

Vargha András

Többváltozós statisztikai elemzések
pszichológiai kutatásokban
ROP-R-rel

Pólya Kiadó

Budapest, 2023

Utolsó javítás: 2026.07.01.

A könyv megjelenését
a Károli Gáspár Református Egyetem támogatta

Szakmai lektor:
Rózsa Sándor

ISBN 978-615-81228-5-6

© Vargha András
© Pólya Kiadó

Tartalom

Előszó	7
Bevezetés	9
B1. Mit kell tudni a ROP-R-ről?	10
B1.1. Általános tudnivalók	10
B1.2. Az R szoftver telepítése	10
B1.3. A ROP-R szoftver telepítése	12
B1.4. A ROP-R menüpontjai	13
B1.5. A ROP-R statisztikai moduljai	14
B2. A könyvben használt adatminták	15
B2.1. Gyermek és szülei antropometriai adatai (ANTR1980)	15
B2.2. Adatok egy kötődéskutatásból (KÖT2016)	16
B2.3. Magyarország boldogságtérképe 2022 (Btérkép2022)	18
I. RÉSZ: Regressziós modulok	21
1. fejezet: A hierarchikus regresszió (HierR) modul	23
1.1. Elméleti alapok	23
1.1.1. Többszörös lineáris regresszió	23
1.1.2. Hierarchikus regresszió	24
1.1.3. Mediációs regresszió	24
1.2. A HierR menüablak	28
1.2.1. Blokkok megadása HierR-ben	29
1.2.2. Influenza esetek	29
1.2.3. Mediációs elemzés kijelölése	29
1.3. Mit tartalmaz a HierR eredménylistája?	30
1.4. A HierR szemléltetése valódi adatokon	30
1.4.1. Hierarchikus regresszióelemzés antropometriai adatokkal	30
1.4.2. Hierarchikus regresszióelemzés a pszichológiai kötődés vizsgálatában	33
1.4.3. Mediációs elemzések a globális jóllét bejósolásában	35
2. fejezet: A polinomiális regresszió (PolR) modul	40
2.1. Elméleti alapok	40
2.2. A PolR menüablak	42
2.3. Mit tartalmaz a PolR eredménylistája?	43
2.4. A PolR szemléltetése valódi adatokon	43
2.4.1. Az extraverzió nemlineáris hatása a szubjektív jóllétre	43
2.4.2. Az érzelmi instabilitás nemlineáris hatása a szubjektív jóllétre	44
2.4.3. A szubjektív jóllét nemlineáris hatása az érzelmi instabilitásra	45

3. fejezet: A bináris logisztikus regresszió (BLR) modul	48
3.1. Elméleti alapok	48
3.2. A BLR menüablak	51
3.3. Mit tartalmaz a BLR eredménylistája?	52
3.4. A BLR szemléltetése valódi adatokon	52
3.4.1. A globális jóllét magas szintjének bejósítása szociodemográfiai változók segítségével	52
3.4.2. A globális jóllét magas szintjének bejósítása szociodemográfiai változók, a PIK16 skálái és a Caprara-féle Pozitivitás skála segítségével	55
II. RÉSZ: Dimenziócsökkentő modulok	58
4. fejezet: A főkomponens-analízis (FKA) modul	59
4.1. Elméleti alapok	59
4.2. Az FKA menüablak	61
4.3. Mit tartalmaz az FKA eredménylistája?	62
4.4. Az FKA szemléltetése valódi adatokon	63
4.4.1. Főkomponens-analízis antropometriai adatokkal	63
4.4.2. A Caprara-féle Pozitivitás skála tételeinek főkomponens-analízise	65
5. fejezet: A feltáró faktoranalízis (EFA) modul	68
5.1. Elméleti alapok	68
5.2. Az EFA menüablak	70
5.3. Mit tartalmaz az EFA eredménylistája?	71
5.4. Az EFA szemléltetése valódi adatokon	72
6. fejezet: A konfirmatív faktoranalízis (CFA) modul	76
6.1. Elméleti alapok	76
6.1.1. Elsőrendű faktormodell több korreláló faktorról	79
6.1.2. A másodrendű faktormodell	79
6.1.3. A bifaktoros faktormodell	81
6.1.4. A CFA elemzés lényege, adekvációs mutatók	82
6.1.5. A CFA modellbecslési módszerei	84
6.1.6. A modifikációs indexek	85
6.2. A CFA menüablak	86
6.2.1. Az input változók CFA-ban	87
6.2.2. Skálahovatartozás megadása	87
6.2.3. Modifikációs index küszöbök megadása	87
6.2.4. Faktornevek megadása és a ROP-R által elkészített faktorábrák	88
6.3. Mit tartalmaz a CFA eredménylistája?	88
6.4. A CFA szemléltetése valódi adatokon	89
6.4.1. A Caprara-féle Pozitivitás skála egydimenziós faktorszerkezetének tesztelése	89

6.4.2. A Caprara-féle Pozitivitás skála háromskalás faktorszerkezetének tesztelése	93
6.4.3. A PIK16 kérdőív négyskalás faktorszerkezetének tesztelése	97
6.4.4. A PIK16 kérdőív másodrendű faktorszerkezetének tesztelése	101
6.4.5. A PIK16 kérdőív bifaktoros faktorszerkezetének tesztelése	104
III. RÉSZ: Klaszterező modulok	115
7. fejezet: Az összevonó hierarchikus klaszteranalízis (AHKA) modul	117
7.1. Elméleti alapok	117
7.1.1. Személyek távolsága	117
7.1.2. Klaszterek távolsága	118
7.1.3. Klaszterstruktúrák jóságának mérése	119
7.2. Az AHKA menüablak	121
7.2.1. Input változók kiválasztása	121
7.2.2. Személytávolság megválasztása	122
7.2.3. Klaszterösszevonás módszerének megválasztása	122
7.2.4. Segítő ábrák	122
7.2.5. Egyéb lehetőségek	122
7.2.6. Az elemzés után létrejövő fájlok	123
7.3. Mit tartalmaz az AHKA eredménylistája?	123
7.4. Az AHKA szemléltetése valódi adatokon	123
8. fejezet: Az osztódó hierarchikus klaszteranalízis (OHKA) modul	132
8.1. Elméleti alapok	132
8.2. Az OHKA menüablak	133
8.3. Mit tartalmaz az OHKA eredménylistája?	133
8.4. Az OHKA szemléltetése valódi adatokon	134
9. fejezet: A k -középpontú klaszteranalízis (KKA) modul	137
9.1. Elméleti alapok	137
9.1.1. A klasztercentrum kiválasztása	137
9.1.2. A k -közép módszer	137
9.1.3. A k -medián módszer	138
9.1.4. A k -medoid módszer	138
9.2. A KKA menüablak	138
9.3. Mit tartalmaz a KKA eredménylistája?	140
9.4. A KKA szemléltetése valódi adatokon	140
9.4.1. k -közép elemzés két kötődésváltozóval	140
9.4.2. k -medián és k -medoid elemzés két kötődésváltozóval	146
9.4.3. k -közép elemzés a PIK16 kérdőív négy skálájával	149
10. fejezet: A modell-alapú klaszteranalízis (MKA) modul	154
10.1. Elméleti alapok	154

10.2. Az MKA menüablak	157
10.3. Mit tartalmaz az MKA eredménylistája?	158
10.4. Az MKA szemléltetése valódi adatokon	159
10.4.1. MKA elemzés két kötődésváltozóval	159
10.4.2. MKA elemzés a PIK16 kérdőív négy skálájával	164
11. fejezet: ROPstat és ROP-R	168
Záró gondolatok	
Irodalom	170

Előszó

A pszichológiai kutatások és a pszichológiai modellek egyre bonyolódnak, s ezzel párhuzamosan egyre magasabb szintű statisztikát igényelnek. Ennek következtében a pszichológia nivósabb szakfolyóiratai egyre magasabb módszertani követelményeket állítanak a tanulmányaikat megjelentetni szándékozó szerzők számára. Erre a kihívásra a pszichológusok által használt profi statisztikai szoftverek (pl. a SAS vagy az SPSS) évről-évre komoly fejlesztéseket építenek be rendszerükbe. Ennek persze ára van, amit magánszemélyek alig, de még az egyetemi intézmények is olykor nehezen tudnak megfizetni.

Ez a helyzet vezetett nemzetközi összefogással egy ingyenesen használható új statisztikai szoftverrendszer, az R (R Core Team, 2021) kialakításához, mely szinte mindent tud, amire ma a nemzetközi kutatótársadalomnak szüksége van. Az R CRAN nevű archívumában¹ több mint 16 000 szabadon letölthető adatfeldolgozási és egyéb számolási eljárás áll a felhasználók rendelkezésére széleskörű dokumentációval együtt, speciális programok/programcsomagok formájában (az R-ben ezek neve *package* vagy *library*).

Az R használatát azonban nehezíti, hogy az R-package-ek megfelelő futtatásához szintaktikailag precízen összeállított utasításcsomagokra, ún. *scriptekre* van szükség, ami majdhogynem programozási jártasságot igényel a felhasználó részéről. Szerencsére a legújabb időkben – egyetemi támogatásoknak is köszönhetően – olyan szoftverek születtek, amelyek a kissé körülményesen használható R-package-eknek tetszetős menürendszerrel felhasználóbarát keretet alakítottak ki. Ilyen a jamovi (The jamovi project, 2022; Şahin és Aybek, 2019), a JASP (JASP Team, 2023) és ezekhez sorakozott fel újabban a Windows operációs rendszerben futtatható ROP-R (Vargha és Bánsági, 2022; Vargha, Bánsági és Jantek, 2024) is, melyek közös alapelve, hogy bárki által ingyen használhatók, hiszen a szintén ingyenes R-package-ekre építenek.

A jelen könyv célja, hogy bemutassa: miként lehet a ROP-R moduljai segítségével magas szintű többváltozós statisztikai elemzéseket végrehajtani. A ROP-R tíz modulja három csoportba sorolható:

- regressziós elemzések,
- dimenzió redukciók (főkomponens- és faktorelemzések),
- klaszteranalízisek.

ROP-R különlegessége, hogy benne több olyan statisztikai módszer is megtalálható (pl. a polinomiális regresszió, a konfirmatív faktoranalízis egyes komplexebb modelljeinek a vizsgálata vagy a modell-alapú klaszteranalízis), amelyek sem jamoviban, sem JASP-ban nem elérhetők.

Néhány szó a könyv szerkezetéről. A Bevezetés írja le, hogy miképpen kell a ROP-R-t és előtte az R szoftvert telepíteni, hogy utána a ROP-R moduljait gond nélkül használhassuk. A Bevezetés ismerteti továbbá azokat a valódi kutatásból származó adatállományokat is, amelyeket a ROP-R moduljainak működését szemléltető mintapéldák során felhasználunk. Ezt követően az I. rész fejezetei a regressziós elemzések, a II. rész fejezetei a dimenzió redukciók, a III. rész fejezetei pedig a klaszteranalízisek moduljait mutatják be, minden fejezet egy-egy modult. Egy fejezeten belül először mindig az adott modul használatához

¹ https://cran.r-project.org/web/packages/available_packages_by_name.html

szükséges többváltozós statisztikai ismereteket foglaljuk össze. Ezután a modul menüablakát és használatának módját taglaljuk. Ezt követően valódi pszichológiai adatelemzésekkel szemléltetjük a modul működését és az eredmények értelmezésének módját.

A könyv megértéséhez elegendő a magyarországi egyetemek pszichológia alap- és mesterszakán oktatott statisztika ismeretanyaga (lásd pl. Vargha, 2007, 2019), így nagy hasznára lehet a felsőbb éves pszichológia szakos hallgatóknak. Ugyanakkor ez a könyv a pszichológiai kutatásokat végzők számára is íródott, olyan eszközöket adva a kezükbe, ami emeli kutatásaik, a statisztikai adatkiértékelés, valamint publikációik színvonalát. A könyvet haszonnal forgathatják a társszakmák (pedagógia, szociológia, biológia, orvostudomány stb.) többváltozós statisztikai elemzések iránt érdeklődő kutatói is.

Itt szeretnék köszönetet mondani Bánsági Péter barátomnak, aki önzetlen segítséggel lett társam a ROP-R szoftver kifejlesztésében és végső alakjának megformálásában. Szeretném megköszönni a könyv szakmai lektorának, Rózsa Sándornak is alapos, sok kis részletre kiterjedő, figyelmes munkáját és értékes javaslatait. Köszönettel tartozom Gergely Bencének, Jakab Zoltánnak és Takács Szabolcsnak is, akik a könyv több fejezetével kapcsolatban fogalmaztak meg olyan megjegyzéseket, amelyek hatására a kifejtéseken és magyarázatokon számottevő mértékben javítottam. Végül köszönöm Oláh Attilának és a Jobb Veled a Világ Alapítványnak, hogy a Magyarország 2022-es Boldogságtérképe kutatás adatait, valamint Jantek Gyöngyvérnek, hogy kötődéskutatási adatait a könyv pszichológiai szemléltető példáiban felhasználhattam.

Végül köszönet illeti a Károli Gáspár Református Egyetem Bölcsész- és Társadalomtudományi Karát, hogy a könyv elkészültét a *Pszichológiai kutatások módszertani platformja* című és 20754B800/2022 témaszámú kutatói pályázat keretében támogatta.

FONTOS ÚJDONSÁG!!!

A klaszterelemzéseket leíró III. részt – ezen belül is különösen a 9. fejezetet – 2025 és 2026 során jelentősen javítottam a ROP-R megfelelő négy moduljával együtt. Új alapelv (Moeller, 2025) szerint érdemes arra törekedni, hogy a változókat közös skálára hozzuk, amivel elkerülhető a klaszterelemzés előtt a változók standardizálása (z -standardizálás). Az elemzések lefutása **után** azonban ilyenkor hasznos, ha az eredmények grafikonon történő ábrázolásához standardizáljuk a klaszterátlagokat a változók összátlagával és együttes szórásával (z_T -standardizálás), és a centroidok mintázatát a nyersátlagok mellett ily módon is szemléltetjük (Vargha, 2025; Vargha & Postáné, 2026).

Bevezetés

Itt mondjuk el, hogy hogyan kell a ROP-R-t és előtte az R szoftvert telepíteni, hogy utána a ROP-R különböző moduljait gond nélkül használhassuk (B1). Itt mutatjuk be továbbá azokat a valódi kutatásból származó adatállományokat is, amelyeket a ROP-R moduljainak működését szemléltető mintapéldák során a könyv különböző fejezeteiben felhasználunk (B2).

B1. Mit kell tudni a ROP-R-ről?

Ebben a pontban adunk tájékoztatást a ROP-R, valamint a ROP-R megfelelő működéséhez szükséges R szoftver telepítéséhez.

B1.1. Általános tudnivalók

ROP-R első verziója 2022 tavaszán készült el, szerzői a jelen könyv szerzője, valamint Bánsági Péter matematikus mérnök. Tekintve, hogy ROP-R a ROPstat statisztikai szoftver többváltozós bővítése, a www.ropstat.hu weboldaltól tölthető le. A ROPstattal kapcsolatos minden információ megtalálható Vargha (2016), illetve Vargha (2020, Melléklet) művében, illetve a www.ropstat.com weboldalon. ROP-R legfontosabb általános jellemzői az alábbiak.

1. Windows operációs rendszerben futtatható.
2. Kétnyelvű (magyar és angol) többváltozós statisztikai szoftver, amelynek jelenlegi 10 modulja a többváltozós statisztika alábbi három témakörében kínál teljeskörű statisztikai elemzéseket: regresszióelemzés, dimenzióredukció (főkomponens- és faktoranalízis), valamint klaszteranalízis.
3. A kiválasztott statisztikai elemzés minden modul esetében egy átlátható, egyszerű ROP-R menüablakban (feladatablakban) paraméterezhető és futtatható.
4. Az elemzés elindítása után a ROP-R létrehoz egy R számára olvasható adatfájlt és egy vagy több megfelelő R-scriptet, amelyeket lefuttatva ROP-R a kapott R-outputot tetszetős formára hozza és elhelyezi azt a ROP-R nézőkében.
5. Ezeket a scripteket ROP-R a felhasználó által elérhető szövegfájlokba írja, amelyek hasznosak lehetnek az R szoftvert tanulók számára az R-scriptek megértésében, az R-ben már korábbi tapasztalatokkal rendelkezők számára pedig a ROP-R-belinél komplikáltabb elemzések R-beli elvégzéséhez.

Fontos, hogy a ROPstattal való szoros kapcsolata ellenére a ROP-R önálló szoftver, amely a ROPstat nélkül is futtatható, sőt, még az sem szükséges, hogy a ROPstat telepítve legyen a gépen. Mivel a ROP-R a programból futtat R-scripteket, működésének feltétele, hogy az R szoftver (speciálisan annak Rcmd.exe programja) installálva legyen. Az ezzel kapcsolatos teendőket az alábbiakban részletezzük.

B1.2. Az R szoftver telepítése

Mivel a ROP-R a programból futtat R-scripteket, működésének feltétele, hogy az R szoftver (speciálisan annak Rcmd.exe programja) installálva legyen. Az ezzel kapcsolatos teendőket az alábbiakban részletezzük. A ROP-R megfelelő működéséhez először az R szoftvert kell telepíteni (2026 elején a legújabb verziója: R-4.5.2). Megjegyezzük, hogy a ROP-R megfelelően működik a korábbi verziók közül számos (pl. az R-4.4.2) verzióval is, de pl. az R-4.1.3 verzió esetén már problémák léphetnek fel egyes R-package-ek installálásakor. A Windows alapú R-4.5.2 – annak 64 bites futásokat is lehetővé tevő x64 moduljával

kiegészített – telepítéséhez erre a weboldalra kell ellátogatni: <https://mirror.metanet.ch/cran/bin/windows/base/>. Innen tölthető le a telepítő program (R-4.5.2-win.exe). Számítógépünk fájlkezelő programjával keressük meg ezt a letöltött programok között és indítsuk el. A felajánlott angol kommunikációs nyelv és a „GNU GENERAL PUBLIC LICENSE” licenc szerződés elfogadásával a program standard körülmények között a „C:\Program Files\R\R-4.5.2” mappát ajánlja fel az R szoftver telepítéséhez. Ha a „C:\Program Files” mappához nincs hozzáférésünk, ehelyett nyugodtan megadhatjuk a „C:_vargha\R\R-4.5.2” mappát is, mert a „C:_vargha” mappa a ROPstat és/vagy a ROP-R szoftver telepítése során amúgy is létrejön, s ezért ehhez mindenképpen lesz hozzáférésünk. Ezután mindig fogadjuk el a felajánlott opciót a „Next” gombra kattintva, csak a „Select additional tasks” ablakban kérjük még a „Create a Quick Launch shortcut” opciót is, melynek hatására létrejön az R számára egy parancsikon az asztalon (R 4.4.2 névvel). Az utolsó ablakban a „Finish” gombra kattintva befejeződik az R installálása.

Az R szoftver sikeres telepítése után létrejön (ha nem, akkor manuálisan létrehozandó) az asztalon egy R parancsikon, mely az RGui.exe keretprogram segítségével képes R-package-eket futtatni. Ha az R-4.5.2 korábban már telepítve volt számítógépünkön, akkor a fenti lépések átugorhatók. A további lépéseket az alábbiak szerint foglaljuk össze.

1. Indítsuk el a letöltött R-verzióhoz tartozó RGui.exe programot.
2. Ha a gépen korábban már használták az ehhez az R-verzióhoz tartozó RGui-t és installáltak benne R-package-et, akkor be kell másolni az alábbi utasításokat egy csomagban az RGui konzoljába (pl. a Ctrl-C, Ctrl-V billentyűkombináció segítségével), majd lenyomni meg az Enter billentyűt.

```
install.packages("cluster", dependencies = TRUE)
install.packages("jmv", dependencies = TRUE)
install.packages("psych", dependencies = TRUE)
install.packages("olsrr", dependencies = TRUE)
install.packages("GPArotation", dependencies = TRUE)
install.packages("lavaan", dependencies = TRUE)
install.packages("lavaanPlot", dependencies = TRUE)
install.packages("factoextra", dependencies = TRUE)
install.packages("ggplot2", dependencies = TRUE)
install.packages("ClusterR", dependencies = TRUE)
install.packages("Gmedian", dependencies = TRUE)
install.packages("mclust", dependencies = TRUE)
install.packages("MBESS", dependencies = TRUE)
install.packages("MASS", dependencies = TRUE)
install.packages("haven", dependencies = TRUE)
```

Ennek hatására az RGui installálja az utasításokban megadott package-eket (több perc kell hozzá), ami után az RGui-ből ki lehet lépni. Ha az RGui rákérdez a CRAN mirrors ablakban az országra, adjuk meg mondjuk a Czech Republic-ot, de erre jó a legtöbb más európai ország is.

3. Ha a gépen ebben az RGui-ben korábban még nem installáltak R-package-et, akkor a fenti 2. pont installáló utasításai közül először csak az elsőt célszerű bemásolni az RGui konzoljába, majd futtatni azt az Enter billentyűvel. Fogadjuk el a program által felkínált mentési helyet, majd a sikeres installálás után másoljuk be a többi utasítást is egy csomagban és futtassuk őket. Néha RGui-ben problémás az új package-ek egy csomagban történő installálása. Ilyenkor érdemes a package-

eket egyenként installálni (azaz a fenti lista sorait RGui-be egyenként bemásolni és az Enter billentyű megnyomásával futtatni). A CRAN mirrors ablakban ilyenkor is meg kell adni egy országot (vagy azon belül egy várost).

4. Előfordulhat, hogy egyes package-ek (*lavaanPlot*, *jmv* és *psych*) installálása nem sikerül (ilyenkor *Error* üzenet jelenik meg az RGui konzolján), jelezve, hogy az adott package installálásához szükséges mely package-et nem sikerült betölteni. Tapasztalataim szerint segíthet, ha a sikertelen installálások után négy speciális segéd package-et (*stringi*, *gtable*, *fansi* és *utf8*) installálunk az alábbi utasítások segítségével (s utána megismételjük az először sikertelen installálásokat):

```
install.packages("stringi", dependencies = TRUE)
install.packages("gtable", dependencies = TRUE)
install.packages("fansi", dependencies = TRUE)
install.packages("utf8", dependencies = TRUE)
```

5. A package-ek installálása akkor tekinthető sikeresnek, ha a

```
library("cluster")
library("jmv")
library("psych")
library("olsrr")
library("GPArotation")
library("lavaan")
library("lavaanPlot")
library("factoextra")
library("ggplot2")
library("ClusterR")
library("Gmedian")
library("mclust")
library("MBESS")
library("MASS")
library("haven")
```

utasításokat az RGui konzoljában egy csomagban lefuttatva minden package sikeresen betöltődik.

B1.3. A ROP-R szoftver telepítése

A ROP-R telepítésével kapcsolatos általános tudnivalókat az alábbiak szerint foglaljuk össze.

1. A ROP-R szoftver a www.ropstat.com weboldalról tölthető le (ott a ROP-R ikonra kattintva). A program sikeres installálás után a „c:_vargha\ropstat” mappában lesz elhelyezve (ez egyébként a ropstat.exe program helye is). A ROP-R.exe program innen is futtatható, de gyakori használat esetén célszerű a ROP-R számára parancsikont elhelyezni az asztalon vagy a tálcán. Figyeljünk arra, hogy a ROP-R.exe fájl (ahogy a ropstat.exe is) feltétlenül maradjon a „c:_vargha\ropstat” mappában.
2. Az R telepítése során felmásolt Rcmd.exe program elérési útját a ROP-R többváltozós moduljainak első használata előtt be kell állítani a ROP-R Beállítások/R-path menüpontjában. Az R-4.5.2 R-verzió standard telepítése során ez az elérési út: c:\Program Files\R\R-4.5.2\bin\x64\Rcmd.exe.

3. ROP-R-ben az adatállományok ugyanúgy olvashatók be benne, mint a ROPstatban. A beolvasáskor az alapértelmezett fájl típus a ROPstat msw típusa. Ezen kívül a ROP-R elfogad Excel fájlokat² (xls vagy xlsx kiterjesztéssel), szövegfájlokat tabulátorral formattálva vagy csv formátumban, valamint SPSS sav és por adatfájlokat.
4. Mindezek után, ha beolvasunk egy adatfájlt, akkor a ROP-R többváltozós moduljai a „Többváltozós_elemezések_R_segítségével” menüpont segítségével futtathatók.
5. Egy modul elindítása során a ROP-R mindig elkészít és futtat egy vagy több R scriptet (futás alatt jelezve ezt a képernyőn). Ezután az eredmények tetszetős táblázatokba rendezve megtekinthetők a ROP-R nézőkéjében, ahonnan a táblázati formák megtartásával Excelbe vagy Wordbe átküldhetők, illetve egyszerűen átmásolhatók.
6. A ROP-R fontos tulajdonsága, hogy az elemzésekhez elkészített R-scriptek *.r formátumú (pl. EFA.r, CFA.r, PolReg.r, MBCA.r stb.) szövegfájlokba íródnak, egy speciális mappában (c:_vargha\ropstat\aktualis), amelyek az elemzések után a felhasználó által elérhetők, és a ROP-R-ből való kilépés után önállóan is futtathatók R-ben (pl. RGui vagy RStudio segítségével).
7. Ha a futtatás során grafikus ábrák is készülnek (pl. mediációs elemzés vagy konfirmatív faktoranalízis során útvonalábra, vagy hierarchikus klaszteranalízisben dendrogram), akkor a ROP-R ezeket ugyanebben az „aktualis” mappában helyezi el, pdf vagy jpg kiterjesztésű fájlokban.
8. Egyes elemzések (pl. az összes regressziós elemzés) során a nyers R-outputok is megőrződnek ebben a mappában, egy oo.txt (vagy oo1.txt vagy oo2.txt stb.) nevű szövegfájlban, továbbá a ROP-R az R-package-ek futtatása során kapott futási visszajelzéseket is összegyűjti és ugyanitt elmenti egy Rreport.txt nevű szövegfájlban. Ha az R nem találja a futtatni kívánt valamelyik R-package-et vagy problémát talál a ROP-R által elkészített R-scriptben, arról ebben a fájlban tájékozódhatunk.

B1.4. A ROP-R menüpontjai

ROP-R statisztikai moduljai a „Többváltozós_elemezések_R_segítségével” menüponttal futtathatók. A ROP-R többi menüpontja (Fájl, Szerkesztés, Esetek, Változók, Transzformációk stb.) ugyanúgy használható, mint ROPstatban (vö. Vargha, 2016; Vargha, 2020, Melléklet), illetve a legtöbb más statisztikai szoftverben.

A „Fájl” menüpontban lehet különböző típusú adatfájlokat (ROPstat és ROP-R közös msw fájljait, Excel-féle xls és xlsx, SPSS-féle sav és por, valamint csv vagy tabulált szövegfájlokat) beolvasni, új msw fájlt nyitni, msw fájlokat SPSS vagy szövegfájl formátumban elmenteni stb.

A „Szerkesztés” menüpontban a szokásos lehetőségek (Kivág, Másol, Beilleszt stb.) mellett a *Keres*, *cserél* parancs alkalmas arra, hogy egy adatoszlopban kilistázza azokat az eseteket, amelyek ebben az adatoszlopban (változóban) egy megadott karaktersorozatot

² Excel fájlok beolvasásakor az adatfájlt tartalmazó munkalapot kell aktív lapnak beállítani.

tartalmaznak, s ezt egy másikra cserélik. Például így lehet a legegyszerűbben egy változó valamelyik értékét átkódolni (pl. minden 6-os értéket 1-esre átírni egy oszlopban).

Az „Esetek” menüpont alkalmas például bizonyos kijelölt sorok törlésére, új sorok beillesztésére, vagy a sorok átrendezésére valamely változó vagy változók értékei szerint növekvő vagy csökkenő sorrendben.

A „Változók” menüpont alkalmas például bizonyos kijelölt változók törlésére vagy új változók beillesztésére. Változók törléséhez csak annyit kell tennünk, hogy a képernyő bal felső sarkában látható *Kiválaszt* ikonra kattintunk, ezután az egérrel kijelölünk egy-egy cellát a törlendő változók oszlopában, majd végül a *Változók* menüpont *Változók törlése* parancsára kattintunk.

A „Transzformációk” menüpont gazdag lehetőségeket kínál a változók átalakítására vagy új változók létrehozására egy- és kétváltozós műveletek vagy statisztikai függvények segítségével. Itt lehet új random változókat is létrehozni 11-féle eloszlás alapján, vagy átkódolni egyes változókat övezetek vagy kódértékek szerint.

A „Beállítások” menüpontban lehet a ROP-R számára megadni a többváltozós modulok futtatásához nélkülözhetetlen Rcmd.exe program elérési útját, továbbá itt lehet betűméretet is megadni az adatcellák számára.

B1.5. A ROP-R statisztikai moduljai

A ROP-R tíz modulja három csoportba sorolható:

I. Regressziós elemzések

1. Hierarchikus regresszió (HierR)
2. Polinomiális regresszió (PolR)
3. Bináris logisztikus regresszió (BLR)

II. Dimenzió redukciók

4. Főkomponens-analízis (FKA)
5. Feltáró faktoranalízis (EFA)
6. Konfirmatív faktoranalízis (CFA)

III. Klaszteranalízisek

7. Összevonó hierarchikus klaszteranalízis (AHKA)
8. Osztódó hierarchikus klaszteranalízis (OHKA)
9. K-centrumú klaszteranalízis (KKA)
10. Modell-alapú klaszteranalízis (MKA)

Minden modulnak van egy saját menüablaka. Könyvünk tíz fejezete részletesen ismerteti ezeket a modulokat, menüablakuk használatát, valamint a futási eredmények olvasását és értelmezését. A modulokban közös, hogy a bennük kijelölt elemzések végrehajtása után a „c:_vargha\ropstat\aktualis” mappában mindig megtalálunk egy *Rreport.txt* nevű szövegfájlt, mely visszajelez, hogy miként futott le a ROP-R által elkészített R-script. Ha a ROP-R bárhol leáll vagy nem teljes értékű eredménylistát közöl, ebben a fájlban nézhetünk utána a részleteknek, hogy kiderüljön a hiba oka. Keressük itt a „Warning message” és az „Error” kezdetű sorokat!

B2. A könyvben használt adatminták

B2.1. Gyermek és szüleik antropometriai adatai (ANTR1980)

Ez az adatminta, melyre könyvünkben ANTR1980 névvel fogunk a továbbiakban hivatkozni, 4128 magyar gyerek (2147 fiú és 1981 lány) születéskori és 10 éves korban mért testsúlyát, illetve testhosszát (testmagasságát), valamint szüleik hasonló adatait tartalmazza, 1980-as évek elején Magyarországon született gyerekekre vonatkozóan (Darvay, 1997). A minta része egy 6720 személyt tartalmazó gyermekegészségügyi longitudinális vizsgálat teljes adatállományának³. Egy ebből a mintából véletlenszerűen kiválasztott 500 fős adatfájl egyébként megtalálható a ROP-R és a ROPstat demonstrációs adatfájlokat tartalmazó közös „demodat” mappájában⁴.

Az adatállomány változói a következők: *Tsúly0* (születéskori testsúly), *Thossz0* (születéskori testhossz), *Tsúly10* (testsúly 10 éves korban), *Tmag10* (testmagasság 10 éves korban), *Anyasúly* (anya átlagos testsúlya), *Anyamag* (anya testmagassága), *Apasúly* (apa átlagos testsúlya), *Apamag* (apa testmagassága). A testsúlyt mindig kg-ban mérték (születéskor két tizedes pontossággal, minden más viszonylatban egész kg-ra kerekítve), a testhosszt, illetve testmagasságot pedig egész cm-re kerekítve. E változók alapstatisztikáit a B.1. táblázat tartalmazza⁵, ahol a ferdeség és a csúcsosság szignifikanciája a normalitás nullhipotézisének elutasíthatóságát jelzi. A változónkénti minimum és maximum értékek nem jeleznek outliereket, a ferdeség és a csúcsosság értékek pedig arra utalnak, hogy a normális eloszlástól való eltérés a gyermek 10 éves kori testsúlya (*Tsúly10*) esetén a legmarkánsabb. A születéskori testhossz 1,68-as (második legnagyobb értékű) csúcsossága ugyan nem kirívóan magas, mégis azt jelzi, hogy ennél a változónál is az átlagtól mindkét irányban előfordulnak szélsőséges adatok. Leginkább normálisnak a testmagasság változói tűnnek.

B.1. táblázat. Az ANTR1980 minta változóinak alapstatisztikái

Változó	Érvényes esetek	Átlag	Szórás	Minimum	Maximum	Ferdeség	Csúcsosság
Tsúly0	4128	3,21	0,49	1,11	5,00	-0,18***	0,57***
Thossz0	4128	50,3	2,5	36	61	-0,23***	1,68***
Tsúly10	3859	33,7	7,4	19	79	1,36***	2,89***
Tmag10	3860	138,8	6,5	112	169	0,07	0,35***
Anyasúly	3620	65,4	11,9	35	121	0,86***	1,04***
Anyamag	3619	162,1	6,3	140	195	0,14***	0,36***
Apasúly	3185	79,7	11,7	48	130	0,51***	0,30***
Apamag	3181	174,1	6,7	150	200	0,08	0,24**

Jelölés: *: $p < 0,05$ **: $p < 0,01$ ***: $p < 0,001$

³ Terhesek és csecsemők egészségügyi és demográfiai vizsgálata. Témafelelős: dr. Ágfalvi Rózsa (OCSGYI) és dr. Joubert Kálmán (KSH). Az adatok személyes közlésből származnak, melyekért ezúton mondok köszönetet.

⁴ Elérési út: c:_vargha\ropstat\dat\demodat

⁵ ROPstattal kiszámítva

Tekintve, hogy az antropometriai adatoknak van egy erős nemi függése, az alapstatisztikákat a gyermek neme szerinti bontásban is kiszámítottuk (lásd B.2. táblázat). Itt a születési adatok tekintetében láthatunk számottevő nemi különbségeket.

B.2. táblázat. Az ANTR1980 minta gyermekre vonatkozó változóinak nemenkénti alapstatisztikái

Változó	Érvényes esetek	Átlag	Szórás	Minimum	Maximum	Ferdeség	Csúcsosság
Fiúk							
Tsúly0	2147	3,28	0,50	1,20	5,00	-0,16**	0,53***
Thossz0	2147	50,7	2,5	37	61	-0,25***	1,60***
Tsúly10	2012	34,1	7,5	19	79	1,51***	3,26***
Tmag10	2013	138,9	6,5	114	169	0,14*	0,46***
Lányok							
Tsúly0	1981	3,13	0,47	1,11	4,70	-0,27***	0,62***
Thossz0	1981	49,8	2,4	36	60	-0,28***	2,00***
Tsúly10	1847	33,2	7,0	20	73	1,16***	2,17***
Tmag10	1847	138,7	6,5	112	163	0,01	0,23*

Jelölés: *: $p < 0,05$ **: $p < 0,01$ ***: $p < 0,001$

A B.1. táblázatból kiolvasható az érvényes esetek száma is minden változóra. Ez azt mutatja, hogy a szülőktől – különösen az apáktól – rendelkezésre álló adatok határozottan foghíjasabbak, mint a gyermekektől származók, akiknél egyébként – a 10 éves kori adatok vonatkozásában – szintén tekintélyes az adathiányzás. Összességében 3072 személy (1593 fiú és 1479 lány) esetében volt minden változónak érvényes értéke.

B2.2. Adatok egy kötődéskutatásból (KÖT2016)

Egy kötődéssel kapcsolatos pszichológiai kutatásban az ECR–RS (Experiences in Close Relationships – Relationship Structures) kérdőív (Fraley, Heffernan, Vicary és Brumbaugh, 2011) magyar populációra való adaptációja volt a fő feladat. A pszichometriai elemzésekhez egy 336 fős felnőtt magyar minta (124 férfi és 212 nő) állt rendelkezésre, ahol a személyek mind heteroszexuális kapcsolatban éltek (Jantek és Vargha, 2016). Életkori átlaguk 30,0 év volt (szórás: 10,7, min = 18, max = 76). Végzettség szerint 2,4% általános iskolát végzett, 23,3% középiskolát, 22,1% főiskolát, 52,2% pedig egyetemet. Az ECR–RS a kötődést két dimenzió (elkerülés és szorongás) mentén méri négy kötődési személy (anya, apa, romantikus partner, barát) tekintetében. A 40 tételes kérdőív minden viszonylatban ugyanazt a 10 tételt alkalmazza, amelyek közül az első hat item az elkerülést, az ez utáni 4 tétel pedig a szorongást méri.

Az ECR-RS teszt 10 alaptétele 7-fokú Likert-skálán mér, mind a négy kötődési személy esetében ugyanaz, ezeket a fordítottság és a skálahovatartozás jelzésével a B.3. táblázat foglalja össze. Fraley és munkatársai (2011), valamint Jantek és Vargha (2016) szerint a 10 tétel közül az utolsó nem illeszkedik egyértelműen saját – szorongás – skálájának faktorára, ezért a szorongás skálát minden kötődési személy viszonylatában a 7–9. tételek

átlagaként definiáltuk, a kötődés skálát pedig – a fordított tételek pontértékének átfordítása után⁶ – az 1–6. tételek átlagaként. Az így kapott skálák:

- Anyai elkerülés (AnyaiElk)
- Anyai szorongás (AnyaiSzor)
- Apai elkerülés (ApaElk)
- Apai szorongás (ApaSzor)
- Romantikus partnerrel kapcsolatos elkerülés (PárElk)
- Romantikus partnerrel kapcsolatos szorongás (PárSzor)
- Baráttal kapcsolatos elkerülés (BarátElk)
- Baráttal kapcsolatos szorongás (BarátSz).

B.3. táblázat. Az ECR-RS kérdőív tételei a fordítottság és a skálahovatartozás jelzésével

Tétel	Fordított?	Skála
1. Segít, ha a bajban hozzá fordulhatok.	Igen	Elkerülés
2. Általában megbeszélem a gondjaimat, problémáimat vele.	Igen	Elkerülés
3. Átbeszéljük a dolgokat vele.	Igen	Elkerülés
4. Nem okoz nehézséget, hogy rá támaszkodjak.	Igen	Elkerülés
5. Nem szívesen nyílok meg neki.	Nem	Elkerülés
6. Inkább nem mutatom ki neki, hogy mit érzek mélyen a szívemben.	Nem	Elkerülés
7. Gyakran aggódom, hogy nem törődik velem igazán.	Nem	Szorongás
8. Félek, hogy esetleg elfordul tőlem.	Nem	Szorongás
9. Aggódom, hogy nem törődik annyira velem, mint én vele.	Nem	Szorongás
10. Nem bízom benne teljesen.	Nem	Szorongás

Fraley és munkatársai (2011) modelljében biztonságos, jó kötődéssel rendelkeznek azok, akik az elkerülés és a szorongás tekintetében egyaránt alacsony szinten vannak, míg a félelemteli, elkerülő kötődésűek mindkét dimenzióan magas értékűek. A magas elkerülés – alacsony szorongás kombinációja az elutasító–elkerülő, a magas szorongás – alacsony elkerülés kombinációja pedig az elárasztott–megszállott típusra jellemző (vö. Jantek és Vargha, 2016, 1. ábra). A vizsgálatban az ECR–RS kérdőív mellett felvételre került többek között három másik kérdőív is.

1. A Rövidített Személyiségvonás Kérdőív (Big Five Inventory, BFI-44, vö. John, Donahue és Kentle, 1991; John, Naumann és Soto, 2008; magyar adaptáció: Rózsa, Tárnok és Nagy, 2020). A tételeik átlagolásával kiszámított skálák:
 - Extraverzió (Extrav)
 - Érzelmi instabilitás (Érzinst)
 - Barátságosság (Barátság)
 - Lelkiismeretesség (Lelkiism)
 - Nyitottság (Nyitott)
2. A szubjektív jóllétet mérő 5-tételes WHO Jóllét Skála (WBI-5 Well-Being Index; vö. Bech, 1996, 2012; magyar adaptáció: Susánszky, Konkoly Thege, Stauder és Kopp, 2006).

⁶ Az átfordítás igen egyszerűen elvégezhető a ROPstat *Itemanalízis* moduljában.

3. A Beck-féle Depresszió Kérdőív (BDI) rövidített változata (Beck és Beck, 1972; magyar adaptáció: Kopp, Skrabski és Czakó, 1990). A 9 darab 4-fokú Likert skálájú tétel átlagolásával nyert, depressziót mérő skálára a továbbiakban Beck-9 névvel hivatkozunk.

E kötődéskutatás skálaváltozóinak alapstatisztikáit a B.4. táblázat tartalmazza. Ebből a ferdeség és a csúcsosság alapján jól látható, hogy legnagyobb mértékben az ECR-RS szorongás skálái és Beck-9 esetében, legkevésbé pedig a Big Five skálák esetében sérül a normalitás.

B.4. táblázat. A KÖT2016 minta skálaváltozóinak alapstatisztikái

Változó	Érvényes esetek	Átlag	Szórás	Min.	Max.	Ferdeség	Csúcsosság
AnyaElk	329	3,4	1,5	1,0	7,0	0,256	-0,667*
AnyaSzor	323	1,6	1,1	1,0	7,0	2,499***	6,102***
ApaElk	313	4,0	1,5	1,0	7,0	0,069	-0,718**
ApaSzor	309	1,8	1,3	1,0	7,0	2,187***	4,858***
PárElk	336	2,0	1,0	1,0	5,5	1,349***	1,282***
PárSzor	329	2,0	1,3	1,0	7,0	1,749***	2,855***
BarátElk	325	2,4	1,1	1,0	6,5	0,724***	0,259
BarátSz	320	1,8	1,1	1,0	6,7	1,639***	2,589***
Extrav	330	3,4	0,7	1,6	4,9	-0,144	-0,590*
Érzinst	330	2,8	0,8	1,0	5,0	0,296*	-0,109
Barátság	330	3,7	0,6	2,0	5,0	-0,182	-0,309
Lelkiism	330	3,6	0,7	2,0	5,0	-0,141	-0,651*
Nyitott	330	3,8	0,6	1,6	5,0	-0,545***	0,354
WBI-5	335	2,9	0,9	0,2	5,0	-0,418**	0,17
Beck-9	335	1,4	0,4	1,0	3,1	1,694***	3,387***

Jelölés: *: $p < 0,05$ **: $p < 0,01$ ***: $p < 0,001$

B2.3. Magyarország boldogságtérképe 2022 (Btérkép2022)

Az ELTE PPK Pozitív Pszichológiai Kutatócsoportja prof. Oláh Attila irányításával – a Bagdi Bella vezette Jobb Veled a Világ Alapítvánnyal együttműködve – 2016 óta évente elkészíti Magyarország boldogságtérképét egy több ezer fős mintára kiterjedő, pszichometriailag hiteles és validált kérdőíveket alkalmazó internetes vizsgálat eredményei alapján⁷. E kutatás keretében 2022 tavaszán 1003 fő (26% férfi, 74% nő) került be a mintába. Életkori átlaguk 36,2 év volt (szórás: 11,3, min = 17, max = 81). Végzettség szerint a minta 5,6%-a általános iskolát végzett, 44,8%-a középiskolát, 21,7%-a főiskolát, 27,9%-a pedig egyetemet. Családi állapot szerint a minta 25,8%-a egyedül élőnek, 28,2%-a kapcsolatban élőnek, 38,3%-a házasnak, 1%-a özvegynek és 6,7%-a elválnak mondta magát.

Könyvünk elemzéseibe e kérdőíves kutatás változói közül a fentebb említetteken kívül még az alábbiakat vontuk be.

