x_t^* &= w_t^* + u_t^{x*} \end{aligned} \quad (6.13)$$

vagyis y_t^* és x_t^* éppúgy közös trendet tartalmaz, mint a nem aggregált y_t és x_t . A hibakorrekciós egyenlet paraméterei a szezonális differenciák következtében nyilván módosul(hat)nak, de az ECM létezése triviális.

Haug (2002) rendkívül alapos, számos modellvariánsra kiterjedő Monte Carlo vizsgálattal elemezte a kointegráció megállapítására kifejlesztett tesztek hatékonyságát aggregált idősorok esetén. A tanulmány összehasonlítja a közös trendre vonatkozó Engle-Granger-féle (a kointegráló regresszió hibatagjának ADF-tesztjére alapozó) EG-próbát, a Philips-Ouliaris által javasolt $\hat{Z}_\alpha$ -próbát, az Andrews által kifejlesztett $\hat{Z}_t$ -tesztet és továbbfejlesztett változatait, a Stock és Watson által elsőként bemutatott SW , SW_{nopw} , SW_{AR} próbákat, valamint a legszélesebb körben használt (a dolgozatomban is részletesen tárgyalt) Johansen-féle maximális saját-érték ($\lambda_{\max}$) tesztet.

A cikk gondolatmenete szerint a szerző 47 év, azaz 188 negyedév, tehát 564 hónap adatait generálja, majd a havi adatokat aggregálja az említett negyedéves, illetve éves szintre. A Monte Carlo vizsgálat eredményei igazolták korábbi megállapításainkat, vagyis az eredetileg kointegrált (értsd kointegráltként generált) idősorok az aggregálást követően is tartalmaztak 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 ennél az idősor-hosszúságnál az $\lambda_{\max}$ -próba ereje mintegy harmadára csökken, ha a havi bontású idősor helyett az éves szintre (12-ed rendben) aggregált idősorok között keressük az együttmozgást. Haug tanulmánya ugyanakkor számos esetben számol be nem várt eredményekről, elsősorban olyankor, amikor az idősorok az aggregálás következtében túlságosan röviddé válnak. Mindennek alaposabb körüljárása indokolja, hogy a fejezet szimulációs részében némiképpen eltérő identifikáció mellett újra megvizsgáljam az irodalomban ismertett szimulációs eredményeket.

6.5 Az időbeli aggregálás modellre gyakorolt hatásának összegzése

Számos kitűnő tanulmány található, amely konkrét gazdasági idősorokon vizsgálja végig az időbeli aggregálás következményeit. Érdemes elolvasni Rossana-Seater (1995) cikkét, melyben a szerzők 19 makrogazdaságot jellemző idősor (kibocsátás, fogyasztás, foglalkoztatás, kamatok, árindexek, tőzsdeindexek) havi bontású (illetve a reál GNP esetében negyedéves) idősorait aggregálják negyedéves, illetve éves szintre (a GNP esetén természetesen csak az utóbbi az érdekes), majd az egységgyök-tesztek elvégzése után versengő modellek közül megkeresik az optimálisnak tűnőt.

A tanulmány mondanivalóját semmiképpen sem bagatellizálva kijelenthetjük, hogy az idősorok mindegyikénél jelentősen megváltozik az optimálisnak tűnő identifikáció, ráadásul az éves bontású (aggregált) idősorok jelentős része egyszerűen véletlen bolyongásnak tűnik, miközben az eredeti (értsd nagyobb frekvencián mért) idősor ennél lényegesen összetettebb identifikációval bírt. A szerzőpáros azt is megmutatta, hogy az időbeli aggregálás nem csak a szezonális hullámzást simítja ki, hanem – részben ennek következményeként – a rövid távú üzleti ciklusok során felfedezhető együttmozgásokat, illetve kölcsönhatásokat is elfedheti.

Az elméleti összefüggéseket tartalmazó részben nem tértem ki valamennyi olyan kérdésre, amely az időbeli aggregálás hatásait vizsgálja. Összefoglalás helyett álljon itt Marcellino (1999, 131. pp.) táblázata, melyben az aggregálás különböző kitüntetett idősori tulajdonságokkal kapcsolatos invariánságát mutatja be, feltüntetve azon szerzőket/tanulmányokat¹⁰¹, akik elsőként írták le a vizsgált jelenséget:

6-1. táblázat: Az aggregálás hatása néhány fontos idősori tulajdonságra

<i>Tulajdonság/jellemző</i>	<i>Szerző(k)</i>	<i>Aggregálás invariáns-e?</i>
Egységgyök	Pierse-Snell (1995)	igen
Kointegráltság	Granger (1990)	igen
Szezonális egységgyök	Granger-Siklos (1995)	nem
Exogenitás	Hendry (1992)	nem
Okság	Sims (1971)	nem
Impulzus-válasz függvény	Granger-Swanson (1992)	nem
Trend-ciklus dekompozíció	Lippi-Reichlin (1991)	nem
Perzisztencia (tartós hatás)	Rossana-Seater (1995)	nem
Előrejelzések	Lütkepohl (1987)	nem

A fenti táblázatból is jól látható, hogy az aggregálás képes jelentősen megváltoztatni az idősor, illetve az idősorok közötti összefüggések alapvető tulajdonságait, így körültekintő kezelése feltétlenül kívánatos. A következőkben az aggregálás okozta anomáliák kialakulási valószínűségeire vonatkozó Monte Carlo vizsgálatokat mutatok be.

6.6 Trend, illetve okság az aggregált idősorokban

Az aggregálás hatását az adatgeneráló folyamatok, illetve az ezek között kimutatható legfontosabb összefüggések vonatkozásában ismét szimulációkkal vizsgálom. Külön foglalkozom a trend-megítélése körül fellépő anomáliákkal, illetve az ok-okozati összefüggések elfedése, vagy a hamis okság okozta problémákkal. Mindkét esetben ugyanazt az eljárást követem (ez nyilvánvalóan hasonlít a korábbi fejezetek logikájára), az ismert DGP-vel rendelkező, eredetileg 1000 elemű idősorokból aggregálással rövidebb idősorokat hozok létre, majd az új (alacsonyabb frekvenciájú) idősorokon végzem el a szokásos specifikációs teszteket (ADF, Granger-okságra vonatkozó Wald-

¹⁰¹ Annyiban módosítottam Marcellino táblázatát, hogy csak azokat a tanulmányokat (szerzőket) tüntettem fel, amelyek dolgozatomban irodalomjegyzékében (amúgy is) szerepelnek.

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.

Az aggregálás eredményeképpen nyert új idősorok hossza rendre 333, 250, 200, 83 és 47, melyek közül az utolsó már nem teljesíti a sztochasztikus idősor-modellezésben Box-Jenkins-i elvárásként ismert¹⁰² minimum 80-100 megfigyelés kritériumot, ugyanakkor mégis alkalmaztam, mert szimulációs eredményeim ezáltal összehasonlíthatóvá válnak a már említett Haug (2002) tanulmány megállapításaival.

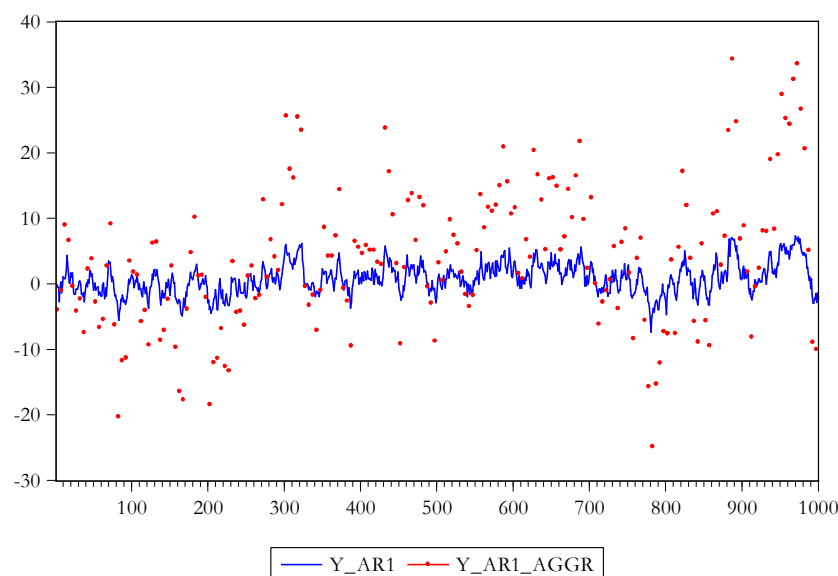
6.6.1 *Az aggregálás hatása a trend megítélésében*

Elsőként azt vizsgáltam, hogy milyen hatással van az időbeli aggregálás a trend megítélésére, vagyis milyen gyakran fordul elő fals pozitivitás¹⁰³ stacionárius idősornál, vagy fals negativitás egységgyököt tartalmazó folyamatokban. Noha a szakirodalom, elsősorban Pierse-Snell (1995) cikkére hivatkozva az egységgyök-jelenséget az aggregálásra nézve invariánsnak tartja, viszonylag kevés eredményt lehet találni, amely az elsőrendű autoregresszív folyamat együtthatója, vagy eltolásos véletlen bolyongás esetén a drift nagyságának függvényében vizsgálná a hibás döntések előfordulási gyakoriságát.

A 6-1. ábra illusztrálja a vizsgálatok logikáját: az eredetileg elsőrendű autoregresszív folyamatot (Y_AR1) ebben az esetben harmadrendben aggregáltam és így kaptam az Y_AR1_AGGR idősort. Ez utóbbi az eredeti megfigyelési gyakoriságnál értelemszerűen ritkább (ha úgy tetszik hiányos), emellett azonnal szembeötlő, hogy magasabb varianciájú. (Ez triviális, hiszen az eredeti értékek összeadásával az idősor varianciája is háromszorozódik.)

¹⁰² Lásd a klasszikus Box-Jenkins (1970).

¹⁰³ A fals pozitivitás itt is, akárcsak korábban, az ADF-próbával elkövetett másodfajú hibát jelenti.



6-1. ábra: 1000 elemű szimulált AR1-folyamat és harmadrendű aggregáltja

A kérdés ezt követően mindössze annyi, hogy a variancia növekedése okozhatja-e az eredetileg stacionárius adatgeneráló folyamat félrespecifikálását, illetve milyen autoregresszív együtthatók esetén kell óvatosabbnak lennünk?

A 6-2. táblázat az ADF-próba alapján hibásan egységgyököt tartalmazónak tekintett idősorok számát mutatja 1000 szimuláció esetén, különböző modellparaméterek mellett.

6-2. táblázat: Hibásan egységgyököt tartalmazónak vélt AR1-folyamatok száma 1000 szimuláció esetén

Aggregálás rendje	ϕ					
	0,9	0,5	0,1	-0,1	-0,5	-0,9
$k=3$	0	0	0	0	0	0
$k=4$	0	0	0	0	0	0
$k=5$	0	0	0	0	0	0
$k=12$	2	0	0	0	0	1
$k=21$	12	6	10	7	6	10

Az eredményekből láthatjuk, hogy az aggregálás ebben a vonatkozásban szinte semmilyen veszéllyel nem jár, az időszori értékek alacsonyabb frekvenciára transzformálása, az autoregresszív együttható nagyságától (praktikusan a folyamat várható értékétől) függetlenül minimális másodfajú hibával jár. (A $k=21$ esetben előforduló 1% körüli másodfajú hiba arány valószínűleg sokkal

inkább az ADF-próbával tesztelendő idősor rövidségének, mint a jelenség drasztikus változásának tudható be.)

A szimuláció elvégezhető fordított kérdésfeltevés mellett is: egységgyököt (0-tól különböző mértékű eltolás esetén sztochasztikus trendet) tartalmazó folyamatok esetén milyen arányban fordul elő fals negativitás, vagyis elsőfajú hiba, az ADF-próba alkalmazása esetén. (Az eredmények értékelése során nem feledkezhetünk meg arról, hogy a döntéseket 5%-os szignifikancia-szinten hoztam, vagyis az 1000 esetből előforduló 50 hibás döntés a próba természetes velejárója!)

6-3. táblázat: Hibásan stacionáriusnak vélt RW-folyamatok száma
1000 szimuláció esetén

Aggregálás rendje	μ			
	0,00	0,01	0,02	0,05
$k = 3$	64	63	54	67
$k = 4$	66	70	65	68
$k = 5$	67	73	51	63
$k = 12$	61	80	57	70
$k = 21$	75	90	72	79

A 6-3. táblázatból viszonylag könnyen leolvasható, hogy az első fajú hiba előfordulási gyakorisága minimálisan meghaladja az elméleti 5%-os arányt, ám a jelenség korántsem olyan mértékben zavaró, hogy nagy erőfeszítéseket kellene tennünk a kiküszöbölése érdekében.

Az egyváltozós idősor-elemzés számára általánosan megfogalmazható megállapítás tehát az, hogy az időbeli aggregálás ez esetben *nem hordoz jelentős veszélyt*.

6.6.2 Okság megítélésének torzulása aggregált idősorok esetén

Az elméleti fejtegetések során már említettem, hogy az ok-okozati viszonyok megítélése tekintében már nem ennyire kedvező a kép, a szakirodalom több helyen is arról értekezik, hogy az időbeli aggregálás jelentősen gyengíti az okság-teszteket. Ebben az alpontban előbb stacionárius, egymással kölcsönösen Granger-okságban álló idősorok, majd kointegrált rendszerek esetén vizsgálom meg, hogy a különböző rendű aggregálás mennyiben befolyásolja a valós összefüggések felismerhetőségét. (Bevallom, némiképp furcsállom, hogy a 6.1 táblázat alapján Granger szerint az kointegráltság invariáns az időbeli aggregálásra, ugyanakkor Sims szerint az okság nem... A fejezetben természetesen nem megkérdőjelezni, csak árnyalni kívánom a Nobel-díjas tudósok megállapításait!)