⁷ Lásd pl. <http://boldogsagprogram.hu/magyarorszag-boldogsagterkepe-2019/>

1. Az anyagi helyzetre rákérdező 5 fokozatú tétel [1 = szegény (2,4%), 2 = átlag alatti (11,5%), 3 = átlagos (61,3%), 4 = jómódú (23,1%), 5 = gazdag (1,7%); Anyagi].
2. A testi és fizikai állapotra rákérdező 6 fokozatú tétel (TestFiz).
3. Az általános lelki állapotra rákérdező 6 fokozatú tétel (Áltlelki).
4. A globális jóllétet mérő 17 tételes Globális Jóllét Kérdőív (Oláh, 2021; Vargha et al., 2020) összpontszáma, melyet a teszt érzelmi, pszichológiai, szociális és spirituális jóllét skálájának átlagaként definiálunk (GJóllét).
5. Az élethez való pozitív hozzáállást és általános pozitív attitűdöt mérő Caprara-féle Pozitivitás Skála (Pozitiv; vö. Caprara et al., 2012; Oláh, Vargha & Caprara, 2025). A főskálát az 5-fokozatú tételek átlagolásával kapjuk. A teszt alskálái⁸ az alábbiak (vö. B.5. táblázat):
 - Optimizmus (Optim)
 - Élettel való elégedettség (Élelég)
 - Önbecsülés (Önbecs)
6. A 16-tételes Rövidített Pszichológiai Immunrendszer Kérdőív (PIK16; Oláh, 2005). A 4-fokozatú, Likert-skálájú tételek átlagolásával kiszámított skálák (vö. B.6. táblázat):
 - Megközelítő-monitorozó viselkedés (Mmonit)
 - Alkotó-végrehajtó hatékonyság (AVhat)
 - Reziliencia (Rezil)
 - Önreguláció (Önreg).

B.5. táblázat. A Caprara-féle Pozitivitás Skála tételei a fordítottság és a skálahovatartozás jelzésével

Tétel	Fordított?	Skála
1. Hiszem, hogy a jövőm jól alakul	Nem	Optim
2. Elégedett vagyok az életemmel	Nem	Élelég
3. Általában mások ott vannak számomra, ha szükségem van rájuk	Nem	Élelég
4. Lelkesedéssel és reményekkel telve tekintek a jövőmre	Nem	Optim
5. Összességében elégedett vagyok magammal	Nem	Önbecs
6. Időnként a jövő nem tűnik számomra teljesen egyértelműnek	Igen	Optim
7. Úgy érzem, hogy sok dolog van, amire büszke lehetek	Nem	Önbecs
8. Általában magabiztosnak érzem magamat	Nem	Önbecs

A Btérkép2022 boldogságkutatás kiválasztott változóinak alapstatisztikáit a B.7. táblázat tartalmazza. Ebből a ferdeség és a csúcsosság alapján jól látható, hogy a normalitás ugyan minden változó esetén szignifikánsan sérül, de nem komoly mértékben, mert minden ferdeség érték abszolút értéke jóval kisebb 1-nél, továbbá a csúcsosság értékek abszolút értékének maximuma is éppen hogy csak meghaladja az 1-et (a szubjektív anyagi helyzet esetében a csúcsosság: 1,056).

⁸ szintén tételeik átlagolásával képezve őket

B.6. táblázat. A PIK16 teszt tételei a fordítottság és a skálahovatartozás jelzésével

Tétel	Fordított?	Skála
1. Nagyon örülök magamnak és annak, amit az életben elértem	Nem	M
2. Gyakran vagyok ideges	Igen	Ö
3. Mikor olyan helyzetben voltam, hogy volt valami problémám, megtaláltam a megfelelő embert, aki segített.	Nem	M
4. Gyakran vannak olyan ötleteim, amelyekhez mások eredményesen tudnak kapcsolódni és továbbgondolkodásra készíteti őket.	Nem	AV
5. Könnyen válok türelmetlenné.	Igen	Ö
6. Ha az életemet nézem, úgy látom, hogy az értelmes és következetesen alakul.	Nem	M
7. Gyakran jók a megsejtéseim arról, hogy hogyan gondolkoznak és éreznek az emberek.	Nem	AV
8. Mások szerint is jó problémamegoldó vagyok.	Nem	AV
9. Sikeresen el tudom érni a magam elé kitűzött célokat.	Nem	AV
10. Gyakran van olyan érzésem, hogy a világ csak úgy elmegy mellettem.	Igen	R
11. Ha a dolgok nem terv szerint mennek, könnyen elmegy a kedvem attól, hogy folytassam őket.	Igen	R
12. Olyan ember vagyok, aki nagyon derűlátóan tekint az életre.	Nem	M
13. A fontos dolgok többségét, amelyek velem történnek, előre látni és ellenőrizni tudom.	Nem	M
14. Más emberek úgy tűnik, változnak, magamról úgy érzem, hogy egy helyben állok.	Igen	R
15. Még a váratlan helyzeteket is úgy veszem, hogy azok izgalmas kihívások számomra	Nem	M
16. Hirtelen természetű vagyok (gyakran előbb cselekszem és csak utána gondolkodom)	Igen	Ö

B.7. táblázat. A Btérekép2022 minta kiválasztott változóinak alapstatisztikái ($N = 1003$)

Változó	Átlag	Szórás	Minimum	Maximum	Ferdeség	Csúcsosság
Anyagi	3,10	0,71	1	5	-0,267	1,056***
TestFiz	3,76	1,18	1	6	-0,206**	-0,160
Áltlelki	3,80	1,38	1	6	-0,397***	-0,514***
GJóllét	3,84	1,18	1	6	-0,289***	-0,596***
Pozitiv	3,38	1,01	1	5	-0,368***	-0,820***
PozOpt	3,35	1,02	1	5	-0,337***	-0,618***
PozElég	2,57	0,89	0,75	3,75	-0,416***	-0,887***
PozÖnb	3,41	1,14	1	5	-0,418***	-0,789***
Mmonit	2,75	0,71	1	4	-0,417***	-0,467**
AVhat	3,03	0,65	1	4	-0,630***	0,204
Rezil	2,94	0,84	1	4	-0,570***	-0,565***
Önreg	2,79	0,79	1	4	-0,483***	-0,475**

Jelölés: *: $p < 0,05$ **: $p < 0,01$ ***: $p < 0,001$

I. RÉSZ

Regressziós modulok

A pszichológiai kutatásokban kiemelt fontosságú a változók közötti kapcsolatok vizsgálata. A statisztika a korreláció és a regresszió módszerével áll az ilyen elemzések rendelkezésére. Ebben a részben három regressziós modult ismertetünk. Az 1. fejezet a hierarchikus, a 2. fejezet a polinomiális, a 3. fejezet pedig a bináris logisztikus regresszió módszerével foglalkozik.

Míg a korrelációs együttható két változó közti lineáris típusú kapcsolat szorosságát méri, a regresszió kulcsfogalma az előrejelzés, predikció. Például regresszióval vizsgálható, hogy hogyan függ az egyén globális jólléte (GJóllét) a testi, fizikai állapotától (TestFiz). Előbbit függő, utóbbit független változónak nevezzük. Ha mindkét változót valamilyen 6-fokú kvantitatív skálán mérjük, egy *egyszerű lineáris regressziós függvény* így nézhet ki például:

$$\text{GJóllét} = 2 + 0,5 \cdot \text{TestFiz}.$$

Ez a függvény azért lineáris, mert a GJóllét regressziós becslését egy alapszint (itt 2) és a független változó valahányszorosa (itt 0,5-szöröse) összegeként állítjuk elő. A regressziós függvény segítségével a független változó minden lehetséges értékéhez tudunk egy becslést adni a függő változó értékére. Például egy igen magas, 5-ös értékű TestFiz értékhez a regressziós becslés: $\text{GJóllét} = 2 + 0,5 \cdot 5 = 4,5$, egy igen alacsony, 1-ös értékű TestFiz értékhez pedig a regressziós becslés: $\text{GJóllét} = 2 + 0,5 \cdot 1 = 2,5$.

Ha a globális jóllét előrejelzéséhez a személy 6-fokú skálán mért szubjektív anyagi helyzetét (Anyagi) is figyelembe akarjuk venni, így nézhet ki például egy *többszörös lineáris regressziós függvény*:

$$\text{GJóllét} = 0,95 + 0,44 \cdot \text{TestFiz} + 0,40 \cdot \text{Anyagi}.$$

Egy predikciós hatás nem feltétlenül jelent ok-okozati hatást, hanem pusztán azt, hogy a független változó (X) vagy változók ($X_1, X_2, \dots$) értékének ismeretében képesek vagyunk a függő változó (Y) értékére kisebb hibájú, vagyis jobb előrejelzést (jóslást, predikciót) tenni, mintha a független változók értékét nem ismernénk. A hiba (pontosabban a hibavariancia) csökkenésének mértéke a *többszörös korreláció négyzet* vagy *megmagyarázott varianciaarány*, melyet szokásosan R^2 -tel jelölünk.

Ha arra vagyunk kíváncsiak, hogy a testi, fizikai állapot és a szubjektív anyagi helyzet ismeretében a globális jóllét előrejelzéséhez hozzá tud-e tenni valamit érdemben az életkor ismerete, a *hierarchikus regresszió* módszeréhez jutunk. Statisztikai fogalmakkal ezt a kérdést így szokták megfogalmazni: megnő-e érdemben (szignifikánsan és szakmailag értelmezhető mértékben) R^2 , ha a többszörös lineáris regressziós modellben a TestFiz és az Anyagi mellé bevesszük független változónak az életkor változóját (Életkor) is?

Ha egyetlen független (X) és egyetlen függő (Y) változó kapcsolatára fókuszálunk, számítanunk kell arra, hogy erre a kapcsolatra más változó (M) is hatással van. A mediációs elemzés éppen erre kíváncsi, vagyis arra, hogy M jelenlétében hogyan módosul X és Y kapcsolata, ha ezt a kapcsolatot az $X \rightarrow Y$ lineáris predikció regressziós együtthatójával mérjük. A mediációs elemzés abból áll, hogy lineáris regressziós egyenletek segítségével becslést adunk ezekre a hatásokra, megvizsgáljuk szignifikanciájukat és alkalmas mutató segítségével mérjük az M mediátor változó által közvetített hatás nagyságát. A mediációs elemzést is az 1. fejezetben mutatjuk be.

A többszörös lineáris regressziós modellben a függő változót mindig a független változók súlyozott összegeként⁹, vagyis lineáris függvényeként keressük. Ha megengedjük, hogy a regressziós függvényben a független változók 1-nél magasabb hatvánnyal is szerepeljenek, akkor *polinomiális regresszióról* beszélünk. Például egy egyszerű másodfokú polinomiális regressziós függvény így nézhet ki, ha a szubjektív anyagi helyzet (Anyagi) ismeretében szeretnénk a globális jóllétre előrejelzést tenni:

$$\text{GJóllét} = 3,9 + 0,406 \cdot \text{Anyagi_z} - 0,066 \cdot (\text{Anyagi_z})^2,$$

ahol Anyagi_z az Anyagi változó standardizáltját jelöli. A polinomiális regresszió módszerét a 2. fejezetben ismertetjük.

Ha a hierarchikus regresszió modelljében a függő változó kétértékű, mint például a személy neme, és esetleg a független változók közé kategoriális változókat is be szeretnénk venni, akkor a 3. fejezetben ismertetett *bináris logisztikus regresszió* módszerét alkalmazhatjuk.

A regressziós modulok mindegyikénél fontos, hogy a független változók között ne legyen olyan változó, mely a többi változó determinisztikus lineáris függvénye, vagyis ne álljon fönn a *multikollinearitás*. Ezt a helyzetet minden modul külön teszteli, és ha a ROP-R egy elemzés során ilyet tapasztal, akkor visszajelez az eredménylistán. Tipikusan multikollinearitást okoz, ha egy skála összes tétele mellé bevesszük a tételek összegét vagy átlagát is¹⁰, vagy ha a független változók között szerepelnek egy kategoriális változó összes értékének dummy változói¹¹.

⁹ az összegbe beleértve az alapszintet is

¹⁰ mert az összeg és átlag egyaránt a tételek lineáris függvénye

¹¹ mert bármelyik érték dummy változója (az érték előfordulásának bináris változója) lineárisan kifejezhető a többi érték dummy változója segítségével

1. fejezet

A hierarchikus regresszió (HierR) modul

Ezzel a modullal hierarchikus regresszióelemzés végezhető kvantitatív változókkal (de Jong, 1999; Tabachnick és Fidell, 2013, 5. fejezet; Vargha, 2019, 1. fejezet). A független változók egy blokkindex segítségével különböző blokkokba sorolhatók, így külön tesztelhetők az egyes blokkok hatásai. Minden függő változót a kijelölt független változók együttese segítségével regresszálunk. Ugyanitt lehetőség van a szélsőséges (influenster) esetek azonosítására, valamint mediációs elemzésre is. Ez utóbbira csak akkor, ha pontosan két blokk van megadva, ahol az első (legfeljebb 2 változóval) a független, a második pedig (létszámkorlát nélkül) a mediátor változókat tartalmazza. Ha ez utóbbiak száma 2-nél több, akkor a program a mediátor és a függő változók minden kombinációjára végrehajt egy mediációs elemzést az első blokkal kijelölt 1 vagy 2 független változóra.

Az 1.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, az 1.2. alfejezet a HierR menüablak használatát mutatja be, az 1.3. alfejezet az eredménylista szerkezetét, az 1.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

1.1. Elméleti alapok

1.1.1. Többszörös lineáris regresszió

Tételezzük fel, hogy egy Y kvantitatív függő változóra szeretnénk az $X_1, X_2, \dots, X_p$ kvantitatív független változók (itt p tetszőleges egész szám) segítségével regressziós becslést tenni. Az $\hat{Y}$ *lineáris regressziós függvény* a független változók egy súlyozott összege (lineáris függvénye), mely így néz ki:

$$\hat{Y} = a + b_1X_1 + b_2X_2 + \dots + b_pX_p. \quad (1.1)$$

Az (1.1) formulát *regressziós egyenletnek* is szokták nevezni. Ebben a regressziós egyenletben az a együttható a regressziós becslés alapszintjét beállító, ún. *Y-tengelymetszet* vagy *regressziós konstans*, $b_1, b_2, \dots, b_p$ a *regressziós együtthatók*, ezek együtt a *lineáris regressziós modell paraméterei*. Kicsit más formában (1.1)-et így is fel lehet írni:

$$Y = \hat{Y} + E = a + b_1X_1 + b_2X_2 + \dots + b_pX_p + E, \quad (1.2)$$

ahol $E = Y - \hat{Y}$ az a hibatag (vagy reziduális), amivel ki kell egészíteni az $\hat{Y}$ becslést, hogy megkapjuk a valódi Y -t. E egy változó, melynek értéke személyenként a valódi Y és $\hat{Y}$ regressziós becslésének a különbsége. Az (1.2) regressziós modellben tehát az Y változót egy regressziós becslés és egy regressziós hibatag összegeként írjuk fel (állítjuk elő).

Konkrét adatminta esetén a regressziós paramétereket matematikai módszerekkel úgy határozzuk meg, hogy az $\hat{Y}$ regressziós becslés a lehető legjobban hasonlítson Y -ra. Azt mindig el tudjuk érni, hogy $\hat{Y}$ és Y átlaga megegyezzen, értékeik a minta személyein azonban kisebb-nagyobb mértékben eltérnek egymástól. Ezt a fajta különbözőséget átlagos négyzetes eltérésük, a *hibavariancia* méri. A *legkisebb négyzetes regresszió* egy olyan módszer, melynél az $a, b_1, b_2, \dots, b_p$ regressziós paramétereket úgy határozzuk meg, hogy a hibavariancia, vagyis E^2 mintabeli becslése a lehető legkisebb legyen. A hibavariancia négyzetgyökét a regressziós becslés *standard hibájának* (SH) nevezik. Ez méri a regressziós becslés átlagos pontatlanságát vagy becslési hibáját.

Ha az Y változót minden személy esetén a változó saját átlagával becsüljük, e triviális becslés hibavarianciája az Y változó mintabeli varianciája, $\text{Var}(Y)$. Ha a független változókban van némi információ a független változóra nézve, akkor a regressziós becslés jobb, vagyis kisebb hibavarianciájú lesz, mint a triviális becslés, így ekkor

$$\text{SH}^2 < \text{Var}(Y).$$

Az SH^2 regressziós hibavariancia csökkenésének mértéke a triviális becslés hibavarianciájával való összevetés során az R^2 -tel jelölt *megmagyarázott varianciaarány*:

$$R^2 = (\text{Var}(Y) - \text{SH}^2) / \text{Var}(Y). \quad (1.3)$$

R^2 tehát azt mutatja, hogy SH^2 milyen mértékben, milyen arányban kisebb, mint $\text{Var}(Y)$. Ha a független változókban semmilyen információ nincs az Y változóra vonatkozóan, akkor $R^2 = 0$, ha pedig Y tökéletesen, vagyis hiba nélkül meghatározható a független változók valamilyen lineáris függvénye segítségével, akkor $R^2 = 1$.

1.1.2. Hierarchikus regresszió

Megnő-e érdemben R^2 , ha a többszörös lineáris regressziós modellben a független változók adott csoportjába beveszünk egy vagy több további független változót is? Erre a kérdésre úgy adhatunk választ, ha a független változókat különböző csoportokba, ún. blokkokba soroljuk, majd megnézzük, hogy az első blokk után lépésenként a többi blokk változóit is beemelve a regressziós modellbe hogyan nő meg (milyen mértékben és szignifikánsan-e) az R^2 . A növekedés szignifikanciáját egy F -statisztikával lehet tesztelni. Végső modellnek azt a regressziós modellt fogadjuk el, ami után már egyetlen blokk beemelésével sem nő meg szignifikánsan az R^2 . Ha a modellbe bevont blokkok független változóira felírjuk a regressziós egyenletet, a regressziós együtthatók szignifikanciája változónként tesztelhető, és ennek alapján tovább egyszerűsíthető a modell a nem szignifikáns hatású független változók elhagyásával.

1.1.3. Mediációs regresszió

Ha egy regressziós modellben egy X független változónak van predikciós hatása az Y függő változóra, és Y -nal egy M -mel jelölt másik független változó is kapcsolatban van, felmerülhet a kérdés, hogy az $X \rightarrow Y$ összefüggésben van-e M -nek számottevő közvetítő, mediáló hatása. M mediációs hatásáról akkor beszélhetünk, ha az M változó rögzítése (értékének fix szinten tartása) számottevően csökkenti az X változó Y -ra vonatkozó mérhető predikciós hatását.

Egyszerű lineáris regressziós modellben X teljes hatását Y -ra az

$$Y = i_1 + c_t X + E_1 \quad (1.4)$$

regressziós modell c_t regressziós együtthatója mutatja (i_1 a regressziós konstans, E_1 pedig a hibatag). Ha arra vagyunk kíváncsiak, hogy ugyanez a hatás hogyan alakul M jelenlétében, akkor többszörös lineáris regressziót kell végrehajtanunk, ahol Y a függő, X és M pedig a független változó. Az így adódó

$$Y = i_2 + cX + bM + E_2 \quad (1.5)$$

modellben a c regressziós együttható mutatja az X változó direkt hatását Y -ra, ami M jelenlétében is megmarad. Az M által közvetített, ún. indirekt hatás az, amit M elvesz (átvesz) c_t -től:

$$c_{ind} = c_t - c. \quad (1.6)$$

A mediációs elemzés során döntünk a c_{ind} mediációs hatás szignifikanciájáról (vagyis arról, hogy az elméleti indirekt hatás 0-tól különbözik-e), valamint arról, hogy nagysága szakmailag releváns-e. De hogyan is értelmezhető pontosan szakmailag ez a hatás? A teljes hatás vizsgálatokor c_t mint regressziós együttható azt mutatja, hogy az X változó értékének 1 egységnyi növekedésekor várhatóan hány egységnyi változik Y a saját skáláján. A direkt hatás vizsgálatokor kapott c érték ugyanezt mutatja, de azon feltétel mellett, hogy M -et rögzítjük, vagyis nem engedjük meg, hogy hasson Y -ra. Így a c_{ind} mennyiség jelentése: mennyivel csökken az X változó Y -ra vonatkozó – lineáris – hatása, ha a teljes hatásból elhagyjuk az M által közvetített – szintén lineáris – hatást. Ha az X , Y , M változókat standardizáljuk (z -értékekre térve át), akkor az értelmezés csak annyiban módosul, hogy X és Y értékskáláján 1 egység 1 szórásnyi távolságnak felel meg.

A mediátor változó által átvett rész (indirekt hatás) arányát százalékosan is ki lehet fejezni:

$$\text{Mediáció\%} = 100(\text{indirekt hatás})/(\text{teljes hatás}) = 100c_{ind}/c_t. \quad (1.7)$$

Mediáció% értelmes voltához szükséges, hogy a direkt és az indirekt hatást mérő regressziós együttható (c és c_{ind}) előjele megegyezzen. Ökölszabály szerint a mediációs hatást akkor szokták szakmailag értékelhetőnek tekinteni, ha a Mediáció% eléri a 10-es szintet.

A mediációs elemzés egy speciális útelemzés (Füstös, Kovács, Meszéná és Simonné Mosolygó, 2004, 105. o.; Hunyadi és Vita, 2008, 210. o.), mely maga is a strukturális egyenletekkel történő modellezés (structural equation modeling = SEM, vö. Füstös et al., 2004, 109. o.; Koltai, 2014) körébe tartozik. A klasszikus mediációs elemzés egy maximum likelihood (ML) elemzéstípus, melynek során a becsült paraméterek standard hibáit a változók normalitásának feltételezésével számítják ki (Baron és Kenny, 1986). A modernebb módszerek robusztus standard hibákkal dolgoznak vagy a modellparaméterekre bootstrap módszerrel készítenek intervallumbecslést (MacKinnon, 2012; Saunders és Blume, 2018; Shrout és Bolger, 2002; Hayes és Scharkow, 2013).

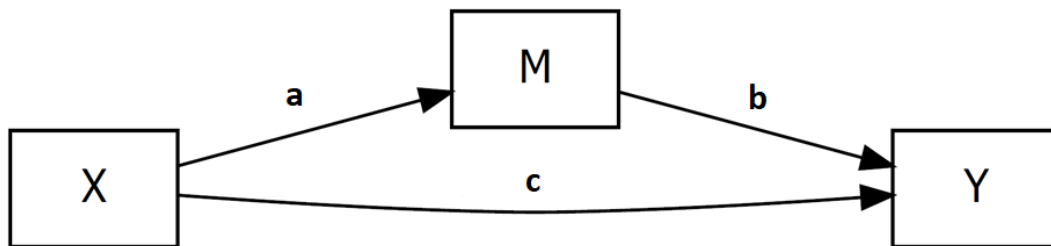
A mediációs elemzés és a parciális korreláció viszonyát a következőképpen írhatjuk le. Előbbiben az a fő kérdés, hogy egy X és Y kapcsolatában van-e egy Z változónak közvetítő hatása? Ha Z hatását parciális korrelációval kiszűröm (vö. Vargha, 2020, 2.3. alfejezet), a

parciális korreláció szignifikanciája a maradék, vagyis a direkt hatás létezését ítéli meg. A mediációs modell annyival több, hogy a mediációs hatás szignifikanciáját, valamint nagyságát (Mediáció%) is tudja vizsgálni. A parciális korreláció a kiszűrés utáni maradék hatást méri.

A mediációs összefüggéseket az útelemzés szokásos *útvonalábrájával* (*path diagram*) szokták ábrázolni. Egy ilyen diagramon az elemzésben résztvevő változók neve szögletes keretben jelenik meg, az őket összekötő nyilakon pedig az útegyütthetők láthatók. Például az 1. 1. ábra diagramja az M változó mediációs hatásának ábrázolása az $X \rightarrow Y$ kapcsolatban. Itt a az $X \rightarrow M$ hatás (a sima $X \rightarrow M$ predikció lineáris regressziós egyenletének regressziós együtthetője), b az $M \rightarrow Y$ hatás (lásd fentebb az (1.5) modellt), c pedig az $X \rightarrow Y$ direkt hatás. Ezek egyben az útvonalábra együtthetői. Az 1.1. ábrán látható mediációs modellt röviden így jelöljük: $X \rightarrow M \rightarrow Y$. Megjegyezzük, hogy a diagram a és b útegyütthetőjének szorzata megegyezik a fentebb már említett indirekt hatással:

$$c_{ind} = ab, \quad (1.8)$$

vagyis az útvonalábráról a direkt (c) és az indirekt (ab) hatás egyaránt leolvasható, illetve kiszámítható. Ez a legegyszerűbb mediációs modellben mindig érvényes $c_{ind} = ab$ összefüggés az általános esetben csak akkor igaz, ha a modell minden változója folytonos és a modellben csak egyetlen mediátor változó szerepel (MacKinnon, Warsi és Dwyer, 1995).



1.1. ábra. Az M változó mediációs hatásának ábrázolása az $X \rightarrow Y$ kapcsolatban útvonalábra segítségével ($X \rightarrow M \rightarrow Y$ modell)

A fentebb ismertetett (1.6) összefüggéssel definiált hagyományos módszer (különbség módszer – difference method) helyett ma már gyakran eleve az (1.8) összefüggéssel (szorzat módszer – product method) definiálják az indirekt hatást (Biesanz, Falk és Savalei, 2010; Nevo, Liao és Spiegelman, 2017), s újabban ezt tekintik a mediációs elemzésekben standardnak (Cheng, Spiegelman és Li, 2021).

Általános esetben egy mediációs modell tetszőleges számú független változót ($X_1, X_2, \dots$) és mediátor változót ($M_1, M_2, \dots$) tartalmazhat. Ennek az általános modellnek a jelölése:

$$(X\text{-ek}) \rightarrow (M\text{-ek}) \rightarrow Y.$$

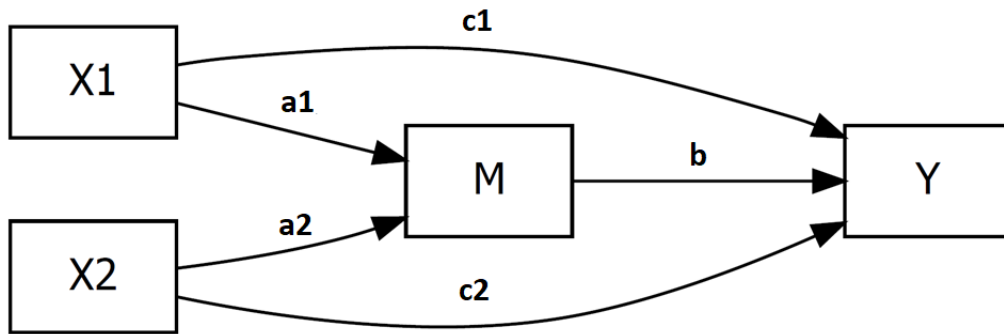
A $c_1, c_2, \dots$ direkt hatások mindig abból a többszörös lineáris regresszióelemzésből adódnak (az X -változók regressziós együtthetőiként), ahol a független változók az X -ek és az M -ek, a függő változó pedig Y , a teljes c_{ij} hatások pedig (minden X változóra külön értékként) abból a regresszióelemzésből (szintén az X -változók regressziós együtthetőiként), ahol a független változók csupán az X -ek (M -ek nélkül), a függő változó pedig Y .

Az útvonalábra $b_1, b_2, \dots$ együtthetői az egyszerű $M_1 \rightarrow Y, M_2 \rightarrow Y, \dots$ lineáris regressziós modellekből adódnak (az M -változók regressziós együtthetőiként), az $a_1, a_2, \dots$

együtthatók pedig az egyszerű $X \rightarrow M$ lineáris regressziós modellekből (az X -ek regressziós együtthatóiként). A mediációs arányokat az így kapott teljes és direkt hatásokból – az (1.6) és a (1.7) formulával összhangban – az alábbi képlet alapján kapjuk minden X_j ($j = 1, 2, \dots$) független változóra:

$$\text{Mediáció}\%(X_j) = 100(c_{tj} - c_j)/c_{tj} = 100(\text{indirekt hatás})/(\text{teljes hatás}) \quad (1.9)$$

Ha a vizsgált hatótényezőt például egyidejűleg két független változóval, $X1$ -gyel és $X2$ -vel mérjük (pl. a kötődést az elkerüléssel és a szorongással), egyetlen M mediátor bevonásával, akkor a külön-külön elemzett egyedi hatások mellett (vagy helyett), érdemes együttes hatásukat is megvizsgálni egyazon mediációs elemzésben ($X1\&X2 \rightarrow M \rightarrow Y$ modell). Ebben az esetben az útvonalábra az 1.2. ábrán látható módon néz ki. Ezen az ábrán $a1$ az $X1 \rightarrow M$, $a2$ az $X2 \rightarrow M$ hatás, b az $M \rightarrow Y$ hatás, $c1$ és $c2$ pedig az $X1 \rightarrow Y$, illetve $X2 \rightarrow Y$ direkt hatás útegyütthatója.



1.2. ábra. A mediációs elemzés útvonalábrája két független változó ($X1$ és $X2$) és egyetlen mediátor változó (M) esetén ($X1\&X2 \rightarrow M \rightarrow Y$ modell)

Míthogy ebben a modellben egyetlen mediátor változó van, a két indirekt hatás különbség módszerrel számított értéke (c_{ind1} és c_{ind2}) megegyezik az $a1$ és b , illetve $a2$ és b útegyütthatók szorzatával (szorzat módszer):

$$c_{ind1} = a1b \text{ és } c_{ind2} = a2b.$$

A teljes hatások ekkor (1.6) segítségével is kiszámíthatók, így a Mediáció% értékek:

$$\begin{aligned} \text{Mediáció}\%(X1) &= 100c_{ind1}/c_{t1}. \\ \text{Mediáció}\%(X2) &= 100c_{ind2}/c_{t2}. \end{aligned}$$

Az (indirekt hatás)/(teljes hatás) arány segítségével definiált Mediáció% [lásd (1.7) és (1.9) formula] néha fura eredményekre vezet. Normál esetben Mediáció% értéke 0 és 100 közé esik. Ha mégsem esik ide, akkor alaposan át kell gondolni a mediációs modell elemeit. Ezt a problémát részletesen tárgyalja Vargha (2023a), ezért az alábbiakban csak összefoglaljuk, hogy milyen esetekben fordul elő ilyen anomália.

- Ha az indirekt hatás ellentétes irányú, mint a direkt hatás, vagyis ha a mediátor változó jelenlétében az X változó hatása nem hogy nem csökken, hanem megnő, akkor Mediáció% értéke negatív lesz. Mediáció% < 0 esetén tehát mindig ilyen szituáció áll fenn.

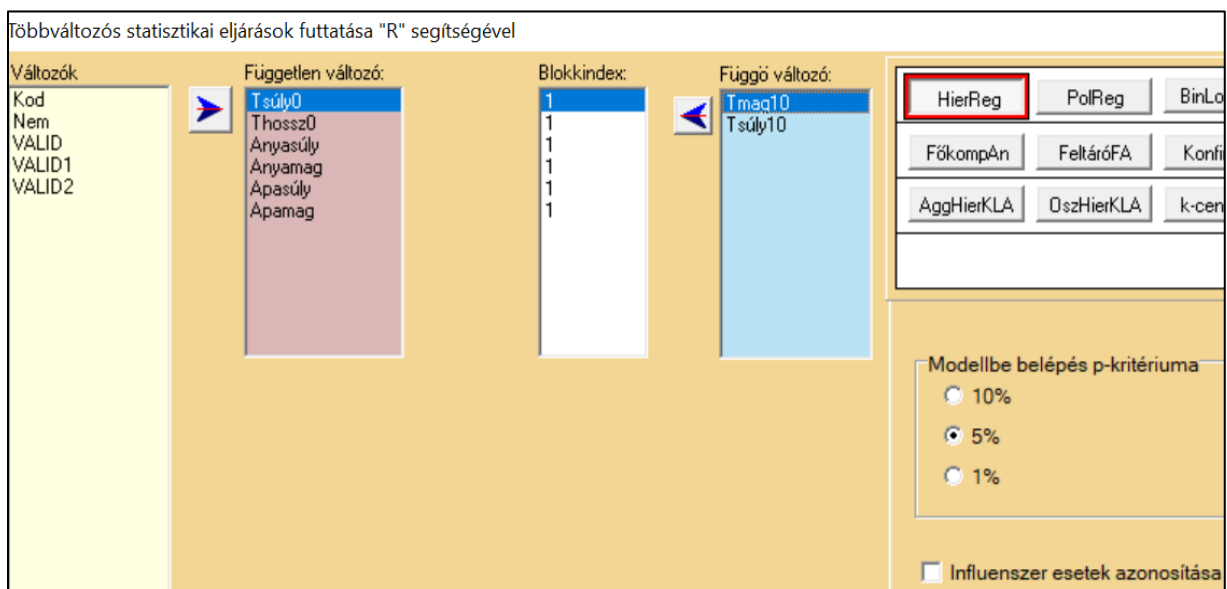
- Ha a mediátor változó modellbe lépésével az X változó hatása oly mértékben lecsökken, hogy még az iránya is megváltozik (pozitívból negatívba), ez 100 fölé viszi a Mediáció% értékét. Ilyen szituáció tehát akkor lép fel, ha a mediátor változó jelenlétében az X változó hatásának iránya megváltozik.
- Előfordulhat az is, hogy van egy egyformán nagy, de ellentétes irányú direkt és indirekt hatás, ami miatt a teljes hatás alacsony szintű lesz, a Mediáció% pedig 100-nál nagyobb (ha a direkt hatás abszolút értéke kisebb, mint az indirekt hatásé) vagy -100-nál kisebb (ha a direkt hatás abszolút értéke nagyobb, mint az indirekt hatásé).

A komplexebb mediációs modellek matematikáját és képleteit nem részletezzük, a gyakorlat szempontjából fontos tudnivalókat az illusztratív példákkal együtt ismertetjük. További információk megtalálhatók Vargha (2023a) cikkében, mely részletesen bemutatja a JASP és a jamovi szoftver mediációs modulját is.

1.2. A HierR menüablak

Többszörös lineáris regresszióelemzést, hierarchikus regresszióelemzést vagy mediációs elemzést ROP-R-ben a hierarchikus regresszió (röviden HierReg vagy HierR) modul segítségével végezhetünk el. HierR a jamovi szoftver (The jamovi project, 2021; Şahin és Aybek, 2019) által használt *jmv*, valamint a *lavaan* (Rosseel, 2012) R-package-re épül.

Ha ROP-R-ben belépünk a HierR modulba, mindössze az a dolgunk, hogy az elemzéshez szükséges változókat a képernyő bal oldalán megjelenő változólistából áttegyük jobbra, a független változókat a „Független változó”, a függő változókat pedig a „Függő változó” ablakba, majd megadjuk a megfelelő blokkindex értékeket (lásd 1.3. ábra). Minden függő változót a kijelölt független (magyarázó) változók együttesével magyarázunk egy többszörös lineáris regressziós modellben.



1.3. ábra. A HierR modul menüablakának fő része

1.2.1. Blokkok megadása HierR-ben

A független változókat a blokkindex segítségével különböző blokkokba sorolhatjuk. A blokkindexek 1 és 9 közötti egész számok, az azonos blokkindexű változók ugyanabba a blokkba tartoznak. HierR segítségével mérhető és tesztelhető, hogy a független változók egymás utáni blokkjai mennyire járulnak hozzá a függő változó bejósolásához. Egy blokk akkor lép be a regressziós modellbe, ha szignifikánsan megemeli az R^2 értéket. Ehhez a szignifikancia szintje a *Modellbe belépés p-kritériuma* panelen állítható be. Persze HierR-rel lehet sima többszörös lineáris regresszióelemzést is végrehajtani, ez esetben a blokkindexeket hagyjuk meg az alapértelmezett 1-es értéken.

1.2.2. Influenster esetek

Külön lehetőség HierR-ben a regressziós elemzés eredményére nagy hatást gyakoroló, ún. influenster esetek – többnyire outlierok – azonosítása és az ezt mérő Cook-féle D-távolság (Marzjarani, 2015) elmentése. Rendszerint a $D_{min} = 4/n$ (itt n az elemzésre kijelölt változók mintájában a komplett esetek száma) küszöbnél nagyobb D-értékű személyeket tekintik outlieroknak, bár ez nem köbe vésett szabály. A menüablakban a „Standard influenster küszöb ($4/n$) szorzója” rovat alapértelmezés szerinti 1 értékének megváltoztatásával lehet ezt a küszöböt igény szerint módosítani. Ha ugyanitt bejelöljük az „Influenster változó elmentése” opciót, akkor a ROP-R minden személyre kiszámítja és az aktuális msw adatállományhoz illeszti a D/D_{min} hányadossal mért, ún. Cook-féle relatív távolságot.

1.2.3. Mediációs elemzés kijelölése

Ha mediációs elemzést akarunk végrehajtani, akkor a következőképpen kell eljárni. HierR akkor kérdez rá a mediációs elemzés végrehajtására, ha a független változóknak pontosan két blokkja van és az 1. blokk egy vagy két változót tartalmaz. Az első blokk változói lesznek a valódi független változók (X -ek), a 2. blokk változói pedig a mediátor változók (M -ek). Az alábbi esetek lehetségesek.

- Ha a független változók ablakában az 1. és a 2. blokk, valamint a függő változók ablaka egyaránt 1-1 változót tartalmaz, akkor $X \rightarrow M \rightarrow Y$ mediációs elemzés kerül végrehajtásra – standard hierarchikus regresszióelemzés mellett (lásd 1. ábra).
- Ha az 1. blokk 2 és a 2. blokk 1 változót tartalmaz, akkor $X1 \& X2 \rightarrow M \rightarrow Y$ mediációs elemzés kerül végrehajtásra (lásd 2. és 5. ábra).
- Ha az 1. blokk 1 és a 2. blokk 2 változót tartalmaz, akkor $X \rightarrow M1 \& M2 \rightarrow Y$ mediációs elemzés kerül végrehajtásra (lásd 6. ábra).
- Ha az 1. és a 2. blokk egyaránt 2 változót tartalmaz, akkor $X1 \& X2 \rightarrow M1 \& M2 \rightarrow Y$ mediációs elemzés kerül végrehajtásra (7. ábra).

Ha az M változók száma 1-nél nagyobb, akkor ROP-R a kijelölt mediátor változók hatását külön-külön modellekben is megnézi. Mindezeket az elemzéseket ROP-R a függő változók ablakában megadott minden Y változóra külön-külön végrehajtja.

Röviden: a ROP-R menürendszerében – a független változók ablakával kapcsolatban – az „1. blokk” elnevezés az 1-2 db független változó, míg a „2. blokk” a tetszőleges számú mediátor változó jelölésére szolgál. Ha utóbbiból csak 1 vagy 2 van, akkor mindezen változókat egyetlen mediációs modellben elemezzük, 2-nél több mediátor változó esetén viszont mediátor változót külön modellben.

Ha minden opciót igény szerint beállítottunk a menüablakban, kattintsunk a „Futtat” gombra a képernyő jobb alsó sarkában és elindul a ROP-R kijelölt elemzése. A modul elemzéseinek végrehajtása után a „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzésekhez elkészített ideiglenes adatfájlt (tmpdat.txt), a futtatott R-scriptet (HierReg.r), a mediációs elemzésekhez tartozó egy vagy több útvonalábrát (path diagramot) modplot*.pdf alakú nevekkal, valamint a részletes R eredménylistát (oo.txt).

1.3. Mit tartalmaz a HierR eredménylistája?

A HierR eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
 - A független változók blokkindexek szerinti csoportosítása (változók besorolása a definiált blokkokba)
- Minden függő változóval elvégzett elemzésre:
- Az egyes blokkok bővülő regressziós modelljeinek főbb jellemzői
 - A végső modell regressziós együtthatói
 - Multikollinearitás diagnosztika
 - Mediációs hatások regressziós együtthatói (csak mediációs elemzésekben)
 - Mediált (indirekt) hatások (csak mediációs elemzésekben)

1.4. A HierR szemléltetése valódi adatokon

1.4.1. Hierarchikus regresszioelemzés antropometriai adatokkal

Tételezzük fel, hogy a gyermek 10 éves kori testmagassága (Tmag10) érdekel bennünket, amit szeretnénk a gyermek születésekor megjósolni a születési adatok (testsúly és testhossz), valamint a szülők magasságának és súlyának ismeretében. Ehhez először többszörös lineáris regresszioelemzést végzünk a 4128 fős ANTR1980 minta alapján, ahol 3072 személy rendelkezett minden antropometriai változó esetén érvényes értékkel (vö. B2.1. alpont).

A változók kiválasztását az 1.3. ábrának megfelelően kell elvégezni, azzal az egyetlen eltéréssel, hogy a Tsúly10 változót nem szükséges a függő változók oszlopába betenni (hacsak nem akarunk a 10 éves kori testsúllyal is hasonló elemzést elvégezni). A regresszioelemzést tehát a Tmag10 függő változóval és a Tsúly0, Thossz0, Anyasúly, Anyamag, Apasúly, Apamag független változókkal végezzük.

A ROP-R outputján először is megismerhetjük az összes változó alapstatisztikáját (ezt itt most nem mutatjuk be), majd a független változók egyetlen definiált blokkjára vonatkozó regressziós eredményeket (lásd 1.1. táblázat).

1.1. táblázat. A többszörös lineáris regresszioelemzés eredménye ROP-R-ben (függő változó: Tmag10; Független változók: gyermek születési testsúlya és testhossza, valamint anyja és apa testsúlya és magassága)

Modell	R	R ²	R ² kor	St.hiba	F	f1	f2	p
1	0,4997	0,2497	0,2483	5,6562	170,04	6	3065	< 0,001

Az 1.1. táblázatból az alábbi következtetéseket vonhatjuk le.

- Az R többszörös korrelációs együttható 0,50, ennek négyzete, R^2 és R^2_{kor} pedig 0,25 körüli¹², ami a független változók erős prediktív erejét jelzi.
- A regressziós becslés standard hibája (St.hiba) 5,7 körüli, vagyis az adott független változók ismeretében kb. 6 cm pontossággal lehet a gyerek 10 éves kori magasságát előrejelezni.
- Ez a predikciós kapcsolat természetesen erősen szignifikáns ($F = 170,04$; $f_1 = 6$, $f_2 = 3065$; $p < 0,001$).

Az output következő fontos eleme a végső modell regressziós együtthatóinak táblázata (lásd 1.2. táblázat). Mivel csak egyetlen blokkot definiáltunk, a végső modell a sima többszörös lineáris regresszióelemzés modellje.

1.2. táblázat. A végső modell regressziós együtthatói ROP-R-ben (a függő és a független változók ugyanazok, mint az 1.1. táblázat esetében)

Prediktor	Együttható	St.hiba	St.eh.	t	p
Konstans	35,098	3,925	-	8,94	< 0,0001
Tsúly0	-0,316	0,332	-0,024	-0,95	0,341
Thossz0	0,605	0,066	0,226	9,14	< 0,0001
Anyasúly	0,048	0,010	0,087	5,03	< 0,0001
Anyamag	0,205	0,019	0,195	11,00	< 0,0001
Apasúly	0,062	0,010	0,111	6,19	< 0,0001
Apamag	0,190	0,018	0,196	10,58	< 0,0001

Az 1.2. táblázatból az alábbi következtetéseket vonhatjuk le. A születés kori testsúly (Tsúly0) kivételével minden független változónak szignifikáns ($p < 0,0001$) hatása, vagyis prediktív információja van Tmag10-re nézve. A hatás mértékét a béta (β) standardizált együttható (St.eh.) jelzi, mely 0,10-es szinten már szakmailag értelmezhető, 0,20 felett pedig már erős hatásról beszélhetünk. Jelen esetben Tmag10-re nézve – nem meglepő módon – Thossz0, vagyis a születés kori testhossz rendelkezik a legerősebb prediktív hatással, de közel ekkora az anya ($\beta = 0,195$) és az apa ($\beta = 0,196$) magasságának prediktív hatása is.

A változók közötti túl erős lineáris összefüggés, amit *multikollinearitásnak* hívunk, megbízhatatlanná teszi a regressziós együtthatók becslését és szignifikanciájuk vizsgálatát. Ezt tehát ellenőrizni, kell, ezt segíti a ROP-R outputján a független változók multikollinearitás diagnosztikájáról tájékoztató táblázat az R^2 , TOL, VIF értékekkel. Minden független változó esetén az R^2 -érték a többi független változó által megmagyarázott, a $TOL = 1 - R^2$ (az ún. *tolerancia*) a többi független változó által meg *nem* magyarázott varianciaarány, a VIF ¹³ pedig a TOL reciproka: $VIF = 1/TOL$. A multikollinearitást a nagyon magas R^2 és VIF, illetve a nagyon alacsony TOL értékek jelzik. A küszöbök tekintetében nincs általános egyetértés (O'Brien, 2007). Van, aki szerint már 5 és 10 közötti VIF (azaz 0,1 és 0,2 közötti TOL) érték esetén is problémás lehet a regresszióelemzés, és 10 fölötti VIF esetén gyakran megbízhatatlan a regressziós együtthatók becslése (Akinwande, Dikko és Samson, 2015). Saját tapasztalataim alapján 20 alatti VIF értékek esetén nincs ok az aggodalomra, 20 fölötti

¹² A ROP-R outputján R^2 -et R^2 jelöli.

¹³ variance inflating factor

VIF értékű független változókat *érdemes*, 100 fölötti VIF értékűeket pedig ki *kell* hagyni az elemzésekből. A jelen esetben (lásd 1.3. táblázat) a legnagyobb VIF-értékek 2,5 körüliek, így ezzel most minden rendben van.

1.3. táblázat. A független változók multikollinearitás diagnosztikája

Változó	R ²	TOL	VIF=1/TOL
Tsúly0	0,599	0,401	2,496
Thossz0	0,599	0,401	2,497
Anyasúly	0,175	0,825	1,212
Anyamag	0,222	0,778	1,285
Apasúly	0,242	0,758	1,319
Apamag	0,290	0,710	1,408

Következő szakmai kérdésünk, hogy összességében a gyermek saját adatai, vagy pedig a szülők adatai rendelkeznek nagyobb prediktív információval Tmag10 bejósolásában? E kérdés megválaszolására hierarchikus regresszióelemzést végzünk, két blokkot definiálva. Elsőként a gyermek két változója (Tsúly0 és Thossz0) mellé tesszük az 1-es, a szülők négy változója (Anyasúly, Anyamag, Apasúly, Apamag) mellé pedig a 2-es blokkindexet, majd egy második elemzésben megfordítjuk a sorrendet.

ROP-R outputján legfontosabb az a táblázat, amely az egyes blokkok bővülő (az alsóbb szinteket mindig tartalmazó) regressziós modelljeinek főbb jellemzőit foglalja össze. Ennek bal oldali része a bővülő modellek szignifikanciáját, R^2 értékét (R^2), R^2_{kor} értékét (R^2_{kor}), valamint standard hibáját (St.hiba) tartalmazza (lásd 1.4. táblázat), jobb oldali része pedig a blokkok belépésének hatását mutatja R^2 -re, valamint az R^2 növekedésének szignifikanciáját jelzi (lásd 1.5. táblázat).

1.4. táblázat. HierR eredménye Tmag10-re, két blokk definiálásával (1. blokk: Tsúly0 és Thossz0; 2. blokk: Anyasúly, Anyamag, Apasúly, Apamag)

Modell	R	R ²	R ² kor	St.hiba	F	f1	f2	p
1	0,3053	0,0932	0,0926	6,2183	157,73	2	3069	< 0,001
2	0,4997	0,2497	0,2483	5,6562	170,04	6	3065	< 0,001

1.5. táblázat. A blokkok belépésének hatása R^2 -re és a növekmény szignifikanciája

Modell	R ² +	F	f1	f2	p-érték
1					
2	0,1565	159,864	4	3065	< 0,0001

Az 1.4. és az 1.5. táblázat alapján most az alábbi következtetések vonhatók le.

- A gyermek születési változói a Tmag10 varianciájának 9,3%-át magyarázzák (1. modell, R^2 oszlop). Ez a hatás $p < 0,001$ szinten szignifikáns.
- Ha az első blokk mellé belép a 2. blokk is (2. modell), R^2 értéke 25%-ra nő. Ez természetesen ugyanaz, mint amit a blokkosítás nélküli többszörös lineáris regresszióelemzés során korábban már kaptunk (lásd 1.1. táblázat).

- A 2. blokk belépésének hatására R^2 értékének növekedése 0,1565 (1.5. táblázat, 2. modell, R^2+ oszlop), azaz csaknem 16%, ami igen tetemes és természetesen erősen szignifikáns ($F = 159,864$; $f_1 = 4$, $f_2 = 3065$; $p < 0,0001$).

Ha a hierarchikus elemzésben a két blokkot úgy definiáljuk, hogy a gyermek két változója mellé a 2-es, a szülők négy változója mellé pedig az 1-es blokkindexet tesszük, akkor az 1.6. és az 1.7. táblázatban látható eredményekhez jutunk.

1.6. táblázat. HierR eredménye Tmag10-re, két blokk definiálásával (1. blokk: Anyasúly, Anyamag, Apasúly, Apamag; 2. blokk: Tsúly0 és Thossz0)

Modell	R	R^2	R^2 korrr	St.hiba	F	f1	f2	p
1	0,4572	0,2090	0,2080	5,8075	202,652	4	3067	< 0,001
2	0,4997	0,2497	0,2483	5,6562	170,04	6	3065	< 0,001

1.7. táblázat. A blokkok belépésének hatása R^2 -re és a növekmény szignifikanciája

Modell	R^2+	F	f1	f2	p-érték
1					
2	0,0407	83,113	2	3065	< 0,0001

Az 1.6. táblázat azt mutatja, hogy a szülők változói érezhetően gazdagabb információt tartalmaznak a gyermek 10 éves magasságára ($R^2 = 20,1\%$), mint a gyermek születési adatai ($R^2 = 9,3\%$; vö. 1.4. táblázat), ami egyben azt is jelenti, hogy a 10 éves kori magasság bejósolásában a születési adatok jóval kevesebbet tesznek hozzá a szülők adataihoz (4,1%, vö. 1.7. táblázat), mint amit a szülők adatai tesznek hozzá a gyermek saját születési adataihoz (15,7%, vö. 1.5. táblázat). Mivel mindkét R^2 -növekedés szignifikáns, a végső regressziós modell ugyanaz a mindkét blokkot tartalmazó modell lesz, amelynek változónkénti kiértékelése az 1.2. táblázatban látható. Ez a regressziós modell már csak úgy egyszerűsíthető, ha a független változók közül elhagyjuk a nem szignifikáns hatású születéskori testsúly (Tsúly0) változót.

1.4.2. Hierarchikus regresszioelemzés a pszichológiai kötődés vizsgálatában

Számos adat utal arra, hogy a depresszió jelentős részben öröklött (Alshaya, 2022). Kendall és munkatársai (2021) szerint például a major depresszió kialakulásában az öröklési arány 30-50%. Kérdés, hogy ezt a viszonylag magas öröklött hajlandóságot mennyire tudják befolyásolni a személy kötődési jellemzői? A problémát a KÖT2016 minta változói (lásd B2.2. alpont) segítségével tudjuk körüljárni, ahol egyaránt találunk kötődési és a depressziót mérő változót.

A megfogalmazott szakmai kérdés megválaszolására egy hierarchikus regresszioelemzést hajtottunk végre, ahol a depressziót a Beck-féle Rövidített Depresszió Kérdőív Beck-9 skálájával, az öröklött hajlandóságot a BFI-44 Big Five teszt 5 skálájával (Extrav, Érzinst, Barátság, Lelkiism, Nyitott), a kötődést pedig az ECR-RS teszt 8 skálájával (Anyaeik, AnyaSzor, ApaEik, ApaSzor, PárEik, PárSzor, BarátEik, BarátSz) mérjük. Az elemzésben Beck-9 lesz a függő változó, a független változók 1. blokkja az 5 Big Five skálát, 2. blokkja pedig a 8 kötődési változót tartalmazza.

Első elemzésünkben kértük az influenszer esetek azonosítását (a Cook-féle D távolság alapján) és az influenszer változó mentését. A kilistázott lehetséges influenszer esetek közül

az első tízet az 1.8. táblázatban foglaltuk össze. Ebből megállapíthatjuk, hogy egyetlen (a 229 sorszámu) személynek van kiemelkedően magas hatása a regressziós elemzés eredményére. Ahhoz, hogy ezt a személyt kihagyjuk az elemzésből, az első futás során létrehozott CookRD változó övezetes csoportkijelölésénél 0–4 övezetet állítottunk be az elfogadható „normál” CookRD értékekre (a Változók deklarációi ablakban), majd a HierR menüablakban CookRD-t (a változólista felújítása után) feltételes csoportosító változónak jelöltük ki.

1.8. táblázat. Influenzser esetek HierR-ben

D/Dmin	Esetsorszám
10,451	229
3,844	90
3,298	193
2,874	305
2,754	200
2,682	157
2,667	311
2,658	307
2,412	20
2,236	72

A fentiek szerinti módosítások után újra futtattuk HierR-t, melynek legfontosabb eredményeit az 1.9. és az 1.10. táblázatban foglaltuk össze.

1.9. táblázat. HierR eredménye Beck-9-re, két blokk definiálásával (1. blokk: BFI Big Five teszt 5 skálája; 2. blokk: ECR-RS kötődésteszt 8 skálája)

Modell	R	R ²	R ² korrr	St.hiba	F	f1	f2	p
1	0,5468	0,2990	0,2821	0,2365	17,743	5	208	< 0,001
2	0,6776	0,4592	0,424	0,2077	13,062	13	200	< 0,001

1.10. táblázat. A blokkok belépésének hatása R²-re és a növekmény szignifikanciája

Modell	R ² +	F	f1	f2	p-érték
1					
2	0,1602	7,405	8	200	< 0,0001

Az 1.9. és az 1.10. táblázat alapján most az alábbi következtetéseket vonhatjuk le.

- A BFI-44 – nagyrészt örökletesnek tekintett – személyiségvonás változói a depressziót mérő Beck-9 változó varianciájának 29,90%-át magyarázzák (1. modell, R² oszlop), ami igen magas, megerősítve a depresszióra való hajlandóság örökletes eredetét.
- Ami viszont igen meglepő: ha az első blokk mellé belép a 2. blokk is (2. modell), R² értéke még számottevően megnő (45,92%-ra).

- A 2. blokk belépésének hatására R^2 értékének növekedése 0,1602 (1.10. táblázat, 2. modell, R^2 oszlop), vagyis 16 százalékpont, ami tetemes és természetesen erősen szignifikáns ($F = 7,405$; $f_1 = 8$, $f_2 = 200$; $p < 0,0001$).

1.11. táblázat. A végső modell regressziós együtthatói (függő változó: Beck-9)

Prediktor	Együttható	St.hiba	St.eh.	t	p
Konstans	1,25760	0,2216	-	5,676	< 0,0001
Extrav	-0,02482	0,0272	-0,0578	-0,912	0,3628
Érzinst	0,12002	0,0251	0,3065	4,782	< 0,0001
Barátság	-0,07894	0,0334	-0,1550	-2,363	0,0191
Lelkiism	-0,04596	0,0261	-0,1032	-1,758	0,0803
Nyitott	-0,01356	0,0294	-0,0267	-0,462	0,6446
AnyaElk	-0,00881	0,0123	-0,0470	-0,715	0,4752
AnyaSzor	0,00735	0,0180	0,0255	0,409	0,6830
ApaElk	-0,00750	0,0115	-0,0408	-0,651	0,5156
ApaSzor	0,02801	0,0163	0,1102	1,722	0,0867
PárElk	0,04828	0,0191	0,1618	2,532	0,0121
PárSzor	0,02662	0,0149	0,1084	1,786	0,0757
Barátelk	0,07667	0,0156	0,2925	4,915	< 0,0001
Barátsz	0,00038	0,0147	0,0015	0,026	0,9795

Összegezve azt mondhatjuk, hogy az igen erős genetikai meghatározottság ellenére a depresszióra való hajlandóság szintjére komoly hatást gyakorolnak az egyén kötődési jellemzői is. Hogy pontosan melyek, azt a végső (mindkét blokk változóit tartalmazó) regressziós modell együtthatóinak táblázatából olvashatjuk ki (lásd 1.11. táblázat), mely azt mutatja, hogy a kötődési jellemzők közül nincs szignifikáns hatása az anyai és az apai kötődésváltozóknak, viszont erősen szignifikáns ($p < 0,0001$) és szakmailag is értelmezhető hatása van a baráttal kapcsolatos elkerülésnek ($\beta = 0,2925$), továbbá szignifikáns ($p = 0,0121$) és értelmezhető mértékű hatása van a romantikus partnerrel kapcsolatos elkerülésnek ($\beta = 0,1618$).

1.4.3. Mediációs elemzések a globális jóllét bejólásában

Ép testben ép lélek, tartja a híres ókori közmondás. Annál inkább igaz lehet, hogy aki mind a testi-fizikai (TestFiz), mind általános lelki állapota (Áltlelki) szerint jól érzi magát, annak a globális jóllét (GJóllét) tekintetében se lehet gondja. Ezt az összefüggést egyszerűen meg tudjuk vizsgálni az 1003 fős Btérkép2022 minta (vö. B.2.3. alpont) változói segítségével, sima többszörös lineáris regresszióelemzést végezve (független változók: TestFiz és Áltlelki, függő változó: GJóllét). Ennél azért szeretnénk kicsit többet is megtudni. Vajon a testi jó állapot és a globális jóllét közti nyilvánvaló pozitív kapcsolatból mennyi köszönhető az élethez való pozitív hozzáállásnak, az általános pozitív attitűdnek, amit például a Caprara-féle Pozitivitás Skálával (Pozitiv) mérhetünk? És hogyan alakul a mediáció, ha az Áltlelki változót is bevesszük TestFiz mellé független változónak?

Első lépésként influenszer elemzést végzünk HierR-ben, GJóllétet függő, a másik 3 említett változót pedig független változóként kijelölve. A regresszió eredményére legnagyobb hatású hat influenszer esetet az 1.12. táblázatban foglaltuk össze, az első négy esetben a 4 regresszióba bevont változó értékét is feltüntetve.

1.12. táblázat. Influenster esetek GJóllét bejósolásában

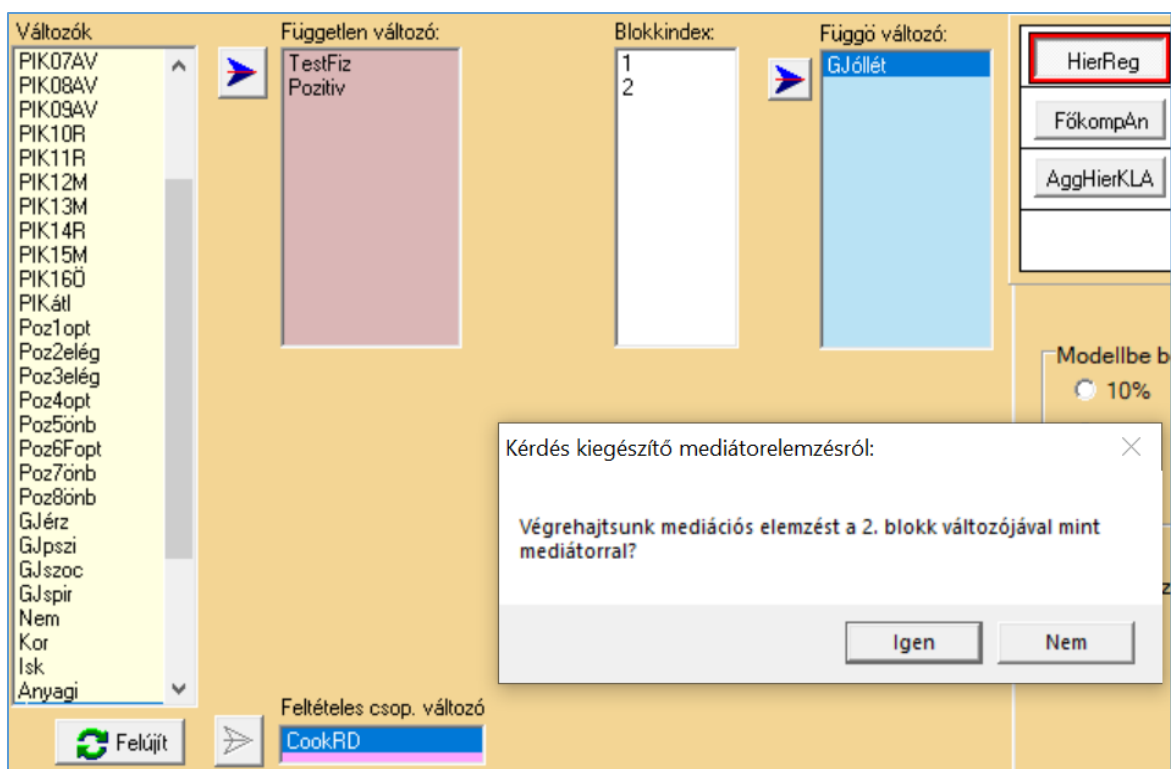
D/Dmin	Esetsorszám	TestFiz	Áltlelki	Pozitiv	GJóllét
10,043	501	4	1	3,75	1,45
9,803	68	5	6	4,63	1
8,342	428	4	6	2,88	1,59
6,666	950	1	1	3	1
4,739	24				
4,456	508				

Az 1.12. táblázat alapján az alábbiakat állapíthatjuk meg.

- A D/Dmin értékek közül az első 4 emelkedik meredeken a többi fölé, ezek az esélyes outlierok.
- Ha ennek a 4 személynek a konkrét adatait is megnézzük, akkor minden esetben találunk olyan inkonzisztenciát, ami kérdésessé teszi, hogy érvényes adatokkal állunk-e szemben.
- Például a 68. és a 428. sorszámú személy esetében nem érthető, hogy maximális (6-os) aktuális lelki állapot mellett miként lehet a globális jólléte minimális (1), vagy ahhoz igen közeli (1,59).
- A 950. sorszámú személy esetében nem érthető, hogy miként lehet a globális jólléte minimális (1), miközben az élethez való pozitív hozzáállása átlagos (5-fokú skálán 3).
- Végül az 501. sorszámú személy esetében az a gyanús, hogy az élethez való átlagosnál jobb (3,75) pozitív hozzáállása mellett hogyan lehet minimális szintű (1) lelki állapotban.

Mindezt figyelembe véve célszerűnek látszik ezt a 4 személyt a mediációs elemzésből kihagyni. Ehhez még egy regressziós futtatással elmentjük az influenszer változót (CookRD) néven, majd ennek csoportkijelölésénél 0–5 övezetet állítunk be az elfogadható „normál” CookRD értékekre (a Változók deklarációi ablakban). Végül a HierR menüablakban CookRD-t (a változólista felújítása után) a mediációs elemzéshez feltételes csoportosító változónak jelöljük ki.

Első mediációs elemzésünkben TestFiz és Pozitiv lesz a függő, GJóllét a független változó, CookRD-t pedig a megbeszéltek szerint feltételes csoportosító változónak jelöljük ki. A két független változó közül TestFiz blokkindexét meghagyjuk az alapértelmezett 1-esnek, Pozitiv mellé azonban 2-es blokkindexet jelölünk ki. Ekkor a „Futtat” gombra kattintva ROP-R megkérdezi, hogy végrehajtsunk-e mediációs elemzést a 2. blokk változójával (ez most Pozitiv), mint mediátorral (lásd 1.4. ábra). Ebben a modulban a ROP-R mindig felteszi ezt a kérdést, ha a független változóknak két blokkja van és az első blokk egy vagy két változót tartalmaz. Megjegyezzük, hogy ha ROP-R két különböző blokkot talál, de nem az 1-es és 2-es (hanem pl. a 4-es és a 7-es) blokkindexszel kijelölve, akkor a program futtatása során ROP-R a kisebbik indexet (a 4-est) automatikusan átjelöli 1-esre, a nagyobbikat (a 7-est) pedig 2-esre.



1.4. ábra. A HierR modul menüablaka mediációs elemzés kijelölése esetén

HierR eredménylistáján elsőként a hierarchikus regresszió eredményét érdemes megtekinteni, melyet az 1.13. és az 1.14. táblázatban foglaltunk össze.

1.13. táblázat. HierR regressziós eredménye az 1.4. ábrán látható kijelölések esetén

Modell	R	R ²	R ² korrr	St.hiba	F	f1	f2	p
1	0,5026	0,2526	0,2518	1,0076	336,94	1	997	< 0,001
2	0,8373	0,7010	0,7004	0,6373	1167,76	2	996	< 0,001

1.14. táblázat. A 2. blokk belépésének hatása R²-re az 1.4. ábrán látható kijelölések esetén

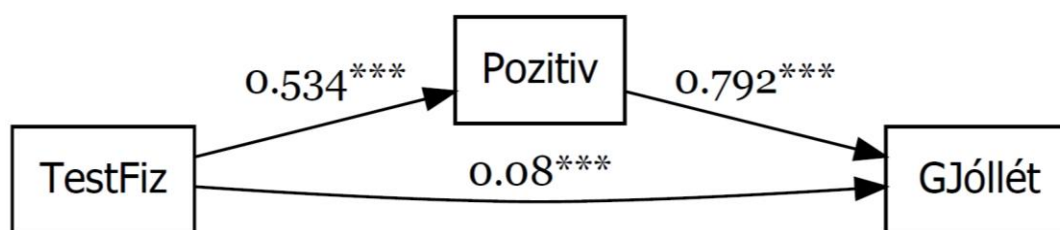
Modell	R ² +	F	f1	f2	p-érték
1					
2	0,4484	1494,01	1	996	< 0,0001

Ebből a két táblázatból kiolvashatjuk, hogy az 1. blokkot (vagyis TestFiz-et) tartalmazó 1. modell a GJóllét varianciájának 25,26%-át magyarázza, ami igen magas érték. A 2. blokk (Pozitiv) belépésével ez a 25,26% megnő további 44,84 százalékponttal (70,1%-ra), vagyis Pozitiv még TestFiz-nél is nagyobb hatást gyakorol GJóllétre. A kérdés csupán az, hogy a TestFiz → GJóllét hatásból mekkora részt mediál Pozitiv (mediációs hatás) és mekkora rész tulajdonítható csupán magának TestFiz-nek (direkt hatás).

Az output további részén (a mediációs elemzések első táblázataként) megtaláljuk a mediációs modell egyes hatásainak nagyságát és szignifikanciáját (lásd 1.15. táblázat), illetve a számítógép „c:_vargha\ropstat\aktualis” mappájában modplot1_TestFiz.Pozitiv.GJóllét.pdf néven elmentett útvonal diagramot (lásd 1.5. ábra).

1.15. táblázat. A mediációs modell regressziós együtthatói (hatások és szignifikanciájuk) egy független változó (TestFiz) és egy mediátor változó (Pozitiv) esetén

Komponens	St.eh.	Reg.eh.	St.hiba	z-érték	p-érték	Hatás
GJóllét ~						
(c)	0,080	0,079	0,021	3,775	< 0,001	TestFiz → GJóllét
Pozitiv ~						
(a)	0,534	0,460	0,022	20,681	< 0,001	TestFiz → Pozitiv
GJóllét ~						
(b)	0,792	0,911	0,023	39,056	< 0,001	Pozitiv → GJóllét



1.5. ábra. A TestFiz → Pozitiv → GJóllét mediációs modell útvonalábrája ROP-R-ben
(Jelölés: ***: $p < 0,001$)

Az 1.15. táblázatból látható, hogy a TestFiz → GJóllét standardizált direkt hatás: $\beta = 0,08$ (St.eh. oszlop), ami ugyan szignifikáns ($< 0,001$), de mivel 0,10-nél kisebb, nem tekinthető szakmailag értelmezhetőnek. Ennél sokkal nagyobb és jelentősebb a TestFiz → Pozitiv és a Pozitiv → GJóllét komponens hatása ($\beta = 0,534$, ill. $\beta = 0,792$), melyek szorzata kiadja a standardizált ab mediációs hatást:

$$ab = 0,534 \cdot 0,792 = 0,423.$$

Ezt az értéket az outputon is megtaláljuk, a „Mediált (indirekt) hatás” című táblázatban, az ab komponens sorában. Ugyanitt láthatjuk, hogy ez a hatás erősen szignifikáns ($p < 0,001$), és a Med% oszlopban megtaláljuk a Mediáció% értékét is (84,1). Mindez azt jelenti, hogy TestFiz GJóllétre gyakorolt hatását 84,1%-ban a Pozitiv változó közvetíti, vagyis a jó fizikai állapot önmagában mit sem ér, ha nem társul egy általános pozitív attitűddel, az élethez való pozitív hozzáállással.

Ha a mediációs modellbe bevonjuk a testi-fizikai állapot mellé az általános lelki állapot (Áltlelki) változót is, komplexebb és pontosabb következtetéseket vonhatunk le. Ehhez a HierR menüablakban a kijelöléseket úgy módosítjuk, hogy betesszük Áltlelki-t is a független változók közé (1-es blokkindexszel). A legfontosabb eredmények az outputon a mediációs modell regressziós együtthatói (lásd 1.16. táblázat) és a mediált (indirekt) hatások táblázatából (lásd 1.17. táblázat) olvashatók ki. Ennek alapján megállapíthatjuk, hogy az Áltlelki változó belépésével a TestFiz változó direkt hatása jelentéktelen szintűre csökkent ($\beta = 0,03$, $p = 0,155$), miközben Áltlelki direkt hatása jelentős és szignifikáns ($\beta = 0,226$, $p < 0,001$) és mindkét független változó esetén domináns Pozitiv mediáló hatása (78,2%, ill. 65,5; vö. 1.17. táblázat).

1.16. táblázat. A mediációs modell regressziós együtthatói (hatások és szignifikanciájuk) két független változó (TestFiz és Áltlelki) és egy mediátor változó (Pozitiv) esetén

Komponens	St.eh.	Reg.eh.	St.hiba	z-érték	p-érték	Hatás
Direkt $X \rightarrow Y$ hatások						
(c1)	0,030	0,029	0,021	1,421	0,155	TestFiz $\rightarrow$ GJóllét
(c2)	0,226	0,191	0,023	8,379	< 0,001	Áltlelki $\rightarrow$ GJóllét
$X \rightarrow M$ hatások						
(a1)	0,165	0,142	0,021	6,685	< 0,001	TestFiz $\rightarrow$ Pozitiv
(a2)	0,659	0,484	0,019	26,048	< 0,001	Áltlelki $\rightarrow$ Pozitiv
$M \rightarrow Y$ hatás						
(b)	0,649	0,747	0,031	24,233	< 0,001	Pozitiv $\rightarrow$ GJóllét

1.17. táblázat. A mediált (indirekt) hatások két független változó (TestFiz és Áltlelki) és egy mediátor változó (Pozitiv) esetén

Komponens	St.eh.	Becslés	St.hiba	z-érték	p-érték	Med%
a1b	0,107	0,106	0,017	6,387	< 0,001	78,2
a2b	0,428	0,362	0,020	17,788	< 0,001	65,5
total	0,790	0,688	0,022	31,829	< 0,001	

2. fejezet

A polinomiális regresszió (PolR) modul

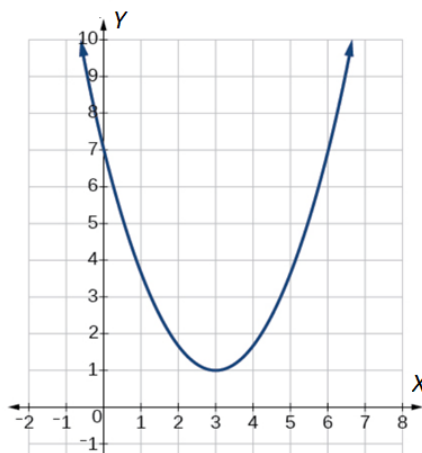
Ebben a modulban egy vagy több függő változót a standardizált független változók legfeljebb első 5 hatványa segítségével lehet külön-külön regresszálni. Az kijelölt számú hatványt egymás után, egyenként léptetjük be a regressziós modellbe. Végző modellként azt a legegyszerűbb modellt fogadjuk el, amely után már nem nő szignifikánsan az R^2 . A polinomiális regresszióelemzés elemzés (PolR) fő célja, hogy nemlineáris kapcsolatokat találjunk a független és a függő változók között. A 2.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 2.2. alfejezet a PolR menüablak használatát mutatja be, a 2.3. alfejezet az eredménylista szerkezetét, a 2.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

2.1. Elméleti alapok

A polinomiális regresszióelemzés (PolReg) alapproblémája ugyanaz, mint az egyszerű lineáris regresszióé. Keressük egy kvantitatív független változó (X) olyan egyszerű függvényét, amellyel jól becsülhetünk egy kvantitatív függő változót (Y). Az egyetlen különbség: X függvényeként nemcsak lineáris függvényt alkalmazhatunk, hanem bonyolultabbat is, X különböző hatványainak tetszőleges súlyozott összegét. Ilyenkor tehát a regressziós függvény a következőképpen néz ki, ha a legmagasabb X -hatvány foka p :

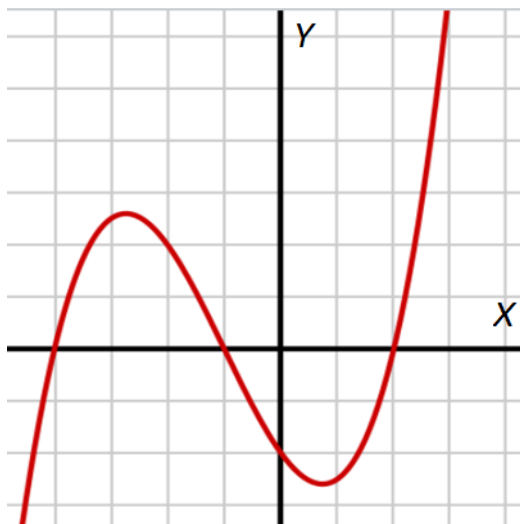
$$\hat{Y} = a + b_1X + b_2X^2 + \dots + b_pX^p. \quad (2.1)$$

Ebben az egyenletben X képviseli a lineáris komponenst, X^2 a négyzettest (kvadrátikus), X^3 a harmadfokú komponenst stb.

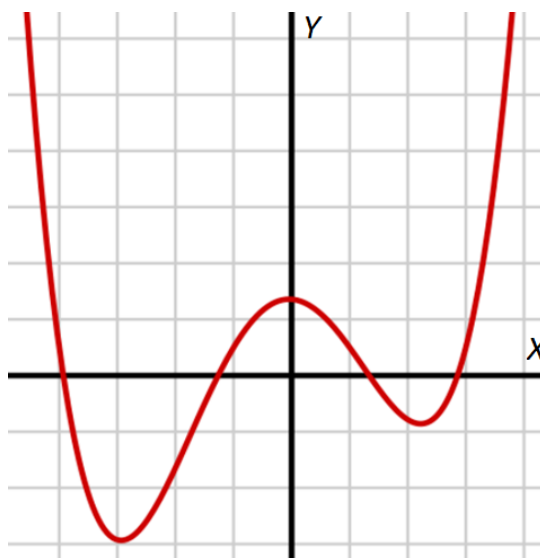


2.1. ábra. Másodfokú (kvadrátikus) polinomiális függés

A polinomiális függvény jó tulajdonsága, hogy minden folytonos függvény jól közelíthető valamilyen polinomiális függvénnyel. Minél nagyobb a legmagasabb hatvány (p), annál komplexebb a kapcsolat. A 2.1., 2.2., 2.3. ábrákon láthatunk egy-egy tipikus másod-, harmad-, illetve negyedfokú polinomiális függést.



2.2. ábra. Harmadfokú polinomiális függés



2.3. ábra. Negyedfokú polinomiális függés

Ha X és Y együttes eloszlása kétdimenziós normális, akkor közöttük csak lineáris típusú kapcsolat lehetséges, hacsak nem függetlenek. Technikailag PolR a többszörös lineáris regresszió egy speciális alkalmazása, ahol a független változók az X változó egymást követő hatványai egy megadott p foksámg:

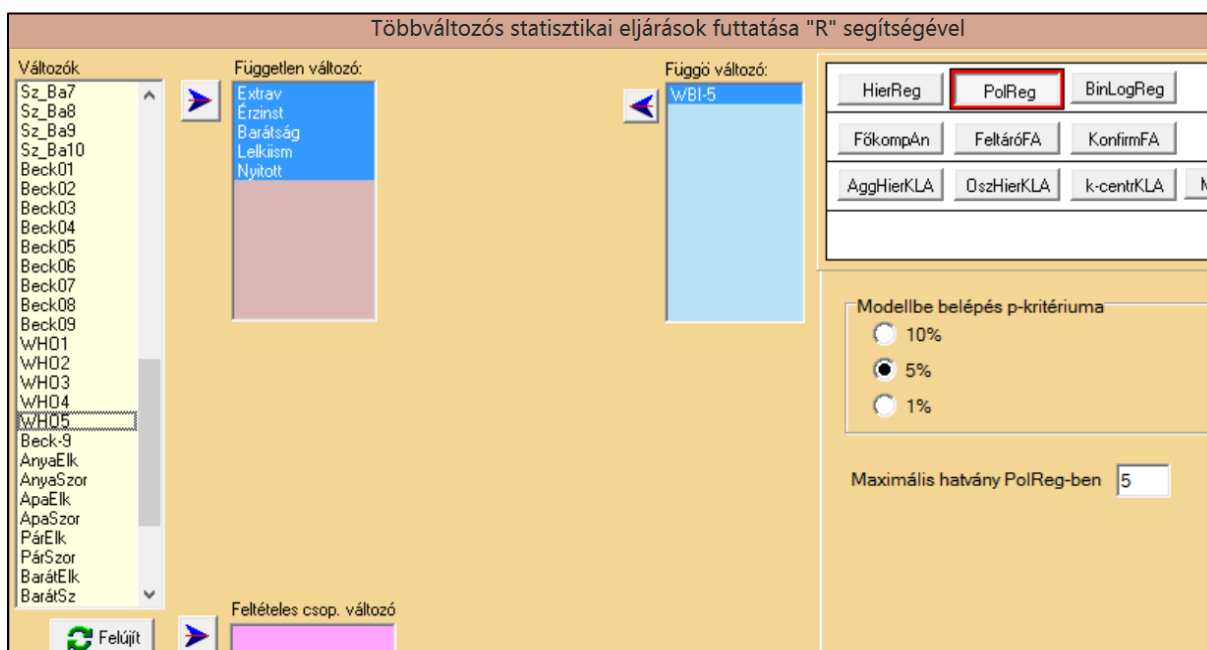
$$X_1 = X, X_2 = X^2, X_3 = X^3 \text{ stb.} \quad (2.2)$$

Emiatt PolR esetében is értelmes az R^2 magyarázott varianciaarány, ami mutatja, hogy az Y függő változót a (2.1) polinomiális regressziós függvény mennyivel becsüli jobban, mint a saját átlaga, vagyis hogy ennek a regressziónak a hibavarianciája mennyivel kisebb, mint $\text{Var}(Y)$.

2.2. A PolR menüablak

Polinomiális regresszioelemzést ROP-R-ben a polinomiális regresszió (röviden PolReg vagy PolR) modul segítségével végezhetünk el. PolR a jamovi szoftver (The jamovi project, 2021; Şahin és Aybek, 2019) által használt *jmv* R-package-re épül.

Ha ROP-R-ben belépünk a PolR modulba, mindössze az a dolgunk, hogy az elemzéshez szükséges változókat a képernyő bal oldalán megjelenő változólistából áttegyük jobbra, a független változókat a „Független változó”, a függő változókat pedig a „Függő változó” ablakba. Minden egyes függő változót minden független változóval külön-külön magyarázunk egy polinomiális regressziós modellben. PolR-ben a beállítható maximális hatványkitevő 2 és 5 közötti érték lehet (alapértelmezés: 5), melyet a menüablakban a „Maximális hatvány PolReg-ben” opció segítségével állíthatunk be (lásd 2.4. ábra).



2.4. ábra. A PolR menüablak

Formálisan X hatványai játsszák a HierR-beli blokkok szerepét (minden blokkban egy-egy ilyen hatvány) és az eredménylistáról azt kell kiolvasni, hogy lépésenként rendre áttérve egy magasabb hatványra, amely egy bonyolultabb nemlineáris kapcsolatot képvisel, szignifikánsan és szakmailag értékelhető mértékben megnő-e az R^2 -tel mért megmagyarázott varianciaarány.

PolR elemzés végrehajtása után a „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzéshez elkészített ideiglenes adatfájlt (tmpdat.txt), a futtatott R-scriptet (PolReg.r), valamint a részletes R eredménylistát (oo.txt).

2.3. Mit tartalmaz a PolR eredménylistája?

A PolR eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
Minden PolR elemzésre:
- A polinomiális regressziós modellek főbb jellemzői
- A végső modell regressziós együtthatói.

2.4. A PolR szemléltetése valódi adatokon

A 336 fős KÖT2016 mintában (lásd B2.2. alpont) a BFI-44 Big Five teszt 5 skálája (Extrav, Érzinst, Barátság, Lelkiism, Nyitott) és a szubjektív jóllétet mérő WBI-5 skála között keresünk nemlineáris kapcsolatokat. Ennek érdekében polinomiális regresszió-elemzést hajtunk végre a ROP-R PolR modulja segítségével, első lépésben független változóként a Big Five teszt 5 skáláját, függő változóként a WBI-5 skálát (lásd 2.4. ábra), második lépésben pedig fordított szereposztásban, független változóként a WBI-5 skálát, függő változóként pedig a Big Five teszt 5 skáláját kiválasztva. A maximális hatványt mindkét esetben a lehetséges legnagyobb 5-ös értékre állítottuk be. Pszichológiai ismereteink szerint egyébként ez előbbi tűnik logikusabb modellnek, amelyben a temperamentum jellegű Big Five vonások hatnak a WBI-5 skálával mért, viselkedés jellegű jóllétre.

Az első lépésben 2, míg a fordított szereposztásban 3 esetben kaptunk szignifikáns nemlineáris hatást, amelyek közül hármat az alábbi alpontokban részletezünk.

2.4.1. Az extraverzió nemlineáris hatása a szubjektív jóllétre

Ez esetben a független változó a KÖT2016 minta Extrav, a függő változó pedig a WBI-5 változója volt. PolR eredménylistáján elsőként a polinomiális regressziós modellek főbb jellemzőinek táblázatát érdemes megtekinteni, amit a 2.1. és a 2.2. táblázatban foglaltunk össze.

2.1. táblázat. A PolR regressziós modelljeinek főbb jellemzői
(független változó: Extrav, függő változó: WBI-5)

X-hatv.	R	R ²	R ² kor	St.hiba	F	f1	f2	p
1	0,3543	0,1255	0,1228	0,8118	47,077	1	328	< 0,001
2	0,3546	0,1258	0,1204	0,8117	23,519	2	327	< 0,001
3	0,3846	0,1479	0,1401	0,8014	18,860	3	326	< 0,001
4	0,3866	0,1495	0,1390	0,8007	14,279	4	325	< 0,001
5	0,3933	0,1547	0,1416	0,7982	11,858	5	324	< 0,001

A 2.1. táblázat szerint már a lineáris hatás is figyelemre méltó [$R^2 = 0,1255$; $F(1, 328) = 47,077$, $p < 0,001$], s ez az R^2 az 5. hatvány belépésével már 0,1547-re nő. A 2.2. táblázat szerint a 3. hatvány belépése 0,0221-es szignifikáns R^2 növekedést eredményez ($p = 0,0039$), de R^2 ez után már egyetlen hatvány belépésével sem nő szignifikánsan, a végső polinomiális modell tehát harmadfokú.

2.2. táblázat. A magasabb X -hatványok belépésének hatása R^2 -re PolR-ben (független változó: Extrav, függő változó: WBI-5)

X-hatv.	R^2	F	f1	f2	p-érték
1					
2	0,0002	0,091	1	327	0,7634
3	0,0221	8,469	1	326	0,0039
4	0,0016	0,605	1	325	0,4373
5	0,0052	1,997	1	324	0,1586

ROP-R a PolR elemzést ténylegesen a standardizált X változóval (Z_x) végzi el, amelyre vonatkozóan a végső polinomiális modell – ROP-R outputjáról leolvasott – regressziós együtthatói a 2.3. táblázatban láthatók. Ez azt mutatja, hogy a végső harmadfokú modellben csak a harmadfokú hatvány (Z_x^3) együtthatója szignifikáns ($p = 0,0039$), jelentősnek mondható standardizált együtthatóval ($\beta = 0,2933$). Érdekes, hogy bár a 2.1. táblázatból kiolvasható egy egyértelmű lineáris hatás, a regressziós táblázatban az elsőfokú hatvány (Z_x) együtthatója közel sincs a szignifikanciához ($p = 0,2602$). Ennek az az oka, hogy a harmadfokú polinomiális függvény tartalmaz egy markáns lineáris trendet is, így a harmadfokú tag belépése már nélkülözhetővé teszi a lineáris hatást képviselő elsőfokú tagot.

2.3. táblázat. A végső polinomiális modell regressziós együtthatóinak táblázata (független változó: Extrav, függő változó: WBI-5)

Prediktor	Együttható	St.hiba	St.eh.	t	p
Konstans	2,86650	0,0585	-	48,959	< 0,0001
Z_x	0,09676	0,0858	0,1113	1,128	0,2602
Z_x^2	0,04075	0,0392	0,0554	1,040	0,2993
Z_x^3	0,09070	0,0312	0,2933	2,910	0,0039

2.4.2. Az érzelmi instabilitás nemlineáris hatása a szubjektív jóllétre

Ez esetben a független változó a KÖT2016 minta Érzinst, a függő változó pedig a WBI-5 változója volt. PolR eredménylistáján elsőként a polinomiális regressziós modellek főbb jellemzőinek táblázatát érdemes megtekinteni, amit a 2.4. és a 2.5. táblázatban foglaltunk össze.

2.4. táblázat. A PolR regressziós modelljeinek főbb jellemzői (független változó: Érzinst, függő változó: WBI-5)

X-hatv.	R	R^2	R^2 korrr	St.hiba	F	f1	f2	p
1	0,4368	0,1908	0,1883	0,7810	77,332	1	328	< 0,001
2	0,4373	0,1912	0,1862	0,7808	38,649	2	327	< 0,001
3	0,4375	0,1914	0,1840	0,7806	25,730	3	326	< 0,001
4	0,4376	0,1915	0,1815	0,7806	19,240	4	325	< 0,001
5	0,4540	0,2061	0,1938	0,7735	16,821	5	324	< 0,001

2.5. táblázat. A magasabb X -hatványok belépésének hatása R^2 -re PolR-ben (független változó: Érzinst, függő változó: WBI-5)

X-hatv.	R^2+	F	f1	f2	p-érték
1					
2	0,0004	0,163	1	327	0,6866
3	0,0003	0,105	1	326	0,7465
4	0	0,004	1	325	0,9476
5	0,0146	5,970	1	324	0,0151

A 2.4. táblázat szerint már a lineáris hatás is figyelemre méltó [$R^2 = 0,1908$; $F(1, 328) = 77,332$, $p < 0,001$]. A 2.5. táblázat szerint csupán az 5. hatvány belépése eredményez szignifikáns R^2 növekedést ($R^2+ = 0,0146$, $p = 0,0151$), a végső polinomiális modell tehát ötödfokú.