Elsőként tekintsük a sokszor alkalmazott kétváltozós, konstans nem tartalmazó VAR-modellt. Emlékeztetőül

$$\mathbf{y}_t = \begin{bmatrix} y_{1t} \\ y_{2t} \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} \beta_{11} & \beta_{12} \\ \beta_{21} & \beta_{22} \end{bmatrix} \quad \boldsymbol{\varepsilon}_t = \begin{bmatrix} \varepsilon_{1t} \\ \varepsilon_{2t} \end{bmatrix}$$

$$\mathbf{y}_t = \mathbf{B}\mathbf{y}_{t-1} + \boldsymbol{\varepsilon}_t$$

ahol a két véletlen változó egymástól független fehér zaj (ez utóbbi feltételezés egyébiránt sokszor sérül, ami identifikációs problémákhoz vezethet). A szimulációk során három különböző együtt-ható-mátrixszal dolgoztam, hogy képet alkothassak arról is, vajon függ-e a hibás döntések aránya attól, hogy milyen értékeket vesznek fel a modell együtt-hatói. A 6-4. táblázatban a fals negatív esetek száma látható, a szokásos 1000 szimulációt követően.

6-4. táblázat: Hibásan okság hiányát találó Wald-próbák száma
1000 szimuláció esetén

Aggregálás rendje	B		
	$\begin{bmatrix} 0,9 & 0,4 \\ -0,4 & 0,9 \end{bmatrix}$	$\begin{bmatrix} 0,6 & -0,3 \\ -0,3 & 0,6 \end{bmatrix}$	$\begin{bmatrix} 0,7 & 0,2 \\ 0,2 & 0,7 \end{bmatrix}$
$k=3$	0	7	18
$k=4$	0	119	137
$k=5$	0	385	334
$k=12$	337	932	927
$k=21$	804	944	936

Az eredmények nemcsak szembeötlőek, hanem rendkívül fontosnak is tűnnek:

- az aggregálás *drasztikusan csökkenti a Wald-próba erejét*, már viszonylag alacsonyrendű össze-vonások (például negyedévről évre áttérés) esetén is gyakran előfordulhat, hogy az eredetileg meglévő ok-okozati összefüggés elfedődik,
- a szimulációk alapján kijelenthető, hogy hónapról évre, vagy munkanapról hónapra történő aggregálás *szinte törvénytzerűen vezet az ok-okozati összefüggés fel nem ismeréséhez*,
- mellékes eredmény, de kijelenthető, hogy a VAR-modell együtt-ható-mátrixában szereplő *paraméterektől függ*, hogy milyen mértékben csökken a Wald-próba ereje.

Ezt követően megvizsgáltam, hogy mennyire érzékeny a kointegráltságot vizsgáló Johansen-próba az idősorok aggregálására. A vizsgált kétváltozós modell az alábbi

$$\begin{aligned}
 x_t &= x_{t-1} + \varepsilon_t \\
 y_{1t} &= x_t + \varepsilon_{1t} \\
 y_{2t} &= \alpha x_{1t} + \varepsilon_{2t}
 \end{aligned}$$

ahol a véletlen változók specifikálása a szokásos módon történt. A 2. fejezetben ismertetett okok miatt y_{1t} és y_{2t} között hibakorrekció írható fel, tehát kointegráltak. A szimulációk során változtattam α paraméter értékét, illetve a hibakorrekciós modellben megjelenített késleltetés értékét valamennyi esetben hozzáigazítottam az aggregálást követően kialakult idősorok hosszához.¹⁰⁴

A fals negatív döntések számát mutatja a 6-5. táblázat:

6-5. táblázat: Hibásan közös trend hiányát találó Johansen-próbák száma
1000 szimuláció esetén

Aggregálás rendje	α			
	0,5	1	2	5
$k=3$	5	3	5	10
$k=4$	9	16	12	12
$k=5$	18	16	17	18
$k=12$	133	146	150	137
$k=21$	199	229	229	220

A közös trend vizsgálata során egyértelmű megállapítások tehetők:

- nincs jelentős különbség az eredményekben, abban a tekintetben, hogy a kointegráló regresszióban milyen értéket vesz fel a regressziós együttható,
- minél magasabb rendben zajlik, az aggregálás annál inkább fenyeget annak a veszélye, hogy az eredeti adatgeneráló folyamatban meglevő kointegráltságot a Johansen-teszt alapján nem ismerjük fel, mindez (a hibázás veszélye) már a hónapról évre áttérés esetén is aggasztó mértékű lehet.

¹⁰⁴ Mindez azt jelenti, hogy míg – akárcsak korábban – 1000 elemű szimulált idősorok esetében 40-es késleltetéssel, addig a harmadrendű aggregálást követően 12-es, a negyedrendű aggregálás után 10-es, s.i.t. késleltetést használtam.

Összefoglalva az időbeli aggregálásra vonatkozó szimulációs eredményeimet három markáns megállapítás tehető:

1. Az idősorokban meglevő trend felismerését nem zavarja meg az időbeli aggregálás, és stacionárius idősorok esetén sincs jelentős valószínűsége annak, hogy hibásan egységgyököt 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.

A következő alfejezetben konkrét példával illusztrálom az előbbi téziseket.

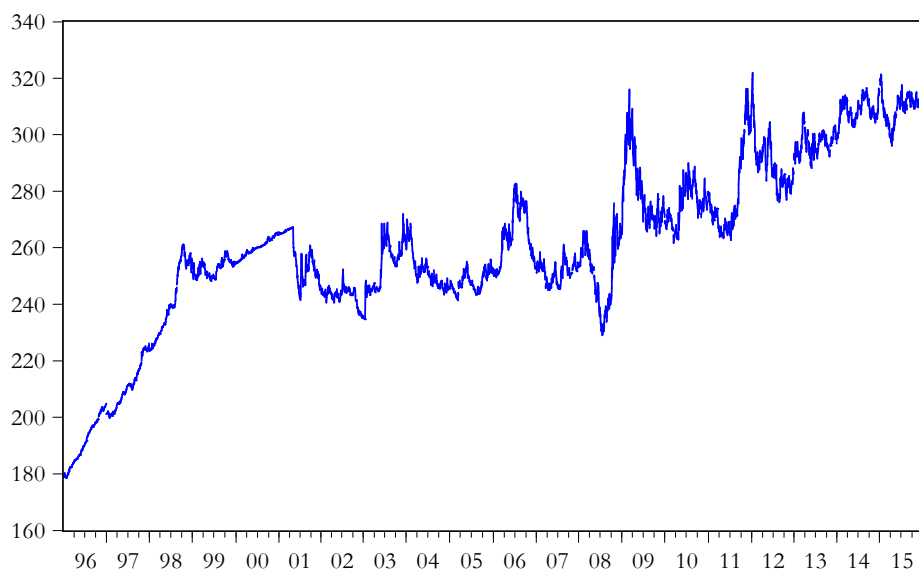
6.7 Árfolyam-idősorok és összefüggéseik az elmúlt 2 évtizedben

Illusztratív példámban az MNB által közölt devizaárfolyamokat fogom vizsgálni az elmúlt 20 év viszonylatában (forrás: www.mnb.hu). Alapvetően három árfolyam érhető el minden különösebb nehézség nélkül:

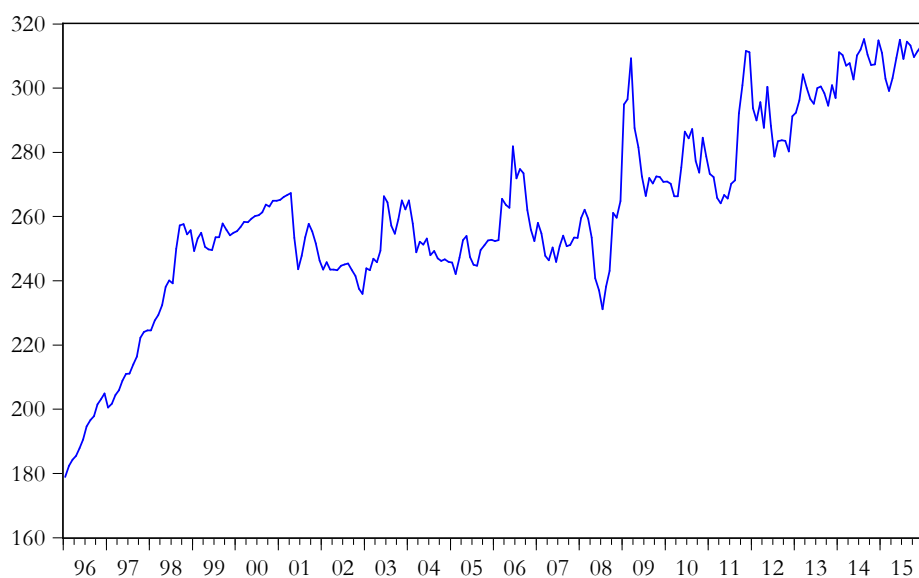
- napi MNB középárfolyam,
- havi záróárfolyam, illetve
- havi átlagos árfolyam.

Valamennyi árfolyam-adat forintban mért és azt mutatja, hogy a vonatkozó deviza (svájci frank, euró, stb.) egy egysége¹⁰⁵ hány forintot ér. A vizsgálat fókuszában az áll, hogy a különböző megfigyelési gyakoriságok, illetve – a havi átlagárfolyam esetében – az aggregálás milyen hatással van az árfolyamok alapvető idősoros tulajdonságaira, illetve az együttmozgásukra. A problémát illusztráló tekintsük az alábbi, 6-2.a-c. ábrákat, melyek a három euró-árfolyam idősort mutatják:

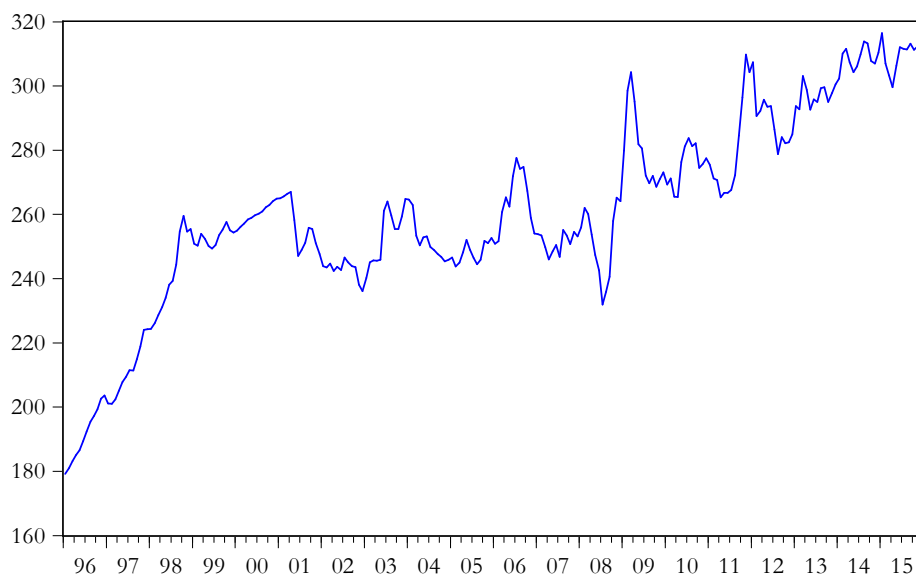
¹⁰⁵ Öt árfolyamot vizsgálok, melyek közül a svájci frank (CHF), az euró (EUR), az angol font (GBP) és az USA dollár (USD) esetében valóban egy egységről van szó, az ötödik árfolyamnál (japán yen, JPY) a viszonyítási egység tradicionálisan 100 JPY.



6-2.a ábra: Az euró forintban mért árfolyama (1996-2015, napi záróárfolyamok)



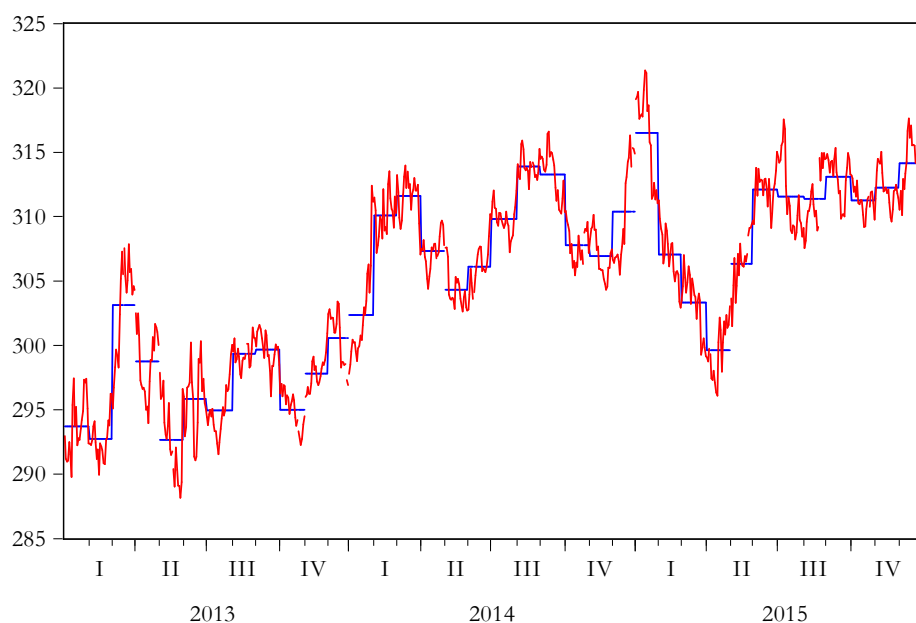
6-2.b ábra: Az euró forintban mért árfolyama (1996-2015, havi záróárfolyamok)



6-2.c ábra: Az euró forintban mért árfolyama (1996-2015, havi átlag-árfolyamok)

Valószínűleg nem elhamarkodott gondolat, hogy pusztán az ábrák alapján jelentős különbség nem látszik az egyes idősorok között, még nagy rutinnal rendelkező idősor-modellezők sem vállalkoznának arra, hogy a grafikonok alapján különböző adatgeneráló folyamatokat vizionáljanak az idősorok mögé.

Valóban csak az illusztráció kedvéért nézzük meg egy rövidebb, 3 éves periódust, amikor ugyanazon az ábrán szerepel a napi, illetve a havi átlag-árfolyam:



6-3 ábra: Az euró forintban mért árfolyama (2013-2015, napi, illetve havi átlag-árfolyamok)

Láthatjuk, amit tudtunk is, hogy a napi árfolyamok volatilitása lényegesen magasabb, a havi átlag-árfolyamok kisebb terjedelemben szóródnak.

Kézenfekvőnek tűnik a kérdés, vajon előfordulhat-e, hogy a napi árfolyamokban megjelenő nagyobb variancia elfedi az idősor(ok)ban esetlegesen meglevő sztochasztikus trendet? Különösen érdekes a felvetés abban az esetben, ha a napi árfolyamok közül csak egy-egy kitüntetettet, a hónap utolsó napjának árfolyamát vizsgáljuk, ilyenkor ugyanis sokszor még a szaksajtó szóhasználat sem tesz különbséget a két idősor között, és a deviza-árfolyam havi bontású idősoráról ír.

A 6-6. táblázat az öt kiválasztott (egyébiránt a hazai devizakereskedelem túlnyomó hányadát kitevő) deviza-árfolyamra vonatkozóan közli a kiterjesztett Dickey-Fuller-próba eredményeit:

6-6. táblázat: Néhány vezető deviza forintban mért árfolyamára vonatkozó ADF-próba eredmények (1996-2015)

Deviza	Havi záró		Havi átlag		Napi	
	τ	p -érték	τ	p -érték	τ	p -érték
CHF	-1,698	0,750	-1,184	0,911	-2,336	0,414
EUR	-3,584	0,033	-3,869	0,015	-3,758	0,019
GBP	-2,771	0,209	-2,526	0,316	-2,834	0,185
JPY	-2,327	0,417	-2,465	0,345	-2,497	0,330
USD	-1,893	0,655	-1,932	0,635	-2,044	0,576

Láthatjuk, hogy a fejezet korábbi részében taglalt elméleti összefüggések igazolást nyertek, vagyis az idősorok megfigyelési gyakorisága *nem módosítja az egységgyök meglétére vonatkozó ítéletünket*. Az euró esetében az ADF-próba nem igazolta az egységgyök jelenlétét, vagyis az idősort stacionáriusnak tekinti, a másik négy deviza esetében sztochasztikus trendet találtam, függetlenül attól, hogy a napi, a havi záró, illetve a havi átlagos értékeket tartalmazó idősort vizsgáltam.

Az euró árfolyam esetében, különösen a korábbi ábráink alapján némiképp meglepőnek tűnik, hogy a próba elvetette az egységgyök létét, ezért elvégeztem a vizsgálatot 10 éves periódusokra is.¹⁰⁶ Az alkalmanként 120 megfigyelést tartalmazó idősorokra vonatkozó ADF-próba eredményeket a 6-7. táblázat tartalmazza:

¹⁰⁶ Természetesen nem feledkezhettünk meg arról, hogy a mintanagyság változása befolyásolja a próba erejét!

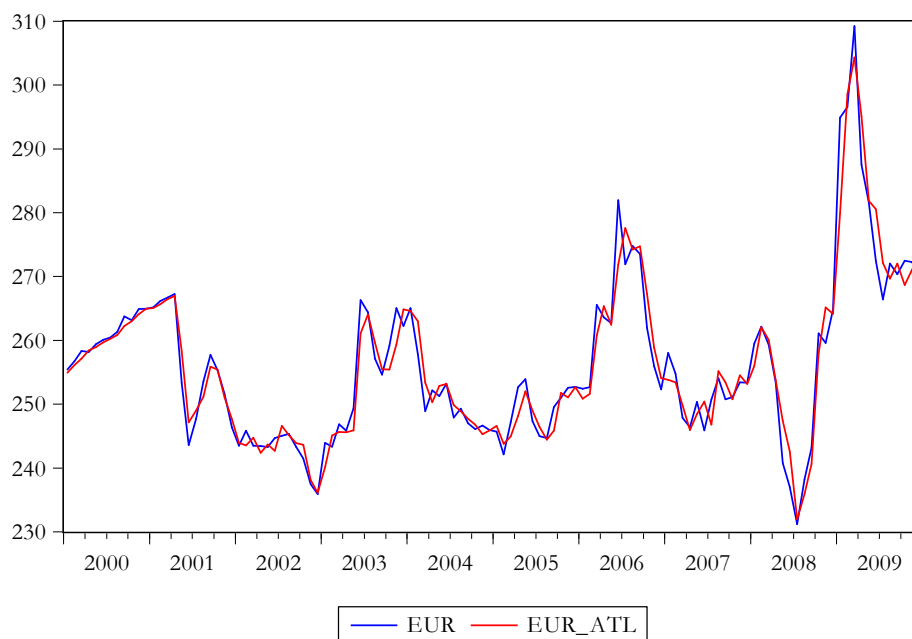
6-7. táblázat: Az euró forintban mért árfolyamára vonatkozó ADF-próba eredmények
(különböző 10 éves periódusok, havi bontás)

Időszak	Havi záró		Havi átlag	
	τ	p -érték	τ	p -érték
1996-2005	-2,221	0,473	-2,137	0,520
1997-2006	-2,439	0,358	-2,369	0,394
1998-2007	-3,444	0,051	-3,605	0,034
1999-2008	-3,188	0,092	-4,044	0,010
2000-2009	-2,897	0,167	-3,667	0,028
2001-2010	-3,300	0,071	-4,029	0,010
2002-2011	-3,886	0,016	-4,114	0,008
2003-2012	-3,368	0,061	-4,016	0,011
2004-2013	-3,629	0,031	-4,239	0,005
2005-2014	-3,503	0,044	-4,112	0,008
2006-2015	-3,589	0,035	-4,175	0,007

A próbák eredménye alapján két, talán némiképp meglepő megállapítást tehetünk:

- a 10 éves idősorok közül jó néhány esetben az ADF-próba alapján nem vethetjük el az egységgyök-hipotézist, a '90-es évek végén bekövetkezett árfolyam-emelkedés periódusában feltétlenül trendet vélelmezhetünk,
- mondanivalóm szempontjából sokkal érdekesebb, hogy a 10 éves időhorizonton végzett elemzések esetén már többször előfordul, hogy a havi záró, illetve a havi átlag-árfolyamok alapján végzett egységgyök-teszt *azonos* időszakon *különböző* eredményre vezet (különösen szembeötlő ez a 2000-2009-es periódusban).

Érzékeltetendő, hogy az előbbi kijelentés cseppet sem triviális, tekintsük a 2000-2009-es időszak árfolyamainak grafikonját:



6-4 ábra: Az euró forintban mért havi záró (EUR), illetve havi átlag-árfolyama (EUR_ATL) a 2000. január és 2009. december közötti időszakban

A 6-4. ábrából egyetlen dologra hívom fel a figyelmet: a két idősor szinte teljesen azonosnak tűnik, a rájuk vonatkozó egységgyök-teszt eredményekből mégis homlokegyenest más következtetések vonhatók le. (Valóban csak zárójelben jegyzem meg: a havi átlagárfolyam idősora alapján a vizsgált időszakban az euró/forint árfolyam stacionáriusnak tűnik, miközben a havi záróárfolyamok alapján trendet vélünk kirajzolódni. Mindez, összevetve például a devizahiteles probléma felelősségi kérdéseivel merőben meglepő!)

Az egyváltozós idősor-elemzés (stacioner, illetve egységgyököt tartalmazó folyamatok megkülönböztetése) ráadásul az elméleti fejtegetések szerint viszonylag egyszerűen megoldható. Minden korábban bemutatott teória szerint sokkal zavarosabb szituációra kell számítanunk, ha idősorok együttmozgását vizsgáljuk. Illusztratív példában mindössze az euró-árfolyam és a többi fontos deviza árfolyama közötti együttmozgást teszteltem, mintaidőszakként az előbb már bemutatott 2000 és 2009 közötti 10 évet választottam, ugyanis a frank, font, yen és dollár idősorok gyakorlatilag az egész 20 éves időtartamban integrálnak tűnnek, miközben az euró (mint korábban láthatuk) egyes szakaszokban stacionárius, így a kointegráltság tesztje ilyenkor aggályos.

A 6-8. táblázat a kointegráció hiányát jelentő nullhipotézist tesztelő Johansen-próba szignifikancia értékeit mutatja, egyrészt a havi záró, másrészt a havi átlag-árfolyamok között:

6-8. táblázat: Az euró és más fontos devizák kointegráltságát tesztelő Johansen-próba szignifikancia értékei (2000-2009, havi bontás)

<i>Deviza</i>	<i>Havi záró</i>	<i>Havi átlag</i>
CHF	0,026	0,072
GBP	0,506	0,521
JPY	0,018	0,073
USD	0,014	0,081

Az eredmények mind modellezési, mind pénzügyi szempontból meglepőek:

- a Johansen-próba alapján nem található együttmozgás az euró és a font között, miközben az euró a másik három árfolyammal *lehet*, hogy közös trendet tartalmaz,
- az előbbi állításban a bizonytalanságot az indokolja, hogy pl. 5%-os szignifikancia-szint mellett másképp ítéljük meg a teszt eredményét a záróárfolyamok, illetve az átlagárfolyamok alapján.

Anélkül, hogy mélyebb pénzügytani fejtegetésekbe bonyolódnék, kimondható, hogy a közös trend megléte nagyjából azt jelenti, hogy nem alakul ki arbitrázs-lehetőség. Mindez megfordítva: kointegráció hiánya esetén nem zárható ki az arbitrázs lehetősége. A fenti eredmények alapján adott napon (a hónap utolsó napján) nincs arbitrázs-lehetőség, ugyanakkor a havi átlagárfolyamok nem feltétlenül bolyonganak együtt.

Dolgozatom szempontjából a példából levonható szintén fontos összefüggés, hogy azonos (vagy közel azonos) elnevezésű, de némiképp eltérő tartalmú idősorok merőben különböző adatgeneráló folyamatból származhatnak, vagyis az aggregálás következtében az adekvát DGP is változhat. Ezzel a kérdéssel a következő, etikus modellezést körbe járó fejezetben még foglalkozom.

7 Konklúziók helyett az etikus idősor-modellezésről

A doktori értekezés végén szokásos konklúziók helyett az utolsó fejezetben inkább a dolgozat írása közben, pontosabban az elmúlt 30 év során szerzett modellezői tapasztalataimról írok. Úgy gondolom, hogy az első fejezetben már említett, közgazdaságtan versus közgazdaságtudomány vita napjainkra némiképp más hangvételt vett: egyesek kifejezetten ellenségként tekintenek az ökonometriai modellezésre, mint olyan szükséges rosszra, ami csak a publikálás során elvárt, de valós eredménnyel nem jár, mások alapvetően tudománytalannak tartanak bármilyen empirikus adatelemzést nélkülöző eredményt, miközben a – reményeim szerint – nagy többség hallgat, vagy talán nem eléggé exponálja magát.¹⁰⁷ 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 megfontolásait.

Nehéz megmondani, hogy mi az oka annak, hogy a statisztikai modellezés eredményeit ilyen mértékű fenntartásokkal kezelik. Nyilván sok esetben az adatokkal szemben megnyilvánuló hitetlenkedés, vagy a bonyolult módszerek meg nem értéséből fakad a statisztikai modellezéssel szembeni ellenérzés, az ilyen averziók valószínűleg kiküszöbölhetetlenek. Dolgozatom zárófejezetében azonban nem ezzel foglalkozom, hanem olyan kérdéseket vettem számba, amelyek az egyébként *módszertanilag* korrektül dolgozó modellezőnek róhatók fel.¹⁰⁸ 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 fejezetben tárgyalt problémák *illusztrálására* alapvetően a magyar államháztartás egyenlegének idősorát használom, vagyis továbbra is az idősor-elemzés témakörében maradok, ugyanakkor előre kívánom bocsájtani, hogy az itt vizsgált etikai vétségek egy része keresztmetszeti, vagy paneladatoknál éppígy elkövethetők. A tárgyalt kérdéskörök:

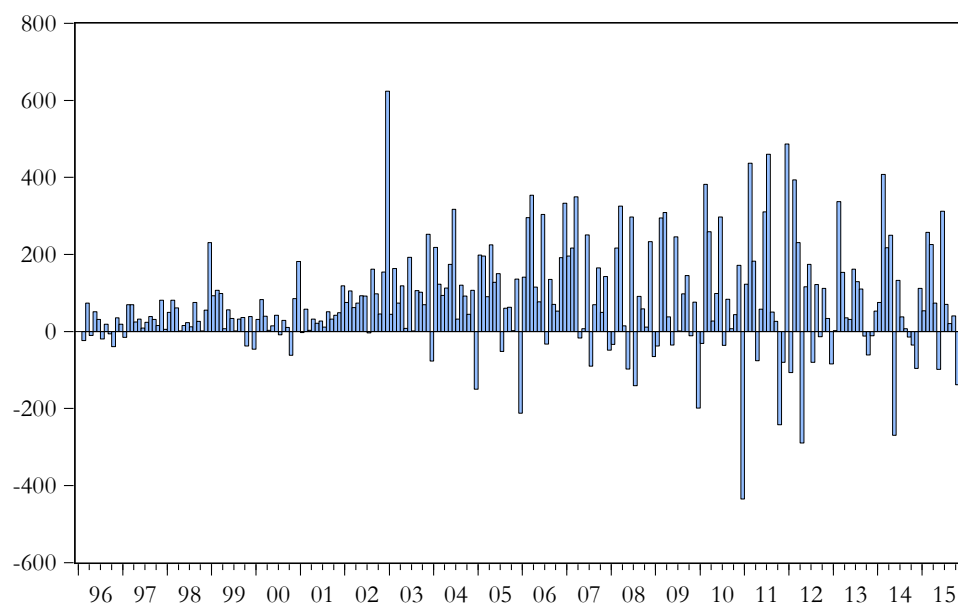
- a vizsgált jelenség tartalmának leírása,
- az ergodicitás és a stacionaritás viszonyának tisztázása,

¹⁰⁷ A kérdésről lásd pl. Mellár (2016) tanulmányát, illetve az erre reflektáló Kőrösi (2016) cikket. Noha nem az említett vitához kapcsolódik (hiszen évekkal megelőzte azt), de érdemes a statisztika (adatgyűjtés, adatelemzés) és az előrejelezhetőség (előrelátás) kapcsolatáról érdemes elolvasni Besenyei már-már filozofikus mélységű tanulmányát (lásd Besenyei, 2011)!