A végső polinomiális modell ROP-R outputjáról leolvasott regressziós együtthatói a 2.6. táblázatban láthatók. Ez azt mutatja, hogy az ötödfokú végső modellben minden páratlan fokszámú tag (Zx , Zx^3 , Zx^5) regressziós együtthatója szignifikáns, jelentősnek mondható standardizált együtthatókkal ($\beta_1 = -0,6318$, $\beta_3 = 0,7216$, $\beta_5 = -0,6994$). Ami feltűnő: a harmadfokú tag első és ötödfokúval ellentétes előjelű regressziós együtthatója, hiszen a páratlan fokszámú harmad- és ötödfokú hatványkomponens a lineáris trend különbözőképpen „elbonyolított” képviselője. A jelenség oka, hogy az adott polinomiális modellben a sima lineáris tag hatása a legdominánsabb ($p < 0,001$ szinten), ezt erősíti az azonos előjelű ötödfokú tag jelenléte is, amit a harmadfokú tag ellentétes előjelű regressziós együtthatóval korrigál.

2.6. táblázat. A végső polinomiális modell regressziós együtthatóinak táblázata (független változó: Érzinst, függő változó: WBI-5)

Prediktor	Együttható	St.hiba	St.eh.	t	p
Konstans	2,93570	0,0642	-	45,741	$< 0,0001$
Zx	-0,54936	0,1136	-0,6318	-4,836	$< 0,0001$
Zx^2	-0,11651	0,0983	-0,1832	-1,185	0,2369
Zx^3	0,17980	0,0815	0,7216	2,207	0,0280
Zx^4	0,03289	0,0229	0,2828	1,440	0,1510
Zx^5	-0,03156	0,0129	-0,6994	-2,443	0,0151

2.4.3. A szubjektív jóllét nemlineáris hatása az érzelmi instabilitásra

Ez esetben a független változó a KÖT2016 minta WBI-5, a függő változó pedig az Érzinst változója volt. PolR eredménylistáján elsőként a polinomiális regressziós modellek főbb jellemzőinek táblázatát érdemes megtekinteni, amit a 2.7. és a 2.8. táblázatban foglaltunk össze.

A 2.7. táblázat szerint már a lineáris hatás is figyelemre méltó [$R^2 = 0,1908$; $F(1, 328) = 77,332$, $p < 0,001$], s ez az R^2 az 5. hatvány belépésével már 0,2238-ra nő. A 2.8. táblázat szerint egyedül a 4. hatvány belépése eredményez szignifikáns R^2 növekedést ($R^2+ = 0,0266$, $p = 0,0009$), így a végső polinomiális modell negyedfokú.

2.7. táblázat. A PolR regressziós modelljeinek főbb jellemzői
(független változó: WBI-5, függő változó: Érzinst)

X-hatv.	R	R ²	R ² kor	St.hiba	F	f1	f2	p
1	0,4368	0,1908	0,1883	0,7007	77,332	1	328	< 0,001
2	0,4397	0,1933	0,1884	0,6996	39,184	2	327	< 0,001
3	0,4429	0,1962	0,1888	0,6984	26,522	3	326	< 0,001
4	0,4720	0,2228	0,2132	0,6867	23,294	4	325	< 0,001
5	0,4730	0,2238	0,2118	0,6863	18,681	5	324	< 0,001

2.8. táblázat. A magasabb X-hatványok belépésének hatása R²-re PolR-ben
(független változó: WBI-5, függő változó: Érzinst)

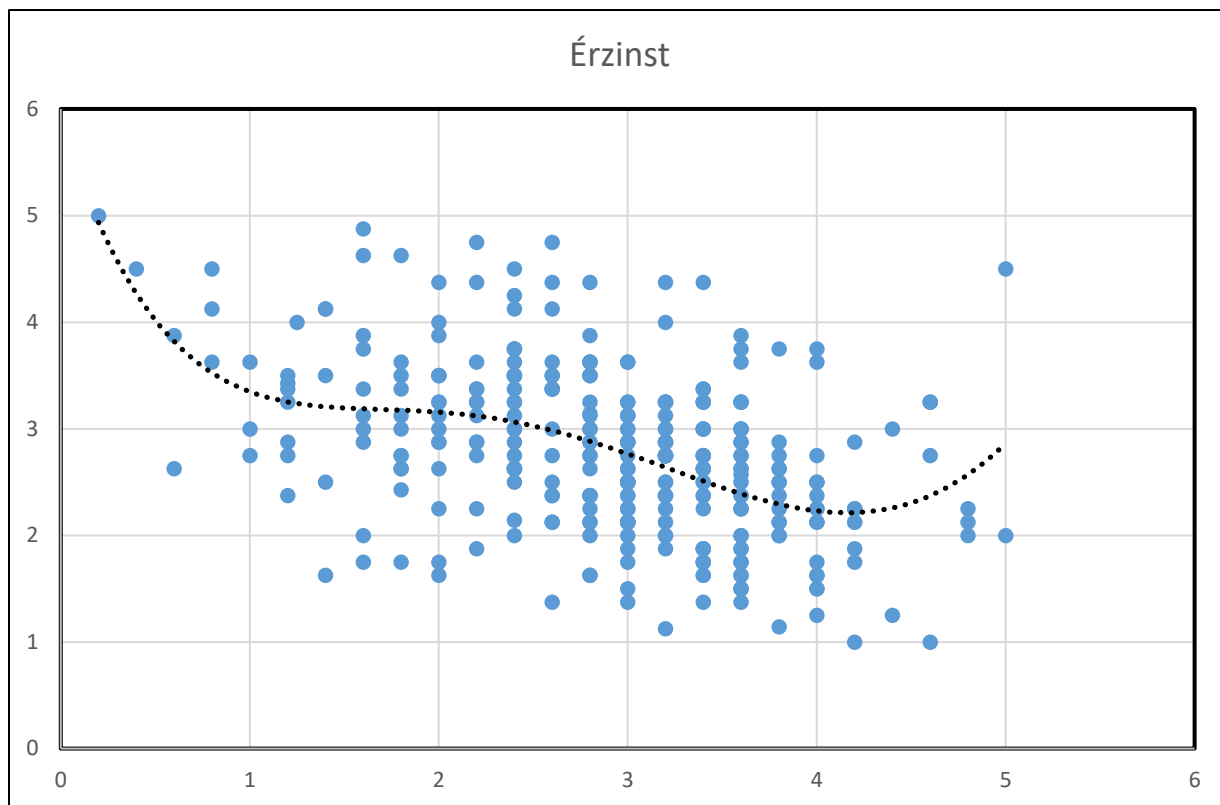
X-hatv.	R ² +	F	f1	f2	p-érték
1					
2	0,0025	1,029	1	327	0,3111
3	0,0029	1,159	1	326	0,2825
4	0,0266	11,136	1	325	0,0009
5	0,0010	0,402	1	324	0,5266

A végső polinomiális modell ROP-R outputjáról leolvasott regressziós együtthatói a 2.9. táblázatban láthatók. Ez azt mutatja, hogy a negyedfokú végső modellben minden tag regressziós együtthatója szignifikáns, értelmezhető szintű (0,20 fölötti abszolút értékű) standardizált együtthatókkal. Legnagyobb hatású a lineáris ($\beta_1 = -0,6051$), illetve a negyedfokú tag ($\beta_1 = 0,4766$).

2.9. táblázat. A végső polinomiális modell regressziós együtthatóinak táblázata
(független változó: WBI-5, függő változó: Érzinst)

Prediktor	Együttható	St.hiba	St.eh.	t	p
Konstans	2,83410	0,0543	-	52,168	< 0,0001
Zx	-0,46922	0,0693	-0,6051	-6,77	< 0,0001
Zx ²	-0,15871	0,0665	-0,2994	-2,387	0,0176
Zx ³	0,05836	0,0212	0,2951	2,759	0,0061
Zx ⁴	0,03970	0,0119	0,4766	3,337	0,0009

A 2.4.2. és a 2.4.3. alpont elemzéseiben az az érdekes, hogy *WBI-5 és Érzinst között mindkét irányban találtunk értelmezhető szintű nemlineáris összefüggést*. Ha ezt ábrázolni szeretnénk, érdemes a két változót ROP-R-ből Excelbe másolni, a két változó adatszoportát kijelölve pontdiagramot készíteni (Beszúrás/Pontdiagram), majd ráhelyezve az egeret egy pontra, a jobb egérgombot megnyomva a Trendvonal felvétele menüpontban a Polinomiális opciót kiválasztani 5-ös fokszámmal. Az eredmény a 2.5. ábrán látható.



2.5. ábra. A 2.9. táblázatban látható ötödfokú hatás (Érzinst függése WBI-5-től) ábrázolása Excelben

3. fejezet

A bináris logisztikus regresszió (BLR) modul

Ebben a modulban minden kétértékű függő változót a kijelölt független változók együttese segítségével regresszálunk, azaz készítünk rá regressziós becslést. A független változókat, amelyek között lehetnek *kategoriálisak* (nominális vagy ordinális skálájúak) is, a blokkindex segítségével különböző blokkokba sorolhatjuk. Végző modellként azt a legegyszerűbb modellt fogadjuk el, amely után már nem nő szignifikánsan a predikciós hatás, nem javul szignifikánsan a modell illeszkedése. BLR tehát valójában egy olyan hierarchikus regresszióelemzés, ahol semmilyen megkötést nem teszünk a független változókra, viszont a függő változó viszont csak kétértékű (dichotóm) lehet. Egy kétértékű változót mindig át tudunk fogalmazni, illetve kódolni úgy, hogy annak értéke 1 legyen, ha egy bizonyos T tulajdonság teljesül (pl. ha a személy kék szemű, ha a személy nő, ha a személy beteg, ha a személy fiatalos, ha a személy rendelkezik egy bizonyos végzettséggel stb.) és értéke 0 legyen, ha a T tulajdonság nem teljesül. Az ilyen 0, 1 értékekkel kódolt változókat nevezzük *bináris* változóknak, erre utal BLR nevében a bináris jelző is. BLR-rel kapcsolatban az a legfőbb szakmai kihívás, hogy találjunk olyan független változókat, amelyek segítségével minél pontosabban, nagy megbízhatósággal be tudjuk jósolni, hogy a T tulajdonság fennáll-e.

BLR egyben *nemlineáris diszkriminancia-analízis* is, ahol a függő változó két értéke által definiált két csoportot próbáljuk minél jobban elkülöníteni, illetve azonosítani a kijelölt független változók segítségével (vö. Tabachnick és Fidell, 2013, 10. fejezet).

A 3.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 3.2. alfejezet a BLR menüablak használatát mutatja be, a 3.3. alfejezet az eredménylista szerkezetét, a 3.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

3.1. Elméleti alapok

A BLR és a HierR mint regressziós elemzések nagyon hasonlóak, a különbségek a következőképpen fogalmazhatók meg. A BLR függő változója csak kétértékű változó lehet. A független változók lehetnek folytonosak (ez az alapértelmezés) vagy kategoriálisak (nominális vagy ordinális skálájúak). Kategoriális beállításnál ROP-R a változó minden talált értékére automatikusan készít egy bináris *dummy* (indikátor) változót¹⁴.

A BLR modelljébe a blokkok egymás után, egyenként lépnek be. Végző modellként azt a legegyszerűbb modellt fogadjuk el, amely után már nem nő szignifikánsan a független változók predikciós hatása, vagyis nem javul szignifikánsan a regressziós modell.

¹⁴ A dummy vagy indikátor változó értéke egy személynél 1, ha a személyt az adott érték jellemzi, egyébként 0.

Statisztikai algoritmusát tekintve a BLR egy olyan többszörös lineáris regresszió, ahol a függő változó a kijelölt függő változó nagyobbik értékéhez tartozó valószínűség¹⁵ *logit* értéke, a független változók pedig az elemzéshez kiválasztott X_j ($j = 1, 2, \dots$) független változók. A *logit* az $\text{odds}(p)$ *esély* mutató segítségével definiálható a következőképpen. Ha p egy tetszőleges szám 0 és 1 között (valószínűség mindig ilyen), akkor

$$\text{odds}(p) = p/(1 - p) \quad (3.1)$$

és

$$\text{logit}(p) = \ln(\text{odds}(p)), \quad (3.2)$$

ahol $\ln$ a természetes (e alapú) logaritmus. A BLR-ben p a függő változó nagyobbik értékének valószínűségét jelöli, és ennek logitjára készítünk regressziós becslést. Ekkor az $Y = \text{logit}(p)$ jelöléssel a regressziós becslés ilyen alakú lesz:

$$\hat{Y} = a + b_1X_1 + b_2X_2 + \dots + b_pX_p = L(\underline{X}), \quad (3.3)$$

azzal a kiegészítéssel, hogy minden kettőnél több értékű kategoriális X_j független változó esetén a változóhoz tartozó b_jX_j komponens maga is egy összeg, mely az X_j változó különböző értékeinek dummy változóiból építkezik (azoknak egy súlyozott összege). BLR esetében tehát a független változó formálisan, vagyis statisztikai szempontból a fejezet bevezetőjében említett T tulajdonság előrejelzési valószínűségének a logitja.

Ha bármely személy esetén ismerjük az $L = L(\underline{X})$ regressziós becslés értékét, akkor ebből p -re is vissza lehet következtetni, mert $\text{odds}(p) = \exp(L)$, amiből p értéke is egyértelműen meghatározható¹⁶. Ez a regressziós becslésből személyenként meghatározható p az ún. *posterior valószínűség*, mely arról tájékoztat, hogy a független változók ismeretében mekkora az esélye, hogy az adott személy a függő változó nagyobbik értékével jellemezhető (vagyis hogy a személy nő, ha a függő változó – pl. férfi = 1, nő = 2 kódolással – a személy neme).

A BLR eredménylistáján minden belépő blokk esetén az alábbi statisztikai mutatók láthatók.

- AIC: a kumulatív BLR modellhez tartozó Akaike-féle információs kritérium (Akaike, 1974). Jobb regressziós modellben (ugyanazon mintában ugyanazon függő változókkal) AIC értéke kisebb lesz.
- R2-McF: a kumulatív BLR modellhez tartozó McFadden-féle pszeudo R^2 (Tjur, 2009).
- R2-CS: a kumulatív BLR modellhez tartozó Cox és Snell-féle pszeudo R^2 (Mbachu, Nduka és Nja, 2012).
- R2-Nag: a kumulatív BLR modellhez tartozó Nagelkerke-féle pszeudo R^2 (Nagelkerke, 1991).
- A kumulatív BLR modell szignifikanciája (χ^2 statisztika, szabadságfok, p -érték), amellyel azt teszteljük, hogy a modellbeli független változóknak van-e 0-nál erősebb hatása a függő változóra.
- A belépő blokk szignifikanciája (χ^2 statisztika, szabadságfok, p -érték), amely azt mutatja meg, hogy a belépő blokk szignifikánsan javítja-e a BLR modell illeszkedését.

¹⁵ A fejezet bevezetőjében említett T tulajdonság felléptének valószínűsége.

¹⁶ $p = \exp(L)/(1 + \exp(L))$

Az R^2 -McF, R^2 -CS és R^2 -Nag pszeudo R^2 statisztikai mutatók a lineáris regresszió R^2 megmagyarázott varianciaarányára hasonlító együttthatók, amelyek 0 és 1 közötti értékeket vehetnek fel és minél közelebb vannak az 1-hez, annál jobb regressziós becslést jeleznek. A három mutató közül sokan a legkevésbé konzervatív Nagelkerke-féle pszeudo R^2 (R^2 -Nag) mutatót preferálják, mert ennek lehetséges értékei kitöltik a teljes $[0-1]$ értéktartományt¹⁷.

Míg a többszörös lineáris regresszió a legkisebb négyzetek módszerével keresi a független változók együttthatóit a regressziós egyenletben (a hibavarianciát minimalizálva, BLR egy maximum likelihood (ML) becslést alkalmaz több iterációs lépésben egy log-likelihood kritérium (LL) segítségével. A kiinduló LL-értéket L_0 -al jelölik. Minél nagyobb a regressziós modellben az adott minta valószínűsége, annál nagyobb LL, ezért az ML-becslés által kapott regressziós megoldás LL-t próbálja maximalizálni. A fenti három pszeudo R^2 mutató mind azt nézi, azt építi be képletébe, hogy a végső LL milyen mértékben nagyobb, mint a kezdeti L_0 ¹⁸. Ennek alapján ugyanúgy egy alaphelyzethez viszonyított relatív javulási mutatók, ahogy R^2 , mely a hibavariancia javulási aránya a sima átlaggal való becsléshez viszonyítva.

Jó modell esetén LL nagy, $LL^* = -2LL$ pedig kicsi lesz. Az LL^* mennyiség mindig pozitív. Az egyes blokkok belépésével LL^* kisebb-nagyobb mértékben nő, a változás egy χ^2 -statisztikával tesztelhető. Ennek segítségével állapítható meg, hogy egy blokk belépésekor szignifikánsan javul-e a függő változó predikciója.

A BLR végső modelljének regressziós táblázatában megtaláljuk az esélyhányados [$odds(p) = \exp(B)$] értékét is minden X prediktorra, amely folytonosként kezelt változók esetén $odds(p)$ arányos változását jelzi, ha X értéke 1 egységgel megnő. Ha ez 1-nél nagyobb, akkor X növekedésével $p/(1-p)$ – és akkor egyben p értéke – is várhatóan nő, vagyis nő a függő változó nagyobbik értékének várható előfordulása a kisebbik értékkel szemben, 1-nél kisebb esélyhányados esetén pedig csökken. Kategóriális X változók esetén az esélyhányados a következőképpen értelmezhető.

- Ha X kétértékű változó, akkor az esélyhányados értéke azt mutatja, hogy milyen lesz $odds(p)$ arányos változása, ha áttérünk a változó kisebb értékéről a nagyobbik értékre (bináris változó esetén a 0-ról az 1-re).
- Ha X kategóriálisnak jelölt és kettőnél több értékű változó, akkor ROP-R referencia értéknek tekinti a változó legkisebb értékét. Ez esetben az esélyhányados értéke minden más X -érték esetén azt mutatja, hogy milyen lesz $odds(p)$ arányos változása, ha áttérünk a változó legkisebb értékéről erre a nagyobb értékre.

Az eredménylista tartalmazza a végső modellre vonatkozóan a független változók multikollinearitás-diagnosztikáját (TOL- és VIF-értékek; vö. 1.4.1. alpont), valamint a végső modellre vonatkozó klasszifikációs táblázatot is. Ez utóbbiból kiolvasható, hogy a végső BLR modell alapján történő besorolás rendre hány személyt sorol be helyesen a saját csoportba, illetve helytelenül a másik csoportba a függő változó két értéke esetén. Végül a *predikciós mutatók* minden kumulatív modellre megadják az átlagos helyes besorolási arányt (Pontos%), illetve a függő változó kisebb és nagyobb értékére a helyes besorolás arányát (*specificitás* és *szenzitivitás*¹⁹ százalék). Természetesen egy regressziós becslés annál jobb (vagyis a független

¹⁷ A McFadden- és a Cox és Snell-féle pszeudo R^2 mutató értéke mindig kisebb, mint a Nagelkerke-féle mutatóé.

¹⁸ vö. <https://en.wikipedia.org/wiki/Pseudo-R-squared>

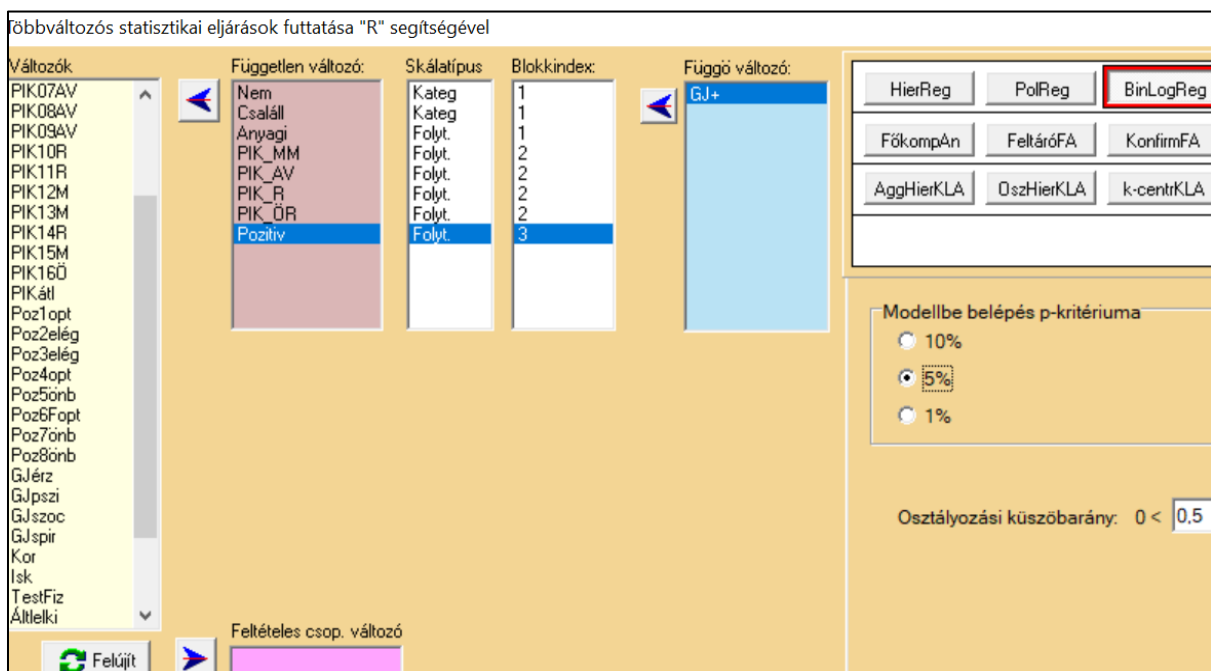
¹⁹ Ha a bináris függő változó nagyobbik értéke bármilyen tulajdonság (pl. valamilyen betegség) meglétét, kisebbik értéke pedig annak hiányát jelzi, akkor a szenzitivitás (érzékenység) arra utal, hogy a teszt (az X változó alapján tett becslés) milyen megbízhatósággal képes például a betegségben szenvedő egyén betegségét

változók annál informatívabbak a függő változóra nézve), minél nagyobb a Pontos%, az átlagos helyes besorolási arány, vagyis minél pontosabban lehet előrejelezni a függő változó értékét a független változók értékének ismeretében. Megjegyzendő, hogy mindezen besorolások erősen függenek az osztályozási/klasszifikációs küszöbaránytól is, amely a BLR menüablakában adható meg (0 és 1 között, az alapértelmezés 0,50; vö. 3.1. ábra).

3.2. A BLR menüablak

BLR-t ROP-R-ben a bináris logisztikus regresszió (röviden BinLogReg vagy BLR) modul segítségével végezhetünk el. A BLR a jamovi szoftver (The jamovi project, 2021; Şahin és Aybek, 2019) által is használt *jmv* R-package-re épül.

Mivel a BLR ugyanúgy képes hierarchikus regresszióra, mint HierR, ha ROP-R-ben belépünk a BLR modulba, ugyanúgy jelölhetjük ki a független és a függő változókat, valamint a független változók különböző blokkjait, mint a HierR modulban. Minden egyes függő változót a kijelölt független változók együttesével magyarázunk egy BLR modellben. HierR-hez viszonyítva BLR-ben plusz lehetőség, hogy a független változók között szerepelhetnek kategoriálisak is. Ezt a kijelölést a „Skálatípus” ablakban állíthatjuk be. Az alapértelmezés szerint minden változó folytonos (Folyt) beállítású, ezt lehet változónként kategoriális (Kateg) típusra átállítani (vö. 3.1. ábra).



3.1. ábra. A BLR menüablak

A BLR modul elemzéseinek végrehajtása után a „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzésekhez elkészített ideiglenes adatfájlt (tmpdat.txt), a futtatott R-scriptet (BinLogReg.r), valamint a részletes R eredménylistát (oo.txt).

kimutatni. Nagy érzékenységgű teszt esetében kevés hamis negatív eredmény születik. A teszt specificitása pedig azt méri, hogy milyen százalékos arányban nem jelez a teszt betegséget, ha a személy nem beteg.

3.3. Mit tartalmaz a BLR eredménylistája?

A BLR eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
 - A független változók blokkindexek szerinti csoportosítása (változók besorolása a definiált blokkokba)
- Minden függő változóval elvégzett elemzésre:
- Az egyes blokkok bővülő regressziós modelljeinek főbb jellemzői
 - A végső modell regressziós együtthatói
 - Multikollinearitás diagnosztika
 - A végső modell klasszifikációs táblázata
 - A predikciós mutatók összefoglaló táblázata

3.4. A BLR szemléltetése valódi adatokon

Az 1003 fős Btérkép2022 minta (vö. B.2.3. alpont) egyik fontos pozitív pszichológiai változója a globális jóllét (GJóllét), mely a jóllét bio-pszicho-szocio-spirituo modelljét operacionalizálja, hangsúlyozva, hogy a teljes jóllét feltétele, hogy az ember humán természetének minden aspektusában jól működjön, miközben jól érzi magát a bőrében (Oláh, 2021).

Izgalmas kérdés, hogy mitől kerülhet valaki a globális jóllét legmagasabb szintjére, mondjuk a GJóllét eloszlásának felső 10%-ába? Ennek tisztázására először is létre kell hozni egy olyan bináris változót, melynek 1 az értéke, ha az egyén a GJóllét eloszlásának felső 10%-ába esik és 0, ha nem. Ehhez először is a GJóllét eloszlásában meg kell keresni azt a küszöbértéket, mely alatt az 1003 fős minta 90%-a, fölötte pedig a 10%-a található. Ezt nem nehéz megcsinálni (pl. a ROPstat *Gyakorisági eloszlás* modulja segítségével), de van erre egy egyszerű eljárás ROP-R-ben is. A főmenü Esetek / Esetek sorbarendezeése menüpontban egyszerűen át kell tenni a GJóllét változót a „Sorbarendeze” ablakba és a „Futtat” gombra kattintani, melynek hatására a ROP-R a GJóllét változó értéke szerint sorbarendezi a teljes mintát. Mivel 1003-nak a 90%-a 902,7, már csak annyi a dolgunk, hogy a sorbarendezezt fájlban a 903. személynek megnézzük a GJóllét értékét. Ez a jelen esetben 5,375 (901-től 904-ig mindenkinek ekkora). Ezután a binarizálást egyszerűen végrehajthatjuk a ROP-R Transzformációk / Átkódolás menüpontjában, egy új GJ+ nevű változót létrehozva, a [0–5,37] övezettel definiálva GJ+ 0-s értékét és az [5,375–6] övezettel GJ+ 1-es értékét. Az így módon létrehozott GJ+ lesz a BLR bináris függő változója. Ezzel kapcsolatban két BLR elemzést hajtunk végre. Először csak a szociodemográfiai változókat vonva be prediktorként (3.4.1. alpont), majd a PIK16 Rövidített Pszichológiai Immunrendszer Kérdőív skáláit és a Caprara-féle Pozitivitás skálát is (3.4.2. alpont).

3.4.1. A globális jóllét magas szintjének bejósolása szociodemográfiai változók segítségével

A független változóknak először egyetlen blokkját hozzuk létre, a Nem, Csaláll (családi állapot) és Anyagi (szubjektív anyagi helyzet) szociodemográfiai változók felhasználásával, az első két (Nem és Csaláll) változót kategoriálisnak jelölve a menüablakban.

A blokk globális hatása szignifikáns predikciót jelez GJ+-ra, 8,5%-os Nagelkerke-féle R^2 értékkel (lásd 3.1. táblázat), ami szakmailag értelmezhető, de alacsony szintű kapcsolat. Részletesebb eredményeket a regressziós együtthatók táblázatából olvashatunk ki (lásd 3.2. táblázat).

3.1. táblázat. A Nem, Csaláll, Anyagi változók hatása BLR-ben a GJ+ függő változóra

Modell	AIC	R2-McF	R2-CS	R2-Nag	Khi2	f	p
1	635,6	0,0636	0,0413	0,0852	42,26	6	< 0,001

3.2. táblázat. A Nem, Csaláll, Anyagi változók hatása BLR-ben a GJ+ függő változóra

Prediktor	Együtth.	St.hiba	Z	p	Esélyhányados
Konstans	-5,723	0,6588	-8,687	< 0,0001	0,003
Nem:					
2 – 1	0,503	0,2803	1,795	0,0727	1,654
Csaláll:					
2 – 1	0,212	0,3425	0,618	0,5362	1,236
3 – 1	0,670	0,3054	2,195	0,0282	1,955
4 – 1	0,203	1,1260	0,180	0,8571	1,225
5 – 1	0,733	0,4669	1,569	0,1166	2,081
Anyagi	0,843	0,1654	5,095	< 0,0001	2,323

A 3.2. táblázat azt mutatja, hogy az előrejelzésben legnagyobb predikciós hatása a szubjektív anyagi helyzetnek van ($p < 0,0001$), 2,323 esélyhányados értékkel. Ez azt jelenti, hogy ha az 5-fokú Anyagi változó értéke 1 egységgel megnő, akkor várhatóan több mint kétszeresére nő a GJ+ odds-a, a $p/(1-p)$ esély, vagyis nagyobb szubjektív anyagi helyzetben érezhetően nagyobb az egyén esélye, hogy bekerüljön GJóllét tekintetében a top10%-ba.

Szignifikáns még 5%-os szinten a családi állapot 3 értékének a hatása is ($p = 0,0282$) az 1 referencia értékkel összevetve. Ez a 3.3. táblázatból kiolvasható címke nevek fényében úgy fogalmazható meg, hogy a házasságban élő felnőttek (3-mas érték) nagyobb eséllyel élnek a globális jóllét legmagasabb szintjét ($GJ+ = 1$), mint azok, akik egyedül élnek (1-es érték), 1,955-ös esélyhányados értékkel. Emellett még egy tendenciát ($p = 0,0727$) is látunk arra, hogy a nők könnyebben kerülnek be a top10%-ba, mint a férfiak, 1,654-es esélyhányados értékkel.

3.3. táblázat. A Nem és a Csaláll kategoriális változó értékcímkei

Csaláll		Nem	
Érték	Címke	Érték	Címke
1	egyedül	1	Férfi
2	kapcsolatban	2	Nő
3	házas		
4	özvegy		
5	elvált		

Ezek után érdekes lehet, hogy a blokk változói segítségével milyen bizonyossággal lehet előrejelezni GJ+ értékét. Ezzel kapcsolatos adatokat az eredménylistán a végső modell klasszifikációs táblázata (lásd 3.4. táblázat) és a predikciós mutatók összefoglaló táblázata (lásd 3.5. táblázat) tartalmaz.

3.4. táblázat. A végső modell klasszifikációs táblázata a GJ+ változó 0 és 1 értékére vonatkozóan (osztályozási küszöbarány = 0,50)

Csoport (GJ+ értéke)	0	1	Korrekt%
0	900	0	100
1	103	0	0

3.5. táblázat. A predikciós mutatók összefoglaló táblázata (osztályozási küszöbarány = 0,50)

Modell	Pontos%	Specif%	Szenzitivitás%
1	89,73	100	0

A 3.4. és a 3.5. táblázat alapján azt mondhatjuk, hogy 0,50-es osztályozási küszöbarány esetén a független változók ismeretében egyetlen top10% személyt sem sikerült helyesen besorolni, mert az átlagos helyes besorolási arány (lásd a 3.5. táblázatban a Pontos%-ot) maximalizálása érdekében az algoritmus senkit se sorolt be a GJ+ = 1 csoportba. Ennek az az oka, hogy a top10%-ba tartozó személyek a teljes mintának mindössze a 10%-át képezik. Ha ehhez közelebbi, például 20%-os küszöbarányt választunk, máris megnő a top10%-ba besorolt személyek Korrekt% értéke 0-ról 22,33%-ra (lásd 3.6. táblázat), de ezzel párhuzamosan a Pontos% lecsökken a korábbi 89,73%-ról 83,85%-ra (lásd 3.7. táblázat).

3.6. táblázat. A végső modell klasszifikációs táblázata a GJ+ változó 0 és 1 értékére vonatkozóan (osztályozási küszöbarány = 0,20)

Csoport (GJ+ értéke)	0	1	Korrekt%
0	818	82	90,89
1	80	23	22,33

3.7. táblázat. A predikciós mutatók összefoglaló táblázata (osztályozási küszöbarány = 0,20)

Modell	Pontos%	Specif%	Szenzitivitás%
1	83,85	90,89	22,33

3.8. táblázat. A végső modell klasszifikációs táblázata a GJ+ változó 0 és 1 értékére vonatkozóan (osztályozási küszöbarány = 0,10)

Csoport (GJ+ értéke)	0	1	Korrekt%
0	517	383	57,44
1	28	75	72,82

Ha a küszöbarányt még lejjebb visszük, a top10% arányával megegyező 10%-ra, még inkább megnő a top10%-ba besorolt személyek Korrekt% értéke (immár 72,82%-ra, lásd 3.8. táblázat), de ezzel párhuzamosan a Pontos% lecsökken 59,02%-ra (lásd 3.9. táblázat).

3.9. táblázat. A predikciós mutatók összefoglaló táblázata
(osztályozási küszöbarány = 0,10)

Modell	Pontos%	Specif%	Szenzitivitás%
1	59,02	57,44	72,82

Végső tanulságként azt állapíthatjuk meg, hogy az osztályozási küszöbarány beállítása előtt el kell döntenünk, hogy a besorolási döntésben mi a legfontosabb számunkra. Az átlagos helyes besorolási arány (Pontos%)? Mert ekkor az 50%-os küszöbarány választása tűnik a legjobbnak. A bináris függő változó 0 vagy 1 értékének a bejósolása? Ekkor viszont úgy kell a küszöbarányt megválasztani, hogy ezen értékek helyes besorolási aránya (a kisebbik érték esetén Specif%, a nagyobb érték esetén Szenzitivitás%) legyen minél nagyobb. Ha a predikciós erő gyenge (mert a Nagelkerke-féle R^2 alacsony), akkor nem lehet egyidejűleg mindkét Korrekt% értéket magasan tartani. Ha a függő változó nagyobbik értékének a pontossága a fontosabb, akkor a küszöbarányt közelíteni kell ezen érték mintabeli arányához. Ha pedig a kisebbik érték bejósolása fontosabb (Specif%), akkor a küszöbarányt ezen érték mintabeli arányához kell közelíteni.

3.4.2. A globális jóllét magas szintjének bejósolása szociodemográfiai változók, a PIK16 skálái és a Caprara-féle Pozitivitás skála segítségével

Annak érdekében, hogy a PIK16-skálák és a Pozitiv skála hatását is megvizsgálhassuk a GJ+ bináris változóra, ezeket is bevettük 2 újabb blokkban a 3.4.1. alpontban leírt elemzésbe (vö. 3.1. ábra), s emellett két apró változtatást is tettünk: 1) a ROP-R „Változók deklarációi” ablakában a GJ+ változó értékeihez címkét, azaz nevet rendeltünk (0 = normál, 1 = top10%); 2) az osztályozási küszöböt a 3.4.1. alpontbeli elemzés tanulságaként átállítottuk 0,5-ről 0,1-re, a top10% érték mintabeli arányához igazodva. Az elemzés eredményeit az alábbiak szerint foglaljuk össze.

Az egymás után beléptetett blokkok hatását a 3.10. táblázat foglalja össze. Mivel az első blokk ugyanaz, mint a 3.4.1. alpontban, hatása is pontosan ugyanaz lesz, mint amit a 3.1. táblázatból már megismertünk, ezért ezt nem kommentáljuk. Érdekes viszont a 3.10. táblázatban a 2. sor (2. modell), amely az 1. és a 2. blokkot (a szociodemográfiai mutatókat és a PIK16 skálákat) egyaránt tartalmazza.

3.10. táblázat. A szociodemográfiai változók (1. blokk), a PIK16-skálák (2. blokk) és a Pozitiv skála (3. blokk) hatása BLR-ben a GJ+ függő változóra

Modell	AIC	R2-McF	R2-CS	R2-Nag	Khi2	f	Khi2+	f
1	635,6	0,0636	0,0413	0,0852	42,26***	6		
2	380,8	0,4596	0,2623	0,5417	305,10***	10	262,85****	4
3	355,6	0,5006	0,2820	0,5825	332,32***	11	27,22****	1

Jelölés: ***: $p < 0,001$; ****: $p < 0,0001$

Ez nemcsak hogy erősen szignifikáns ($p < 0,001$), hanem még a hatás mértékében is kiemelkedő, amit az AIC zuhanásszerű csökkenése, illetve a három pszeudo R^2 ugrásszerű megemelkedése is jelez. A 3. blokk (Pozitív) belépése további szignifikáns és szakmailag is értelmezhető mértékű modelljavulást eredményez (lásd 3. sor). Részletesebb eredményeket a regressziós együtthatók táblázatából olvashatunk ki (lásd 3.11. táblázat).

3.11. táblázat. A végső BLR modell regressziós együtthatói (függő változó: GJ+)

Prediktor	Együtth.	St.hiba	Z	p	Esélyhányados
Konstans	-23,797	2,3100	-10,301	< 0,0001	0
Nem:					
2 – 1	0,319	0,3876	0,823	0,4105	1,376
Csaláll:					
2 – 1	-0,424	0,4649	-0,912	0,3617	0,654
3 – 1	-0,056	0,4199	-0,132	0,8948	0,946
4 – 1	-0,917	1,4367	-0,638	0,5232	0,400
5 – 1	-0,170	0,6487	-0,262	0,7935	0,844
Anyagi	0,101	0,2435	0,415	0,6782	1,106
Mmonit	2,980	0,5655	5,270	< 0,0001	19,694
AVhat	0,566	0,4115	1,375	0,1692	1,761
Rezil	-0,055	0,2792	-0,198	0,8432	0,946
Önreg	0,271	0,2309	1,176	0,2398	1,312
Pozitív	2,086	0,4376	4,765	< 0,0001	8,049

A 3.11. táblázat alapján az alábbi következtetéseket vonhatjuk le.

- Bár az 1. blokk önmagában szignifikáns hatású, a 2. és a 3. blokk belépése után már egyik változójuknak sem szignifikáns a hatása a regressziós modellben, még a kezdetben legerősebb hatású anyagi helyzetnek sem (lásd a 3.2. táblázat utolsó sorát).
- A 2. blokk változói közül a legnagyobb predikciós hatása a Mmonit skálának van ($p < 0,0001$), 19,694 esélyhányados értékkel. Ha Mmonit értéke 1 egységgel megnő, akkor tehát várhatóan csaknem húszszorosára nő a GJ+ odds-a, vagyis akinek a viselkedésére jellemző a fizikai és a szociális környezet megismerésére, megértésére és kontrollálására vonatkozó aktív tendencia (vö. Oláh, 2005), az nagyobb eséllyel kerül be GJóllét tekintetében a top10%-ba. Mmonit jelenlétében más PIK16-beli skálának nincs szignifikáns hatása a modellben.
- A 3. blokk változói közül a legnagyobb predikciós hatása a Pozitív skálának van ($p < 0,0001$), 8,049 esélyhányados értékkel. Ha Pozitív értéke 1 egységgel megnő, akkor kb. nyolcszorosára nő a GJ+ odds-a, vagyis akinél erősen pozitív az élethez való hozzáállás, az nagyobb eséllyel kerül be GJóllét tekintetében a top10%-ba.

Ezek után ismét izgalmas kérdés, hogy ebben a végső modellben a független változók segítségével milyen bizonyossággal lehet előrejelezni GJ+ értékét. Az ezzel kapcsolatos eredményeket a 3.12. és a 3.13. táblázatban foglaltuk össze. Ezekből azt látjuk, hogy 10%-os osztályozási küszöbarány esetén már egészen pontos előrejelzést tehetünk a független változók ismeretében GJ+ értékére. A top10% csoport helyes azonosítási aránya már 95%

feletti (Korrekt% = 96,12), miközben a normál csoportot is 85% körüli pontossággal lehet bejósolni. A helyes azonosítás aránya tekintetében a 2. blokk belépése a legfontosabb, de még a 3. blokkra is szükség volt ahhoz, hogy a top10% csoport Korrekt% értéke 95% fölé menjen (lásd a Szenzitivitás% oszlopot a 3.13. táblázatban).

3.12. táblázat. A végső modell klasszifikációs táblázata a GJ+ változó normál és top10% értékére vonatkozóan (osztályozási küszöbarány = 0,10)

Csoport	normál	top10%	Korrekt%
normál	761	139	84,56
top10%	4	99	96,12

3.13. táblázat. A predikciós mutatók összefoglaló táblázata (osztályozási küszöbarány = 0,10)

Modell	Pontos%	Specif%	Szenzitivitás%
1	59,02	57,44	72,82
2	84,45	83,56	92,23
3	85,74	84,56	96,12

Elgondolkodtató, hogy miért hagyunk benn a modellben olyan sok független változót, miközben közülük csak kettőnek van szignifikáns hatása? Nem lenne célszerű csak ezzel a két változóval (Mmonit és Pozitiv) végezni az elemzést? Végre is hajtottunk egy ilyen BLR-t, melynek illeszkedése alig volt gyengébb, mint a fentebb taglalt modellé (Nagelkerke-féle R^2 csökkenése 0,0084, Szenzitivitás% csökkenése 1,95, miközben a Specifitás% értéke még nőtt is kb. 1 százalékponttal). Jogosnak tűnik tehát ezen egyszerűbb BLR modell elfogadása.

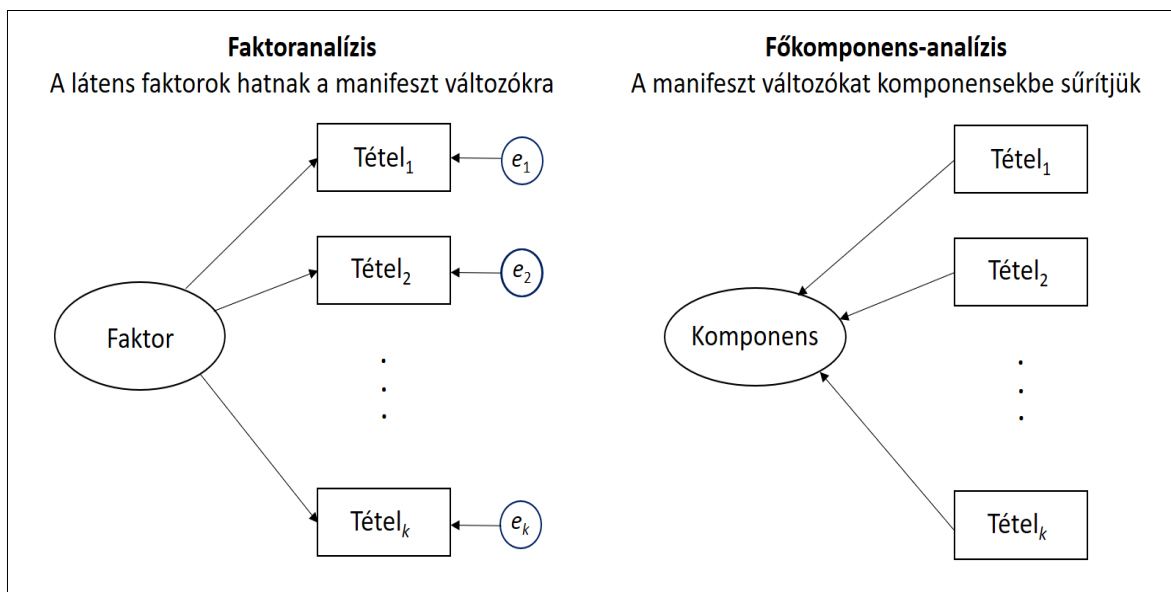
II. RÉSZ

Dimenziócsökkentő modulok

A főkomponens-analízis (FKA) célja a vizsgált kvantitatív változókat jól helyettesítő olyan kevés számú új változó létrehozása az eredeti (megfigyelt, manifeszt) változók súlyozott összegeiként, amelyek egymástól függetlenek (páronként korrelálatlanok) és a lehető legjobban magyarázzák az eredeti változókat. FKA-t és az FKA végrehajtására alkalmas ROP-R modult a 4. fejezetben ismertetjük.