¹⁰⁸ Nem állítom, hogy sikerült minden ilyen problémát összegyűjtenem (bár törekedtem rá), éppen ezért kifejezetten kérem az Olvasót, hogy ha tudja, bővítse a felsorolást!

- a mintaidőszak megválasztása,
- az eredeti adatsor „kilúgozásával” keletkező problémák,
- az eredmények értelmezési korlátainak (pl. véletlen szerepe, vagy ceteris paribus elv) nem kellő súllyal történő hangsúlyozása.

Az alábbi ábrán a fejezetben mindvégig illusztrációként szereplő államháztartási egyenleg idősora látható a szokásos deficit szemléletben, vagyis a pozitív számok a deficitet (negatív egyenleg), míg a negatív számok a szufficitet mutatják:

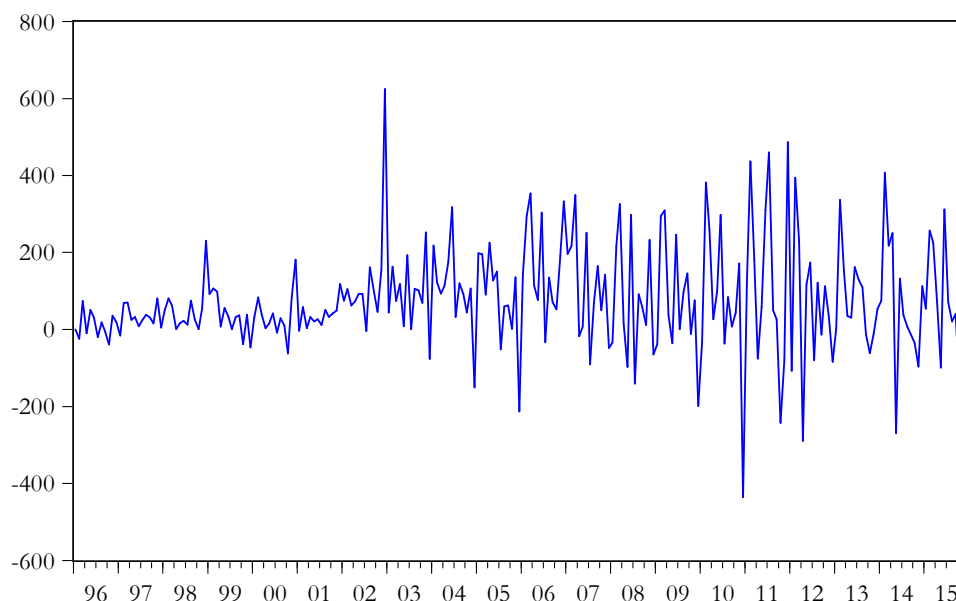


7-1. ábra: Az államháztartás deficit nagysága 1996. január – 2015. december
(havi bontás, folyó áron, Mrd Ft)

Ebben a zárófejezetben alapvetően egyetlen kérdést vizsgálok: vajon stacionárius-e az egyenleg, vagy trendet tartalmaz. Ismételten nem az a célom, hogy eddig ismeretlen gazdaságpolitikai összefüggéseket, vagy rezsimváltásokat tárjak fel, ehelyett megmutatom, hogy milyen etikai kérdések merülhetnek fel egy triviálisnak tűnő kiterjesztett Dickey-Fuller-próba, vagy KPSS-teszt alkalmazása esetén.

A fenti 7-1. ábrán egyértelműen egy tartam-idősor látható (hiszen a deficit az adott hónapban elsejétől a hó utolsó napjáig keletkezik) ezért ábrázoltam – Hunyadi (2002) útmutatásának megfelelően – oszlop-diagrammal a jelenséget. Ugyanakkor annak érdekében, hogy a dolgozat valamennyi korábbi ábrájával összehasonlítható grafikonokat kapjunk az ábrát újraserkesztettem, majd ezt és a későbbi ábrák mindegyikét vonaldiagramként mutatom be, ami akár már az első etikai vétség kategóriájába tartozhat...

Tehát az államháztartási egyenleg deficit-szemléletben az alábbi módon alakult az elmúlt két évtizedben:



7-1.a ábra: Az államháztartás deficit nagysága 1996. január – 2015. december
(havi bontás, folyó áron, Mrd Ft)

Az ábra alapján nehéz egyértelműen állást foglalni, a várható érték stacionaritás többé-kevésbé valószínűsíthető, ugyanakkor a deficit nagyságának ingadozása (az idősor volatilitása) a mintaidőszak végéhez közeledve növekedni látszik. Elvégezve a szokásos próbákat (tehát kiterjesztett Dickey-Fuller-tesztet az egységgyök létére, illetve KPSS-tesztet a stacionaritásra) az alábbi eredményeket¹⁰⁹ kapjuk:

7-1. táblázat: ADF és KPSS teszt az államháztartási deficit idősorára

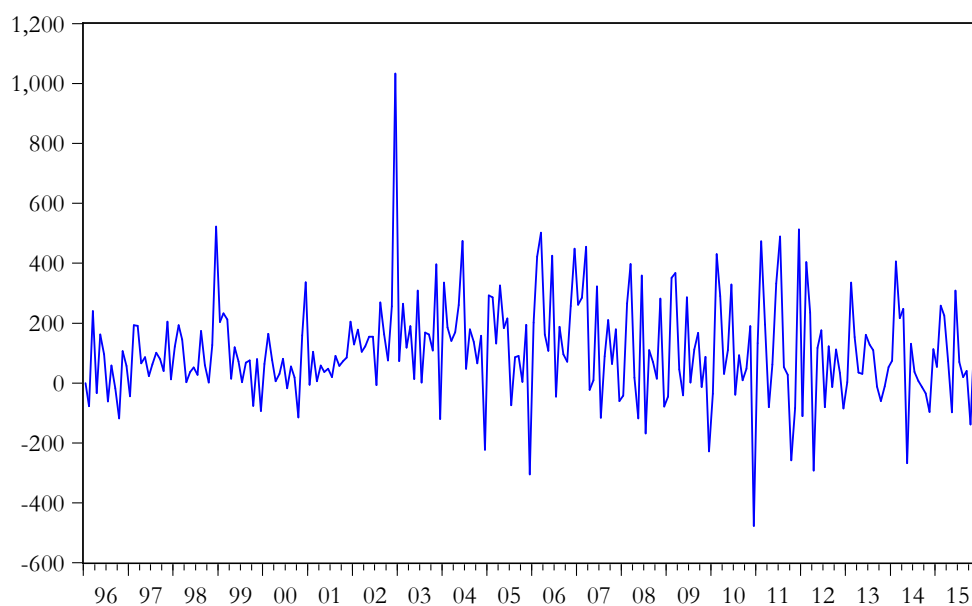
Modell-feltevés	ADF			KPSS	
	próbafüggvény	p-érték	lag	próbafüggvény	p-érték
eltolást tartalmaz	-15,349	0,000	0	0,631	$0,01 < p < 0,05$
eltolást és determinisztikus trendet is tartalmaz	-15,539	0,000	0	0,275	$p < 0,01$

Láthatjuk, hogy korábbi várakozásunkat a próbák is alátámasztották: az egységgyök megléte az ADF-próba alapján elvethető, ugyanakkor – valószínűsíthetően a variancia nem konstans volta miatt – az

¹⁰⁹ Az alábbi, illetve a fejezet további táblázataiban a kiterjesztett Dickey-Fuller próbánál a modellben szereplő késleltetés optimális mértékét a Schwarz-féle információs kritérium (BIC) alapján határoztam meg. A stacionaritás, illetve sztochasztikus trend létének tesztelését csak eltolást, illetve eltolást és determinisztikus trendet is tartalmazó alapmodell esetén is vizsgáltam (lásd (2.28), illetve (2.29) modell).

idősor stacionaritása nem igazolható megnyugtatóan. Érdekes felfigyelnünk arra, hogy az optimális késleltetés mértékét mindkét esetben 0-nak választotta az eljárás, vagyis a DF-próba kiterjesztése nem tűnt indokoltnak. A determinisztikus trendet is tartalmazó modelleknél a trendváltozóhoz tartozó paraméterre vonatkozó parciális t-próba szignifikancia értéke rendre 0,065, illetve 0,058, vagyis a modellbecslések alapján *a determinisztikus trend megléte nem zárható ki* egyértelműen. Rögtön felmerülhet az első etikai jellegű kérdés: hibázik-e a modellező, ha pl. csak az ADF-próba eredményét közli és veszi figyelembe, és a későbbiekben az idősorra, mint stacionárius változóra tekint? A 2. fejezetben írottak alapján – megítélésem szerint – elvárható a korrekt modellezőtől, hogy az óvatosság elve alapján ne nyugodjon meg, és a deficit-idősor stacionárius jellegének megállapítása érdekében további vizsgálatokat végezzen!

Kézenfekvőnek látszik, hogy egy 20 éves időhorizontot felölelő pénzügyi idősor esetén az árváltozással valamit kezdeni kellene, ezért létrehoztam a *deflált* (változatlan áras idősort is). Deflátorként – jobb híján – a kumulált fogyasztói árindexet használtam, miközben az államháztartás felhasználási oldalán megjelenő kiadási tételek nem feltétlenül csak a fogyasztói árváltozásoktól terheltek. A 2015-ös árakon számított, változatlan áras idősor képe az alábbi:



7-2. ábra: *Az államháztartás deficit nagysága 1996. január – 2015. december*
(havi bontás, 2015-ös változatlan áron, Mrd Ft)

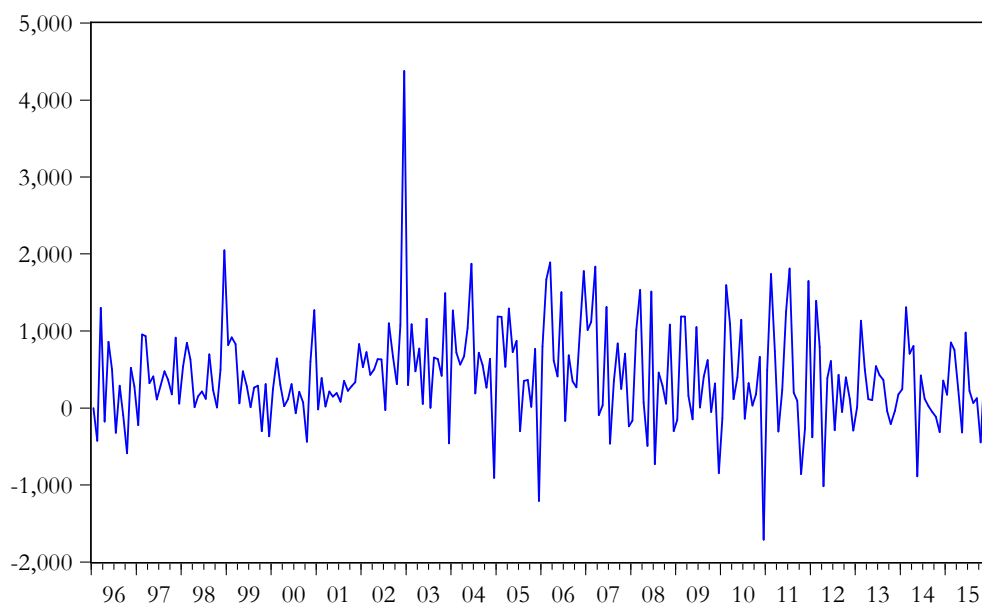
Első ránézésre nyilván nehéz különbséget felfedezni a folyó, illetve változatlan áras idősor között, ám a vonatkozó próbák némi tulajdonságbeli eltérést mutatnak.

7-2. táblázat: ADF és KPSS teszt a változatlan áras államháztartási deficit idősorára

Modell-feltevés	ADF			KPSS	
	próbafüggvény	p-érték	lag	próbafüggvény	p-érték
eltolást tartalmaz	-15,268	0,000	0	0,285	$0,1 < p$
eltolást és determinisztikus trendet is tartalmaz	-15,238	0,000	0	0,284	$p < 0,01$

A második modellvariáns mindkét próba esetén a determinisztikus trend hiánya mellett foglal állást ($p > 0,1$), vagyis – a próbák alapján – kijelenthető, hogy a *változatlan áras deficit* idősora *nem tartalmaz sem sztochasztikus, sem determinisztikus trendet és stacionáriusnak tekinthető*.

Folytassuk az adatsor elméleti megfontolásokon alapuló transzformálását! Kis, nyitott gazdaságok esetén szokásos modellfeltevés, hogy az államháztartási egyenleg nagysága nem független a hazai deviza árfolyamától, ezért célszerű a makropénzügyi adatokat valamilyen kevésbé volatilis devizában – Európában, illetve az Európai Unióban kézenfekvően *euróban* – nyilvántartani. A változatlan áras, euróban elszámolt államháztartási deficit nagysága a 7-3. ábrán látható módon¹¹⁰ alakult:



7-3. ábra: Az államháztartás deficit nagysága 1996. január – 2015. december
(havi bontás, 2015-ös változatlan áron, millió euró)

¹¹⁰ Ügyeljünk a léptékre! Eddig havi „néhány száz” milliárd forintos, most havi „egy-két ezer” millió eurós tételekről van szó, amiben – figyelembe véve, hogy az €/Ft árfolyam a vizsgált időszakban a 230-320 Ft-os sávban ingadozott – semmi meglepő nincs.

Az immár megszokottnak tekinthető próbák eredménye:

7-3. táblázat: ADF és KPSS teszt a változatlan áras, euróban mért államháztartási deficit idősorára

Modell-feltevés	ADF			KPSS	
	próbafüggvény	p-érték	lag	próbafüggvény	p-érték
eltolást tartalmaz	-15,168	0,000	0	0,376	$0,05 < p < 0,1$
eltolást és determinisztikus trendet is tartalmaz	-15,191	0,000	0	0,295	$p < 0,01$

A determinisztikus trend meglétét egyik teszttel sem sikerült igazolni, vagyis a próbák egy *trendet* (növekvő tendenciát) *nem tartalmazó jelenséget* mutatnak, amelyben a *variancia-stacionaritás nem feltétlenül igazolható*.

Ne feledjük, hogy céloom nem az államháztartás egyenlegének alakulását legjobban leíró időszori modell megtalálása volt, hanem annak bemutatása, hogy a vizsgált jelenség operacionalizálása során alkalmazott idősor(ok) nem kellően körültekintő definiálása akár jelentősen különböző megállapításokat is eredményezhet, egy még oly egyszerű kérdésben, mint a trend megléte is. Miben rejlik a modellező etikussága ebben a tárgykörben? Véleményem szerint két dologban:

1. egészen *pontosan definiálni* kell a vizsgálandó jelenségek (idősorok) tartalmát, az adatok forrása mellett kitérve az árazás, a mértékegység, az esetlegesen elvégzett transzformációk során alkalmazott segédváltozók körére is;
2. amennyiben kézenfekvően több „első ránézésre” azonos tartalmú és főleg azonosan elnevezhető változó áll rendelkezésre, akkor fel kell hívni a felhasználó figyelmét az esetlegesen kimutatható tulajdonság-különbségekre, és *részletesen le kell írni* a további elemzés céljára kiválasztott *változó pontos tartalmát és a kiválasztás okát*.

A fenti második szabály értelmében jelzem, hogy a dolgozat hátralevő részében a *változatlan áras, euróban mért államháztartási deficit* idősorával dolgozom, amely tehát – lásd korábban – havi bontást vizsgálva *nem tartalmaz sem determinisztikus, sem sztochasztikus trendet és a szokásos 5%-os szignifikanciaszint mellett stacionáriusnak tekinthető*.¹¹¹

A sztochasztikus idősor-modellezés módszertanának, főképpen az ilyen modellek alapján készülő előrejelzések ellenzői legtöbbször azzal érvelnek, hogy a jelenségek nem ergodikusak, így egy múltbeli időperiódusból nem lehet következtetni az idősor jövőbeli alakulására. Az ergodicitás definíciója szerint egy sztochasztikus folyamat *ergodikus*, ha az egyetlen realizált idősorból (mintá-

¹¹¹ Noha a dolgozatban korábban már jeleztem, annak érdekében, hogy az első szabály se sérüljön, megismétlem, hogy az adatok forrása a Központi Statisztikai Hivatal által kiadott, a „KSH jelenti” című korábban havi bontású, ma már negyedéves kiadvány, melyet „gördülő” módon használtam, ügyelve arra, hogy valamennyi adatrevíziót visszamenőlegesen is átvezessek.

ból) becsült momentumok az idősor hosszának növelésével konvergálnak a folyamat (alapsokaság) momentumaihoz. A várható értékében (átlagában) ergodikus adatgeneráló folyamatra igaz, hogy

$$\lim_{T \rightarrow \infty} E \left[\left(\frac{1}{T} \sum_{t=1}^T y_t - \mu \right)^2 \right] = 0$$

valamint egy varianciájában (is) ergodikus DGP esetén teljesül a

$$\lim_{T \rightarrow \infty} E \left[\left(\frac{1}{T} \sum_{t=1}^T (y_t - \mu)^2 - \sigma_y^2 \right)^2 \right] = 0$$

feltétel.

Az ergodicitás fogalmával az egyetlen „probléma”, hogy a hagyományos statisztika eszközeivel nem tesztelhető. A kritikusok (lásd pl. Mellár, 2016) itt megállnak, és kijelentésüket, miszerint idősorok múltbeli értékeiből prognózis nem készíthető bizonyítottan vélik. Ezzel szemben a modern idősor-kutatás alaptétele, miszerint amennyiben egy idősor *stacionárius*, akkor *ergodikus is*, az előző okfejtést felülírja. A stacionáriuság ugyan a kívántnál erősebb feltétel, vagyis elképzelhető lenne, hogy egy nem stacionárius idősor ergodikus, azaz lehet, hogy a stacionaritás megkövetelésével elveszítünk előrejelzésre alkalmas folyamatokat, ugyanakkor kijelenthető, hogy amennyiben egy idősor stacionárius, akkor abból hatékony előrejelzés készíthető. A gazdaságmodellezők körében elfogadott megállapítás, hogy az általunk leggyakrabban vizsgált jelenségek, vagy ezek logaritmusának első-, esetleg második differenciája stacionáriusnak bizonyul, tehát ergodikus, azaz alapját képezheti jól teljesítő statisztikai modelleknek.

Ebből következően talán a legtöbb, amit a stacionaritás, vagy ergodicitás témakörében a felhasználó számára megmutathatunk, az egy adott empirikus idősor különböző hosszúságú *részperiódusa*-ra vonatkozó várható érték, illetve variancia (esetleg az értelmezhetőség kedvéért szórás), valamint a különböző részmintákra vonatkozó egységgyök-tesztek. Amennyiben az előbbi mutatók, illetve próbastatisztikák állandónak (konstansnak) tűnnek, akkor plauzibilissé tettük az ergodicitás meglétét.

Tekintsük az alábbi, 7-4. táblázatot, amely a teljes 20 éves idősor, és annak rövidebb¹¹² szakaszaira vonatkozóan mutatja meg az előbbieken említett statisztikákat:

¹¹² A szakaszok mindig egész éveket (januártól decemberig) tartalmaznak, és a szakaszhatárok kijelölésekor törekedtem arra, hogy a kormányzati ciklusok, illetve a válság okozta esetleges strukturális törés nyomon követhető legyen.

7-4. táblázat: Az államháztartási deficit havi bontású (2015-ös változatlan áras, euróban mért) idő-sorára vonatkozó alapstatisztikák különböző időintervallumokra

Időszak	Átlagos havi deficit (millió €)	Havi deficit szórása (millió €)	ADF-próba (eltolással és determinisztikus trenddel)			
			empirikus érték	p-érték	lag	trendhez tartozó p-érték
1996 - 2015	422,9	647,0	-15,191	0,000	0	0,360
1996 - 2007	495,3	639,0	-11,902	0,000	0	0,016
1996 – 2002	403,1	614,8	-5,082	0,004	0	0,091
2001 – 2007	613,0	716,3	-9,429	0,000	0	0,555
2005 – 2011	496,0	744,6	-8,435	0,000	1	0,060
2009 - 2015	307,6	633,2	-8,847	0,000	1	0,217

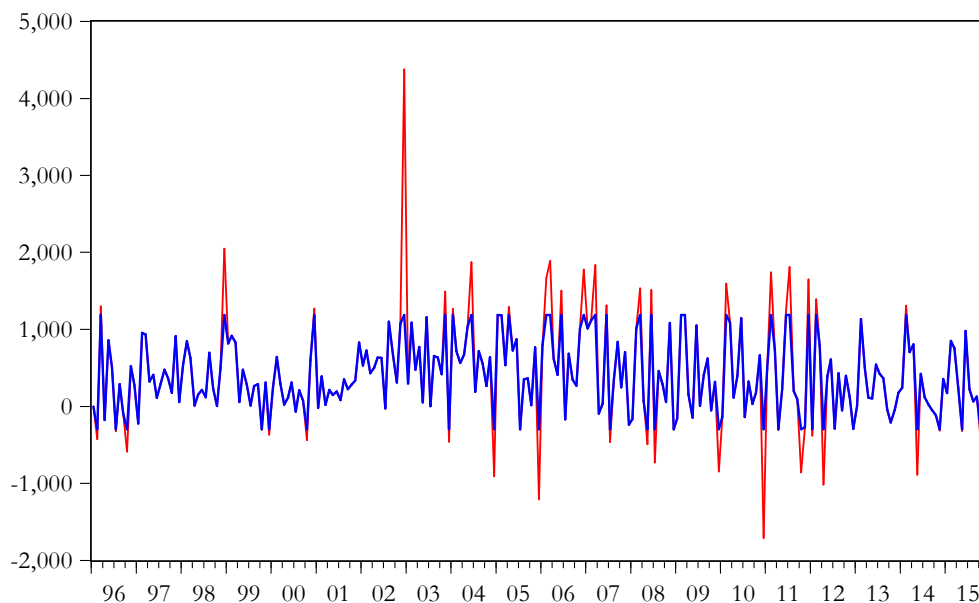
A fenti táblázatban olvasható eredmények nem hiszem, hogy meglepőek lennének bármely empirikus modellező számára, hiszen mit látunk? *Az alapvető momentumok ingadoznak ugyan, de egyáltalán nem tűnik úgy, hogy a rövidebb periódusokra számított átlagok, vagy szórások ne konvergálnának a teljes időszakra számított alapstatisztikákhoz.* A sztochasztikus trend (egységgyök) megítélését tekintve az eredmények egybehangzóak: a folyamat stacionernek tűnik, a determinisztikus trend megítélése ugyanakkor változó, egyes rövidebb időszakokban ennek megléte nem zárható ki. A legvalószínűbb, hogy egy-egy kiugró érték (2002 augusztusában a 100 napos csomag idején, illetve 2008-ban a globális gazdasági és pénzügyi válság kitörését követően), illetve néhány erőteljesebb árfolyam-ingadozás rövidebb periódusokban eltéríti az empirikus idősorra vonatkozó alapvető statisztikákat a DGP-t jellemzőktől, ám ezek az eltérések nem eredményezik az ergodicitás megkérdőjelezhetőségét.

Összességében körülbelül az mondható ki, hogy *amennyiben az idősor hossza ezt lehetővé teszi, akkor végezzük el a legfontosabb próbákat rövidebb szakaszokra is,* és ügyeljünk arra, hogy a szakaszhatárok gazdaságtörténeti (gazdaságelméleti) szempontok alapján védhetőek legyenek.

Az előbbieken felmerült, hogy a szakaszokra bontott empirikus idősor esetén az okozhatta az egyes szakaszokra vonatkozó eltérő teszteredményeket, hogy bizonyos – dolgozatomban részletesen tárgyalt – *időszaki anomáliák torzították* az elemzést. Mivel markáns strukturális törés nem rajzolódik ki a deficit alakulásában, ezért csak arra gyanakodhatunk, hogy néhány kiugró érték befolyásolja az eredményeket. (Logikus gondolat, hogy például a növekvő tendenciát az fedi el, hogy az időszak elején néhány kiugróan magas deficit is előfordult, ami tompítja a várható értékben meg-

levő tendenciát.¹¹³) Az idősor outlierektől való megtisztítása érdekében winsorizáltam (lásd (4.31) képlet), $\alpha = 0,1$ paraméterrel.

Az outlierektől tisztított idősort a 7-4. ábrán vastag kék vonallal jelöltem:



7-4. ábra: Az államháztartás deficit nagysága (outlier-szűrés után) 1996. január – 2015. december
(havi bontás, 2015-ös változatlan áron, millió euró)

Az ábra ismét jól mutatja az outlier-szűrés logikáját: a kiugró értékeket a megfelelő kvantilisokkal helyettesítve az idősor belesimul egy előre kijelölt csatornába, tendencia szinte nem is alakulhat ki... Ezt követi a meglepetés, ugyanis ha megvizsgáljuk az outlier-szűrt idősorra vonatkozó egységgyök-, illetve stacionaritás teszt eredményeket, az alábbi értékeket kapjuk:

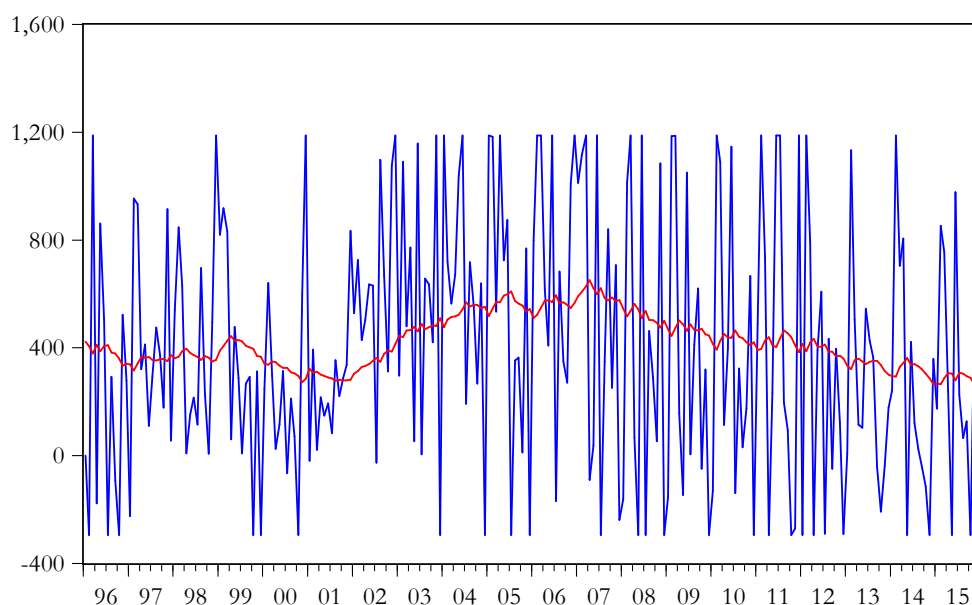
7-5. táblázat: ADF és KPSS teszt az outlierektől tisztított, havi bontású, változatlan áras, euróban mért államháztartási deficit idősorára

Modell-feltevés	ADF			KPSS	
	próbafüggvény	p-érték	lag	próbafüggvény	p-érték
eltolást tartalmaz	-2,167	0,219	11	0,383	$0,05 < p < 0,1$
eltolást és determinisztikus trendet is tartalmaz	-2,333	0,414	11	0,341	$p < 0,01$

¹¹³ Mindemellett nem kívánom elhallgatni, hogy sok mértékadó statisztikus kifejezetten hibásnak tartja ezt a megközelítést. Véleményük szerint nemhogy kiszűrni kellene az outliereket, hanem megbecsülni, hiszen éppen ritkaságuknál fogva sokkal több információt hordoznak, mint a folyamatba jól illeszkedő megfigyelések.