A faktoranalízis (FA) alapgondolata ezzel szemben az, hogy létezik egy közvetlenül nem megfigyelhető, *látens faktorstruktúra*, mely valamilyen módon kifejezi, képviseli vizsgált változóink lényegét – és lényegében ezen látens faktor(ok) hatásaként alakulnak megfigyelt változóink értékei. FA egyik fő célja a *látens faktorstruktúra* feltárása (*feltáró, exploratív FA*, röviden EFA; lásd 5. fejezet), másik fő célja pedig ennek tesztelése, ha van egy konkrét feltételezésünk a változók korrelációs struktúrájáról, vagyis a faktormodellről (*modelltesztelő, konfirmatív FA*, röviden CFA; lásd 6. fejezet).

Az FA és az FKA modellje közti alapvető különbséget szemlélteti a 4.0. ábra, melyen a megfigyelt, manifeszt változókat egy tesztskála k számú tétele képviseli. FA modelljében ezen tételeket a tételek közös faktora (vagy faktorai) határozzák meg, kiegészítve ezek egyediségével ($e_1, e_2, \dots, e_k$). FKA modelljében a tételekből kiszámítható főkomponensek (röviden komponensek) pusztán arra valók, hogy a sok tételt kevés számú új változóval helyettesítsük, azaz kevés számú komponensbe sűrítsük. FKA-ban tehát a fő hatásirány: manifeszt változók $\rightarrow$ főkomponensek, FA-ban pedig: faktorok $\rightarrow$ manifeszt változók.



4.0. ábra. A faktoranalízis és a főkomponens-analízis modellje

4. fejezet

A főkomponens-analízis (FKA) modul

Ezzel a modullal kvantitatív változók főkomponens-analízise (FKA) végezhető el, amelyet ortogonális (Varimax) vagy ferde (Promax vagy Oblimin) forgatással illeszthetünk jobban az input változóegyütteshez. FKA-ban lehetőség van a szélsőséges esetek azonosítására (outlier detekció), valamint a Cronbach-alfa és a McDonald-féle omega kiszámítására is (vö. T. Kárász, Nagybányai Nagy, Széll és Takács, 2022), ha az input változók mind egyazon skála egyirányú tételeinek tekinthetők.

A 4.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 4.2. alfejezet az FKA menüablak használatát mutatja be, a 4.3. alfejezet az eredménylista szerkezetét, a 4.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

4.1. Elméleti alapok

Ahogy fentebb már említettük, FKA lényege, hogy az eredeti változók súlyozott összegeiként kevés számú korrelálatlan összetevőt, ún. *főkomponenst* hozunk létre, amelyek az eredeti változók teljes varianciájának a lehető legnagyobb részét megmagyarázzák (Tabachnick és Fidell, 2013, 13. fejezet; Vargha, 2019, 4. fejezet). A főkomponenseket egymás után, egyenként hozzuk létre az alábbi algoritmus szerint.

1. Az első főkomponens a változók olyan súlyozott összege, amely a változókval a lehető legszorosabb lineáris kapcsolatban van, azaz amelynél a lehető legnagyobb az egyes változókkal számított korrelációk négyzetösszege.
2. Minden további főkomponensre is ugyanez igaz, csak azon feltétel mellett, hogy az újabb főkomponensek korrelálatlanok az összes előtük kiszűrt főkomponenssel.
3. Adott főkomponens esetén az egyes változókkal számított korrelációk négyzetösszegét a főkomponens *sajátértékének* nevezzük, és ez gyakorlatilag azt mutatja, hogy a főkomponens hány változónyi információt magyaráz meg, illetve fed le.
4. Az egyértelműség kedvéért a főkomponenseket standardizált (0 átlagú és 1 szórású) alakban definiáljuk.

Logikájában FKA hasonlít a többszörös lineáris regresszió módszeréhez. A különbség köztük az, hogy FKA-ban nem egyetlen függő változót, hanem az összes vizsgált változót szeretnénk megmagyarázni, és nem független változók egy csoportja, hanem saját magukból gyártott új változók segítségével, amelyek tömörebben tartalmazzák az általuk lefedett információt. FKA fő célja tehát a változók számának csökkentése a takarékoság jegyében, hogy további elemzésekben már elég legyen ezzel a kevés számú új változóval dolgozni. Elvárjuk, hogy a megtartott főkomponensek a változók teljes varianciájának minimum 50-

60%-át magyarázzák, de jónak csak a 70% fölötti megmagyarázott varianciájú megoldásokat tekintjük (lásd pl. Shaharudin, Ahmad, Zainuddin és Mohamed, 2018).

Az FKA elemzés alapja a vizsgált változók közti lineáris kapcsolatokat összefoglaló korrelációs mátrix, alkalmazási feltételei pedig az alábbiak.

- **Mintanagyság:** A statisztikai elemzésbe bevont személyek száma legyen jó nagy és mindenképpen jóval nagyobb, mint a vizsgált változók száma. Általánosan elfogadott, hogy 100-nál kisebb elemszám szóba se jöhet. Megbízható elemzéshez – az elméleti korrelációk pontos becslése érdekében – legalább 300-400 személyre van szükség és jó, ha minden változóra legalább 10 személy jut. Minimális kikötés, hogy a személyek száma legyen legalább 3-5-szöröse a változók számának.
- **Linearitás:** A változók közti kapcsolat legyen lineáris jellegű, vagyis a Pearson-féle korrelációs együttható legyen alkalmas a köztük lévő összefüggések szorosságának a mérésére. Ezért fontos, hogy a változók kvantitatívak legyenek. Megjegyezzük, hogy a monoton típusú kapcsolatok gyakran elfogadhatóan közelíthetők lineárisal, ezért értelmezhető eredményt kaphatunk ordinális típusú változókkal is. Ha a vizsgált változók együttes eloszlása többdimenziós normális²⁰, az már maga után vonja, hogy közöttük csak lineáris típusú kapcsolatok lehetségesek (vö. Vargha, 2019, 17. o.).
- **Interkorrelációk:** A változók legyenek kapcsolatban egymással! Független vagy kismértékben összefüggő változók esetén nincs esélye a változóredukciónak. Gyakorlatilag ez azt jelenti, hogy a változók korrelációs mátrixában legyenek jócskán a 0,3-0,5-ös szintet elérő vagy meghaladó értékek is.

Bár FKA esetében nem modellt, hanem csupán a változóinkat megfelelő mértékben helyettesíteni képes komponenseket keresünk, ezek értelmes és jelentéssel felruházható volta megkönnyíti a velük végzett későbbi elemzéseket. Ezt az elsődlegesen feltárt struktúra, a megtartott főkomponensek forgatásával (rotációjával) lehet elérni. A *forgatás* egy matematikai művelet valamilyen többdimenziós térben, melynek technikai részletei pszichológiai szempontból nem lényegesek. A lényeg az, hogy a forgatás eredményeképpen létrehozott komponensek²¹ (FKA) jobban illeszkednek az eredeti változókra, mint forgatás előtti megfelelőik, így azok az értelmezést megkönnyítő szakmai jelentéssel is felruházhatók. Fontos, hogy a forgatott főkomponensek (az új komponensek) a változók varianciájából együttesen pontosan akkora részt magyaráznak, mint a nem forgatott főkomponensek, vagyis a forgatással az összes megmagyarázott rész aránya nem változik. A forgatással tehát nem többet, hanem jobban tudunk megmagyarázni a változókból. A forgatott komponensek jobban szeparálódnak egymástól, mint az eredetiek, azaz a változók többsége (esetleg az összes) csak egy-egy komponensre ad erős töltést. Optimális esetben tehát a változók a forgatott komponenseknek megfelelő, nem átfedő csoportokat alkotnak.

Az FKA-ban használt legfontosabb fogalmak az alábbiak.

- **Főkomponenssúly-mátrix:** standardizált regressziós együtthatók, amelyeket a főkomponensek együttese, mint független változók segítségével az FKA-ba bevont változók regressziós előállítására kapunk, ami a standardizálás és a főkomponensek korrelálatlansága (ortogonális szerkezete) miatt megegyezik a változók és a főkomponensek korrelációs mátrixával.

²⁰ Ami igen ritkán fordul elő.

²¹ Forgatás után a főkomponensek már elveszítik maximális megmagyarázó képességüket, ezért használjuk ilyenkor velük kapcsolatban a „fő” előtag nélküli komponens elnevezést.

- **Mintázatmátrix (forgatott komponenssúly-mátrix):** standardizált regressziós együtthatók, amelyeket a forgatott komponensek együttese, mint független változók segítségével az FKA-ba bevont változók regressziós előállítására kapunk. Ez ortogonális (pl. Varimax) forgatás esetén megegyezik a változók és a forgatott komponensek korrelációs mátrixával, de ferde (pl. Promax vagy Oblimin) forgatás esetén a két mátrix különbözik. Elsődlegesen e mátrix alapján szokták értelmezni az FKA eredményét.
- **Struktúramátrix:** a változók és a forgatott főkomponensek korrelációs mátrixa, mely csak ferde forgatás esetén nyer jelentőséget.
- **Forgatott komponensek korrelációk táblázata:** ferde forgatás esetén ebből olvashatjuk ki, hogy a forgatott komponensek milyen szoros kapcsolatban vannak egymással.
- **Kommunalitások:** egy változó 0 és 1 közötti értékű kommunalitása varianciájának azon hányada, amelyet a megtartott főkomponensek megmagyaráznak. A változók kommunalitását a forgatás nem változtatja meg.

4.2. Az FKA menüablak

FKA-t ROP-R-ben a főkomponens-analízis (röviden: FKA) modul segítségével végezhetünk, a menüablak „input változók” ablakába áttett változók együttesén (lásd 4.1. ábra). Ha elfogadjuk a menüablak alapértelmezéseit, akkor ROP-R az 1-nél nagyobb sajátértékű főkomponenseket tartja meg és végez rajtuk ortogonális Varimax rotációt. A *Megtartott főkomponensek* panelen azonban be lehet állítani azt is, hogy a program egy megadott számú főkomponenst tartson meg.

Ezzel a modullal kvantitatív változók főkomponensanalízise (FKA) végezhető el, amelyet ortogonális (Varimax) vagy ferde (Promax vagy Oblimin) forgatással illeszthetünk jobban az input változóegyütteshez. FKA-ban lehetőség van a szélsőséges esetek azonosítására, valamint a Cronbach-alfa és a McDonald-féle omega kiszámítására is, ha az input változók mind egyazon skála egyirányú tételeinek tekinthetők.

4.1. ábra. Az FKA menüablak

A *Forgatás* panelen módosítható az alapértelmezés szerinti Varimax forgatás. Ha úgy akarjuk, nem lesz forgatás (*Nincs forgatás* opció) vagy megengedhetjük, hogy a forgatott komponensek akár korreláljanak is egymással (*Promax* vagy *Oblimin* ferde forgatás).

Az FKA modulban lehetőség van arra is (az *MBESS* R package segítségével; vö. Kelley, 2007), hogy megvizsgáljuk a kiválasztott változók együttesének, mint egy egydimenziós additív skála tételeinek (itemeinek) a belső konzisztenciáját. Ha a „Reliabilitási mutatók kiszámítása” opciót bejelöljük, akkor a program kiszámítja a kiválasztott változókra vonatkozóan a Cronbach- α és a McDonald-féle ω reliabilitásmutató mintabeli értékét és az elméleti értékre vonatkozó 95%-os konfidenciaintervallumot (vö. T. Kárász et al., 2022). Ezzel kapcsolatban fontos tudnivaló, hogy a fordított megfogalmazású negatív tételeket az FKA végrehajtása előtt át kell fordítani. Ez egyszerűen végrehajtható például a ROPstat itemanalízis moduljában, ha bejelöljük ott a „Negatív itemek átfordítása az adatállományban (a régi itemek felülírásával)” opciót. ROP-R-ben ugyanezt csak némileg bonyolultabban, a Transzformációk menüpont segítségével lehet megoldani.

Az FKA modulban lehetőség van a szélsőséges, extrém esetek felderítésére az „Extrém esetek (outlierek) azonosítása” opció bejelölésével, valamint az eset extremitás változó elmentésére (ugyanúgy, ahogy HierR-ben az influenszer esetek esetében; vö. 1.2.2. alpont). A mentéshez csupán az „Eset extremitás elmentése” opciót kell bejelölni. Ez a Mahalanobis-távolság egy GMD²² robusztus változatán alapuló elemzés (vö. Leys, Klein, Dominicy és Ley, 2018), ami a MASS R-package (Venables és Ripley, 2002) segítségével hajtható végre. ROP-R egy beépített küszöbérték segítségével jelöli ki, hogy mely esetek tekintendők outliernek, de ez a küszöb az „Extremitás beépített küszöbének szorzója” rovat segítségével rugalmasan módosítható. Ezen Outli néven elmentett extremitást mérő változó segítségével ezután megbízható többváltozós elemzések (pl. faktor- és klaszteranalízisek) végezhetők, ha az Outli változót feltételes csoportosító változónak jelöljük ki a menüablakban. Sajnos ennek az outlier detekciót végző elemzésnek van egy fontos feltétele: az elemzett változók egyikénél se lehet az alsó és a felső kvartilis (K1 és K3) egyenlő. Ha ugyanis K3-K1 (az interkvartilis terjedelem, angolul IQR) 0-val egyenlő, akkor a MASS package programja hibaüzenettel leáll, amit a ROP-R outputja MASS package problémaként visszajelez.

Külön lehetőségként (a „Komponenseket elment” opció bejelölésével) kérhető még a megtartott komponensek mentése, vagyis az elemzett adatállományhoz való illesztése is. Ha a forgatás opció be van jelölve, akkor a forgatott komponensek²³ lesznek elmentve, ha viszont ezt az opciót nem kérjük, akkor az eredeti, nem forgatott főkomponensek.

Egy FKA elemzés végrehajtása után a szokásos „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzésekhez elkészített ideiglenes adatfájlt (tmpdat.txt), valamint a futtatott R-scripteket (PCA.r, rotate.r).

4.3. Mit tartalmaz az FKA eredménylistája?

Az FKA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- A főkomponensek sajátértékei és magmagyarázott variancia százaléakai

²² Generalized Mahalanobis Distance (általánosított Mahalanobis távolság)

²³ A forgatott főkomponenseket szokták faktoroknak is nevezni.

- Főkomponenssúlyok mátrixa
- Rendezett főkomponenssúlyok mátrixa
- Csak forgatás esetén:
- Forgatás utáni komponenssúly-mátrix
- Forgatás utáni rendezett komponenssúly-mátrix
- Forgatott komponensek korrelációs mátrixa (ferde forgatás esetén)
- Struktúramátrix (ferde forgatás esetén)

4.4. Az FKA szemléltetése valódi adatokon

4.4.1. Főkomponens-analízis antropometriai adatokkal

Első elemzésként a gyermek születéskori és 10 éves kori testsúlya és testhossza, illetve testmagassága (Tsúly0, Thossz0, Tsúly10, Tmag10), valamint a szülők magassága és testsúlya (Anyasúly, Anyamag, Apasúly, Apamag) változókon végzünk FKA-t a 4128 fős ANTR1980 minta alapján, ahol 3072 személy rendelkezett minden antropometriai változó esetén érvényes értékkel (vö. B2.1. alpont). Az elemzés célja, hogy megpróbáljuk ezt a 8 változót lényegesen kevesebb számú, de azért szakmailag értelmezhető változóval helyettesíteni más elemzésekben.

A változók kiválasztását a 4.1. ábrának megfelelően végeztük el. Az érdemi FKA elemzés előtt outlier detekciót kértünk (bejelölve az „Extrém esetek azonosítása” opciót és 0,8-ra állítva a szélsőségeség küszöbének szorzóját), továbbá most még nem kértünk forgatást.

Az első futtatás eredménylistájának elején, közvetlenül az alapstatisztikák táblázata után találjuk az eset extremitásra vonatkozó információt. Ebből megtudhatjuk, hogy a beépített algoritmus által kiszámított küszöb: $D_{min} = 51,749$, emellett külön táblázat adja meg az ezen küszöbérték 80%-át²⁴ elérő GMD általánosított Mahalanobis távolságokat²⁵ a hozzájuk tartozó esetsorszámokkal (lásd 4.1. táblázat).

4.1. táblázat. A D_{min} érték 80%-át elérő GMD távolságok, esetsorszámokkal

GMD	Esetsorszám
51,759	2954
50,451	1638
47,288	2067
45,438	2216
44,628	3807
44,525	2961
42,842	155
42,005	2599
41,518	1539

A legnagyobb GMD érték (51,759) alig haladja meg a beépített algoritmus által kiszámított küszöb (51,749) értékét és nem kiugróan magas, továbbá a három legnagyobb

²⁴ A küszöbszorozót 0,80-ra állítottuk be.

²⁵ vö. Aggarwal, 2017, ill. https://www.cfholbert.com/blog/outlier_mahalanobis_distance/

GMD értékű személy adatait megtekintve nem találtunk érvénytelen, műtermék adatokra utaló jeleket. Ennek figyelembevételével az eredménylistán érdemes megtekinteni az FKA első fontos táblázatát, mely a feltárt főkomponensek sajátértékeit tartalmazza (lásd 4.2. táblázat). A táblázatból láthatjuk, hogy az 1-nél nagyobb sajátértékű főkomponensek együtt az input változók variációjának csupán a 66,45%-át magyarázzák (Kum% oszlop), ezért minimum 4 főkomponenst kell megtartanunk ahhoz, hogy a megmagyarázott varianciaarány 70% fölé emelkedjen. Az újabb FKA elemzés előtt ezért a menüablakban 4-ben rögzítettük a megtartandó főkomponensek számát és Oblimin ferde forgatást kértünk, hogy a forgatott komponensek szakmailag minél értelmezhetőbbek legyenek.

4.2. táblázat. A főkomponensek sajátértékei és magmagyarázott variancia százalécai

Index	Sajátérték	Megmagy. var. (%)	Kum%
1	2,840	35,50	35,50
2	1,431	17,89	53,39
3	1,045	13,06	66,45
4	0,946	11,82	78,27
5	0,787	9,84	88,11
6	0,434	5,42	93,53
7	0,305	3,82	97,35
8	0,212	2,65	100

Az Oblimin ferde forgatás utáni rendezett komponenssúly-mátrix (mintázatmátrix) a 4.3. táblázatban látható, ahol a 0,2-nél kisebb súlyok helyét a jobb áttekinthetőség érdekében üresen hagytuk. A rendezettség azt jelenti, hogy a forgatott komponensek (Komp1, Komp2, ...) oszlopaiban a változók a komponenssúly szerint csökkenő sorrendben vannak feltüntetve.

4.3. táblázat. Oblimin ferde forgatás utáni rendezett komponenssúly-mátrix (mintázatmátrix; a 0,2-nél kisebb súlyok helyét üresen hagytuk)

Változó	Komp1	Komp2	Komp3	Komp4
Thossz0	0,941			
Tsúly0	0,936			
Tsúly10		0,958		
Tmag10		0,797		
Apamag			0,891	
Apasúly			0,752	
Anyasúly				0,872
Anyamag			0,239	0,725

A 4.3. táblázat egy szakmailag jól értelmezhető struktúrát mutat, a komponenseket minden esetben 0,70 feletti súlyokkal kifizető változókkal. Az 1. komponens (Komp1) a gyermek születési, a 2. (Komp2) a 10 éves kori antropometriai adatait képviseli. Ugyanakkor Komp3 az apa, Komp4 pedig az anya antropometriai adatait képviseli. Ez a négy forgatott komponens szakmailag jól értelmezhető, együtt a nyolc input változó variációjának 78%-át

magyarázza (vö. 4.2. táblázat), így más elemzésekben jól helyettesíthetik az eredeti változókat.

Ha tudni szeretnénk, hogy ez a négy komponens milyen szoros kapcsolatban van eredeti változóinkkal, ehhez az outputon szintén megtalálható struktúramátrixot kell megnéznünk (lásd 4.4. táblázat, ahol a komponenseket kifesztítő változók korrelációit félkövérrel kiemeltük). E táblázat alján megtaláljuk azt az információt is, hogy az egyes komponensek rendre az összes változó varianciájából mennyit magyaráznak. Ezek az értékek kiegyenlített struktúrára utalnak.

4.4. táblázat. Oblimin ferde forgatás utáni struktúramátrix
(korrelációk a változók és a forgatott komponensek között)

Változó	Komp1	Komp2	Komp3	Komp4
Tsúly0	0,936	0,210	0,113	0,227
Thossz0	0,942	0,225	0,147	0,210
Tsúly10	0,168	0,935	0,187	0,264
Tmag10	0,310	0,870	0,397	0,317
Anyasúly	0,181	0,307	0,066	0,857
Anyamag	0,241	0,206	0,398	0,770
Apasúly	0,082	0,294	0,776	0,225
Apamag	0,176	0,217	0,889	0,203
MMV:	2,016	1,996	1,782	1,686
MMV%:	25,20	24,95	22,27	21,07

Megjegyzés: MMV = Megmagyarázott variancia

4.4.2. A Caprara-féle Pozitivitás skála tételeinek főkomponens-analízise

Következő elemzésként az élethez való pozitív hozzáállást és általános pozitív attitűdöt mérő Caprara-féle Pozitivitás Skála 8 tételén végzünk FKA-t az 1003 fős Btérkép2022 minta alapján (vö. B2.3. alpont). Az elemzés célja, hogy megnézzük: a 8 tétel valóban egyetlen fő konstruktmra illeszkedik-e. Ha igen, akkor azt várjuk, hogy a sajátértékek között egyetlen 1-nél nagyobb van, ez legalább 5-szöröse az utána következőnek, és minden tétel magas súllyal illeszkedik az első főkomponensre.

4.5. táblázat. A Dmin érték feletti GMD távolságok, a távolsághoz tartozó esetsorszám, valamint a skála 8 tételének az értéke az adott személynél (a tételneveket rövidítettük)

GMD	Esetsorszám	Poz1	Poz2	Poz3	Poz4	Poz5	Poz6	Poz7	Poz8
88,01	15	5	1	1	1	1	1	1	5
79,43	294	1	1	1	5	1	1	4	1
56,22	8	1	1	1	1	5	5	5	5
55,13	658	5	1	2	5	1	4	5	5
51,98	182	2	1	1	5	1	4	1	1
50,15	150	3	1	1	5	1	1	4	2
48,70	716	5	5	5	5	1	2	5	5

Első lépésként, az érdemi FKA elemzés előtt ismét outlier detekciót kértünk (az Outli nevű outlier változó elmentésével) és nem kértünk forgatást. A program 7 személynél jelezte, hogy a minta centrumától szélsőségesen távol esik (mert GMD értéke a küszöbérték felett van). Ezeket az eseteket a GMD távolsággal és a Pozitív skála 8 tételének értékével együtt a 4.5. táblázatban foglaltuk össze. Az értékek minden személy esetén mutatnak olyan inkonzisztenciát, aminek alapján jobbnak láttuk kihagyni őket a további elemzésekből (a létrehozott és elmentett Outli változó feltételes csoportosító változóként való kijelölésével).

Az ily módon 7 fővel csökkentett, de tisztított mintán ($n = 996$) való új futtatás után kapott főkomponensek sajátértékeit a 4.6. táblázat tartalmazza. Ennek alapján megállapíthatjuk, hogy a sajátértékek között valóban egyetlen 1-nél nagyobb van, és ez több mint 5-szöröse az utána következőnek. De vajon minden tétel magas súllyal illeszkedik az első főkomponensre? Ehhez meg kell tekinteni a főkomponenssúlyok táblázatát is erre az első főkomponensre (Komp1) vonatkozóan (lásd 4.7. táblázat).

4.6. táblázat. A főkomponensek sajátértékei és magmagyarázott variancia százalékai

Index	Sajátérték	Megmagy. var. (%)	Kum%
1	5,299	66,23	66,23
2	0,970	12,12	78,36
3	0,564	7,05	85,40
4	0,391	4,88	90,28
5	0,261	3,26	93,55
6	0,237	2,96	96,51
7	0,163	2,04	98,55
8	0,116	1,45	100

4.7. táblázat. Az első főkomponens súlyainak táblázata

Változó	Fkomp1
Poz1Opt	0,888
Poz2Elég	0,910
Poz3Elég	0,779
Poz4Opt	0,913
Poz5Önb	0,906
Poz6OptF	0,231
Poz7Önb	0,847
Poz8Önb	0,805

A 4.7. táblázatból azt látjuk, hogy a 6. számú és fordított megfogalmazású Poz6OptF tétel kivételével minden tétel 0,77 feletti, igen magas súllyal illeszkedik az első főkomponensre. Azt nem tudjuk, hogy Poz6OptF esetében az alacsony érték oka a tétel tartalmilag nem megfelelő volta, vagy az, hogy a tétel fordított megfogalmazású. A teszt megszerkesztője, Gian Vittorio Caprara szerint az utóbbi az igaz. Személyes véleménye alapján éppen azért tartotta benn a skálában ezt a tételt, hogy mérsékelje a többi tétel formai dominanciáját, ugyanis egy pszichológiai konstruktum nem függhet a tételek formai jellemzőitől.

Mindenesetre a Pozitivitás skála ilyen formájú változatának a pszichometriai megfelelőségét támasztják alá azok a reliabilitási adatok, amelyeket FKA-ban a „Reliabilitási mutatók kiszámítása” opció bejelölésekor kaptunk, ugyanis a Poz6OptF tétel átfordítása után ismét futtatva FKA-t, Cronbach- α értékére 0,914 ($CI_{0,95} = [0,906; 0,922]$), a McDonald-féle ω értékére pedig 0,927 ($CI_{0,95} = [0,920; 0,933]$) adódott.

5. fejezet

A feltáró faktoranalízis (EFA) modul

Ezzel a modullal kvantitatív változók *feltáró, exploratív faktoranalízise* (EFA) végezhető el három különböző faktorkiszűrési módszerrel (maximum likelihood módszer, főfaktoranalízis vagy főfaktorelemzés, illetve minimum reziduális módszer). A végső faktorstruktúrát ortogonális (Varimax) vagy ferde (Promax vagy Oblimin) forgatással illeszthetjük jobban az input változóegyütteshez.

Az 5.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, az 5.2. alfejezet az EFA menüablak használatát mutatja be, az 5.3. alfejezet az eredménylista szerkezetét, az 5.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

5.1. Elméleti alapok

FKA-val szemben a faktoranalízis (FA) alapgondolata, hogy létezik egy közvetlenül nem megfigyelhető, *látens faktorstruktúra*, mely valamilyen módon kifejezi, képviseli vizsgált változóink – korrelációs vagy kovariancia-mátrixuk által tartalmazott – közös lényegét, és lényegében ezen látens faktorok hatásaként alakulnak megfigyelt változóink értékei. FA egyik fő célja e faktorstruktúra feltárása, amit EFA segítségével hajthatunk végre (lásd pl. Vargha, 2019, 5. fejezet). FA modelljében a vizsgált változók mindegyikének a varianciája két összetevőre bontható:

- a több változóban is jelen levő, közös faktorok által meghatározott rész;
- a csak erre a változóra jellemző, s ily módon a többi változótól független egyedi rész (reziduális), mely tartalmazza az ezen változó mérésével kapcsolatos hibát is.

FKA-val ellentétben FA csak a változók által közösen lefedett részre alkot modellt, nem célja a változók teljes varianciáját megmagyarázni. FKA esetében a hangsúly az eredeti változókon van, FA esetében viszont a közvetlenül nem megfigyelhető, látens faktorokon, amelyek a dolgok igazi lényegét fejezik ki, megbújnak a háttérben és onnan hatnak mért változóinkra (vö. 4.0. ábra). FKA főkomponensei nem látensek, azok az eredeti változók rögzített szabály szerint képezett transzformáltjai (súlyozott összegei), konkrét mintán személyenként egyértelműen meghatározható konkrét értékűek. Ezzel szemben az FA látens faktorait konkrét mintán is csak becsülni tudjuk, ahogy a köztük lévő korrelációkat vagy faktorsúlyaikat, s ezek értékei a becslési módszertől is függő közelítéseknek tekinthetők.

Mivel EFA-ban a változóknak csak a közös (legalább két változó által lefedett) részét faktorizáljuk, egy jó faktorstruktúra létezésének elengedhetetlen feltétele, hogy a változók közös része, amit a Kaiser-Meyer-Olkin (röviden KMO) mutatóval szoktak leggyakrabban mérni, kellően nagy legyen (vö. Vargha, 2019, 100. o.). Fentiekből következően az EFA-ban feltárt faktorok az eredeti változók varianciájának csak kisebb részét magyarázzák, mint az FKA főkomponensei.

Az EFA-ban használt legfontosabb fogalmak, amelyek részben a konfirmatív FA-ban is szerepet kapnak, az alábbiak.

- **Elsődleges faktorsúlymátrix:** standardizált regressziós együtthatók, amelyeket a választott módszerrel elsődlegesen feltárt látens faktorok együttese, mint független változók segítségével a faktormodellbe bevont változók regressziós előrejelzésére kapunk, s ami a standardizálás és a látens faktorok korrelálatlansága (ortogonális szerkezete) miatt megegyezik a változók és a látens faktorok becsült korrelációs mátrixával.
- **Mintázatmátrix (forgatott faktorsúlymátrix):** standardizált regressziós együtthatók, amelyeket a forgatott látens faktorok együttese, mint független változók segítségével a faktormodellbe bevont változók regressziós előrejelzésére kapunk. Ez ortogonális (pl. Varimax) forgatás esetén megegyezik a változók és a látens faktorok becsült korrelációs mátrixával, de ferde (pl. Promax vagy Oblimin) forgatás esetén a két mátrix különbözik. Elsődlegesen e mátrix alapján szokták értelmezni a feltárt faktorstruktúrát.
- **Struktúramátrix:** a változók és a látens faktorok becsült korrelációs mátrixa.
- **Faktorkorrelációk táblázata:** a forgatott látens faktorok becsült korrelációs mátrixa. Ferde forgatás esetén ebből olvashatjuk ki, hogy a forgatott faktorok milyen szoros kapcsolatban vannak egymással.
- **Kommunalitás becslések:** a végső struktúrában a faktorok által egy változó variáciájából megmagyarázott varianciaarány (a mintázatmátrixszal kapcsolatban említett, adott változóra vonatkozó regressziós becslés R^2 -értéke). Egy változó 0 és 1 közti értékű kommunalitása tehát varianciájának azon hányada, amelyet a közös faktormodell megmagyaráz. *Unicitása (reziduális varianciaarány)* pedig az, ami a kommunalitást 1-re kiegészíti. A változók kommunalitását a forgatás nem változtatja meg. A kommunalitások összege rögzített faktorszám esetén a változókból a faktorok által megmagyarázott összvariancia, amelyet a változók számával leosztva az összvariancia-arányt kapjuk. Ez az arány mindig kisebb, mint ugyanazon a mintán ugyanazokkal a változókkal FKA elemzést végezve, az azonos számú megtartott főkomponenshez tartozó kumulatív százalékarány.

A kapott faktorstruktúra értelmezhetősége EFA esetében is fontos szerepet játszik. A forgatás lényege, hogy a forgatás eredményeképpen létrehozott faktorok jobban illeszkednek az eredeti változókra, mint forgatás előtti megfelelőik, így azok az értelmezést megkönnyítő szakmai jelentéssel is felruházhatók. Jó tudni, hogy forgatással a változókból megmagyarázott teljes variancia aránya nem változik, vagyis a forgatással nem többet, hanem jobban tudunk megmagyarázni a változókból.

EFA-ban több módszer is létezik az *elsődleges faktorstruktúra* feltárására, ezek közül a leggyakoribb a *maximum likelihood* módszer (ML), a *főfaktorelemzés* (principal axis factoring vagy PAF) és a *minimum reziduális* módszer (minimum residual vagy MinRes) (lásd Comrey, 1962; Osborn, 2014; Tabachnick és Fidell, 2013, 13. fejezet; Vargha, 2019, 5. fejezet).

ML alkalmazási feltétele, hogy a vizsgált változók együttes eloszlása többdimenziós normális legyen. E feltétel teljesülése esetén ML elsődleges faktormodellje olyan struktúra, amelynek fennállása mellett a legnagyobb a valószínűsége az adott mintabeli korrelációs mátrixnak. ML iteratív módszerrel becsüli a faktorsúlyokat. ML a normalitás súlyos sérülése esetén gyenge megoldást produkál.

PAF esetén a normalitás nem fontos feltétel és ez a módszer hasonló logikájú, mint FKA. Csupán annyiban különböznek, hogy míg FKA-ban olyan főkomponenseket hozunk létre, amelyekkel a standardizált változók teljes összvarianciájából a lehető legnagyobb részt tudjuk magyarázni, a PAF által feltárt faktorokkal a változók közös összvarianciájának szeretnénk a lehető legnagyobb részét megmagyarázni. Ez is iterációs módszer, melynek első lépésében minden változó kommunalitását a többi változó által magyarázott varianciaarányal becsüljük. Ezek ismeretében meghatározható egy faktormodell a hozzá tartozó factorsúlyokkal. Az ezekből kiszámítható kommunalitások a közös összvariancia jobb becslését szolgáltatják, s ezt az iterációs folyamatot addig folytatjuk, amíg a modell javul (az input korrelációs mátrix és a faktormodellből kiszámítható korrelációs mátrix hasonlóbb lesz).

MinRes esetén sem fontos feltétel a normalitás. MinRes a regresszióból jól ismert legkisebb négyzetek módszerével keres olyan megoldást, amelynek során a lehető legkisebb lesz a faktormodellben a változók reziduálisainak a négyzetösszege. Egy változó reziduálisa a változónak az a része, amelyet a faktormodell nem magyaráz meg. Hasonló logikájuk miatt a MinRes és a PAF hasonló megoldásokra szokott vezetni (Osborn, 2014).

5.2. Az EFA menüablak

EFA-t ROP-R-ben a feltáró faktoranalízis (röviden: EFA) modul segítségével végezhetünk, a menüablak „input változók” ablakába áttett változók együttesén (lásd 5.1. ábra). Ha elfogadjuk a menüablak alapértelmezéseit, akkor ROP-R a *parallel* módszerrel dönt az optimális factorszámról és PAF típusú EFA-t hajt végre ortogonális Varimax rotációval.

5.1. ábra. Az EFA menüablak

A parallel-elemzés lényege, hogy összehasonlítjuk a mintából generált EFA-sajátértékeket egy azonos méretű, véletlenszerű (random) adatokból létrehozott, szimulált mintából kiszámított EFA-sajátértékekkel, majd annyi faktort tartunk meg, amennyi esetében a valódi minta faktorának sajátértéke nagyobb, mint a sorrendben neki megfelelő szimulált minta faktorának sajátértéke (Lim és Jahng, 2019). Számos statisztikai szoftver EFA-moduljában (jamoviban és a ROP-R legújabb változatában is) ez a faktorszám megállapításának alapmódszere.

A *Faktorszám meghatározása* panelen be lehet állítani azt is, hogy a program egy megadott számú faktort tartson meg, de erre a célra alkalmazható még az FKA sajátértékei vagy az EFA-sajátértékek alapján könnyen elkészíthető *lejtődiagram* (Cattell, 1966) is. Míg az FKA-sajátértékek az input változók korrelációs mátrixából határozhatók meg, az EFA-sajátértékeket ugyanebből a mátrixból úgy képezzük, hogy a főátló 1-es értékei²⁶ helyett változónként a változó varianciájából a többi változó által meghatározott varianciaarányt (a kommunalítások első becslését), vagyis a velük közös rész arányát írjuk. Emiatt az EFA-sajátértékek mindig kisebbek, mint a nekik megfelelő FKA-sajátértékek, s így közöttük negatív értékek is előfordulhatnak. A pozitív sajátértékek számánál nagyobb faktorszámot nincs értelme megadni.

Ha a faktorszám meghatározását az FKA sajátértékei alapján végezzük, itt a beállítható küszöb alapértéke 1, vagyis alapértelmezésben annyi faktort tartunk meg, ahány 1 változónál többet megmagyarázni képes főkomponens van. Ha az EFA-sajátértékek alapján döntünk, a küszöb felajánlott értéke 0,1, vagyis alapértelmezésben annyi faktort tartunk meg, ahány a változók közös részéből legalább 0,1 varianciányinál többet meg tud megmagyarázni.

A *Forgatás* panelen az alapértelmezés a Varimax forgatás, de ha úgy akarjuk, nem lesz forgatás (*Nincs forgatás* opció) vagy megengedhetjük, hogy a forgatott faktorok korreláljanak egymással (*Promax* vagy *Oblimin* ferde forgatás).

Külön lehetőségként kérhető még a megtartott faktorok mentése (a „Faktorokat elment” opció bejelölésével), vagyis az elemzett adatállományhoz való illesztése is. Ekkor ha a *Nincs forgatás* opció van bejelölve, akkor az elsődlegesen kiszűrt faktorok becslései lesznek elmentve, egyébként a forgatott faktorbecslések.

Egy EFA elemzés végrehajtása után a szokásos „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzéshez elkészített ideiglenes adatfájlt (tmpdat.txt), a futtatott R-scripteket (Rbegin.r, EFA.r), valamint a parallel-elemzés lejtődiagramját (screep1.jpg). Ez utóbbit ROP-R akkor is elkészíti, ha a faktorszámot nem a parallel módszerrel állítjuk be.

5.3. Mit tartalmaz az EFA eredménylistája?

Az EFA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- Kaiser-Meyer-Olkin adekvációs mutató (KMO)
- A főkomponensek sajátértékei (FKA-sajátértékek) és magmagyarázott variancia százalécai
- Az elsődleges faktorok sajátértékei (EFA-sajátértékek)
- Multikollinearitás diagnosztika

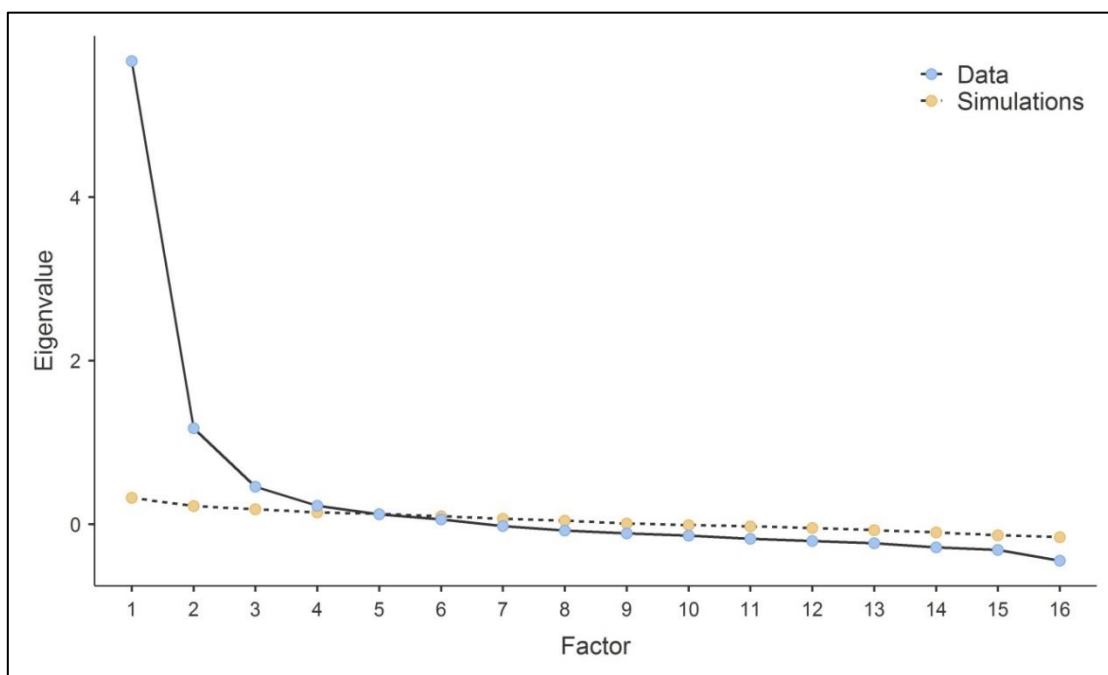
²⁶ Ezek a standardizált változók teljes varianciáját képviselik.

- Rotálatlan faktorsúlyok mátrixa
 - Modellilleszkedés tesztelése khi-négyzet-próbával
 - A feltárt modell néhány kiszámítható adekvációs mutatója
- Csak forgatás esetén:
- Forgatott (rotált) faktorsúlymátrix
 - Forgatás utáni rendezett faktorsúlysúly-mátrix
 - Forgatott faktorok korrelációs mátrixa (ferde forgatás esetén)
 - Struktúramátrix (ferde forgatás esetén)

5.4. Az EFA szemléltetése valódi adatokon

Szemléltetéseként a pszichológiai immunrendszer állapotának és összetevőinek feltérképezésére megszerkesztett PIK16 Pszichológiai Immunrendszer Kérdőív (Oláh 2005) 4-fokozatú, Likert-skálájú tételein (vö. B.6. táblázat) végzünk EFA-t az 1003 fős Btérkép2022 minta alapján (vö. B2.3. alpont). Az elemzés célja, hogy megnézzük: a kérdőív 16 tétele valóban a 4 skálának megfelelő 4 fő konstruktumot feszít-e ki, és ha igen, akkor ez a skálabesorolásokkal összhangban van-e. Az elemzések előtt a negatív tételeket átfordítottuk.

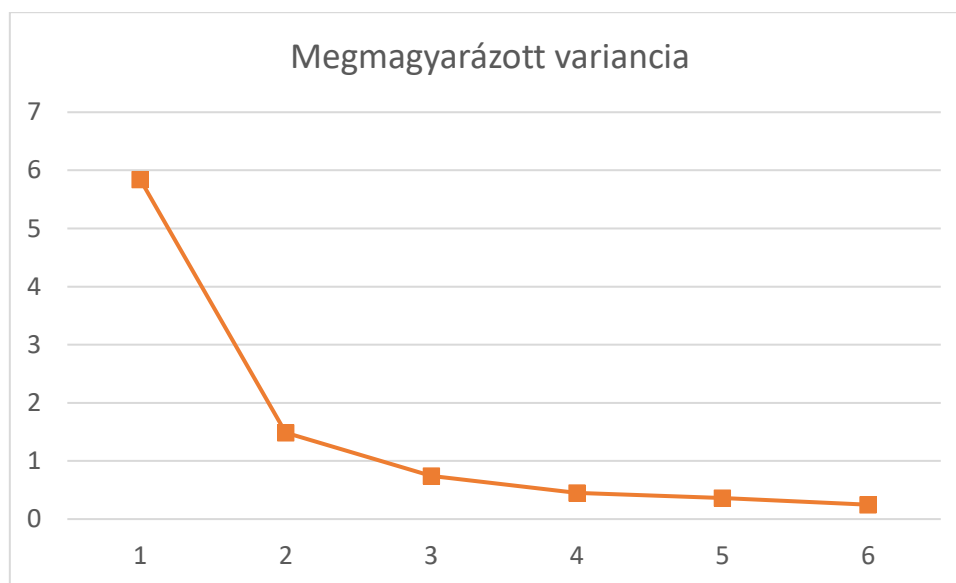
Első lépésként, az érdemi EFA elemzések előtt az FKA modul segítségével outlier detekciót kértünk a 16 PIK16 tételre. A program 6 személynél jelezte, hogy a minta centrumától szélsőségesen távol esik, de közülük egyikük GMD-értéke sem volt kiugróan magas a program által jelzett küszöbhez viszonyítva, másrészt értékeiket megszemlélve nem láttunk olyan inkonzisztenciát, ami fals adatokra, outlier esetre utalt volna. Ennek alapján nem láttuk indokoltnak, hogy bárkit is kihagyjunk a további elemzésekből. Megjegyezzük, hogy az outlier detekciót elemzésenként és változócsopontonként érdemes elvégezni. Az, hogy valaki a Caprara kérdőív alapján outliernek tekinthető (vö. 4.4.2. alpont), még nem jelenti azt, hogy a PIK16 elemzésekor is automatikusan annak kell tekinteni.



5.2. ábra. A PIK16 tételein alkalmazott parallel-elemzés lejtődiagramja

Elfogadva EFA-ban az alapértelmezés szerinti parallel módszert, ROP-R az 5.2. ábrán látható lejtődiagramot készítette el, hogy informálódjunk a 16 tételt kifeszítő látens faktorok számáról. A lejtődiagram azt mutatja, hogy az EFA-sajátértékek csupán az első négy faktor esetén nagyobbak érzékelhetően a szimulációs adatok alapján meghatározott sajátértékeknél, így a parallel módszer alapján 4 faktor megtartását javasoljuk.

Megjegyezzük, hogy a szimulációval létrehozott random adatok elemzésről-elemzésre mások lesznek, így a parallel módszer futásról-futásra kis mértékben eltérő faktorszámhoz is vezethet. Például konkrét adatmintánk esetében három futásból a program egy esetben 4, két esetben 5 faktort javasolt. A lejtődiagram azonban mind a három esetben az 5.2. ábrán láthatóhoz hasonlított, minden esetben 4 faktor esetén jelezve a mintánkból számított EFA-sajátérték dominanciáját a random mintából számított sajátértékkel szemben.



5.3. ábra. Az elsődleges faktorokhoz tartozó megmagyarázott varianciák lejtődiagramja (változók: 16 PIK16 tétel, EFA módszer típusa: PAF)

A faktorszám meghatározásának egy másik módszereként EFA elemzést végeztünk forgatás nélkül, 6 elsődleges faktort beállítva (az első futásnál az EFA-sajátértékek itt nem közölt táblázatában a sajátértékek a 7-ediktől kezdve fordultak pozitívból negatívba), hogy megnézzük az ezen feltárt faktorok által megmagyarázott variancia értékeket. Ezek mindig kisebbek, mint az azonos sorszámmal rendelkező sajátértékek, mert FKA a változók teljes, EFA pedig csak a változók közös varianciáját faktorizálja. Mivel a PIK16 4-fokozatú tételei nehezen fogadhatók el normális eloszlásúnak, EFA módszereként PAF-ot választottuk. Az elsődleges faktorok által megmagyarázott varianciákat az eredménylistán a rotálatlan faktorsúlyok mátrixa alatt találjuk, az ezek alapján elkészített lejtődiagram pedig az 5.3. ábrán látható. Ez az ábra is azt mutatja, hogy a lejtődiagram a 4. faktor után, az 5. faktortól kezd ellaposodni, ami ismét 4 faktor jelenlétére utal.

EFA értelmes voltához ezután az eredménylistán megnéztük a Kaiser-Meyer-Olkin adekvációs mutató értékét is (az alapstatisztikák táblázata alatt). Mivel KMO értéke 0,911 (kitűnő), ez arra utal, hogy a 16 tétel elegendő közös részt tartalmaz egy EFA elvégzéséhez. Így az EFA-t 4 faktorral, PAF módszerrel és Promax ferde forgatással végeztük el.

5.1. táblázat. A feltárt faktormodell néhány kiszámított adekvációs mutatója

RMSEA =	0,057	CI_90 RMSEA elméleti értékére: (0,050; 0,064)
RMSR =	0,023	Root Mean Square of Residuals
CFI =	0,969	Comparative Fit Index
TLI =	0,941	Tucker-Lewis Index

Az eredménylistán érdemes először a feltárt 4-faktoros faktormodell néhány kiszámított adekvációs mutatójának táblázatát megtekinteni (lásd 5.1. táblázat). Ezek majd a 6. fejezetben ismertetésre kerülő konfirmatív FA során nyernek igazán jelentőséget, de azért itt is megállapíthatjuk, hogy ezek az adekvációs mutatók jónak mondhatók, mert $RMSEA < 0,06$, $RMSR < 0,08$, $CFI > 0,95$ és $TLI > 0,90$ (vö. 6.1. táblázat). Ennek ellenére a táblázat felett azt látjuk az outputon, hogy a modellilleszkedés tesztelése χ^2 -próbával erősen szignifikáns ($\chi^2 = 266,039$, $f = 62$, $p < 0,0001$), ami a faktormodell nem megfelelő voltát jelzi. Ez persze nem lep meg, mert egy 1003 fős mintán parányi hatások is szignifikánsak tudnak lenni.

5.2. táblázat. A PIK16 tételeinek Promax ferde forgatás utáni rendezett faktorsúlymátrixa

Változó	Faktor1	Faktor2	Faktor3	Faktor4
PIK12M	0,958			
PIK01M	0,645	0,292		
PIK15M	0,566			
PIK03M	0,501			
PIK06M	0,431	0,366		
PIK13M	0,417		0,309	
PIK14R		0,755		
PIK10R	0,236	0,703		
PIK11R		0,491		0,257
PIK08AV			0,791	
PIK07AV			0,674	
PIK09AV	0,337		0,415	
PIK05Ö				0,887
PIK02Ö	0,215			0,561
PIK16Ö	-0,243			0,524
PIK04AV	0,314		0,275	

EFA eredménylistáján a legfontosabb a forgatás után kapott faktorsúlymátrix (ferde forgatás esetén az ún. *mintázatomátrix*), melynek a faktorok és faktorsúlyok szerint rendezett táblázatát az 5.2. táblázat tartalmazza. Ez a táblázat jól reprodukálja a PIK16 skálabeosztását, ugyanis a 16 tétel között mindössze egyetlen olyat találunk (PIK04AV), amelyik nem illeszkedik 0,40 feletti faktorsúllyal a saját skálájának megfelelő faktorhoz. Megjegyezzük, hogy a MinRes módszerrel is elvégezve ugyanezt az elemzést, nagyon hasonló táblázatot kaptunk, a faktorsúlyok egyetlen esetben sem tértek el 0,01-nél nagyobb mértékben az 5.2. táblázatban látottól. Promax helyett Oblimin ferde forgatást választva hasonló struktúrájú, de a PIK16 skáláit kicsit gyengébben reprodukáló megoldást kaptunk. Ha kíváncsiak vagyunk a

változók és a látens faktorok közti korrelációkra is, mely ferde forgatás esetén különbözik a mintázatmátrixtól, akkor az outputon a struktúramátrixot kell megkeresni.

5.3. táblázat. A forgatott látens faktorok korrelációs mátrixa

	Faktor1	Faktor2	Faktor3	Faktor4
Faktor1	1	0,649	0,692	0,309
Faktor2	0,649	1	0,444	0,463
Faktor3	0,692	0,444	1	0,180
Faktor4	0,309	0,463	0,180	1

Az eredménylistán említésre méltó még a multikollinearitás diagnosztika táblázata (a főkomponensek sajátértékei alatt), valamint a forgatott látens faktorok korrelációs mátrixa (a ferde forgatás utáni rendezett faktorsúlymátrix alatt), mely az 5.3. táblázatban látható, s ahonnan azt olvashatjuk ki, hogy erős korrelációs kapcsolat van a Megközelítő-monitorozó skálát képviselő 1. faktor, valamint a Reziliencia skálát képviselő 2. és az Alkotó-végrehajtó hatékonyság skálát képviselő 3. faktor között.

Végül érdemes megnézni, hogy milyen korrelációs kapcsolat van az EFA által feltárt négy faktor és a PIK16 négy skálája között. Ezt a korrelációs elemzést – Pearson korrelációkat kiszámítva – ROPstatban végeztük el és az eredményeket az 5.4. táblázatban foglaltuk össze. Ebből jól látható, hogy az EFA által feltárt mind a négy faktor legalább 0,95-ös szinten korrelál a neki megfelelő PIK16 skálával (lásd az 5.4. táblázatban a félkövérrel kiemelt korrelációkat).

5.4. táblázat. Az EFA által feltárt és elmentett faktorok korrelációi a PIK16 skálákkal (a faktornév után zárójelben jelezzük, hogy a faktor mely skálát képviseli)

Változó	Mmonit	AVhat	Rezil	Önreg
Faktor1 (MM)	0,980**	0,767**	0,606**	0,375**
Faktor2 (R)	0,719**	0,549**	0,947**	0,526**
Faktor3 (AV)	0,779**	0,964**	0,383**	0,247**
Faktor4 (ÖR)	0,327**	0,223**	0,558**	0,954**

Jelölés: **: $p < 0,001$

Összességében elmondhatjuk, hogy mind a faktorok száma, mind a skálahovatartozás, mind az EFA-val feltárt és a mintában kiszámított skálák közti szoros kapcsolat tekintetében sikerült a PIK16 kérdőív faktorstruktúráját jól reprodukálni.

6. fejezet

A konfirmatív faktoranalízis (CFA) modul

Ezzel a modullal konfirmatív faktorelemzés (CFA) végezhető, mely által tesztelhető egy hipotetikus (pl. más mintán EFA-val feltárt) faktorstruktúra elfogadhatósága, illetve jósága különböző adekvációs mutatók segítségével. Kérésre végezhető másodrendű CFA²⁷ is, mely esetben a ROP-R elkészíti a bifaktoros CFA R-beli futtatására alkalmas scriptet is.

A 6.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 6.2. alfejezet a CFA menüablak használatát mutatja be, a 6.3. alfejezet az eredménylista szerkezetét, a 6.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

6.1. Elméleti alapok

FKA-ban főkomponensekként a változók olyan súlyozott összegeit keressük, amelyek a változók teljes varianciájának a lehető legnagyobb részét megmagyarázzák. EFA-ban olyan faktormodellt keresünk, amelynek faktorai a változók *közös részét* magyarázzák a lehető legnagyobb mértékben. Ennek egy kritériuma lehet például az, hogy a faktormodell mennyire képes reprodukálni a változók közös részét képviselő kovariancia-mátrixot, illetve korrelációs mátrixot (a standardizált változók kovariancia-mátrixát). CFA-ban is az EFA kritériumait alkalmazzuk, csak az alábbi plusz megszorítások mellett.

1. Minden látens faktor csak azokból a változókból (tégelekből) építhető, amelyeket számára kijelölünk. Így a látens faktorok csak saját tételeikkel korrelálhatnak, más faktorok tételeivel nem, vagyis keresztöltéseket nem engedünk meg.
2. A végső modellben a tételek reziduálisai egymással nem korrelálhatnak (sem egyazon faktoron belül, sem különböző faktorok között).

Ezek a megszorítások a modellillesztés javítása céljából olykor enyhíthetünk, ha ezt szakmai érvekkel is alá tudjuk támasztani.

Egy konkrét mintán a CFA elemzés abból áll, hogy keresünk egy olyan faktormodellt, amely eleget tesz a fenti megszorításoknak és a lehető legjobban illeszkedik adatainkra, illetve az azokból kiszámítható korrelációs vagy kovariancia-mátrixra. Az alkalmas faktormodell feltárása nem olyan egyszerű és egyértelmű, mint FKA-ban a főkomponensek meghatározása. A faktormodellt különböző módszerek (pl. ML, MLR, MLMV, ULS, WLS²⁸ stb.) segítségével, minden esetben csak több lépésben, a modellt lépésről-lépésre javítva lehet becsülni.

²⁷ Szokták a hierarchikus CFA elnevezést is használni.

²⁸ Az ML-lel kezdődő módszerek mind maximum likelihood becslés alapúak, az LS végződésűek pedig a legkisebb négyzetes (Least Squares) regressziós becslésen alapulnak.

Egy faktormodell jóságának a mérésére CFA-ban különböző együttthatók (RMSEA, CFI, TLI stb.) állnak rendelkezésre, a tesztelést pedig a más statisztikai eljárásokból is ismert χ^2 -próbás illeszkedésvizsgálattal lehet elvégezni. Például RMSEA²⁹, mint egy abszolút illeszkedési mutató a tételek kovariancia-mátrixát és a CFA modelljéből becsült kovariancia-mátrixot hasonlítja össze. Ehhez kiszámítjuk a két mátrixból az elemenkénti négyzetes eltérések átlagát, RMSEA ennek a négyzetgyöke (Kline, 2010). CFI³⁰ ugyanakkor egy relatív összehasonlító mutató, mely azt méri, hogy a faktormodell tesztelésére alkalmazott χ^2 illeszkedési statisztika mennyivel kisebb (tehát a modell mennyivel jobb illeszkedésű), mint annak a lehető legrosszabb, ún. alapmodellnek a χ^2 -értéke, ahol a változók és a látens faktorok együttesében minden korreláció 0, vagyis ahol semmi nem függ össze semmivel (Vargha, 2019, 121. o.). Az illeszkedés jóságának mérésére használt főbb mutatók összefoglalását lásd például Vargha (2019, 122. o.), Schreiber és munkatársai (2006), illetve Kenny (2020).

Ha egy faktormodell CFA-ban jónak bizonyul, még további komplexebb kérdéseket is meg lehet foglalmazni. Például igaz-e, hogy a teszt skálái egy közös konstrukcióra illeszkednek (pl. a MET 5 skálája a mentális egészség konstrukciójára)? Vagy igaz-e az, hogy egy megerősített faktormodell ugyanúgy érvényes különböző csoportokban? Előbbi kérdést például a másodrendű (Schivinski, 2013) vagy a bifaktoros (Dunn és McCray, 2020) CFA-modellek segítségével lehet megvizsgálni, utóbbit pedig a mérési invariancia módszerével (Chen, 2007; Kim és Yoon, 2011; Vandenberg és Lance, 2000).

Egy faktormodellt a CFA-ban a *faktorábra*³¹ (*faktordiagram*) segítségével szoktak ábrázolni, melynek több formája is létezik. Közös bennük, hogy a diagramon a látens faktort vagy faktorokat ellipszis vagy kör, a megfigyelt (manifeszt) változókat pedig téglalap alakú keretbe helyezik. Például egy sima egydimenziós (egyskálás) faktormodell a 6.1. ábrán látható. Az ábra a 16-tételes PIK16 (vö. B2.3. alpont) egyfaktoros struktúráját ábrázolja, melyen a teszt 16 tétele egy PIK³² elnevezésű közös faktorra illeszkedik.

A 6.1. ábrán a tételek nevének végén – a sorszám jelzete után – látható betűk a skálához tartozásra utalnak, bár az ábrán látható egydimenziós konstrukcióban a négy skálát külön nem vontuk be a modellbe. Az ábrán a PIK látens faktortól a tételekhez vezető egyirányú nyilak (útvonal szakaszok) fölötti számok egy konkrét faktormodell standardizált faktorsúly becslései (egyben a látens faktor és a tételek közti korrelációk becslései), melyeket az 1003 fős Btérkép2022 mintán (lásd B2.3. alpont) elvégzett CFA elemzés során kaptunk. A 6.1. ábrán a tételek között azért nem látunk nyilakat, mert ebben az egydimenziós modellben alapértelmezés szerint feltételezzük, hogy minden, ami a tételekben közös, az benne van a látens faktorban, a tételek egyediségei, reziduális részei pedig egymástól függetlenek, azaz nem korrelálnak egymással. Ezen a megszorításon indokolt esetben enyhíteni szoktak, ami látszólag ellentmond annak a korábbi kijelentésünknek, hogy EFA-ban a változók teljes közös, legalább két változó által lefedett részét faktorizáljuk. A kérdés az, hogy miért nem része a közös faktornak az, ami csak két tételt köt össze? A válasz: EFA megpróbálja a változókban lévő összes közös részt feltérképezni és csoportosítani, CFA pedig a közös szakmai lényegét azonosítani. Ha két tételt az a formai jellemző köt össze, hogy verbális megfogalmazásuk nagyon hasonló és ez inspirál esetükben hasonló válaszokat és eredményez

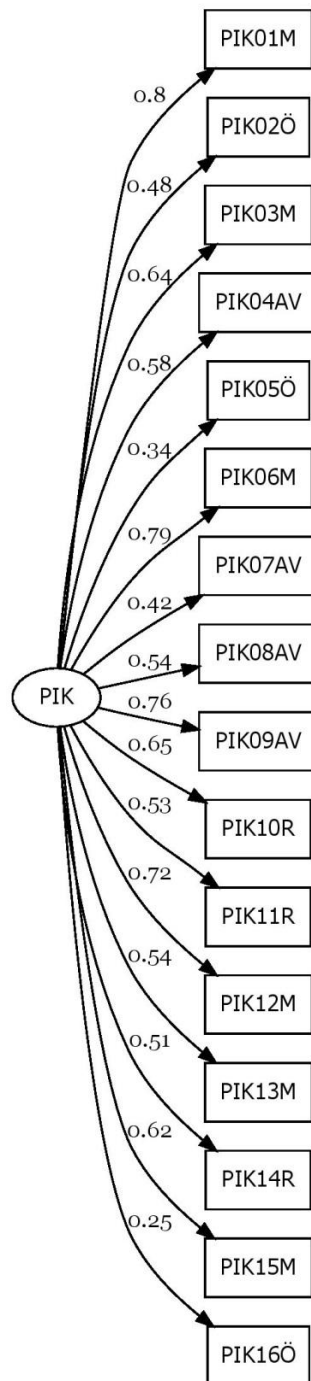
²⁹ Root Mean Square Error of Approximation

³⁰ Comparative Fit Index

³¹ Ez is egyfajta útvonalábra (path diagram)

³² Pszichológiai immunkompetencia

a két tétel között pozitív korrelációt, érdemes kapcsolatukat mint egy egyedi összefüggést bevonni a modellbe, mely független a teszt fő konstruktumától.

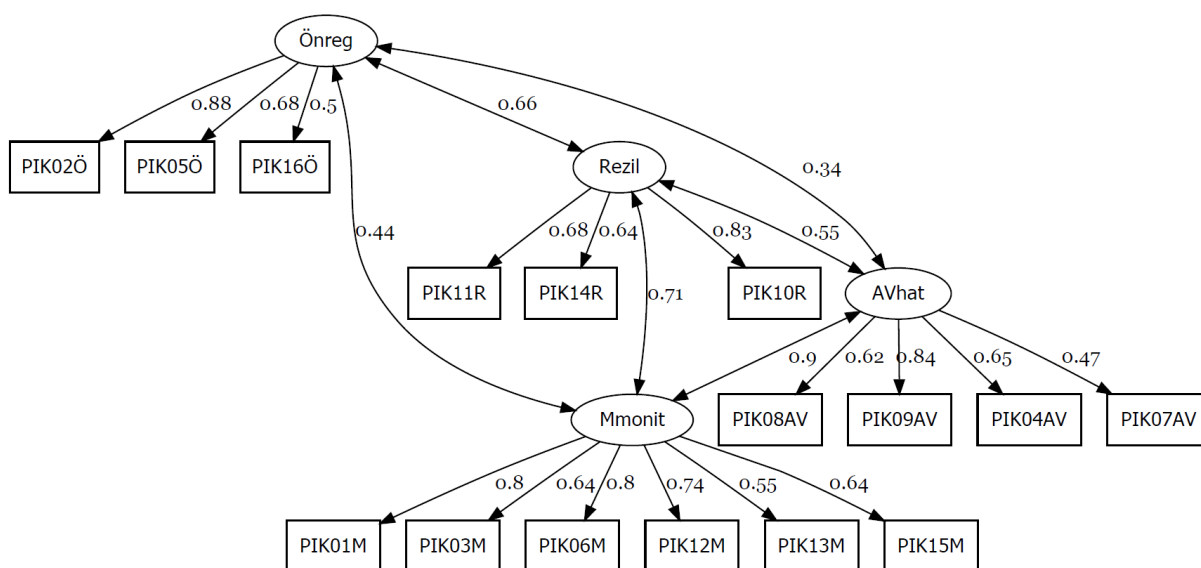


6.1. ábra. A CFA grafikonja egyetlen elsőrendű látens faktor (PIK) esetén

EFA nem hierarchikusan tárja fel a faktorképző tényezőket. CFA-ban azonban lehetőségünk van arra, hogy az adatainkat egy ennél komplexebb faktormodellre illesszük. Már maga az is a komplexitás irányába tett lépés, hogy egyes tételpárok kovarianciáit beépítjük a modellbe. Az alábbiakban ismertetésre kerülő elsőrendű többfaktoros, másodrendű, majd bifaktoros modellek a komplexitás különböző szintjeit és típusait képviselik.

6.1.1. Elsőrendű faktormodell több korreláló faktorra

Ha a PIK16 teszt tételeinek skálabesorolását is figyelembe vesszük, feltételezve, hogy a 16 tétel a négy skálának megfelelő, négy egymással korreláló látens faktorra illeszkedik, akkor a 6.2. ábra modelljéhez, egy *elsőrendű többfaktoros faktormodell*hez jutunk. Ilyen modelleket szoktak standard CFA elemzésekkel megvizsgálni (tesztelni és az illeszkedés jóságát mérni). A 6.2. ábrán a négy látens faktortól (Mmonit, AVhat, Rezil és Önreg) a saját tételekhez vezető egyirányú nyilak fölötti számok itt is egy konkrét faktormodell standardizált faktorsúly becslései, melyeket a fent említett minta adatain végzett elsőrendű CFA elemzés során kaptunk. A látens faktorokat összekötő kétirányú nyilak fölötti számok a látens faktorok közötti páronkénti korrelációk (a standardizált kovarianciák) becslései.



6.2. ábra. A CFA grafikonja négy elsőrendű látens faktor (Mmonit, AVhat, Rezil és Önreg) esetén

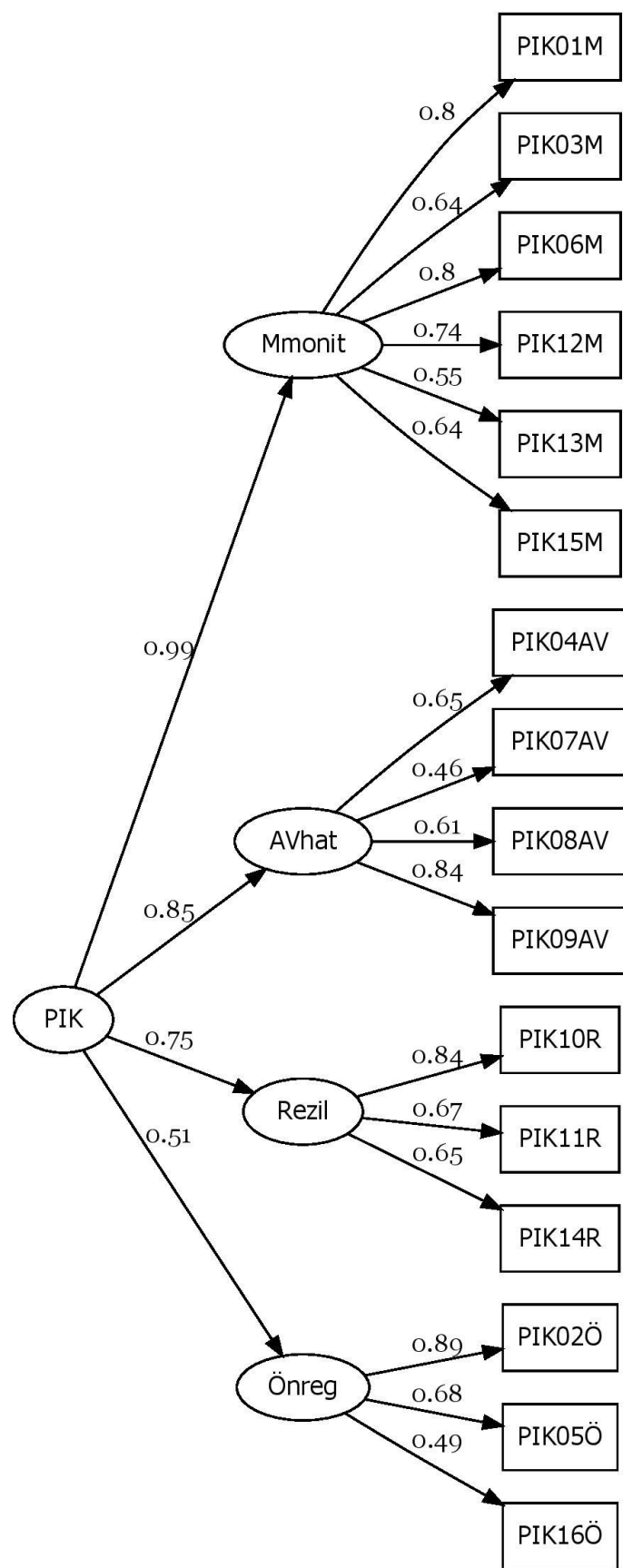
A 6.2. ábrán sem látunk nyilakat a tételek között, mert alapértelmezés szerint ebben a modellben is feltételezzük, hogy minden, ami a tételekben közös, az benne van a négy egymással kisebb-nagyobb mértékben korreláló látens faktorban, a tételek reziduálisai pedig egymástól függetlenek. Nyilakat a látens faktorok és a nem saját tételek között sem látunk. Ez annak a feltételezésnek felel meg, hogy minden tétel csak a saját skálájának megfelelő látens faktorra illeszkedik, a többi faktorra nem korrelál, vagyis nincsenek keresztöltések. Indokolt esetben ezeken a megszorításokon enyhíteni szoktak.

6.1.2. A másodrendű faktormodell

Ha feltételezzük, hogy az elsőrendű látens faktorok egy közös másodrendű látens faktorra illeszkednek, vagyis hogy ezek egy általuk képviselt konstruktum különböző megnyilvánulási formái, akkor egy *másodrendű faktormodell*hez³³ jutunk. A modell értelmes voltához elvárás, hogy az elsőrendű faktorok száma legalább három legyen, de ha bizonyos korlátozásokat teszünk³⁴, a modell két faktor esetén is becsülhető. A másodrendű faktor képviseli az elsőrendű faktorok közös lényegét, ami nem minden elsőrendű struktúra esetén létezik.

³³ Szokták a hierarchikus faktormodell elnevezést is használni.

³⁴ Például úgy, hogy 1-re állítjuk be az elsőrendű látens faktorok varianciáját.



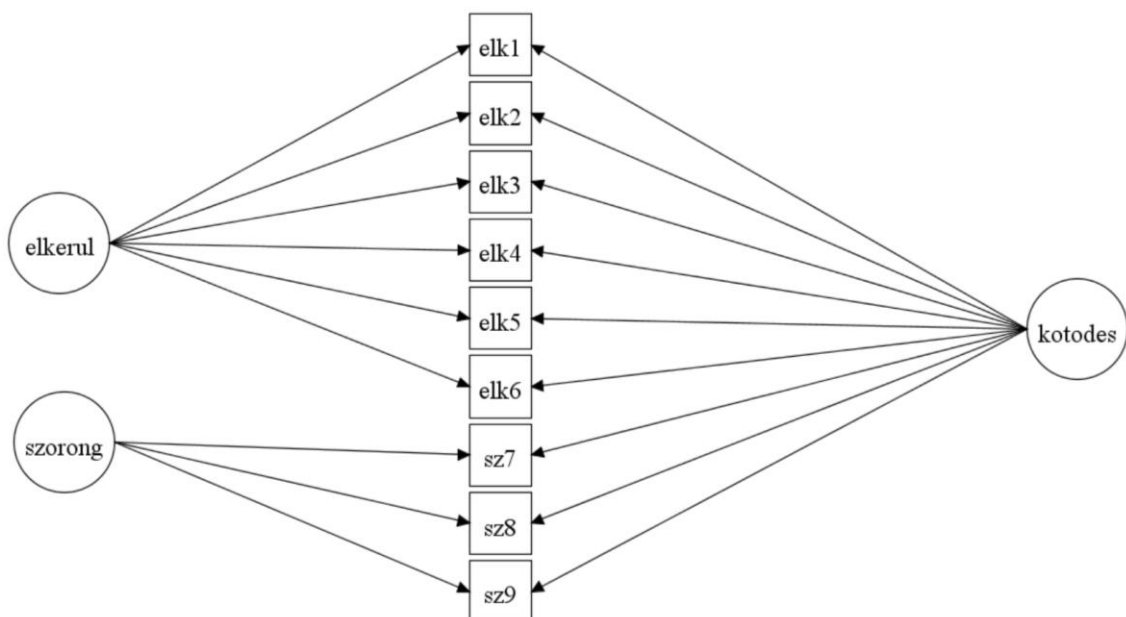
6.3. ábra. A CFA grafikonja négy elsőrendű (Mmonit, AVhat, Rezil és Önreg) és egy másodrendű látens faktor (PIK) esetén

Például a Globális Jólét Kérdőív esetében jogos lehet az a gondolat, hogy a teszt négy skálája (Érzelmi jólét, Pszichológiai jólét, Szociális jólét, Spirituális jólét) által képviselt négy faktor egy magasabb rendű konstrukció, a globális jólét négy megnyilvánulási formája, ezért itt a másodrendű faktormodell feltételezése szakmailag értelmes. Ugyanakkor a Big Five modell öt fő faktora (Extraverzió, Barátságosság, Lelkiismeretesség, Érzelmi stabilitás, Kultúra/Intellektus) esetében értelmes másodrendű faktor lehetősége fel sem merül.

A PIK16 esetén is elképzelhető, hogy a skálák által képviselt négy látens faktor egy pszichológiai immunkompetenciát képviselő másodrendű konstruktumra illeszkedik. Ez a másodrendű modell látható a 6.3. ábrán, ahol a PIK másodrendű látens faktor magyarázza a négy elsőrendű látens faktor (Mmonit, AVhat, Rezil és Önreg) mindegyikét. Ez utóbbiak között azért nem látunk nyilakat, mert a modell szerint mindazt, ami bennük közös, a PIK másodrendű faktor tartalmazza. A 6.3. ábrán a nyilak fölötti számok itt is egy konkrét faktormodell standardizált faktorsúly becslései, amelyeket a fentebb már említett Btérkép2022 mintán végzett másodrendű CFA elemzés során kaptunk. Az ábrán a teszt 16 tétele nincs összekötve a PIK másodrendű faktorról, a modell szerint ugyanis a másodrendű faktor csak az elsőrendű faktorokon keresztül, azok közvetítésével hat az egyes tételekre. Végül az egyes tételek sincsenek egymással összekötve, mert ami bennük közös, azt saját elsőrendű faktoruk képviseli, reziduálisai pedig függetlenek egymástól. A másodrendű modellben az elsőrendű faktorokat a másodrendű faktor alfaktorainak vagy *facetjeinek* tekintjük.

6.1.3. A bifaktoros faktormodell

A bifaktoros modell feltételez egy általános faktort, amelyre minden tétel illeszkedik és több specifikus faktort, amelyek rendre a tételek egy-egy alcsoportját magyarázzák (Dunn & McCray, 2020). Ennek a modellnek a sémáját mutatja be a 6.4. ábra, ahol középen az ECR-RS teszt (vö. B2.2. alpont) 9 tétele látható (az első hat elkerülés, az utolsó három pedig szorongás tétel), bal oldalon az elkerülés (elkerul) és a szorongás (szorong) *specifikus faktora*, a jobb oldalon pedig a kötődés (kotodes) *általános faktora*.



6.4. ábra. A bifaktoros modell sematikus ábrája az ECR-RS kötődésteszt 9 tételével (*elkerul* és *szorong*: specifikus faktorok, *kotodes*: általános faktor)

Az általános faktor fejezi a tételek közös lényegét, míg a specifikus faktorok a skálák egyediségét, amit az általános faktor nem tartalmaz. Alapértelmezés szerint a specifikus faktorok nem korrelálnak sem egymással, sem az általános faktorrall.

Ezt a modellt valószínűsíti, ha a tesztételek FKA elemzése során az első főkomponens sajátértéke legalább háromszorosa az utána következőnek (jelezve egy általános faktor létezését), és ha a tételeknek létezik egy olyan csoportosítása, amelynél az egyes csoportok, mint tesztiskálák létezését alkalmas EFA elemzéssel alá tudjuk támasztani (forgatás után kellően magas faktorsúlyokkal). A „bifaktoros” jelző arra utal, hogy minden tétel két faktorra illeszkedik. Egyrészt a közös általános faktorra, másrészt egy-egy specifikus faktorra (a specifikus faktorok valamelyikére). A másodrendű és a bifaktoros modell közötti legfőbb különbségek az alábbiak.

1. A tesztételek oldaláról abban térnek el egymástól, hogy az a konstruktum, ami az összes tétel közös lényegét képviseli, a másodrendű modellben hierarchikus módon, az elsőrendű faktorok közvetítésével, mint azok közös faktora kapcsolódik a tesztételekhez, a bifaktoros modell esetén pedig közvetlenül, direkt módon.

2. Az elsődleges faktorok szemszögéből pedig az a nagy különbség a két modell között, hogy a másodrendű modellben a másodrendű faktort, mely az elsőrendű faktorokon keresztül a tételek közös tartalmát is hordozza, az elsőrendű faktorok feszítik ki, tehát azokkal szoros kapcsolatban van, míg a bifaktoros modellben a specifikus faktorok és az általános faktor között nem tételezünk fel korrelációt. **Ez a legfontosabb különbség a kétféle modell között!** Matematikailag nem teljesen precízen megfogalmazva, a bifaktoros modell szerint minden tételnek van egy közös T tulajdonsága, s emellett a tételek egy részének A, másik alcsoportjának B, harmadiknak C tulajdonsága is van (pl. 3 alszála esetén), viszont sem A-nak, sem B-nek, sem C-nek nincs köze T-hez, függetlenek tőle. Egy ilyen tulajdonság lehet például egy olyan formai dolog is, hogy a tétel átfordítandó, negatív megfogalmazású. Ugyanakkor a másodrendű modell szerint a tételek egy része A, másik része B, harmadik része C tulajdonságú, és létezik egy T tulajdonság úgy, hogy A, B és C egyaránt implikálja T-t (de egyik sem azonos vele), tehát T egy közös, általánosabb tulajdonság. Ez esetben tehát az A, B, C tulajdonságok nem esnek kívül a fő konstruktumon (mint attól független specifikus faktorok), hanem alkotó részei annak.

3. A harmadik csupán technikai jellegű eltérés az, hogy a másodrendű modellben legalább három elsőrendű faktorra van szükség, hogy értelme legyen az elsőrendű faktorok közös másodrendű faktorának, míg a bifaktoros modell esetén akár két tételcsoport (skála) is elegendő. Mindamellet specális megkötésekkel (pl. 1-ben rögzítve a specifikus faktorok varianciáját) kétfaktoros másodrendű faktormodell is tesztelhető. Ha adott esetben szakmai szempontból mindkét modell értelmes, azt szoktuk elfogadni, amelyiknek a modellje jobb illeszkedést mutat az empirikus adatokhoz.

6.1.4. A CFA elemzés lényege, adekvációs mutatók

A CFA elemzés lényege, hogy a fentiek szerint felállítunk egy faktormodellt, majd egy megfelelően kiválasztott empirikus minta alapján – egy χ^2 -próbát végrehajtva – döntünk arról, hogy a modell érvényességét állító nullhipotézis elfogadható-e (modelltesztelés), s egyben arra alkalmas adekvációs mutatók segítségével megmérjük, hogy a modell mennyire illeszkedik empirikus adatainkra. Mindezt az adott faktormodell kiértékelésének nevezzük. Különböző szóba jöhető modellek közül azt preferáljuk, amelyik a modelltesztelés során a

legkevésbé bizonyul rossznak (a legkevésbé szignifikáns), s amelynek egyben az adekvációs (illeszkedési) mutatói a legjobbak.

6.1. táblázat. A CFA legfontosabb adekvációs mutatói

Mutató	Jelentés	Megfelelőség
RMSEA	Az input változók kovariancia-mátrixát és a CFA modelljéből becsült kovariancia-mátrixot hasonlítja össze. RMSEA az elemenkénti négyzetes eltérések átlaga, majd ennek az átlagnak a gyöke.	0,06 alatt jó, 0,06-0,08 között elfogadható, 0,08 fölött gyenge.
C90 (RMSEA)	90%-os intervallumbecslés az elméleti RMSEA értékre.	Akkor elfogadható, ha tartalmazza a 0,05-ös értéket, vagy ha az egész intervallum 0,08 alatt helyezkedik el.
pClose	Annak a próbának a p -értéke, mely azt teszteli, hogy RMSEA kisebb-e 0,05-nél: $\text{Prob}(\text{RMSEA} \leq 0,05)$.	Akkor jó, ha nagyobb 0,05-nél, vagyis ha nem szignifikáns a próba.
χ^2/f	Takarékossági mutató, melynél χ^2 azt az illeszkedést méri, ami RMSEA-ban is benne van. f a modell szabad paramétereinek a száma (a χ^2 statisztika szabadságfoka), így a hányados az egy paraméterre eső átlagos χ^2 komponenst adja meg.	Akkor jó, ha 3,5-nél kisebb, 3,5-5 között elfogadható.
SRMR	A modellből és a mintából kapott korrelációs mátrix elemei közti átlagos négyzetes eltérés négyzetgyöke. Minél közelebb van a nullához, annál jobban illeszkedik a faktormodell.	0,05 alatt jó, 0,08 alatt elfogadható az illeszkedés.
RMSR	Ugyanaz, mint SRMR, csak más névvel.	Lásd SRMR.
CFI	Egy relatív mutató, mely azt méri, hogy a modell tesztelésére alkalmazott χ^2 statisztika mennyivel kisebb (tehát a modell mennyivel jobb illeszkedésű), mint annak a lehető legrosszabb, ún. alapmodellnek a χ^2 -értéke, ahol a változók és a látens faktorok együttesében minden korreláció 0. Értéke 0 és 1 között lehet (Vargha, 2018, 121. o.).	0,95 felett kiváló, 0,90 felett elfogadható, 0,90 alatt gyenge.
TLI	CFI-hez hasonló módon definiált mutató, mely többnyire kisebb, mint CFI. Értéke legtöbbször 0 és 1 közé esik, ritkán azonban kicsivel 1 fölé is mehet (Hu és Bentler, 1999; Tucker és Lewis, 1973).	0,95 felett kiváló, 0,90 felett jó, 0,85 felett elfogadható, 0,85 alatt gyenge.
AIC, BIC	Akaike ³⁵ és Bayes ³⁶ -féle információs kritérium mutató, amelyekkel ML-típusú becsléseknél ugyanazon adatállomány több modellje hasonlítható össze.	A kisebb AIC, illetve BIC értékű modell a jobb.

³⁵ vö. Akaike, (1974) és https://en.wikipedia.org/wiki/Akaike_information_criterion

³⁶ vö. Schwarz (1978) és https://en.wikipedia.org/wiki/Bayesian_information_criterion

Tekintve, hogy a modellt tesztelő χ^2 -próba szignifikanciája erősen függ a minta nagyságától (pl. 1000 főnél nagyobb minták esetén szinte mindig szignifikáns), a modell kiértékelésében döntő szerepük van az adekvációs mutatóknak. A könnyebb áttekinthetőség érdekében a legfontosabb modell-adekvációs mutatókat külön is összefoglaltuk (lásd 6.1. táblázat), további mutatókkal (GFI, NFI stb.) kapcsolatban lásd Schreiber és munkatársai (2006), illetve Kenny (2020).

6.1.5. A CFA modellbecslési módszerei

A CFA-ban valamely faktormodell elfogadhatóságát szokásosan egy χ^2 -statisztikával teszteljük, illeszkedésének jóságát pedig különböző adekvációs mutatók segítségével mérjük (vö. 6.1. táblázat). A definiált faktormodell becslésére három alpmódszer áll rendelkezésre (Li, 2016; Shi és Maydeu-Olivares, 2020):

- ML = maximum likelihood módszer;
- ULS = unweighted least squares, vagyis a súlyozatlan legkisebb négyzetek módszere;
- DWLS = diagonally weighted least squares, vagyis a diagonálisan súlyozott legkisebb négyzetek módszere.

ML a strukturális egyenletmodellek, és ezen belül a CFA legszélesebb körben használt faktormodell becslési módszere. ML egyaránt alkalmas a modellilleszkedés tesztelésére és az illeszkedés jóságának a kiértékelésére. ML egyik hátránya, hogy alkalmazási feltétele a manifest változók többdimenziós normális eloszlása, ami a gyakorlatban gyakran nem teljesül. Ha ez a feltételezés sérül, akkor a modell eredmények nem feltétlenül megbízhatóak. A torzítás a becsült paraméterek standard hibáiban és a modelltesztelésre használt χ^2 -statisztikában jelenik meg (Li, 2014; Míndrilă, 2010; Schermelleh-Engel, Moosbrugger és Müller, 2003). A normalitás sérülését Curran, West és Finch (1995), illetve Míndrilă (2010) ajánlása nyomán akkor érdemes komolyan venni, ha a ferdeségi együttható 2-nél, a csúcossági együttható pedig 7-nél nagyobb abszolút értékű.

A normalitás súlyos sérülése esetén szokták a regressziós elemzésekből is ismert legkisebb négyzetek módszerét alkalmazni súlyozott (WLS) vagy súlyozatlan (ULS) formában, melynek során olyan modellt keresünk, amelyből a kiszámított korrelációs vagy kovariancia-mátrix a lehető legkisebb mértékben tér el az adatmintából kiszámított korrelációs vagy kovariancia-mátrixtól. A DWLS a WLS módszer olyan változata, amelynél Pearson helyett polichorikus korrelációt számítunk. Itt kategoriális változók esetén a korrelációt azon feltételezés mellett számítjuk ki, hogy a diszkrét értékek egy eredetileg folytonos normális eloszlású változó kerekített értékei és a polichorikus korreláció a feltételezett folytonos eloszlású változók között számított Pearson korreláció (Lee, Poon és Bentler, 1995).

Az ML, ULS, DWLS alpmódszerekre építve ROP-R-ben öt robusztus becslési variáns érhető el, amelyek jól lefedik a CFA elemzésekben használt módszerek széles spektrumát (Muthén, 1994; Li, 2021; Kyriazos és Poga-Kyriazou, 2023):

- MLMV: ML-becslés robusztus standard hibákkal, valamint korrigált átlagot és varianciát alkalmazó próbastatisztikával. Ez folytonosnak tekinthető, de erősen nem normális eloszlású változók esetén jó alternatívája ML-nek (vö. Gao, Shi és Maydeu-Olivares, 2019; Vargha et al., 2020).
- MLR: ML-becslés Huber-White robusztus standard hibákkal, valamint egy Yuan-Bentler-féle statisztikával aszimptotikusan megegyező tesztstatisztikával. Ez a

módszer a normalitás enyhe vagy közepes mértékű sérülése esetén ajánlható ML helyett folytonos változók esetén (Li, 2016).