A próbák eredményei azt mutatják, hogy *nem zárható ki a sztochasztikus trend (egységyök) megléte*, ugyanakkor a további eredményekből látszik, hogy *determinisztikus trend nem mutatható ki* és az idő-sorban viszonylag magasrendű autokorreláció is megjelenik.

Rögtön megértjük a trend megítélésére vonatkozó, a korábbiakkal ellentétes eredményt, ha az outlier-szűrt idősort és a belőle exponenciális simítással előállítható, becsült idősort összevetjük (a 7-4.a. ábrát szemlélve ismét ügyeljünk arra, hogy a lépték megváltozott!):



7-4.a ábra: Az állambájtartás deficit nagysága outlier-szűrés után, valamint ennek exponenciálisan simított értékei (havi bontás, 2015-ös változatlan áron, millió euró)

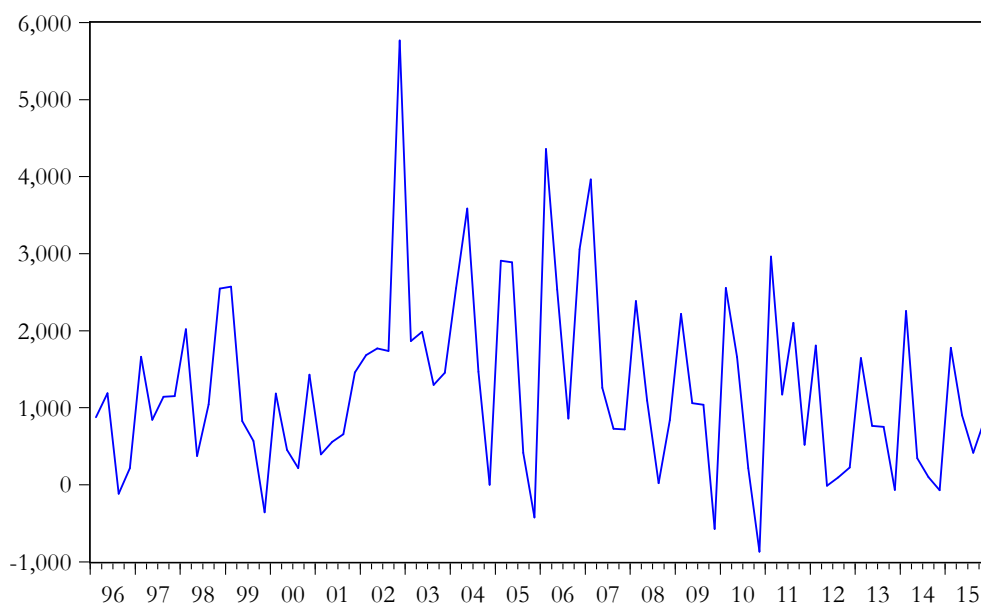
Az ábra alapján kijelenthetjük, hogy a deficit idősorában a körülbelül 2002 közepéig tartó csökkenő trendet 2002 és 2007 között egyértelműen növekvő tendencia váltotta fel, majd a válság kitörésétől kezdve ismét a csökkenő trend vált uralkodóvá.

Ismét nem kívánok a fenti megállapítások gazdaságtörténeti jelentőségével foglalkozni, mindössze arra kívánok rámutatni, hogy az eredeti idősor „kilúgozása” (esetünkben a nem feltétlenül indokolt outlier-szűrés) akár azt is eredményezheti, hogy a tendenciát látunk ott is, ahol eredetileg erre nem utalt semmi. Nagyon fontos tehát kijelentenünk, hogy *az időszori értékeket* (tehát nem a modellt!) *módosító eljárásokkal* csak akkor kísérletezzünk, *ha a statisztikai (módszertani) okon túl más, elsősorban a jelenség természetéből, vagy a jelenséget alakító körülmények változásából eredő outlierekre gyanakszunk!*

A dolgozat elején, a pénzügyi idősorokról szólva, megjegyeztem, hogy modellezés szempontjából a nagy adatsűrűség (frekvencia) jó tulajdonság, ám gyakran – mintegy ellentételezésként – ez magával hozza a nagy varianciát (volatilitást). A különböző időszori anomáliák vizsgálata során többször is tapasztaltuk, hogy az empirikus idősorban meglevő jelentős (az átlagot akár nagyságren-

dekkal meghaladó) szórás sokszor képes elfedni a jelenség legalapvetőbb tulajdonságait is. Joggal merülhet fel ezen a ponton a kérdés (különösen az előző bekezdésekben tárgyalt outlier-szűrést követően), hogy elképzelhető-e, hogy a nagy volatilitás az oka annak, hogy az államháztartási deficit eredeti idősorában nem találtunk trendet.

A volatilitás csökkentésének egyik kézenfekvő módja az idősor aggregálása (vagy az ezzel ekvivalens mozgóátlagolás), amit ebben az esetben is kényelmesen megtehetünk. A 7-5. ábra az államháztartási deficit negyedéves bontású idősorát¹¹⁴ mutatja:



7-5. ábra: Az államháztartás deficit nagysága 1996. I. negyedév – 2015. IV. negyedév
(negyedéves bontás, 2015-ös változatlan áron, millió euró)

Vizsgáljuk meg a szokásos eszközökkel, hogy található-e trend a negyedéves deficit-idősorban!

7-6. táblázat: ADF és KPSS teszt a negyedéves bontású államháztartási deficit idősorára

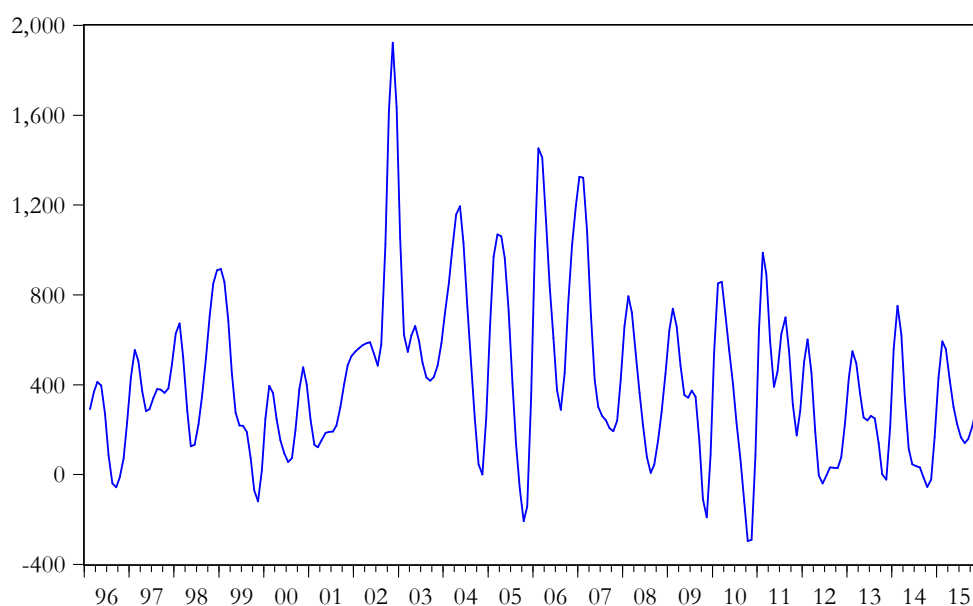
Modell-feltevés	ADF			KPSS	
	próbafüggvény	p-érték	lag	próbafüggvény	p-érték
eltolást tartalmaz	-2,563	0,105	3	0,333	$0,1 < p$
eltolást és determinisztikus trendet is tartalmaz	-2,738	0,225	3	0,259	$p < 0,01$

Eredményeink meglehetősen ellentmondásosak: az ADF-próba a sztochasztikus trend (egységgyök) létét látzik igazolni, miközben a vizsgált modellek egyikében esetben sem szignifikáns a determi-

¹¹⁴ Ügyeljünk rá, hogy a három hónap összegeként keletkező negyedéves adatok más léptékkel rendelkeznek, mint a korábbi havi deficit értékek!

nisztikus trend. A csak konstans tartalmazó modellt vizsgáló *KPSS-próba* mindeközben nem veti el a *stacionaritást feltételező nullhipotézist*. A teszt-eredményeket alaposabban vizsgálva azt mondhatjuk, hogy megfigyelések számának csökkenése (hiszen ebben az esetben mindössze 80 negyedévet vizsgáltam) eredményezheti a trend létének megítélésében megjelenő bizonytalanságot.

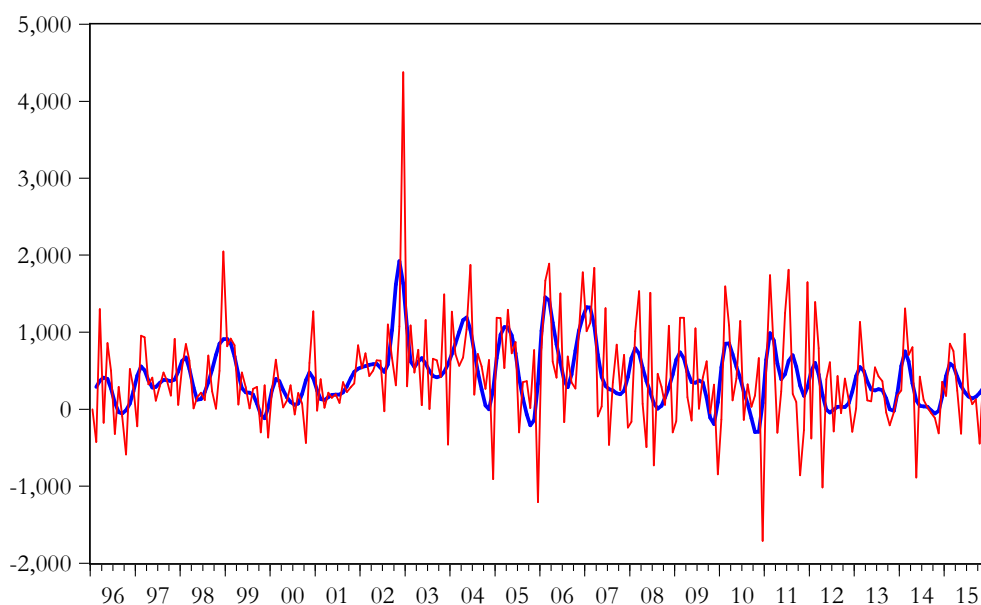
Érdekes gondolkísérlet, ha feltesszük, hogy eredetileg csak negyedéves adatok állnak rendelkezésre, melyekből – annak érdekében, hogy a sztochasztikus idősor-modellezésben megkövetelt minimum 100 megfigyelést elérjük – Catmull-Rom spline (CRS) segítségével¹¹⁵ havi bontású idősort becsülünk:



7-6. ábra: Az állambáztartás deficit becsült havi nagysága 1996. január – 2015. december
(negyedéves bontásból CRS-sel, 2015-ös változatlan áron, millió euró)

Az idősor képe nem nagyon igazít el a stacionaritás tekintetében. Érdekes ugyanakkor a következő ábra, melyen a tényleges (vékony piros), illetve a negyedévesből „visszabecsült” (vastag kék) idősorok szereplnek:

¹¹⁵ Lásd 3.1 alfejezet.



7-6.a. ábra: Az államháztartás deficit havi nagysága 1997. január – 2015. december
(tényleges, illetve negyedéves bontásból CRS-sel becsült, változatlan áron, millió euró)

Látható, hogy az eljárás segítségével a volatilitás jelentősen csökkent, tulajdonképpen az aggregálás-felosztás felfogható outlier-szűrésként is. Tekintsük a szokásos próbákat az így nyert idősorra:

7-7. táblázat: ADF és KPSS teszt a becsült, havi bontású államháztartási deficit idősorára

Modell-feltevés	ADF			KPSS	
	próbafüggvény	p-érték	lag	próbafüggvény	p-érték
eltolást tartalmaz	-2,810	0,058	14	0,359	$0,05 < p < 0,1$
eltolást és determinisztikus trendet is tartalmaz	-2,986	0,138	14	0,277	$p < 0,01$

Az eredmények azt mutatják, hogy a korábbiakkal ellentétben magas rendű autokorreláció is terheli az idősort, valamint érdemes észrevenni, hogy ha a determinisztikus trendet is tartalmazó modellvariánst teszteljük, a KPSS-próbánál a determinisztikus trend megjelenése nem zárható ki minden ésszerű szignifikancia-szinten ($p=0,096$). Mindezt összefoglalva ismét csak azt láthatjuk, hogy az eredeti idősor transzformálása, ha még oly korrektül történik is, jelentősen növelheti az elemzési eredményekkel kapcsolatos bizonytalanságot.

A magyar államháztartás egyenlegének vizsgálata mindössze egy – nem is biztos, hogy a legszembeötlőbb eredményeket szolgáltató – példa arra, hogy milyen modellezői vétségeket követhetünk el, akarva, vagy akaratlanul, ha a standard szakirodalom útmutatásait mechanikusan alkalmazzuk.