- ULSMV: az ULS-becslés robusztus variánsa robusztus standard hibákkal, valamint korrigált átlagot és varianciát alkalmazó próbastatisztikával. Ez a módszer kategoriális változók esetén jó alternatívája lehet ML-nek (Savalei és Rhemtulla, 2013).
- WLSM: a DWLS-becslés robusztus variánsa robusztus standard hibákkal, valamint korrigált átlagot alkalmazó próbastatisztikával. Ez a módszer kategoriális változók esetén szintén jó alternatívája lehet ML-nek (Savalei és Rhemtulla, 2013), bár nem minden esetben (Navruz, 2016).
- WLSMV: a DWLS-becslés robusztus variánsa robusztus standard hibákkal, valamint korrigált átlagot és varianciát alkalmazó próbastatisztikával. Mint DWLS robusztus változata, kategoriális változók esetén ajánlható, de számos más nem normális eloszlástípus esetén is jobb becslésnek tűnik, mint MLR (Li, 2016; Holtmann, Koch, Lochner és Eid, 2016). Mplus futási tapasztalatok alapján többen megjegyzik³⁷, hogy azonos feladatban WLSM illeszkedése gyengébb a χ^2 -statisztika és RMSEA, de jobb CFI és TLI tekintetében. Muthén, DuToit és Spisic (1997) szerint elméletileg WLSMV preferálandó WLSM-mel szemben.

Fontos még, hogy a faktorúlmátrix ugyanúgy néz ki egy alaplómódszer és bármely robusztus variánsa esetén. Amiben különböznek: a faktorsúly-bebecslések standard hibái.

A CFA futtatása során ROP-R-ben mindig megkapjuk a kiválasztott robusztus módszer mellett a megfelelő ML, ULS vagy DWLS alapbecsléshez tartozó eredményeket is. Nehéz pontos eligazítást adni ahhoz, hogy konkrét esetben melyik módszert válasszuk a CFA-ban. A kiindulás az lehet, hogy normális eloszlású tételek esetén ML alapú, folytonos nem normális eloszlású tételek esetén ULS alapú, kategoriális (kétértékű vagy ordinális) változók esetén DWLS alapú becslés ajánlott. Jelenleg még hiányoznak az alapos kutatási eredmények az összes lehetséges módszer előnyeiről és hátrányairól. Érdemes ezért adott esetben több módszert is kipróbálni. Ami a legfontosabb: olyan modellt fogadjunk el, amely a kiértékelés során a lehető legjobb illeszkedést mutatja (vö. 6.1. táblázat) és szakmailag is értelmes.

6.1.6. A modifikációs indexek

Alapértelmezésben a faktorindexek segítségével definiált faktorstruktúra látens faktorai korrelálhatnak egymással, de az egyes faktorokat alkotó tételek reziduálisai (a faktor által meg nem magyarázott részei) nem – sem az azonos faktorba tartozó tételek, sem a különböző faktorokba tartozó tételek esetén. Ugyanígy alapértelmezésben nincs megengedve a keresztöltés sem, vagyis az, hogy egy tétel korreláljon valamely nem saját faktorra. Ez a legegyszerűbb faktorstruktúra azonban nem mindig állja meg a helyét. Néha meg kell engedni, hogy a fenti kapcsolatokat képviselő kovarianciák 0-tól különböző értéket is felvehessenek. Erre utaló jelzést a modifikációs indexek magas értékei adhatnak az elsődleges futás eredménylistáján.

Modifikációs indexe a CFA modell olyan paramétereinek van, amelyekre valamilyen korlátozást teszünk. Például az egy faktorba tartozó tételek reziduálisaira kikötjük, hogy nem korrelálnak (azaz a kovarianciájuk 0). Ugyanez a helyzet a különböző faktorokba tartozó tételekkel, valamint a keresztöltésekkel is. A modifikációs index megadja, hogy milyen

³⁷ lásd pl. <http://www.statmodel.com/discussion/messages/23/121.html?1469548078>

mértékben javul a modell (pontosabban milyen mértékben csökken a modellilleszkedést mérő χ^2 -érték), ha egy-egy ilyen korlátozást megszüntetünk, például ha megengedjük egy kovarianciára, hogy 0-tól különbözzön (MacCallum, Roznowski és Necowitz, 1992). A CFA modul menüablakában ennek alapján beállítható, hogy a faktorokon belüli vagy különböző faktorokba tartozó tételek közötti, illetve a tételek és a nem saját faktorok közötti kovarianciák milyen modifikációs küszöb felett legyenek beépítve egy javított faktormodellbe. Mivel egyetlen paraméter korlátozásának feloldásával a χ^2 -érték változása 1 szabadságfokú χ^2 -eloszlást követ, 3,841 feletti érték már 5%-os szinten szignifikáns. Ugyanakkor azt is meg kell vizsgálni, hogy a legnagyobb modifikációs index értékek kiemelkednek-e a többi közül. Ha ilyen értékek vannak, először csak 1–2 kovariancia bevonásával érdemes próbálkozni és azt is mérlegelni kell, hogy a szóban forgó tételek szakmai tartalma, jelentése feljogosít-e erre a bevonásra. A gyakorlatban – saját tapasztalataim alapján – ritkán érdemes 20 alatti modifikációs értékű kovarianciát bevonni a faktormodellbe.

6.2. A CFA menüablak

CFA elemzéseket ROP-R-ben a *Konfirmatív faktoranalízis* modul (röviden KonfirmFA vagy CFA) segítségével végezhetünk el. Ez a modul egyaránt alkalmas elsőrendű egy- és többfaktoros, másodrendű, valamint bifaktoros elemzések végrehajtására, amelyeket majd a 6.4. alfejezetben részletezünk, valódi pszichológiai adatokkal szemléltetve.

A technikai részleteikkel kapcsolatban megjegyezzük, hogy ROP-R CFA modulja a *lavaan* (Rosseel, 2012) és a *lavaanPlot* (Lishinski, 2021) R-package-re épül, mégpedig az Mplus szoftver (Muthén & Muthén, 1998-2017) CFA elemzéséhez igazított beállítással.

6.5. ábra. A CFA menüablak

A CFA menüablak szolgál többek között az elemzésbe bevonandó változók kiválasztására (input változók ablaka), skálahovatartozásuk jelzésére (Faktorindex ablak), a CFA modellbecslés módszerének megválasztására és a modifikációs index küszöbök beállítására (lásd 6.5. ábra). A „Másodrendű CFA” opció bejelölése esetén nem sima elsőrendű, hanem másodrendű CFA elemzést végez a program (feltéve, hogy a definiált faktorok száma legalább kettő). Ez esetben ROP-R elkészíti a bifaktoros CFA R-beli futtatására alkalmas scriptet is. A CFA menüablak speciális beállításával kapcsolatos tudnivalókat az alábbi alpontokban foglaljuk össze.

6.2.1. Az input változók CFA-ban

A CFA-hoz szükséges változók, a teszttételek kiválasztásához a „Változók” feliratú ablakból kell az elemzéshez szükséges tételeket (kijelölésük után akár egy csomagban) az „input változók” ablakba átküldeni, a két ablak közti nyílra rákattintva. Ezek a változók lesznek a CFA megfigyelt, manifeszt változói. Elméletileg elvárjuk tőlük, hogy folytonos, normális eloszlásúak legyenek, de manapság már léteznek a CFA-ban olyan robusztus módszerek (lásd 6.1.5. alpont), amelyek ezt a szigorú korlátozást enyhítik, s akár kétértékű nominális (pl. férfi - nő, vagy beteg - nem beteg) vagy ordinális változók elemzését is lehetővé teszik. A gyakorlatban a legtöbb kutatás az 5 vagy több kategóriát tartalmazó ordinális változókat folytonosként kezeli, és van némi bizonyíték arra, hogy ez az egyszerűsítés nem gyakorol nagy hatást az eredményekre (vö. Babakus, Ferguson és Jöreskog, 1987; Dolan, 1994; Johnson és Creech, 1983; Hutchinson és Olmos, 1998). A továbbiakban az egyszerűség kedvéért folytonos változóként hivatkozunk mi is az ilyen változókra, míg kategoriális jelzővel a kétértékű vagy a 3-4 értékű ordinális változókat illetjük. E terminológia a CFA módszerének kiválasztása során kap szerepet.

6.2.2. Skálahovatartozás megadása

A Faktorindex ablakban egy faktorindex segítségével jelölhető ki, hogy a kiválasztott változók (tételek) rendre mely faktorokhoz (skálákhoz) sorolandók. A faktorindex értéke 1 és 9 közötti egész szám lehet. Az azonos faktorindexű változók azonos faktorhoz tartoznak. Négy faktor kijelölésekor célszerű az 1–4 faktorindexeket használni, de ha rendre a 2, 3, 5, 7 indexértékeket használjuk a faktorok azonosítására, ROP-R akkor is értelmesen átkódolja ezeket (az indexek sorrendjét megőrizve) az 1–4 faktorindexekké.

6.2.3. Modifikációs index küszöbök megadása

A CFA standard faktormodelljében a változók reziduálisai sem faktoron belül, sem különböző faktorok között nem korrelálnak, vagyis kovarianciájuk nulla. Hasonlóképpen nincs korreláció a változók és az idegen faktorok között, vagyis keresztöltéseket nem feltételezünk. Ha először lefuttatunk egy standard CFA-t, akkor a modifikációs indexek jelzik, hogy ezen korlátozó feltételek oldásával milyen mértékben javulna a faktormodell illeszkedése. A CFA menüablak *Modifikációs index küszöb* panelje szolgál arra, hogy a faktormodellt javítsuk azon kovarianciák modellbe való beépítésével, amelyek modifikációs indexe az első CFA futás során nagyobb a megadott megfelelő küszöbnél. Külön küszöbök állíthatók be

- a faktorokon belüli reziduális kovarianciákra (alapértelmezés: 400),
- a különböző faktorok elemei közötti reziduális kovarianciákra (alapértelmezés: 500),
- a keresztöltések kovarianciáira (itt az alapértelmezés: 600).

ROP-R tehát a CFA modul futtatása során először végrehajt egy standard CFA elemzést a menüablakban megadott faktorbesorolásokkal a kijelölt modellbecslési módszerrel, majd ezen első futás eredménylistáján megnézi, hogy van-e olyan modifikációs index, amelyik nagyobb, mint a saját kategóriájára vonatkozó modifikációs index küszöb. Ha van, akkor annak kovarianciáját a faktorilleszkedés javítása céljából bevonja a faktormodellbe. A küszöbök alapértelmezés szerinti értékei (400, 500, 600) elég nagyok ahhoz, hogy az esetek túlnyomó többségében csak a standard CFA elemzés történjen meg, és ennek eredménylistáját áttanulmányozva a felhasználó mérlegelhesse, nem szeretné-e ezen küszöbök közül valamelyiket lejjebb szállítania, hogy egyes kovarianciák bekerülhessenek a modellbe. A kovarianciák modellbe való bevonásakor nem célszerű egyidejűleg 1-2-nél több kovarianciát bevonni és elvárjuk, hogy a bevont kovariancia a többihez képest legyen kiugró és szakmailag megmagyarázható. Például ilyen szakmai indok lehet két kérdőíves tételnél a nagyon hasonló megfogalmazás, ami arra inspirálhatja a tesztet kitöltőt, hogy a két tételre hasonló választ adjon. 20 alatti modifikációs indexű kovarianciát ritkán szoktunk bevonni a faktormodellbe.

6.2.4. Faktornevek megadása és a ROP-R által elkészített faktorábrák

CFA indítása után ROP-R rákérdez a kijelölt faktorok nevére, a tételnevek nem numerikus közös része alapján kínálva fel neveket az elsőrendű faktoroknak. Ha a menüablakban be van jelölve a „Másodrendű CFA” opció, akkor – legalább két faktor definiálása esetén – ROP-R rákérdez a másodrendű faktor nevére is. Ha az elsőrendű faktorok nevei tartalmazznak közös részt, akkor a program ezt a közös részt kínálja fel másodrendű faktornévnek. Ha az elsőrendű faktornevekben nincs közös rész, akkor az ajánlott másodrendű faktornév: F2rend. Persze nem kötelező a felkínált név, mód van bármilyen alfanumerikus faktornév megadására.

ROP-R-ben a CFA elemzés végrehajtása után a „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzésekhez elkészített ideiglenes adatfájlt (tmpdat.txt), a futtatott R-scripteket (sima modell: CFA.r, javított modell: CFA2.r, 1-nél nagyobb indexű feltételes csoport javított modellje: CFA3.r), a részletes R-eredménylistákat (oo.txt, o2.txt), valamint a faktormodellek elkészített diagramjait (pathplot1.pdf és pathplotR1.pdf). Ha feltételes csoportosító változót is használunk, akkor a feltételes csoportok összefűzött R-eredményeit a javítás nélküli modellekre az oo.txt, a javított modellekre pedig az ooR.txt fájl tartalmazza. A pathplotR1.pdf fájlban látható diagram a kovarianciákkal javított modell faktorábrája, s ez természetesen csak akkor jön létre, ha van olyan modifikációs index, amely nagyobb, mint a hozzá tartozó küszöb, s emiatt az ezzel javított modell CFA elemzése is megtörténik (ehhez tartozik a CFA2.r script-fájl, feltételes csoportok esetén pedig még a CFA3.r script-fájl is).

Ezek a pdf formátumban elmentett *faktordiagram* ábrák alapértelmezés szerint függőleges irányultságúak (ilyen pl. a 6.1. és a 6.3. ábra), de a menüablak *Faktordiagram irányítása* paneljén be lehet állítani, hogy az irányítás vízszintes legyen (lásd pl. a 6.2. ábrát).

6.3. Mit tartalmaz a CFA eredménylistája?

A CFA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- Az input változók faktorindexek szerinti csoportosítása (tételek besorolása a definiált skálákba)
- Megadott modifikációs küszöbök

- A kovarianciákra vonatkozó modifikációs indexek
- A megadott faktormodell tesztelése
- A minden változó és faktor függetlenségét képviselő alapmodell tesztelése
- Illeszkedési mutatók a vizsgált módszerekre/modellekre
- Standardizált faktorsúlyok és kommunalitások (elsőrendű CFA esetén)
- A másodrendű faktormodell standardizált regressziós együttható becslései (másodrendű CFA esetén)
- A látens faktorok páronkénti standardizált kovariancia becslései (csak több faktor esetén)
- A kovarianciákra vonatkozó legnagyobb modifikációs indexek a javított modellben (ha van javított modell)
- Standardizált reziduális kovariancia (korreláció) becslések a javított modellben (ha van javított modell)

6.4. A CFA szemléltetése valódi adatokon

Az alfejezet alábbi alpontjaiban a CFA végrehajtását szemléltetjük különböző faktormodellek kijelölésével, valódi pszichológiai kutatások valódi adataival. Ha másodrendű CFA-t kérünk, akkor két kijelölt faktor esetén ROP-R 1-ben rögzíti a specifikus faktorok kezdeti varianciáját a több iterációs lépésből álló modellillesztés indulásakor, hogy a CFA érvényesen végrehajtható legyen.

6.4.1. A Caprara-féle Pozitivitás skála egydimenziós faktorszerkezetének tesztelése

Első elemzésként az élethez való pozitív hozzáállást és általános pozitív attitűdöt mérő Caprara-féle Pozitivitás Skála 8 tételén végzünk CFA-t az alapértelmezés szerinti MLMV robusztus módszerrel az 1003 fős Btérkép2022 minta alapján (vö. B2.3. alpont). Az elemzés célja, hogy teszteljük a 8 tétel egyfaktoros struktúráját. Ezzel kapcsolatban már kaptunk pozitív eredményeket FKA végrehajtásakor (lásd 4.4.2. alpont), a CFA azonban statisztikailag precízebb módszer egy faktormodell tesztelésére és az illeszkedés jóságának a mérésére.

FKA-ból felhasználhatjuk azt a tudást, hogy a minta 7 személye outliernek tekinthető (lásd 4.5. táblázat), akiket jobbnak láttuk kihagyni a CFA elemzésből (az FKA-ban létrehozott és elmentett Outli változó feltételes csoportosító változóként való kijelölésével), így CFA-ban egy csökkentett, 996 fős mintával dolgoztunk.

Az eredménylistán először kovarianciákra vonatkozó modifikációs indexeket tekintjük meg, amelyek közül a 10 legnagyobb a 6.2. táblázatban látható. Eszerint az 1. és a 4., valamint a 7. és a 8. tétel kovarianciájának modellbe való bevétele javítaná számottevően a faktormodell illeszkedését. Ez a kérdések tartalmát tekintve szakmailag is elfogadható (vö. B.5. táblázat). Ezt figyelembe véve olyan új CFA elemzést végeztünk, amelynél a faktorokon belüli reziduális kovarianciák küszöbét 400-ról 100-ra csökkentettük, így az új (javított) faktormodellt két kovariancia bevétele javítja. A modifikációs indexek listája a javítás után annyiban módosul, hogy a listában ROP-R „+” jellel meg is jelöli a javított modellbe bevont kovarianciák modifikációs indexét.

6.2. táblázat. A 10 legnagyobb modifikációs index

Változó1	Változó2	Mod.ind.
Poz1Opt	Poz4Opt	202,49
Poz7Önb	Poz8Önb	162,12
Poz1Opt	Poz5Önb	55,04
Poz5Önb	Poz8Önb	43,35
Poz2Elég	Poz5Önb	33,51
Poz1Opt	Poz8Önb	26,54
Poz4Opt	Poz8Önb	25,19
Poz2Elég	Poz4Opt	23,61
Poz4Opt	Poz7Önb	19,87
Poz2Elég	Poz3Elég	15,86

E lista után következik a modellillesztés χ^2 -próbas tesztelésére vonatkozó eredmények táblázata (lásd 6.3. táblázat). Ebben megtaláljuk a modellt tesztelő χ^2 -statisztika értékét, annak szabadságfokát (f), valamint a szignifikancia mértékét jelző p -értéket. A szabadságfok alatti sor a χ^2/f takarékosági index értékét tartalmazza (vö. 6.1. táblázat). A táblázat négy oszlopa megfelel az eredeti modell sima ML és robusztus MLMV módszerrel, valamint a bevont kovarianciákkal javított ML (jele: ML+) és MLMV (jele: MLMV+) módszerrel történő tesztelésének.

6.3. táblázat. A kijelölt elsőrendű CFA faktormodell tesztelése

Mutató	ML	MLMV	ML+	MLMV+
χ^2	456,35	310,8	133,52	95,32
f	20	20	18	18
χ^2/f	22,82	15,54	7,42	5,30
p -érték	< 0,001	< 0,001	< 0,001	< 0,001

Jelen esetben azt láthatjuk, hogy a χ^2 -érték jelentősen csökken ugyan a modell javításakor (pl. ML/456,35-ről MLMV+/95,32-re), de még mindig olyan magas, hogy a modell $p < 0,001$ szinten elvethető. Ez az erős szignifikancia a nagy elemszámnak ($N = 996$) is köszönhető. Fontos ezért modellünk jóságát más mutatók segítségével is mérlegre tenni. A 6.3. táblázatban a χ^2/f takarékosági index (= 5,30) az MLMV+ modell esetében már majdnem elfogadható szintű illeszkedést jelez.

A 6.1. táblázat többi mutatójának értékét a ROP-R outputján egy külön táblázat tartalmazza, a 6.3. táblázatban feltüntetett négy modell x módszer kombinációra (lásd 6.4. táblázat). Ebből elsőként azt állapíthatjuk meg, hogy a sima ML-nél mindig jobb illeszkedést produkál a robusztus MLMV módszer, és hogy a két kovariancia bevonásával mindig érezhetően javulnak az illeszkedési mutatók. Ennek következtében a legjobb mutatók az MLMV+ modellt jellemzik, CFI, TLI és SRMR esetében kiváló szinten. RMSEA esetében a 0,080 alatti 0,066-os érték még elfogadható, az elméleti RMSEA értékre vonatkozó $C90 = (0,053; 0,079)$ konfidencia-intervallummal együtt. A pClose p -érték (0,021) is már csak $p < 0,05$ szintű szignifikanciát jelez.

6.4. táblázat. Illeszkedési mutatók a vizsgált modell-módszer kombinációkra

Mutató	ML	MLMV	ML+	MLMV+
AIC	20360,5	20360,5	20041,4	20041,4
BIC	20478,2	20478,2	20168,9	20168,9
RMSEA	0,148	0,121	0,08	0,066
C90	(0,136; 0,160)	(0,109; 0,133)	(0,068; 0,093)	(0,053; 0,079)
pClose	< 0,001	< 0,001	< 0,001	0,021
CFI	0,934	0,940	0,982	0,984
TLI	0,907	0,916	0,973	0,975
SRMR	0,033	0,033	0,019	0,019

Ugyanezt a CFA elemzést az ML alapú MLR robusztus módszerrel, mint alternatív lehetőséggel is elvégeztük, az összehasonlítást segíti a kapott eredményekből elkészített 6.5. táblázat. Az eredmények nagyon hasonlóak, de SRMR kivételével minden mutató tekintetében kissé gyengébb illeszkedést tapasztalunk, mint az MLMV módszer esetében.

6.5. táblázat. A javított faktormodell illeszkedési mutatói két robusztus módszer esetén

Mutató	MLMV+	MLR+
χ^2/f	5,30	5,51
RMSEA	0,066	0,067
C90	(0,053; 0,079)	(0,056; 0,079)
pClose	0,021	0,005
CFI	0,984	0,981
TLI	0,975	0,971
SRMR	0,019	0,019

A CFA outputjának következő fontos eleme a standardizált faktorsúlyok és a kommunalítások táblázata. Ez a robusztus MLMV módszer használata esetén az eredeti (MLMV) és a javított (MLMV+) modellre a 6.6. táblázatban látható.

6.6. táblázat. Standardizált faktorsúlyok és a kommunalítások

Faktor/tétel	MLMV	MLMV	MLMV+	MLMV+
Pozitív	Faktorsúly	Kommunalitás	Faktorsúly	Kommunalitás
Poz1Opt	0,887	0,787	0,856	0,733
Poz2Elég	0,903	0,815	0,916	0,839
Poz3Elég	0,739	0,547	0,743	0,552
Poz4Opt	0,913	0,834	0,887	0,787
Poz5Önb	0,890	0,792	0,902	0,814
Poz6OptF	0,203	0,041	0,200	0,040
Poz7Önb	0,798	0,638	0,791	0,626
Poz8Önb	0,748	0,560	0,741	0,548

A 6.6. táblázatból megállapíthatjuk, hogy a tételek – egyetlen kivétellel – a javított modell minden faktorában igen magas 0,70 feletti faktorsúly és 0,50 feletti kommunalitás becsléssel jellemezhetők, ami az egyfaktoros struktúra konstrukciós validitását erősíti. A kivétel: a negatív megfogalmazású 6. tétel (Időnként a jövő nem tűnik számomra teljesen egyértelműnek).

Logikusan merül fel, hogy ha Poz6OptF ennyire rosszul illeszkedik a látens faktorra, akkor talán ki is lehetne hagyni a skálából. Érdekes módon az eredmények nem javulnak. Ugyanazt a két kovarianciát építve be a modellbe, és ismét MLMV modellbecslést alkalmazva, $\chi^2/f = 6,08$, RMSEA = 0,071, C90 = (0,056; 0,088), pClose = 0,11, CFI = 0,987, TLI = 0,977, SRMR = 0,016, ami RMSEA tekintetében romlás, más mutatók esetében gyakorlatilag stagnálás. Ezért úgy látjuk, hogy nincs sem statisztikai, sem szakmai oka a Poz6OptF tétel elhagyásának.

A fentebb taglalt eredményeken kívül az is érdekes lehet, hogy mekkorák a modifikációs indexek a javított modellben a (lásd 6.7. táblázat). A 6.7. táblázat alapján ezek közül csak a legnagyobb (a Poz5Önb és a Poz8Önb közötti) érdemel figyelmet. Mivel a két tétel (Poz5Önb: *Összességében elégedett vagyok magammal*, ill. Poz8Önb: *Általában magabiztosnak érzem magamat*) nem tűnik egymás szinonímájának, s formailag sem hasonlítanak egymásra, így nem látjuk indokoltnak, hogy megengedjük reziduálisuk korrelációját a modellben. Ez is azt mutatja, hogy érdemes volt az első futtatás során kapott modifikációs indexek közül csupán az első kettőnek megfelelő kovarianciát bevonni a faktormodellbe.

6.7. táblázat. A 8 legnagyobb modifikációs index a javított modellben

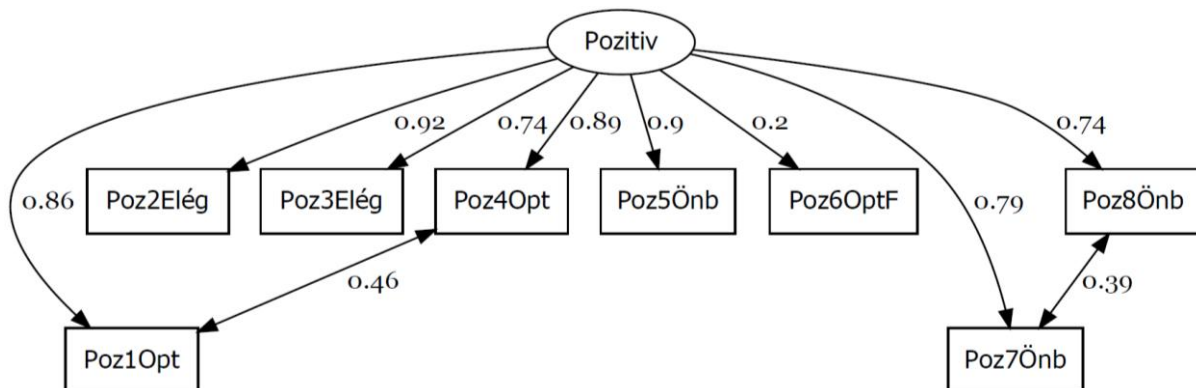
Változó1	Változó2	Mod.ind.
Poz5Önb	Poz8Önb	58,03
Poz1Opt	Poz2Elég	24,64
Poz3Elég	Poz5Önb	17,21
Poz1Opt	Poz5Önb	17,17
Poz2Elég	Poz8Önb	16,98
Poz2Elég	Poz4Opt	14,94
Poz3Elég	Poz4Opt	11,34
Poz3Elég	Poz6OptF	10,63

Végül érdemes megnézni az eredménylista utolsó táblázatát is, mely a modellbe bevont kovarianciák értékeit adja meg a javított modellben. A 6.8. táblázat alapján ezek se nem túlságosan nagyok (tehát nincs redundancia), se nem túlságosan kicsik (tehát nem feleslegesek), vagyis pont megfelelőek.

6.8. táblázat. Standardizált reziduális kovariancia (korreláció) becslések az MLMV+ modellre

Vált_1	Vált_2	Korreláció
Poz1Opt	Poz4Opt	0,461
Poz7Önb	Poz8Önb	0,395

A végső – javított – egydimenziós faktormodell diagramja a 6.6. ábrán látható. Ezt ROP-R a CFA elemzés során (a *Path diagram irányítása* panelen a „Vízszintes” opciót választva) a „c:_vargha\ropstat\aktualis” mappában menti el, pathplotR1.pdf néven. A nyilak melletti számok a faktormodell standardizált regressziós együtthatói (standardizált faktorsúlyok, illetve kovarianciák; vö. 6.6. és 6.8. táblázat).

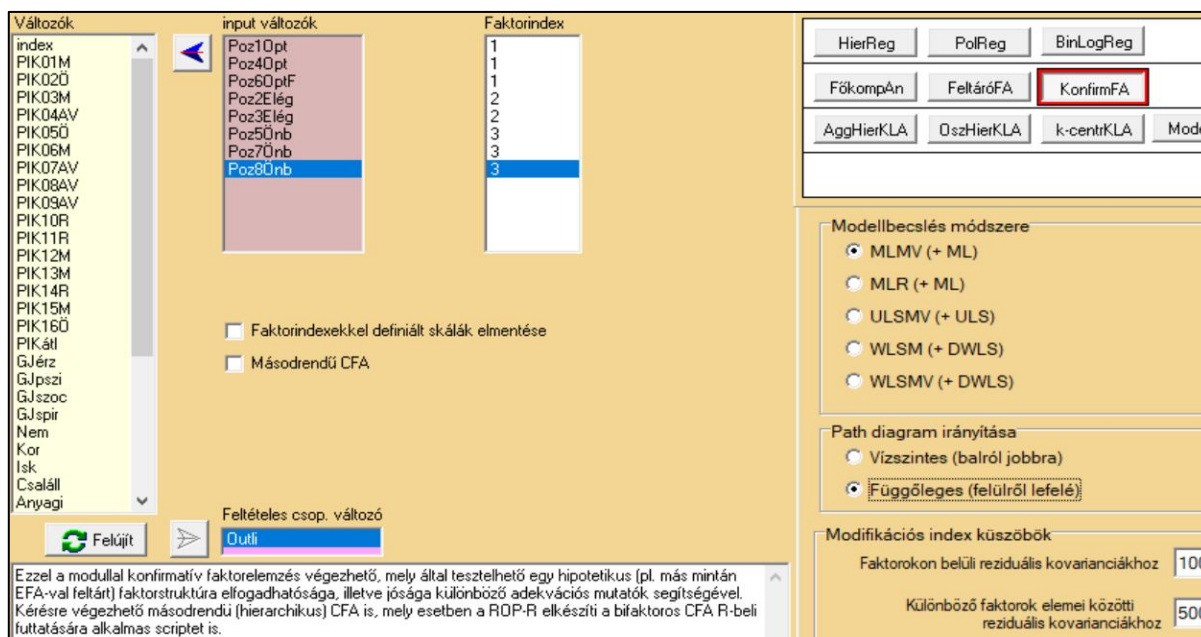


6.6. ábra. A Caprara-féle Pozitivitás kérdőív egyfaktoros diagramja (MLMV modellbecslés)

Az elvégzett egyfaktoros elsőrendű CFA elemzések alapján azt mondhatjuk, hogy a Caprara-féle Pozitivitás skála összpontszámának szerkezeti validitása összességében elfogadható, néhány illeszkedési mutató (pl. CFI, TLI, SRMR) tekintetében pedig egyenesen kiváló. Három alskálájának (Optimizmus: Optim; Élettel való elégedettség: Élelég; Önbecsülés: Önbecs; vö. B5. táblázat) megfelelőségét az alábbi 6.4.2. alponthban vizsgáljuk.

6.4.2. A Caprara-féle Pozitivitás skála háromskálás faktorszerkezetének tesztelése

A Caprara-féle Pozitivitás skála három alskálájának megfelelőségét egy elsőrendű háromfaktoros CFA-val vizsgáljuk meg ugyanazon az outlierekkel csökkentett 996 fős Btérkép2022 mintán, mint az előző alponthban, ismét MLMV robusztus modellbecslési módszert használva.



6.7. ábra. A CFA menüablak beállításai háromfaktoros elsőrendű CFA esetén

A CFA menüablakában mindössze annyi plusz feladatunk van az előző elemzéshez viszonyítva, hogy a három alskálának megfelelő faktorindexeket be kell állítanunk a „Faktorindex” ablakban (lásd 6.7. ábra, illetve B5. táblázat).

Az eredménylistán először ismét a kovarianciákra vonatkozó modifikációs indexeket tekintjük meg, amelyek közül a 10 legnagyobb a 6.9. táblázatban látható. Ez a táblázat eggyel több oszloppal rendelkezik, mint a 6.2. és a 6.7. táblázat, mert 1-nél több faktor esetén nemcsak az egyazon faktorba tartozó tételek reziduálisai között jelentkezhethet a kovariancia igénye, hanem különböző faktorokba tartozó tételek reziduálisai, valamint a tételek és idegen faktorok között (keresztöltés esete) is. Mivel a lehető legkevesebb kovarianciát szeretnénk a CFA modelljébe bevenni, csak a faktorokon belüli modifikációs küszöböt csökkentettük le 100-ra (lásd 6.7. ábra), s ennek megfelelően a táblázat csupán egyetlen kovariancia modellbe való bevitelét jelzi (+ jellel a Poz7Önb és Poz8Önb tétel reziduálisai között), melyet egyébként a 6.4.1. alpont elemzése során is bevontunk a faktormodellbe (lásd pl. a 6.6. ábrát).

6.9. táblázat. A 10 legnagyobb modifikációs index

Változó1	Változó2	Mod.ind.	Kovariancia típusa
Poz7Önb	Poz8Önb	111,12+	itemek egyazon faktorban
ÉlElég	Poz5Önb	110,23	kapcsolat egy idegen faktorrall
ÉlElég	Poz8Önb	80,09	kapcsolat egy idegen faktorrall
Poz5Önb	Poz7Önb	78,29	itemek egyazon faktorban
Poz2Elég	Poz5Önb	32,67	itemek különböző faktorokban
Optim	Poz8Önb	30,79	kapcsolat egy idegen faktorrall
Optim	Poz5Önb	30,58	kapcsolat egy idegen faktorrall
Poz4Opt	Poz2Elég	28,10	itemek különböző faktorokban
Poz2Elég	Poz8Önb	19,84	itemek különböző faktorokban
Poz1Opt	Poz2Elég	16,50	itemek különböző faktorokban

E lista után következik a modellillesztés χ^2 -próbás tesztelésére vonatkozó eredmények táblázata (lásd 6.10. táblázat). Ebben megtaláljuk a modellt tesztelő χ^2 -statisztika értékét, annak szabadságfokát (f), a χ^2/f takarékosági index értékét, valamint a szignifikancia mértékét jelző p -értéket. A táblázat négy oszlopa megfelel az eredeti modell sima ML és robusztus MLMV módszerrel, valamint a bevont kovarianciákkal javított ML+ és MLMV+ módszerrel történő tesztelésének. A táblázat még mindig azt jelzi, hogy a faktormodell illeszkedése nem jó ($p < 0,001$ szinten szignifikánsan rossz), de például a χ^2/f takarékosági index értéke már 5 alatt van, ezért eszerint a modell elfogadható és jobb, mint a 6.3. táblázatban kiértékelt modell.

6.10. táblázat. A kijelölt elsőrendű CFA faktormodell tesztelése

Mutató	ML	MLMV	ML+	MLMV+
χ^2	205,93	145,29	99,48	72,38
f	17	17	16	16
χ^2/f	12,11	8,55	6,22	4,52
p -érték	< 0,001	< 0,001	< 0,001	< 0,001

A többi illeszkedési mutató értékét a 6.11. táblázatban foglaltuk össze. Ebből elsőként azt állapíthatjuk meg, hogy az elsőrendű háromfaktoros modell illeszkedése minden tekintetben jobb, mint az elsőrendű egyfaktorosé (vö. 6.4. táblázat). Itt most az MLMV+ modell esetében CFI, TLI és SRMR értéke kiváló szintű, RMSEA esetében a 0,060-as érték jó, ugyancsak az elméleti RMSEA értékre vonatkozó és 0,050-et tartalmazó C90 = (0,046; 0,074) konfidencia-intervallum, és az is jó, hogy a pClose érték már nem szignifikáns ($p = 0,119$). Mindez azt jelenti, hogy az Optim, Élelég, Önbecs alskálák definiálása a kérdőív 8 tételéből javítja a faktorstruktúrát, s egyben lehetőséget ad arra, hogy értelmes szakmai tartalommal bővítsük a Pozitivitás kérdőív értelmezési spektrumát.

6.11. táblázat. Illeszkedési mutatók a vizsgált modell-módszer kombinációkra

Mutató	ML	MLMV	ML+	MLMV+
AIC	20115,8	20115,8	20011,3	20011,3
BIC	20248,2	20248,2	20148,6	20148,6
RMSEA	0,106	0,087	0,072	0,060
C90	(0,093; 0,119)	(0,074; 0,100)	(0,059; 0,086)	(0,046; 0,074)
pClose	< 0,001	< 0,001	0,003	0,119
CFI	0,971	0,974	0,987	0,988
TLI	0,953	0,957	0,978	0,980
SRMR	0,023	0,023	0,016	0,016

Megjegyezzük, hogy ugyanezt a CFA elemzést az MLR robusztus módszerrel, mint alternatív lehetőséggel is elvégeztük és egyetlen mutató tekintetében sem tapasztaltunk jobb illeszkedést (a legtöbb esetben egy picit gyengébbet).

6.12. táblázat. Standardizált faktorsúlyok és a kommunalítások

Faktor/tétel	MLMV	MLMV	MLMV+	MLMV+
Optim	Faktorsúly	Kommunalitás	Faktorsúly	Kommunalitás
Poz1Opt	0,918	0,843	0,917	0,841
Poz4Opt	0,947	0,897	0,948	0,898
Poz6OptF	0,210	0,044	0,210	0,044
Élelég	Faktorsúly	Kommunalitás	Faktorsúly	Kommunalitás
Poz2Elég	0,925	0,855	0,927	0,859
Poz3Elég	0,751	0,563	0,749	0,561
Önbecs	Faktorsúly	Kommunalitás	Faktorsúly	Kommunalitás
Poz5Önb	0,918	0,843	0,927	0,859
Poz7Önb	0,825	0,681	0,798	0,637
Poz8Önb	0,793	0,629	0,762	0,581

A standardizált faktorsúlyok és a kommunalítások táblázata a robusztus MLMV módszer használata esetén az eredeti (MLMV) és a javított (MLMV+) modellre a 6.12.

táblázatban látható. Ebből megállapíthatjuk, hogy a tételek – egyetlen kivétellel – a javított modell minden faktorában igen magas 0,70 feletti factorsúly és 0,60 feletti kommunalitás becsléssel jellemezhetők, ami a háromfaktoros struktúra konstrukciós validitását erősíti. A kivétel itt is a fordított megfogalmazású 6. tétel.

Az elsőrendű többfaktoros CFA faktormodellje megengedi, hogy a látens faktorok korreláljanak egymással. Hogy milyen szoros kapcsolat van a faktorok között, arról az output súlymátrix alatti táblázata informál (lásd 6.13. táblázat). Ebből megállapíthatjuk, hogy a robusztus MLMV becslést alkalmazva igen erős korrelációkat láthatunk a faktorok között mind az eredeti, mind a javított modell esetén. A leggyengébb korreláció (Optim és Önbecs között) is már 0,90 körüli, a legnagyobb (Élelég és Önbecs között) pedig 0,95 fölötti, ami meglehetősen magas.

6.13. táblázat. A látens faktorok páronkénti korreláció becslései

F(i)	F(j)	MLMV	MLMV+
Optim	Élelég	0,934	0,932
Optim	Önbecs	0,896	0,904
Élelég	Önbecs	0,945	0,955

A fentebb taglalt eredményeken kívül az is érdekes lehet, hogy mekkorák a modifikációs indexek a javított modellben a (lásd 6.14. táblázat). A 6.14. táblázat alapján azt mondhatjuk, hogy ezek jelentősen kisebbek, mint a javítás előtti modell legnagyobb modifikációs indexei (lásd 6.9. táblázat) és nem utalnak arra, hogy újabb kovarianciát be kéne vonni a modellbe.

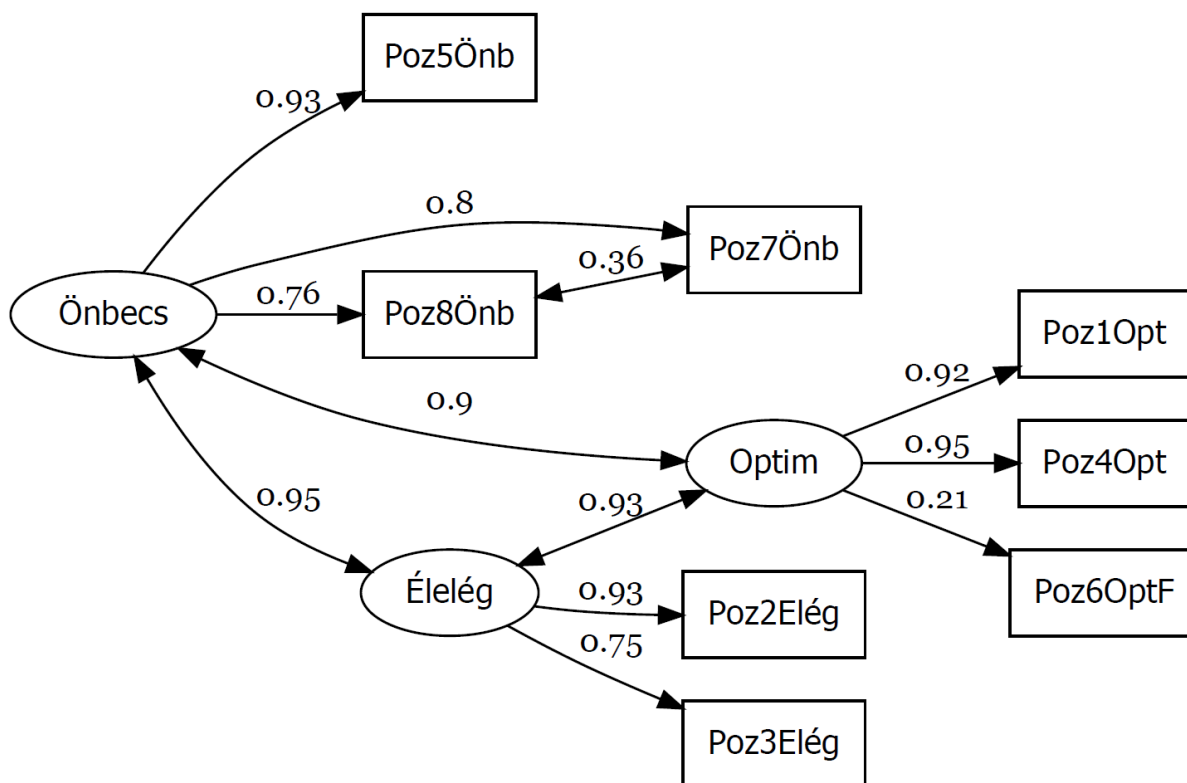
6.14. táblázat. A 8 legnagyobb modifikációs index a javított modellben

Változó1	Változó2	Mod.ind.
Élelég	Poz8Önb	34,93
Poz5Önb	Poz8Önb	33,12
Poz5Önb	Poz7Önb	33,12
Élelég	Poz7Önb	32,66
Poz4Opt	Poz2Elég	31,92
Poz1Opt	Poz2Elég	18,07
Poz1Opt	Poz5Önb	14,05
Poz4Opt	Poz3Elég	13,72

A végső – javított – háromfaktoros modell diagramja a 6.8. ábrán látható. Ezt ismét a pathplotR1.pdf nevű fájlban találjuk. Az ábráról leolvashatók a javított faktormodell standardizált regressziós együtthatói, a standardizált factorsúlyok (vö. 6.12. táblázat) és az egyetlen bevont – standardizált – kovariancia (Poz7Önb és Poz8Önb között).

Az elvégzett elsőrendű háromfaktoros CFA elemzés alapján kijelenthető, hogy a Caprara-féle Pozitivitás kérdőív három alskálás struktúrájának validitása összességében jónak, néhány illeszkedési mutató (pl. CFI, TLI, SRMR) tekintetében pedig egyenesen kiválónak mondható. A 4.4.2. és a 6.4.1. alpont eredményeit is figyelembe véve azt mondhatjuk, hogy a

teszt mind az összpontszáma (Pozitiv), mind a három alskálája (Optimizmus, Élettel való elégedettség és Önbecsülés) tekintetében használható, ha pszichológiai érvényességüket más vizsgálatokkal alá tudják támasztani.