A Magyar Statisztikai Társaság Etikai Testületének elnökeként az ET első ülésén azt tanultam a nálam sokkal tapasztaltabb, esetenként jogvégzett Kollégáktól, hogy az számít etikai vétségnek, amit sem jogszabály, sem egyéb szabályzat nem tilt, de a jó ízlés, az erkölcs normáit sérti. Ha az idősor-kutatók számára „etikai kódexet” kellene egyszer készítenem, valószínűleg szerepelnének benne az alábbi intelmek:

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. Sohase feledkezzünk meg arról, hogy a modell a valóság egyszerűsített változata, ahol az operacionalizálhatóság oltárán mindig feláldozunk valamennyit a valósághűségből! Körültekintően hívjuk fel erre a felhasználók figyelmét, jelezzük a mintából történő következtések korlátait, ismertessük a ceteris paribus elv következményeit!

Első önálló publikációm (Rappai, 1989) bevezetőjében, egy modellezésről, becslési módszer és teszteljárás választásról szóló tartalmas szakmai vita részeként azt írtam: „*Úgy véljük azonban, hogy az ökonometriai 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 gazdálkodás pontosabb megalapozását. A modellek építőiben azonban felmerülhet a kétség, hogy e kompromisszumok mely kombinációja « alkímia » és hol kezdődik a « tudomány ».*”

A véleményem ebben 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, miközben a szakirodalom még számos kérdésben adós az egyértelmű válasszal. Dolgozatommal az idősoros anomáliák kezelése tekintetében igyekeztem néhány használható megoldási javaslatot adni.

8 Irodalom

ÁCS BARNABÁS (2012): Was the financial crisis of 2008 forecastable? *Hungarian Statistical Review*, Special Number 16., 85-100. pp.

AGUINIS, H. – GOTTFREDSON, R. K. – JOO, H. (2013): Best-practice recommendations for defining, identifying and handling outliers. *Organizational Research Methods*, Vol. 16. No. 2., 270-301 pp.

AKAIKE, H. (1969). Statistical predictor identification. *Annals of the Institute of Statistical Mathematics*, Vol. 21. No. 2., 203-217. pp.

AMEMIYA, T. – WU, R. Y. (1972): The effect of aggregation on prediction in the autoregressive model. *Journal of the American Statistical Association*, Vol. 67. No. 3., 628-632. pp.

ANDREWS, D. W. K. (1993): Tests for parameter instability and structural change with unknown change point. *Econometrica*, Vol. 61. No. 4., 817-858. pp.

ANDREWS, D. W. K. – PLOBERGER, W. (1994): Optimal tests when a nuisance parameter is present only under the alternative. *Econometrica*, Vol. 62. No. 6., 1383-1414. pp.

ARDELEAN, V. (2012): *Detecting outlier sin time series*. IWQW Discussion Papers 05/2012. University of Erlangen-Nuremberg, 27 pp.

AUE A. – HORVÁTH L. (2012): Structural breaks in time series. *Journal of Time Series Analysis*, Vol. 34. No. 1., 1-16. pp.

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.

BALKE, N.S. – FOMBY, T.B. (1994): Large shocks, small shocks and economic fluctuations: outlier sin macroeconomic time series. *Journal of Applied Econometrics*, Vol. 9. No. 2., 181-200. pp.

BANDYOPADHYAY, P.S. – FORSTER, M. R. (2011): *Philosophy of statistics (Handbook of the Philosophy of Science. Volume 7.)* Elsevier, Oxford UK., 571 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.

BARNETT, V. – LEWIS, T. (1994): *Outliers in statistical data*. 3rd ed. John Wiley, Chichester, 460 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.

BEDŐ ZSOLT – RAPPAI GÁBOR (2004): Van-e értéke az árfolyamokat befolyásoló híreknek? Eseményablak-elemzés magyar részvényárfolyamokra. *Szigma*, 35. évf. 3-4. sz., 107-122. old.

BEDŐ, ZS. – RAPPAL, G. (2006): Is there causal relationship between the value of the news and stock returns? *Hungarian Statistical Review*, Special Number 10., 81–99. 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.

BERGSTROM, A. R. (1985): The estimation of parameters in non-stationary higher-order continuous-time dynamic models. *Econometric Theory*, Vol. 1. No. 2., 369-385. pp.

BESENYEI LAJOS (2011): Múlt és jövő – statisztika és előrejelzés. *Statisztikai Szemle*, 89. évf. 10-11. sz., 1042-1056. old.

BESENYEI LAJOS – GIDAI ERZSÉBET – NOVÁKY ERZSÉBET (1977): *Jövő kutatás, előrejelzés a gyakorlatban*. Közgazdasági és Jogi Könyvkiadó, Budapest, 291 old.

BOX, G.E.P. – JENKINS, G.M. (1970): *Time series analysis; forecasting and control*. Holden-Day, San Francisco, 553 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.

BROCKWELL, P. J. (2001): *Continuous-time ARMA Processes*. (in Shanbhag, D. N. – Rao, C. R. (ed): *Handbook of Statistics 19; Stochastic Processes: Theory and Methods*.) Elsevier, Amsterdam, 249-276. pp.

BROOKS, C. (2002): *Introductory econometrics for finance*. Cambridge University Press, 701 pp.

BROWN, R. L. – DURBIN, J. – EVANS, J. (1975) Techniques for testing the constancy of regression relationship over time. *Journal of Royal Statistical Society, Series B*, Vol. 37., No. 1., 149–163. pp.

BRYN, G. – HUBERT, M. – STRUYF, A. (2004). A robust measure of skewness. *Journal of Computational and Graphical Statistics*, Vol. 13. No. 4., 996-1017 pp.

CAMPOS, J. – ERICSSON, N.R. – HENDRY, D.F. (1996): Cointegration tests in the presence of structural breaks. *Journal of Econometrics*, Vol. 70. No. 2., 187-220. pp.

CARRION-I-SILVESTRE, J.L. – SANSÓ-I-ROSSELLÓ, A. – ORTUÑO, M.A. (2001): Unit root and stationarity tests' wedding. *Economics Letters*, Vol. 70. No. 1., 1-8. pp.

CATMULL, E. – ROM, R. (1974): *A class of local interpolating splines*. (in Barnhill, R. E. – Reisnfeld, R. F. (ed): *Computer Aided Geometric Design*.) Academic Press, New York, 317-356. pp.

CHAMBERLAIN, G. (1982): The general equivalence of Granger and Sims causality. *Econometrica*, Vol. 50. No. 3., 569-582. pp.

CHAN, N. H. (1999): The ET interview: Professor George C. Tiao. *Econometric Theory*, Vol. 15. No. 3., 389-424. pp.

CHEN, K. – YING, Z. – ZHANG, H. – ZHAO, L. (2008): Analysis of least absolute deviation. *Biometrika*, Vol. 95. No. 1., 107-122 pp.

CHOW, G.C. (1960): Test of equality between sets of coefficient in two linear regressions. *Econometrica*, Vol. 28. No. 3., 591-605. pp.

COCHRANE, J. H. (2012): *Continuous-Time Linear Models*. NBER Working Paper 5807, University of Chicago, 33 pp.

COOK, R. D. – WEISBERG, S. (1982): *Residuals and influence in regression*. Chapman and Hall, New York, 230 pp.

CSEREHÁTI ZOLTÁN (2004): Az outlierok meghatározása és kezelése gazdaságstatisztikai felvételekben. *Statisztikai Szemle*, 82. évf. 3. sz., 728-746. old.

DARVAS ZSOLT (2005): *Bevezetés az idősorlemzés fogalmaiba. (Jegyzet)* Budapesti Corvinus Egyetem, Budapest, 110 old.

DICKEY, D.A. – FULLER, W.A. (1979): Distribution of estimators for time series regression with a unit root. *Journal of the American Statistical Association*, Vol. 74. No. 2., 427-431. pp.

DOORNIK, J.A. – OOMS, M. (2005): *Outlier detection in GARCH Models*. Tinbergen Institute Discussion Paper, TI 2005-092/4., 26 pp.

DUFOUR, J-M. – RENAULT, E. (1998): Short-run and long-run causality in time series: Theory. *Econometrica*, Vol. 66. No. 6., 1009-1125. pp.

DUFOUR, J-M. – PELLETIER, D. – RENAULT, E. (2006): Short and long run causality measures: Inference. *Journal of Econometrics*, Vol. 132. No. 2., 337-362. pp.

DUFOUR, J-M. – TAAMOUTI, A. (2010): Short and long run causality measures: Theory and inference. *Journal of Econometrics*, Vol. 154. No. 1., 42-58. 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.

ENGLE, R. F. (2000): The econometrics of ultra-high frequency data. *Econometrica*, Vol. 68. No. 1., 1-22. pp.

ENGLE, R.F. – GRANGER, C.W.J. (1987): Co-integration and error correction: representation, estimation and testing. *Econometrica*, Vol. 55. No. 2., 251-276. pp.

ENGLE, R.F. – RUSSELL, J.R. (1998): Autoregressive conditional duration: a new model for irregularly spaced transaction data. *Econometrica*, Vol. 66. No. 5., 1127-1162 pp.

E. SZABÓ LÁSZLÓ (2008): *Kauzalitás*. <http://phil.elte.hu/leszabo/> (letöltve 2008. december 12.)

FOX, A. (1972): Outliers in time series. *Journal of the Royal Statistical Society, Series B (Methodological)*, Vol. 34. No. 2., 350-363. pp.

GICZI ANDRÁS BÉLA (2004): Az osztályozás és a káoszelmélet. A rendszer-, a káosz-, az osztályozás-, és az információelmélet találkozása a tudományok dobozai között. *Könyvtári Figyelő*, 50. évf. 1. szám, 59-82. old.

GRANGER, C.W.J. (1969): Investigating causal relations by econometric models and cross-spectral methods. *Econometrica*, Vol. 37. No. 3., 424–438. pp.

GRANGER, C.W.J. (1980): Testing for causality. A personal viewpoint. *Journal of Economic Dynamics and Control*, Vol. 2. No. 2., 329-352. pp.

GRANGER, C.W.J. (1988): Some recent developments in a concept of causality. *Journal of Econometrics*, Vol. 39. No. 2., 199-211. 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.

GRANGER, C.W.J. – NEWBOLD, P. (1974): Spurious regression in econometrics. *Journal of Econometrics*, Vol. 2. No. 2., 111-120. pp.

GRANGER, C.W.J. – SIKLOS, P. L. (1995): Systematic sampling, temporal aggregation, seasonal adjustment, and cointegration. Theory and evidence. *Journal of Econometrics*, Vol. 65., No. 1-2., 357-369. pp.

GRANGER, C.W.J. – SWANSON, N. R. (1992): *Impulse response functions based on a causal approach to residual orthogonalization in vector autoregression*. Working paper, 92-50, University of California, San Diego, 32 pp.

GREGORY, A.W. – HANSEN, B.E. (1996): Tests for cointegration in models with regime and time shifts. *Oxford Bulletin of Economics and Statistics*, Vol. 58. No. 4., 555-560. pp.

GREEN, W.H. (2008): *Econometric Analysis 6th ed*. Pearson Education, New Jersey, 1178 pp.

HAJDU OTTÓ – HERMAN SÁNDOR – PINTÉR JÓZSEF – RÉDEY KATALIN (1987): *Ökonometriai alapvetés*. Tankönyvkiadó, Budapest, 159 old.

HAMILTON, J.D. (1994): *Time series analysis*. Princeton University Press, New Jersey, 820 pp.

HANSEN, B.E. (1997): Approximate asymptotic P values for structural-change tests. *Journal of Business & Economics Statistics*, Vol. 15. No. 1., 60-67. pp.

HANSEN, B.E. (2001): The new econometrics of structural change: Dating breaks of U.S. labor productivity. *Journal of Economic Perspectives*, Vol. 15. No. 4., 115-128. 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.

HECKMAN, J.J. (2008): *Econometric causality*. UCD Dublin, Geary WP/26/2008, 53 pp.

HENDRY, D. F. (1992): An econometric analysis of TV advertising expenditure in United Kingdom. *Journal of Policy Modelling*, Vol. 14. No. 3., 281-311. pp.

HICKS, J. (1979): *Causality in economics*. Basil Blackwell, Oxford UK., 124 pp.

HSIAO, C. (1979): Causality test sin econometrics. *Journal of Economic Dynamics and Control*, Vol. 1. No. 4., 321-346. pp.

HUME, D. (2006): *Értekezés az emberi természetről*. Akadémiai Kiadó, Budapest, 656 old.

HUNYADI LÁSZLÓ (1994): Egységgyökök és tesztjeik. *Sigma*, 25. évf. 3. sz., 135-164. old.

HUNYADI LÁSZLÓ (2002): Grafikus ábrázolás a statisztikában. *Statisztikai Szemle*, 80. évf. 1. sz., 22-52. old.

IGLEWICZ, B. – HOAGLIN, D. C. (1993): *How to detect and handle outliers*. ASQC Quality Press., Milwaukee, 87 pp.

JOHANSEN, S. (1991): Estimation and hypothesis testing of cointegration vectors in gaussian vector autoregressive models. *Econometrica*, Vol. 59. No. 6., 1551–1580. 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.

JONES, R. E. (1985): *Time series analysis with unequally spaced data*. (In HANNAN, E.J. – KRISHNAIAH, P.R. – RAO, M.M. (ed): *Handbook of Statistics. Vol. 5. Time Series in the Time Domain*.), Elsevier, North-Holland Amsterdam, 157-177. pp.

KALMAN, R. E. – BUCY, R. S. (1961): New results in linear filtering and prediction theory. *Journal of Basic Engineering*, Vol. 83. No. 1., 95-108. pp.

KĘBŁOWSKI, P. – WELFE, A. (2004): The ADF-KPSS test of the joint confirmation hypothesis of unit autoregressive root. *Economic Letters*, Vol. 85. No. 2., 257-263. pp.

KIM, K. (2003): Dollar Exchange Rate and Stock Price: Evidence from Multivariate Cointegration and Error Correction model. *Review of Financial Economics*, Vol. 12. No. 2., 301-313. pp.

KNORR, E.M. – NG, R.T. (1997): A unified approach for mining outliers. In: Proceedings Conference of CASCON, *CASCON1997* Toronto, 219-222. pp.

KŐRÖSI GÁBOR (2016): A lány továbbra is szolgál... *Közgazdasági Szemle*, 58. évf. 6. sz., 647-667. old.

KŐRÖSI GÁBOR – MÁTYÁS LÁSZLÓ – SZÉKELY ISTVÁN (1990): *Gyakorlati ökonometria*. Közgazdasági és Jogi Könyvkiadó, Budapest, 481 old.

KŐRÖSI GÁBOR – LOVRICS LÁSZLÓ – MÁTYÁS LÁSZLÓ (1993): *Aggregation and the long run behaviour of economic time series*. Working paper 12/93., Monash University 14 pp.

KRIEGL, H-P. – KRÖGER, P. – ZIMEK, A. (2010): *Outlier detection techniques*. Tutorial notes: SIAM SDM 2010 Columbus Ohio. <http://www.dbs.ifi.lmu.de/~zimek/publications/SDM2010/sdm10-outlier-tutorial.pdf> (letöltve 2015. április 12.)

KWIATKOWSKI, D. – PHILLIPS, P. C. B. – SCHMIDT, P. – SHIN, Y. (1992): Testing the null hypothesis of stationarity against the alternative of a unit root. *Journal of Econometrics*, Vol. 54. No. 2., 159–178. pp.

LEDOLTER, J. (1990): Outlier diagnostics in time series analysis. *Journal of the Time Series Analysis*, Vol. 11. No. 2., 317–324. pp.

LEE, J. – HUANG, C.J. – SHIN, Y. (1997): On stationary test sin the presence of structural breaks. *Economics Letters*, Vol. 55. No. 2., 165–172. pp.

LEWIS, D. (1973): *Counterfactuals*. Basil Blackwell, Oxford, 156 pp.

LIN, J. L. (2008): *Notes on testing causality*. Working paper of Department of Economics, National Chengchi University. <http://faculty.ndhu.edu.tw/~jlin/files/causality.pdf> (letöltve 2013. november 3.)

LIPPI, M. – REICHLIN, L. (1991): Trend-cycle decomposition and measures of persistence: Does time aggregation matter? *Economic Journal*, Vol. 101. No. 405., 314–323. pp.

LJUNG, G.M. (1993): On outlier detection in time series. *Journal of the Royal Statistical Society, Series B (Methodological)*, Vol. 55. No. 2., 559–567 pp.

LOMB, R. (1976): Least-squares frequency analysis of unequally spaced data. *Astrophysics and Space Science*, Vol. 39. No. 2., 447–462. pp.

LORENZ, E. (1963): Deterministic nonperiodic flow. *Journal of the Atmospheric Sciences*, Vol. 20. No. 1., 130–141. pp.

LÜTKEPOHL, H. (1987): *Forecasting aggregated Vector ARIMA process*. Springer Verlag, Berlin, 327 pp.

MACKIE (1965): *Causes and Conditions*. American Philosophical Quarterly, Vol. 2. No. 4., 245–264. pp.

MADDALA, G. S. (2004): *Bevezetés az ökonometriába*. Nemzeti Tankönyvkiadó, Budapest, 704 old.

MAHALANOBIS, P. CH. (1936): On the generalized distance in statistics. *Proceedings of the National Institute of Sciences of India*, Vol 2. No. 1, 49–55. pp.

MÁK FRUZSINA (2011): Egységgyöktesztek alkalmazása strukturális törések mellett a hazai benzinár példáján. *Statisztikai Szemle* 89. évf. 5. sz., 545–573. old.

MARCELLINO, M. (1999): Some consequences of temporal aggregation in empirical analysis. *Journal of Business and Economics Statistics*, Vol. 17. No.1., 129–136. pp.

MELLÁR TAMÁS (2003): *Dinamikus makromodellek a magyar gazdaságra*. Statisztikai Kiadó, Budapest, 215 old.

MELLÁR TAMÁS (2016): Szolgálólányból királycsináló – avagy az ökonometria makroökonómiai térhódítása? *Közgazdasági Szemle*, 63. évf. 3. sz., 285–306. old.

MÉSZÁROS JÓZSEFNÉ (2011): *Numerikus módszerek*. Digitális Tankönyvtár, Miskolci Egyetem Földtudományi Kar, GEMAK 6841B.

MICHELBERGER PÁL – SZEIDL LÁSZLÓ – VÁRLAKI VIKTOR (2001): *Alkalmazott folyamatstatisztika és idősor-analízis*. Typotex Kiadó, Budapest, 390 old.

MILLS, T.C. – PATTERSON, K. (SZERK.) (2006): *Palgrave Handbook of Econometrics. Vol. 1. Econometric Theory*. Palgrave, MacMillan, New York, 1097 pp.

MÓCZÁR JÓZSEF (2008, 2009): Közgazdaságtan vagy közgazdaság-tudomány? A 20. század legfontosabb eredményei. I-II. rész. *Competitio*, 7. évf., 2. szám, 5-34. old., illetve 8. évf., 1. szám, 76-97. old.

MOSCONI, R. – SERI, R. (2006): Non-causality in bivariate binary time series. *Journal of Econometrics*, Vol. 132. No. 2., 379-407. pp.

MUNDRUCZÓ GYÖRGY (1998): *Útmutatás a statisztikai modellezés gyakorlatához*. KSH, Budapest, 295 old.

NARAYAN, P. K. – SMYTH, R. (2009): Multivariate granger causality between electricity consumption, exports and GDP: Evidence from a panel of Middle Eastern countries. *Energy policy*, Vol. 36. No. 1., 229-236. pp.

NGUYEN, T. D. – WELSCH, R. (2010): Outlier detection and least trimmed squares approximation using semi-definite programming. *Computational Statistics & Data Analysis*, Vol. 54. No. 12., 3212–3226. pp.

NOVÁKY ERZSÉBET (1995): Káosz és előrejelzés. *Statisztikai Szemle*, 73. évf. 10. sz., 815-823. old.

NOVÁKY ERZSÉBET (2015): A hazai társadalmi-gazdasági mutatók vizsgálata a káoszelmélet eszközeivel. *Statisztikai Szemle*, 93. évf. 1. sz., 25-52. old.

OLEWUEZI, N. P. (2011): Note on the comparison of some outlier labeling techniques. *Journal of Mathematics and Statistics*, Vol. 7. No. 4., 353-355 pp.

PEARSON, K. (1892): *The grammar of science. (Classic reprint.)* A. and Ch. Black, 574 pp.

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.

PERRON, P. (2006): *Dealing with structural breaks*. (In: MILLS, T.C. – PATTERSON, K. (szerk): *Palgrave Handbook of Econometrics. Vol. 1. Econometric Theory*.) Palgrave, MacMillan, New York, 278-352. pp.

PEÑA, D. (1987): Measuring influence in dynamic regression models. *Technometrics*, Vol. 33. No. 1., 93-101. pp.

PHILLIPS, P.C.B. (2005): Challenges of trending time series econometrics. *Mathematics and Computers in Simulation*, Vol. 68. No. 5-6., 401-416. pp.

PIERSE, R.G. – SNELL, A.J. (1995): Temporal aggregation and the power of tests for a unit root. *Journal of Econometrics*, Vol. 65. No. 2., 333-345. 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.

RAJAGURU, G. – ABEYSINGE, T. (2008): Temporal aggregation, cointegration and causal inference. *Economic Letters*, Vol. 101. No. 3., 223-226. pp.

RAPPAI GÁBOR (1989): A fogyasztási javak piacának nem egyensúlyi modellje. *Statisztikai Szemle*, 67. évf. 7. sz., 663-678. old.

RAPPAI GÁBOR (2004): *A hozam-idősorok természetéről.* (In Vita L. (szerk.): *Egy reneszánsz statisztikus.*) KSH, Budapest. 153-165. old.

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 (2013): *Bevezető pénzügyi ökonometria.* Pearson Education, Harlow UK, 148 old.

RAPPAI GÁBOR (2014): Rendszertelen idősorok modellezése spline-interpolációval. *Statisztikai Szemle*, 92. évf. 8-9. sz., 766-791. old.

RAPPAI GÁBOR (2015): Modeling non-equidistant time series using spline interpolation. *Hungarian Statistical Review*, Special number 19., 22-46. pp.

RAPPAI GÁBOR (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.

REICHENBACH, H. (1956): *The direction of time.* University of California Press, Berkeley, 142 pp.

RÉNYI ALFRÉD (2004): *Levelek a valószínűségről.* Neumann Kht. Budapest. <http://www.mek.oszk.hu/05000/05029/html/index.htm> (letöltve 2011. június 21.)

ROSSANA, R. J. – SEATER, J. J. (1995): Temporal aggregation and economic time series. *Journal of Business and Economic Statistics*, Vol. 13. No. 4., 441-451. pp.

RUSSEL, B. (1976): *Miszticizmus és logika és egyéb tanulmányok.* Magyar Helikon, Budapest, 405 old.

SAID, S.E. – DICKEY, D.A. (1984): Testing for unit roots in autoregressive-moving average models of unknown order. *Biometrika*, Vol. 71. No. 3., 599-608. pp.

SCHLITGEN, R. – STREITBERG, B.H.J. (1991): *Zeitreihenanalyse*. Oldenbourg Verlag, München, 502 pp.

SCHMITZ, A. (2000): *Erkennung von Nichtlinearitäten und wechselseitigen Abhängigkeiten in Zeitreihen*. WUB-DIS-2011, Wuppertal, 123 pp.

SCHOENBERG, I. J. (1946): Contributions to the problem of approximation of equidistant data by analytic functions. Parts A and B, *Quarterly Applied Mathematics* Vol. 4. No. 1. 45-99 pp.; 112-141 pp.

SEO, S. (2006): *A review and comparison of methods for detecting outliers in univariate data sets*. University of Pittsburgh, 53 pp. <http://d-scholarship.pitt.edu/7948/1/Seo.pdf> (2013. szeptember 5.)

SHEFFRIN, S. M. – TRIEST, R. K. (1995): *A new approach to causality and economic growth*. Working Paper Series No. 95-12. Federal Reserve Bank of Boston, 22 pp.

SILVESTRINI, A. – VEREDAS, D. (2005): *Temporal aggregation of univariate linear time series models*. Core Discussion Paper 2005/59., 43 pp.

SIMS, C.A. (1971): Discrete Approximations to Continuous Time Distributed lags in Econometrics. *Econometrica*, Vol. 39. No. 3., 545-563. pp.

SIMS, C.A. (1972): Money, income and causality. *American Economic Review*, Vol. 62. No. 4., 540-552. pp.

SLUTZKY, E. (1937): The Summation of Random Causes as the Source of Cyclic Processes. *Econometrica*, Vol. 5. No. 2, 105-146. pp.

SRIVASTAV, A. K. – AGRAVAL, G. – SRIVASTAVA, A. (2010): A study of exchange rate movement and stock market volatility. *International Journal of Business and Management*, Vol. 5. No. 12., 62-73. pp.

SUGÁR ANDRÁS (1999): Szezonális kiigazítási eljárások I-II. *Statisztikai Szemle*, 77. évf. 9. és 10-11. sz., 705-721. old. és 816-832. old.

SZABÓ GÁBOR (2001): A reichenbach-i közös ok eredete. *Magyar Filozófiai Szemle*, 45. évf. 1-2. sz., 83-112. old.

TOLVI, J. (1998): *Outliers in time series: A review*. University of Turku, 30 pp. <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.198.823&rep=rep1&type=pdf> (letöltve 2013. október 10.)

TSAY, R. S. (1986): Time series model specification in the presence of outliers. *Journal of the American Statistical Association*, Vol. 81. No. 1., 132-141. pp.

TSAY, R. S. (1988): Outliers, level shifts and variance changes in time series. *Journal of Forecasting*, Vol. 7. No. 1., 1-20. pp.

TUKEY, J.W. (1977): *Exploratory data analysis*. Addison-Wesley. 688 pp.

WANG, Z. (2013): cts: An R package for continuous time autoregressive models via Kalman filter. *Journal of Statistical Software*, Vol. 53. No. 5., 1-18. pp.

WEISS, A. (1984): Systematic sampling and temporal aggregation in time series models. *Journal of Econometrics*, Vol. 26. No. 2., 271-281. pp.

WHITE, H. – GRANGER, C.W.J. (2011): Consideration of trends in time series. *Journal of Time Series Econometrics*, Vol. 3. No. 1., 38 pp.

WIENER, N. (1956): The theory of prediction. (In: Beckenback, E.R. (ed.): *Modern Mathematics for Engineers.*) McGraw Hill, New York, 165-190. pp.

WITTGEINSTEIN (2004): *Tractatus logico-philosophicus – Logikai-filozófiai értekezés.* Atlantisz Kiadó. Budapest. 131 old.

WU, L.S.-Y. – HOSKING, J.R.M. – RAVISHANKER, N. (1993): Reallocation outliers in time series. *Journal of the Royal Statistical Society, Series C (Applied Statistics)*, Vol. 42. No. 2., 301-313. pp.

XU, Z. (1996): On the Causality between Export Growth and GDP Growth: An Empirical Reinvestigation. *Review of International Economics*, Vol. 4. No. 2., 172-184. pp.

YOHAI, V. J. – STAHEL, W. A. – ZAMAR, R. H. (1991): A Procedure for Robust Estimation and Inference in Linear Regression (In Stahel, W. – Weisberg, S. (ed): *Directions in Robust Statistics and Diagnostics.*) Springer, München, 365-374. pp.

ZIVOT, E. – ANDREWS, D.W.K. (1992): Further evidence on great crash, the oil price shock and the unit root hypothesis. *Journal of Business and Economic Statistics*, Vol 10. No. 1., 251-270. pp.