6.8. ábra. A Caprara-féle Pozitivitás kérdőív háromskálás faktorábrája (MLMV modellbecslés)

6.4.3. A PIK16 kérdőív négyskálás faktorszerkezetének tesztelése

Harmadik CFA illusztrációs példánkban visszatérünk az 5.4. alfejezetben EFA-val elemzett PIK16 Pszichológiai Immunrendszer Kérdőívhez, mely a pszichológiai immunrendszer állapotának és összetevőinek feltérképezésére lett megszerkesztve (Oláh 2005). Elsőrendű CFA-t végzünk az 1003 fős Btérkép2022 minta alapján (vö. B2.3. alpont), hogy megnézzük: a kérdőív 4 skálája (Megközelítő-monitorozó viselkedés = Mmonit, Alkotó-végrehajtó viselkedés = AVhat, Reziliencia = Rezil, Önreguláció = Önreg) megfelelően illeszkedik-e egy 4-faktoros struktúrára a B.6. táblázatban megadott tételbesorolásokkal.

Az 5.4. alfejezetben láttuk, hogy a mintában a 16 tesztétel alapján nincs egyértelműen azonosítható és elhagyandó outlier eset, ezért most is a teljes 1003 fős mintán végezzük az elemzéseket, szokásosan az MLMV becslési módszerrel. A négy skálába tartozást négy különböző skálaindexszel (1–4) jelöltük a menüablakban.

Az eredménylistán először ismét a kovarianciákra vonatkozó modifikációs indexeket tekintjük meg, amelyek közül a 10 legnagyobbat a 6.15. táblázatban helyeztük el. Mivel most is a lehető legkevesebb kovarianciát szeretnénk a CFA modelljébe bevonni, csak a faktorokon belüli kovarianciákból vesszük be a két legnagyobbat a modellbe (PIK07AV és PIK08AV, ill. PIK05Ö és PIK16Ö reziduálisai között, ami szakmailag is elfogadható; vö. B.6. táblázat), ehhez ennek a kategóriának a modifikációs küszöbét lecsökkentjük 40-re. Újra futtatva CFA-t

a javított modellel, a χ^2 -próbával kapcsolatos eredményeket a 6.16. táblázatban helyeztük el. A táblázat azt jelzi, hogy a faktormodell illeszkedése nem jó ($p < 0,001$ szinten szignifikánsan rossz), de például az MLMV+ modell χ^2/f takarékosági indexe már 5 alatt van, melynek alapján a modell akár el is fogadható.

6.15. táblázat. A 10 legnagyobb modifikációs index (PIK16, MLMV módszer)

Változó1	Változó2	Mod.ind.	Kovariancia típusa
PIK07AV	PIK08AV	64,07	itemek egyazon faktorban
Mmonit	PIK08AV	55,40	kapcsolat egy idegen faktorra
Rezil	PIK02Ö	51,09	kapcsolat egy idegen faktorra
PIK05Ö	PIK16Ö	47,63	itemek egyazon faktorban
PIKMM	PIK09AV	45,07	kapcsolat egy idegen faktorra
PIK03M	PIK04AV	45,03	itemek különböző faktorokban
PIK02Ö	PIK16Ö	39,33	itemek egyazon faktorban
Rezil	PIK05Ö	39,03	kapcsolat egy idegen faktorra
PIK04AV	PIK09AV	37,33	itemek egyazon faktorban
AVhat	PIK13M	36,01	kapcsolat egy idegen faktorra

6.16. táblázat. A kijelölt CFA faktormodell tesztelése ML alapú modellbecslésekkel

Mutató	ML	MLMV	ML+	MLMV+
χ^2	709,95	541,24	598,19	459,07
f	98	98	96	96
χ^2/f	7,24	5,52	6,23	4,78
p -érték	< 0,001	< 0,001	< 0,001	< 0,001

A többi illeszkedési mutató értékét a 6.17. táblázatban foglaltuk össze. Ebből azt állapíthatjuk meg, hogy az elsőrendű négyfaktoros modell illeszkedése számos tekintetben elfogadható (pl. RMSEA és C90 0,08 alatti, CFI, TLI > 0,90), de azért összességében az illeszkedés elég gyenge.

6.17. táblázat. Illeszkedési mutatók a vizsgált modell-módszer kombinációkra

Mutató	ML	MLMV	ML+	MLMV+
AIC	37487,4	37487,4	37379,5	37379,5
BIC	37752,6	37752,6	37654,5	37654,5
RMSEA	0,079	0,067	0,072	0,061
C90	(0,074; 0,084)	(0,062; 0,073)	(0,067; 0,078)	(0,056; 0,067)
pClose	< 0,001	< 0,001	< 0,001	< 0,001
CFI	0,909	0,912	0,925	0,928
TLI	0,888	0,892	0,906	0,910
SRMR	0,049	0,049	0,043	0,043

Mit tehetünk a modellbecslés javítása érdekében? Ha figyelembe vesszük, hogy a PIK16 tételei mindössze 4-értékűek, ezért kategoriális jellegűek, megpróbálkozhatunk az ML-alapú helyett egy polichorikus korrelációt alkalmazó, DWLS alapú modellbecsléssel, például WLSMV-vel (vö. 6.1.5. alpont). A kovarianciákra vonatkozó modifikációs indexek közül az 5 legnagyobbat a 6.18. táblázatban helyeztük el.

6.18. táblázat. Az 5 legnagyobb modifikációs index (PIK16, WLSMV módszer)

Változó1	Változó2	Mod.ind.	Kovariancia típusa
Rezil	PIK02Ö	26,05	kapcsolat egy idegen faktorra
PIK07AV	PIK08AV	25,70	itemek egyazon faktorban
PIK05Ö	PIK16Ö	24,05	itemek egyazon faktorban
Rezil	PIK07AV	22,25	kapcsolat egy idegen faktorra
Mmonit	PIK02Ö	21,19	kapcsolat egy idegen faktorra

A 6.18. táblázat modifikációs indexei sokkal kisebbek, mint amiket a 6.15. táblázatban láthattunk az MLMV módszer esetén, így most egyszerű modellre törekedve nem veszünk be a modellbe egyetlen kovarianciát sem. A χ^2 -próbával kapcsolatos eredményeket a 6.19. táblázatban helyeztük el. A táblázat azt jelzi, hogy a faktormodell illeszkedése most se jó ($p < 0,001$ szinten szignifikánsan rossz), de például a DWLS alapmódszer χ^2/f takarékosági indexe 3,5 alatt van, ami jó jel (vö. 6.1. táblázat). Az különösen érdekes, hogy a robusztus modellbecslést alkalmazva az illeszkedés nem javul, hanem érezhetően romlik.

6.19. táblázat. A kijelölt CFA faktormodell tesztelése DWLS alapú modellbecslésekkel

Mutató	DWLS	WLSMV
χ^2	265,42	501,01
f	98	98
χ^2/f	2,71	5,11
p -érték	$< 0,001$	$< 0,001$

A többi illeszkedési mutató értékét a 6.20. táblázatban foglaltuk össze. Ebből azt látjuk, hogy a DWLS alapmódszert alkalmazva minden illeszkedési mutató kiváló szintű (vö. 6.1. táblázat), ami a PIK16 kérdőív négyfaktoros struktúrájának validitását (az Mmonit, AVhat, Rezil, Önreg skálákkal) megerősíti.

6.20. táblázat. Illeszkedési mutatók DWLS alapú modellbecslések esetén

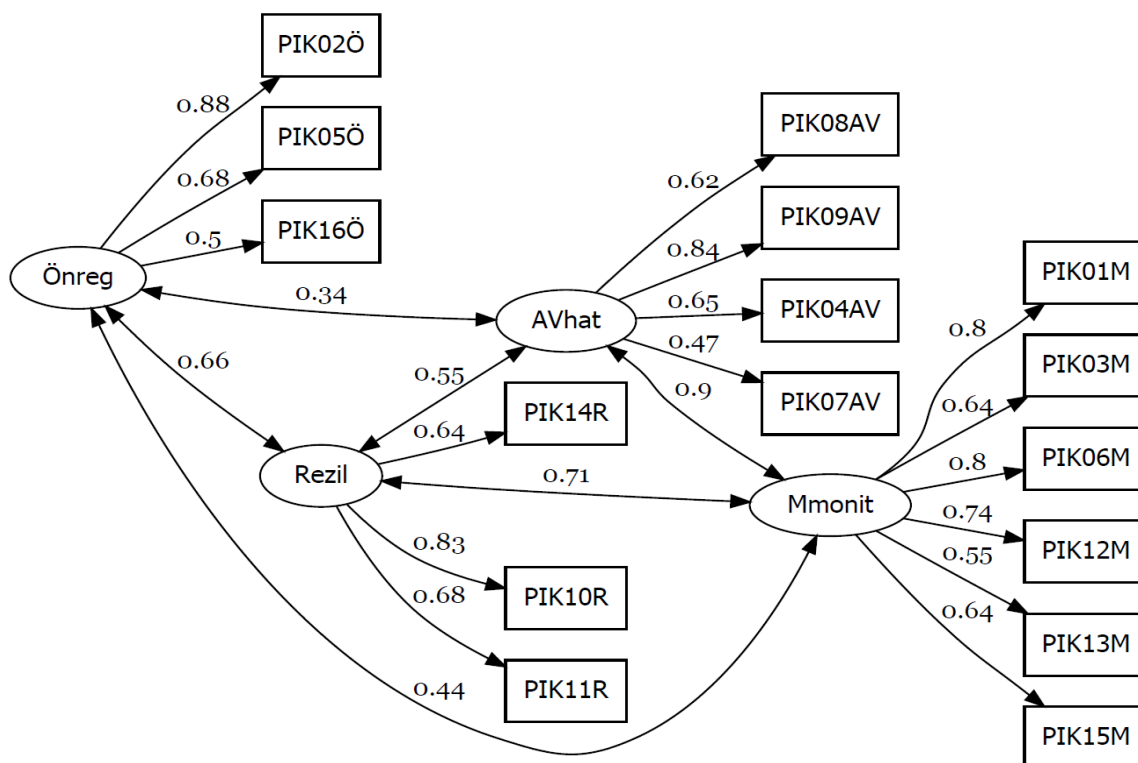
Mutató	DWLS	WLSMV
RMSEA	0,041	0,064
C90	(0,035; 0,047)	(0,059; 0,070)
pClose	0,992	$< 0,001$
CFI	0,988	0,914
TLI	0,985	0,894
SRMR	0,046	0,046

A standardizált faktorsúlyok és a kommunalítások táblázata a DWLS modellbecslési módszer használata esetén a 6.21. táblázatban látható. Ebből megállapíthatjuk, hogy a tételek minden faktor esetén elfogadható (egyetlen 0,47-es kivételtől eltekintve 0,50 feletti) faktorsúly és 0,20 feletti kommunalitás becsléssel jellemezhetők.

6.21. táblázat. Standardizált faktorsúlyok és a kommunalítások (kommun.)
DWLS modellbecslés esetén

Mmonit	Faktorsúly	Kommun.	AVhat	Faktorsúly	Kommun.
PIK01M	0,803	0,645	PIK04AV	0,652	0,425
PIK03M	0,644	0,415	PIK07AV	0,470	0,221
PIK06M	0,799	0,639	PIK08AV	0,615	0,378
PIK12M	0,741	0,550	PIK09AV	0,841	0,707
PIK13M	0,553	0,305			
PIK15M	0,639	0,408			
Önreg	Faktorsúly	Kommun.	Rezil	Faktorsúly	Kommun.
PIK02Ö	0,879	0,773	PIK10R	0,827	0,684
PIK05Ö	0,679	0,462	PIK11R	0,684	0,467
PIK16Ö	0,504	0,254	PIK14R	0,642	0,412

A végső négyfaktoros modell diagramja a 6.9. ábrán látható, ami csak irányításában különbözik a 6.6. ábrán láthatótól (ott függőleges, itt vízszintes). Mivel a modellt nem javítottuk kovarianciákkal, a grafikon a pathplot1.pdf nevű fájlban található. Az ábráról leolvashatók a javított faktormodell standardizált regressziós együtthatói és a standardizált faktorsúlyok (vö. 6.21. táblázat).



6.9. ábra. A PIK16 kérdőív négyeskálás faktorábrája (DWLS modellbecslés)

Az elsőrendű négyfaktoros faktormodell látens faktorai közti korrelációk leolvashatók a 6.9. ábráról, de külön is összefoglaltuk őket a 6.22. táblázatban. Ebből megállapíthatjuk, hogy a látens faktorok közül legszorosabb kapcsolatban Mmonit és AVhat van ($r = 0,900$), ami a kelletténél nagyobb redundanciára utal az Mmonit és az AVhat skála között. A többi faktor közötti korreláció mérsékelt (0,30–0,50 közötti) vagy magas (0,55–0,75 közötti).

6.22. táblázat. A látens faktorok páronkénti korreláció becslései

F(i)	F(j)	DWLS
Mmonit	AVhat	0,900
Mmonit	Rezil	0,710
Mmonit	Önreg	0,445
AVhat	Rezil	0,554
AVhat	Önreg	0,338
Rezil	Önreg	0,656

Az elvégzett elsőrendű négyfaktoros CFA elemzés alapján összefoglalóul azt mondhatjuk, hogy a PIK16 Rövidített Pszichológiai Immunrendszer Kérdőív négyskalás struktúrájának validitása minden tekintetben jónak, néhány illeszkedési mutató (pl. RMSEA, CFI, TLI) tekintetében pedig egyenesen kiválónak mondható, ugyanakkor elgondolkodtató az Mmonit és az AVhat skála között tapasztalt erős redundancia. Mindezek alapján a PIK16 tesztnek mind a négy skálája (Mmonit, AVhat, Rezil, Önreg) használható, ha pszichológiai érvényességüket alkalmas vizsgálatokkal alá tudják támasztani.

6.4.4. A PIK16 kérdőív másodrendű faktorszerkezetének tesztelése

Következő elemzésként másodrendű CFA elemzést végzünk a négyskalás PIK16 kérdőíven, egy újabb szinttel bonyolítva az előző alponthan megvizsgált elsőrendű faktormodell.

ROP-R-ben a „Másodrendű CFA” opciót kell bejelölni a CFA modul menüablakában ahhoz, hogy sima elsőrendű helyett egy másodrendű CFA-t végezhessünk. Indítás után a program itt is rákérdez a kijelölt faktorok nevére, s ugyanúgy kínál fel nevet az elsőrendű faktoroknak, mint az elsőrendű CFA esetén. Ha ezek a faktornevek tartalmaznak közös részt, akkor a program ezt a közös részt kínálja fel másodrendű faktornévnek. Ha az elsőrendű faktornevekben nincs közös rész, akkor az ajánlott másodrendű faktornév: „F2rend”. Persze mód van itt is bármilyen alfanumerikus faktornév megadására. Másodrendű CFA esetén is be lehet vonni kovarianciákat a modellbe, de csak az egyazon faktorba tartozó tételek közöttieket.

Jelen esemben ugyanazokat a skálakijelöléseket és beállításokat használva, mint az előző alponthan az elsőrendű négyfaktoros elemzés során (WLSMV módszer), az elsőrendű faktorok neveit a 6.21. táblázatban, illetve a 6.9. ábrán látható módon adtuk meg, a másodrendű faktornévre pedig az egyszerű PIK nevet.

Az eredménylistán ismét a modifikációs indexeket érdemes először is szemügyre venni, amelyek közül a 15 legnagyobbat a 6.23. táblázatban helyeztük el. Bár a ROP-R másodrendű CFA elemzésében csak az azonos faktorbeli kovarianciákat lehet bevonni a faktormodellbe, a program minden típusra kilistázza a modifikációs indexeket. Például az első két legnagyobb modifikációs index (mindkettő 217,92) arra hívja fel a figyelmet, hogy jelentősen javulna a másodrendű faktormodell illeszkedése, ha megengednénk, hogy az

Mmonit és az AVhat, illetve a Rezil és az Önreg elsőrendű faktor reziduálisai korreláljanak, a harmadik legnagyobb (109,51) pedig arra, hogy további jelentős javulást érjünk el, ha megengednénk, hogy a PIK11R tétel ne csak saját Rezil faktorán, hanem az Önreg faktoron is súlyozódjék. Érdekes, hogy az első faktoron belüli kovarianciával kapcsolatos modifikációs index csak a 15. helyen jelenik meg (a PIK05Ö és a PIK16Ö tétel között). Mivel ez az érték a többihez képest meglehetősen alacsony (27,10), egyszerű modellre törekedve most sem veszünk be a modellbe kovarianciát.

6.23. táblázat. A 15 legnagyobb modifikációs index (PIK16, WLSMV módszer)

Változó1	Változó2	Mod.ind.	Kovariancia típusa
Mmonit	AVhat	217,92	
Rezil	Önreg	217,92	
Önreg	PIK11R	109,51	
Rezil	PIK02Ö	105,68	
AVhat	Rezil	52,60	
Mmonit	Önreg	52,60	
Önreg	PIK10R	47,69	
PIK11R	PIK02Ö	41,56	
PIK11R	PIK05Ö	39,50	
Mmonit	Rezil	35,65	
AVhat	Önreg	35,65	
Rezil	PIK07AV	35,63	
PIK11R	PIK16Ö	28,96	
PIK	PIK02Ö	27,10	
PIK05Ö	PIK16Ö	27,10	itemek egyazon faktorban

A χ^2 -próbával kapcsolatos eredményeket a 6.24. táblázatban helyeztük el. A táblázat azt jelzi, hogy a faktormodell illeszkedése most se jó ($p < 0,001$ szinten szignifikánsan rossz), de például a DWLS alaplómódszer χ^2/f takarékosági indexe 5 alatt van, ami egy elfogadható modell jelzése (vö. 6.1. táblázat). A robusztus modellbecslést alkalmazva az illeszkedés az elsőrendű CFA-hoz hasonlóan most sem javul, hanem érezhetően romlik.

6.24. táblázat. A másodrendű faktormodell tesztelése
DWLS alapú modellbecslésekkel

Mutató	DWLS	WLSMV
χ^2	484,85	805,29
f	100	100
χ^2/f	4,85	8,05
p -érték	$< 0,001$	$< 0,001$

A többi illeszkedési mutató értékét a 6.25. táblázatban foglaltuk össze. Ebből azt látjuk, hogy a DWLS alaplómódszert alkalmazva minden illeszkedési mutató elfogadható (CFI

és TLI esetében kiváló) szintű, de azért érezhetően gyengébb, mint amit az elsőrendű CFA esetén kaptunk (vö. 6.1. és 6.20. táblázat).

6.25. táblázat. A PIK16 másodrendű faktormodelljének illeszkedési mutatói DWLS alapú modellbecslések esetén

Mutató	DWLS	WLSMV
RMSEA	0,062	0,084
C90	(0,057; 0,068)	(0,079; 0,089)
pClose	< 0,001	< 0,001
CFI	0,972	0,849
TLI	0,966	0,819
SRMR	0,061	0,061

6.26. táblázat. Standardizált faktorsúlyok és kommunalitások a DWLS modellbecslés esetén

Mmonit	Faktorsúly	Kommunalitás
PIK01M	0,804	0,646
PIK03M	0,642	0,412
PIK06M	0,799	0,638
PIK12M	0,745	0,555
PIK13M	0,550	0,303
PIK15M	0,638	0,408
AVhat	Faktorsúly	Kommunalitás
PIK04AV	0,654	0,428
PIK07AV	0,464	0,215
PIK08AV	0,615	0,378
PIK09AV	0,842	0,710
Rezil	Faktorsúly	Kommunalitás
PIK10R	0,836	0,699
PIK11R	0,672	0,451
PIK14R	0,646	0,417
Önreg	Faktorsúly	Kommunalitás
PIK02Ö	0,890	0,792
PIK05Ö	0,680	0,462
PIK16Ö	0,492	0,242
PIK	Regressziós eh.	Kommunalitás
Mmonit	0,989	0,978
AVhat	0,851	0,724
Rezil	0,749	0,560
Önreg	0,511	0,261

A faktorsúlyok (standardizált regressziós együtthatók) és a kommunalitások a DWLS modellbecslési módszer használata esetén a 6.26. táblázatban láthatók. Ebből

megállapíthatjuk, hogy a tételek ugyanolyan szinten illeszkednek az elsőrendű faktorokra, mint az elsőrendű CFA modelljében (vö. 6.21. táblázat). A táblázat legalsó blokkja pedig arról informál, hogy az elsőrendű faktorok közül az első három (Mmonit, AVhat és Rezil) magas, Önreg viszont még elfogadható, de náluk érezhetően alacsonyabb súllyal illeszkedik a PIK másodrendű faktorra. A másodrendű faktormodell diagramja a 6.3. ábrán látható, amelyen fel vannak tüntetve a 6.26. táblázatban összefoglalt faktorúlyok (standardizált regressziós együtthatók). Mivel a modellt nem javítottuk kovarianciákkal, a grafikont ismét a pathplot1.pdf nevű fájlban találjuk.

Összefoglalóul megállapíthatjuk, hogy a másodrendű faktormodell illeszkedése elfogadható, vagyis elfogadható egy olyan modell, amely szerint a PIK16 faktorstruktúrájában nemcsak a négy elsőrendű faktornak van helye, hanem a négy elsőrendű faktor még egy közös másodrendű faktorra is illeszkedik. Ebben a struktúrában az Önreguláció skáláját képviselő Önreg faktor illeszkedése a másodrendű faktorra kicsit gyengébb a többinél, ez az oka annak, hogy a másodrendű faktormodell illeszkedése kicsit gyengébbnek bizonyult, mind az elsőrendű modellé (vö. 6.20. és 6.25. táblázat). Mindazonáltal a kapott eredmények alátámasztják nemcsak a teszt négy skálájának, hanem egy összpontszámmal vagy a négy skála átlagával mért általános vagy globális mutató használatának a jogosságát is. Ez a globális PIK mutató a vizsgált személy pszichológiai immunrendszerének általános szintjét méri, vagyis azt, hogy pszichológiai immunrendszere milyen hatékonyan működik (Oláh, 2005).

6.4.5. A PIK16 kérdőív bifaktoros faktorszerkezetének tesztelése

Bifaktoros struktúra segítségével is validálható egy olyan faktormodell, amelyben egyaránt megtalálható egy általános faktor, valamint a négy skálának megfelelő egyedi faktor. De amíg a másodrendű struktúrában az elsőrendű egyedi faktorok inherens összetevői az általános faktort képviselő másodrendű faktornak, s így az egyes tételek ehhez az általános faktorhoz csak az elsőrendű faktorok közvetítésével kapcsolódnak, a bifaktoros modellben a tételek közvetlenül, direkt módon kapcsolódnak a közös lényegüket tartalmazó általános faktorhoz. Ebben a modellben a négy skálának megfelelő elsőrendű faktor, az ún. specifikus faktorok már csak azt tartalmazzák saját tételeik információjából, ami őket a többi skálától megkülönbözteti, s aminek semmi köze a 16 tétel közös tartalmához, közös jelentéséhez.

A bifaktoros CFA elemzést is az 1003 fős Btérkép2022 minta alapján (vö. B2.3. alpont) végezzük el a PIK16 tételein, ahogy tettük ezt a 6.4.3 és 6.4.4. alpontban. A bifaktoros elemzés kicsit komplikáltabb, mint az eddigiek, ugyanis ezt a ROP-R nem végzi el olyan direkt módon, ahogy az eddig bemutatott CFA elemzéseket.

A lényeg a következő. Ha a CFA modulban bejelöljük a „Másodrendű CFA” opciót, akkor létrejön az az R-script, amelyet az R szoftver RGui.exe keretprogramjában (lásd B1.2 alpont) manuálisan lefuttathatunk, s amelynek eredménylistájából kiolvashatjuk a bifaktoros faktormodell jóságára vonatkozó adatokat. Ha emellett a bejelölt „Másodrendű CFA” opció alatt a „Bifaktoros CFA végrehajtása” opciót is bejelöljük, akkor ROP-R maga is lefuttatja a létrehozott bifaktoros scriptet (fájlnév: CFAbif.r), melynek eredménylistája ugyanabban a „c:_vargha\ropstat\aktualis” mappában, egy CFAbif1.txt szövegfájlban kerül elmentésre, a modellhez tartozó faktordiagram pedig a bifplot1.pdf nevű pdf fájlban látható. Ha a javított modell is tesztelésre kerül, akkor ROP-R a scriptet a CFAbifR.r, az outputot a CFAbifR1.txt, a faktordiagramot pedig a bifplotR1.pdf nevű pdf fájlba menti.

A létrehozott script segítségével mi magunk is lefuttathatjuk a bifaktoros CFA-t. Ehhez az kell, hogy elindítsuk az RGui keretprogramot és a scriptben található utasításokat bemásoljuk az RGui konzoljába, majd Enter-rel futtassuk. Ekkor az eredmények a fentebb már említett fájlokban tekinthetők meg.

Jelen esetben a 6.4.4. alpontban leírt másodrendű CFA elemzést elvégezve (a WLSMV modellbecslési módszert választva) a 6.27. táblázatban látható script készült el.

6.27. táblázat. A bifaktoros CFA ROP-R által elkészített R-scriptje a 6.4.4. alpontban leírt másodrendű CFA elemzés beállításával (WLSMV modellbecslési módszert választva)

```
library(lavaan)
library(lavaanPlot)
dat<-read.delim('c:/_vargha/ropstat/aktualis/tmpdat1.txt', header=T)
model <- '
# Specific factors
Mmonit=~ PIK01M+PIK03M+PIK06M+PIK12M+PIK13M+PIK15M
AVhat=~ PIK04AV+PIK07AV+PIK08AV+PIK09AV
Rezil=~ PIK10R+PIK11R+PIK14R
Önreg=~ PIK02Ö+PIK05Ö+PIK16Ö
# General factor
ÁltFak=~ PIK01M+PIK03M+PIK06M+PIK12M+PIK13M+PIK15M+PIK04AV+
PIK07AV+PIK08AV+PIK09AV+PIK10R+PIK11R+PIK14R+PIK02Ö+
PIK05Ö+PIK16Ö
# Set zero correlations between the general factor and each specific factor
ÁltFak~~ 0*Mmonit
ÁltFak~~ 0*AVhat
ÁltFak~~ 0*Rezil
ÁltFak~~ 0*Önreg
# Set zero correlations among specific factors
Mmonit~~ 0*AVhat      # omit "0*" for allowing correlating factors
Mmonit~~ 0*Rezil      # omit "0*" for allowing correlating factors
Mmonit~~ 0*Önreg      # omit "0*" for allowing correlating factors
AVhat~~ 0*Rezil      # omit "0*" for allowing correlating factors
AVhat~~ 0*Önreg      # omit "0*" for allowing correlating factors
Rezil~~ 0*Önreg      # omit "0*" for allowing correlating factors
'

sink('c:/_vargha/ropstat/aktualis/CFAbif1.txt')
fit <- cfa(model, data=dat, mimic="Mplus", estimator = "WLSMV")
summary(fit, fit.measures = T, rsq = T, standardized = T)
modindices(fit, sort = TRUE, maximum.number = 50)
pl <- lavaanPlot(model = fit, graph_options = list(rankdir = "TD", fontsize = "10"),
node_options = list(shape = "box", fontname = "Helvetica"),
edge_options = list(color="black"), coefs=F, stand=T, covs=F, digits=2)
embed_plot_pdf(pl,'c:/_vargha/ropstat/aktualis/bifplot1.pdf')
sink()
```

Ebben a scriptben csupán egyetlen változtatást érdemes tenni, "WLSMV" helyett "DWLS" írandó (alulról a 8. sorban). Ezzel a változtatással a *lavaan* R-package a DWLS alapmódszerrel becsüli a kijelölt CFA faktormodellt.

A 6.27. táblázat R-scriptjének futtatása után a bifaktoros CFA modellre vonatkozó eredmények a CFAbif1.txt szövegfájlba íródtak (lásd az első „sink” névvel kezdődő sorba

beírt fájlnevet). Ezt az outputot értelmes tagolásokkal 8 részre daraboltuk (lásd 6.28-6.35. táblázat).

6.28. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 1. része (modelltesztelés χ^2 -próbával, relatív illeszkedési mutatók)

lavaan 0.6-12 ended normally after 174 iterations	
Estimator	DWLS
Optimization method	NLMINB
Number of model parameters	64
Number of observations	1003
Model Test User Model:	
Test statistic	401.181
Degrees of freedom	88
P-value (Chi-square)	0.000
Model Test Baseline Model:	
Test statistic	13745.462
Degrees of freedom	120
P-value	0.000
User Model versus Baseline Model:	
Comparative Fit Index (CFI)	0.977
Tucker-Lewis Index (TLI)	0.969

Az első, amit érdemes megnézni, a modell tesztelése, mely a „Model Test User Model” című rovatnév után található a 6.28. táblázatban. A χ^2 statisztika értéke a „Test Statistic” sorban olvasható (401.181). E sor alatt látható a szabadságfok (88), melynek segítségével a takarékosági index is kiszámítható:

$$\chi^2/f = 401,181/88 = 4,56,$$

ami picit jobb, mint a másodrendű modellé (4,85), de még mindig határozottan gyengébb, mint az elsőrendű négyfaktoros modellé (2,71).

A CFI és a TLI relatív illeszkedési mutató értéke a „Comparative Fit Index (CFI)” és a „Tucker-Lewis Index (TLI)” kezdetű sorból olvasható ki (lásd a 6.28. táblázat utolsó két sorát). Ezek ugyanúgy kiváló értékek, mint amit az elsőrendű (lásd 6.20. táblázat) és a másodrendű modell (lásd 6.25. táblázat) esetén láttunk.

Az eredménylista következő része az abszolút illeszkedési mutatókkal kapcsolatos eredményeket tartalmazza (lásd 6.29. táblázat). RMSEA (0,060), C90_{RMSEA} (0,054; 0,066) és SRMR (0,055) egyaránt 0,080 alatti, elfogadható illeszkedést jeleznek (lásd 6.1. táblázat). Ezek megint picit jobbak, mint a másodrendű modellé (lásd 6.25. táblázat), de érezhetően gyengébbek, mint az elsőrendű négyfaktoros modellé (lásd 6.20. táblázat).

A pClose érték a „P-value RMSEA <= 0.05” sorban látható (0,004), mely azt jelzi, hogy a 0,060-as RMSEA érték szignifikánsan nagyobb, mint a 0,050-es elméleti felső küszöb.

A 6.29. táblázat alján látunk még egy WRMR nevű abszolút illeszkedési mutatót, melyet ritkábban használnak a gyakorlatban, mint RMSEA-t vagy SRMR-t. Vele kapcsolatban elég annyit tudni, hogy az 1-nél kisebb WRMR-értékek jeleznek jó illeszkedést (DiStefano, Liu, Jiang és Shi, 2018), ami esetünkben láthatóan nem teljesül.

6.29. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 2. része (abszolút illeszkedési mutatók)

Root Mean Square Error of Approximation:	
RMSEA	0.060
90 Percent confidence interval - lower	0.054
90 Percent confidence interval - upper	0.066
P-value RMSEA <= 0.05	0.004
Standardized Root Mean Square Residual:	
SRMR	0.055
Weighted Root Mean Square Residual:	
WRMR	1.625

Igen fontos a „Latent Variables” rovatnév után található táblázat, mely a bifaktoros modell paramétereinek a regressziós becsléseit tartalmazza (lásd 6.30. táblázat). E táblázat fejlécében „Estimate” a nyers becslés, „Std.Err” a standard hiba, „z-value” a becslés szignifikanciáját tesztelő z-érték, „P(>|z|)” e teszteléshez tartozó p -érték, „Std.lv” a becslés standardizált látens faktorok esetén, „Std.all” pedig a becslés, ha minden változót standardizálunk, vagyis a standardizált faktorsúly becslés.

6.30. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 3. része (faktormodell paramétereinek regressziós becslései)

Latent Variables:						
	Estimate	Std.Err	z-value	P(> z)	Std.lv	Std.all
Mmonit =~						
PIK01M	1.000				0.208	0.219
PIK03M	0.434	0.406	1.071	0.284	0.090	0.095
PIK06M	0.917	0.638	1.438	0.151	0.191	0.199
PIK12M	-0.540	0.521	-1.036	0.300	-0.112	-0.111
PIK13M	-0.430	0.421	-1.022	0.307	-0.090	-0.106
PIK15M	-1.373	1.108	-1.240	0.215	-0.286	-0.305
AVhat =~						
PIK04AV	1.000				0.189	0.201
PIK07AV	1.878	0.523	3.591	0.000	0.355	0.451
PIK08AV	2.913	0.943	3.088	0.002	0.550	0.669
PIK09AV	1.212	0.351	3.458	0.001	0.229	0.267
Rezil =~						
PIK10R	1.000				0.569	0.529
PIK11R	0.702	0.152	4.616	0.000	0.399	0.410
PIK14R	0.897	0.209	4.288	0.000	0.510	0.510
Önreg =~						
PIK02Ö	1.000				0.499	0.520
PIK05Ö	1.583	0.245	6.472	0.000	0.790	0.795
PIK16Ö	1.000	0.120	8.365	0.000	0.499	0.507
ÁltFak =~						
PIK01M	1.000				0.769	0.807

PIK03M	0.790	0.026	30.392	0.000	0.608	0.640
PIK06M	0.999	0.032	31.569	0.000	0.768	0.799
PIK12M	0.980	0.033	29.283	0.000	0.754	0.745
PIK13M	0.604	0.024	25.288	0.000	0.465	0.549
PIK15M	0.796	0.032	24.861	0.000	0.612	0.653
PIK04AV	0.699	0.026	26.656	0.000	0.538	0.570
PIK07AV	0.375	0.020	19.085	0.000	0.289	0.367
PIK08AV	0.527	0.023	23.294	0.000	0.405	0.493
PIK09AV	0.811	0.028	29.435	0.000	0.623	0.725
PIK10R	0.874	0.031	28.340	0.000	0.672	0.624
PIK11R	0.640	0.026	24.926	0.000	0.492	0.505
PIK14R	0.613	0.026	23.634	0.000	0.471	0.471
PIK02Ö	0.592	0.024	24.764	0.000	0.455	0.474
PIK05Ö	0.427	0.022	19.516	0.000	0.328	0.330
PIK16Ö	0.298	0.020	14.704	0.000	0.229	0.233

A 6.30. táblázat arról informál, hogy milyen a tételek illeszkedése a specifikus faktorokra (Mmonit, AVhat, Rezil és Önreg), valamint az „ÁltFak” nevű általános faktorra. Ez utóbbi nevet a ROP-R adja, ha ezen változtatni akarunk, ahhoz ezt át kell írni a 6.27. táblázatban látható R-scriptben. A tételek illeszkedése a bifaktoros modellben akkor tekinthető jónak, ha a standardizált faktorsúlyok (lásd az utolsó, Std.all című oszlopot) mind nagyobbak 0,50-nél. Ez ÁltFak esetén bár 16-ból csak 10 esetben teljesül, de legalább mindegyik pozitív és két kivételtől (0,233 és 0,330) eltekintve 0,35-nél nagyobb. A specifikus faktorokra való illeszkedés már érezhetően gyengébb. Önreg, Rezil és AVhat esetében még elfogadható, de az Mmonit faktor esetében már egyértelműen nagyon gyenge az illeszkedés. Itt ugyanis a hat faktorsúly fele negatív, és minden faktorsúly abszolút értéke igen alacsony (a legnagyobb is csak 0,305).

Mindezeket figyelembe véve a kapott eredményeket úgy kell értelmeznünk, hogy a bifaktoros CFA az általános faktor létezését megerősíti, azt viszont nem jelenthetjük ki, hogy a tételek az általuk tartalmazott ezen közös konstruktumon túl rendelkeznének még olyan specifikus információkkal, ami létjogosultságot adna mind a négy tesztskála esetében egy specifikus faktornak. Bár se a másodrendű, se a bifaktoros CFA modellje nem volt kiváló illeszkedésű, mindkettő alátámasztja egy olyan konstrukció (látens faktor) létezését, melyre a PIK16 16 tétele vagy közvetve (másodrendű CFA) vagy direkt módon (bifaktoros CFA) illeszkedik, hitelesítve ezzel a teszt összpontszámának használatát. A négy skálának megfelelő elsőrendű faktor konstrukciója mind az első-, mind a másodrendű CFA eredményei alapján jónak bizonyult (vö. 6.4.3. és 6.4.4. alpont).

A fentiekben áttekintettük a bifaktoros CFA R-beli eredménylistájának legfontosabb elemeit. Az alábbiakban a ritkábban használt egyéb részokről is szólnunk néhány szót. A 6.31. táblázatban például „Covariances” fejléccel találjuk a látens faktorok közti kovarianciákat. Ezek most mind 0-k, mert így állítottuk be őket a modellben. Ez akkor nézne ki másképp, ha megengednénk, hogy a specifikus faktorok korreláljanak egymással, vagy ha a tételek közti kovarianciákkal javítanánk a faktormodellt.

A 6.32. táblázatban „Intercepts” fejléccel találjuk a modell összes változójának regressziós konstans értékét. Ezek a standardizált látens faktorok esetén természetesen 0-k, a teszt tételei esetében pedig mintabeli átlagukkal egyeznek meg.

6.31. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 4. része (kovarianciák)

Covariances:						
	Estimate	Std.Err	z-value	P(> z)	Std.lv	Std.all
Mmonit ~~						
ÁltFak	0.000				0.000	0.000
AVhat ~~						
ÁltFak	0.000				0.000	0.000
Rezil ~~						
ÁltFak	0.000				0.000	0.000
Önreg ~~						
ÁltFak	0.000				0.000	0.000
Mmonit ~~						
AVhat	0.000				0.000	0.000
Rezil	0.000				0.000	0.000
Önreg	0.000				0.000	0.000
AVhat ~~						
Rezil	0.000				0.000	0.000
Önreg	0.000				0.000	0.000
Rezil ~~						
Önreg	0.000				0.000	0.000

6.32. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 5. része (alapszintek/konstansok)

Intercepts:						
	Estimate	Std.Err	z-value	P(> z)	Std.lv	Std.all
.PIK01M	2.863	0.030	95.143	0.000	2.863	3.004
.PIK03M	2.820	0.030	93.993	0.000	2.820	2.968
.PIK06M	2.850	0.030	93.928	0.000	2.850	2.966
.PIK12M	2.775	0.032	86.880	0.000	2.775	2.743
.PIK13M	2.544	0.027	95.274	0.000	2.544	3.008
.PIK15M	2.614	0.030	88.360	0.000	2.614	2.790
.PIK04AV	2.759	0.030	92.724	0.000	2.759	2.928
.PIK07AV	3.269	0.025	131.695	0.000	3.269	4.158
.PIK08AV	3.147	0.026	121.194	0.000	3.147	3.827
.PIK09AV	2.950	0.027	108.757	0.000	2.950	3.434
.PIK10R	2.818	0.034	82.875	0.000	2.818	2.617
.PIK11R	2.831	0.031	92.006	0.000	2.831	2.905
.PIK14R	3.157	0.032	99.947	0.000	3.157	3.156
.PIK02Ö	2.720	0.030	89.687	0.000	2.720	2.832
.PIK05Ö	2.694	0.031	85.788	0.000	2.694	2.709
.PIK16Ö	2.941	0.031	94.597	0.000	2.941	2.987
Mmonit	0.000				0.000	0.000
AVhat	0.000				0.000	0.000
Rezil	0.000				0.000	0.000
Önreg	0.000				0.000	0.000
ÁltFak	0.000				0.000	0.000

A 6.33. táblázatban „Variances” fejléccel találjuk a tesztételek, vagyis a modell manifeszt változói esetében a reziduális, vagyis a látens faktorok által meg nem magyarázott varianciaarányt (lásd a táblázat utolsó oszlopát). Mivel egy manifeszt változót részben saját specifikus faktora, részben az általános faktor magyaráz, ezeket minden tétel esetén úgy kapjuk, hogy négyzetre emeljük a specifikus és az általános faktorra vonatkozó standardizált faktorsúlyát, összeadjuk őket (megmagyarázott rész), majd az összeget 1-ből kivonjuk (nem megmagyarázott rész).

6.33. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 6. része: varianciák

Variances:						
	Estimate	Std.Err	z-value	P(> z)	Std.lv	Std.all
.PIK01M	0.274	0.059	4.681	0.000	0.274	0.302
.PIK03M	0.525	0.039	13.437	0.000	0.525	0.582
.PIK06M	0.298	0.056	5.326	0.000	0.298	0.322
.PIK12M	0.443	0.046	9.527	0.000	0.443	0.433
.PIK13M	0.491	0.033	14.987	0.000	0.491	0.687
.PIK15M	0.422	0.088	4.769	0.000	0.422	0.481
.PIK04AV	0.563	0.036	15.683	0.000	0.563	0.634
.PIK07AV	0.409	0.051	7.948	0.000	0.409	0.661
.PIK08AV	0.209	0.104	2.007	0.045	0.209	0.309
.PIK09AV	0.297	0.037	8.021	0.000	0.297	0.403
.PIK10R	0.384	0.091	4.200	0.000	0.384	0.331
.PIK11R	0.548	0.052	10.471	0.000	0.548	0.577
.PIK14R	0.518	0.079	6.512	0.000	0.518	0.517
.PIK02Ö	0.466	0.055	8.551	0.000	0.466	0.505
.PIK05Ö	0.257	0.114	2.255	0.024	0.257	0.260
.PIK16Ö	0.668	0.055	12.060	0.000	0.668	0.688
Mmonit	0.043	0.042	1.029	0.303	1.000	1.000
AVhat	0.036	0.017	2.114	0.034	1.000	1.000
Rezil	0.324	0.087	3.725	0.000	1.000	1.000
Önreg	0.249	0.044	5.657	0.000	1.000	1.000
ÁltFak	0.591	0.027	22.119	0.000	1.000	1.000

Például a PIK01M tétel esetében

$$1 - (0,219^2 + 0,807^2) = 1 - (0,048 + 0,651) = 1 - 0,699 = 0,301,$$

ami a kerekítési hibától eltekintve megegyezik a 6.33. táblázat legfelső sorában látható 0,302 értékkel (Std.all oszlop).

A 6.34. táblázat a manifeszt változókra vonatkozóan tartalmazza a látens faktorok által megmagyarázott varianciaarányt, melyet értelemszerűen minden változó esetében úgy kapunk, hogy a 6.33. táblázat utolsó oszlopában látható értéket 1-ből kivonjuk (pl. $0,698 = 1 - 0,302$). Ezek a változók kommunalitásai. Ha egy tétel kommunalitása nagyon alacsony (mondjuk kisebb, mint 0,025), ez annak a jele, hogy a tétel kilóg a faktormodellből. Ilyen a jelen esetben szerencsére nem fordul elő.

6.34. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 7. része (R^2 -ek)

R-Square:	
	Estimate
PIK01M	0.698
PIK03M	0.418
PIK06M	0.678
PIK12M	0.567
PIK13M	0.313
PIK15M	0.519
PIK04AV	0.366
PIK07AV	0.339
PIK08AV	0.691
PIK09AV	0.597
PIK10R	0.669
PIK11R	0.423
PIK14R	0.483
PIK02Ö	0.495
PIK05Ö	0.740
PIK16Ö	0.312

Az R segítségével végrehajtott bifaktoros CFA eredménylistáján még megemlítendő az output végén található modifikációs indexek listája, melynek fejléce az „lhs op” szöveggel kezdődik (lásd 6.35. táblázat). Itt a modifikációs indexek (jelzete: mi) az első adatoszlopban láthatók. Konkrét mintánk esetében például az első sorban a Rezil és az Önreg faktor kapcsolatára vonatkozóan a modifikációs index értéke 186,992, mely arra utal, hogy ez a két specifikus faktor olyan szoros kapcsolatban van, hogy a modellben meg kellene engednünk, hogy korreláljanak egymással. Ehhez mindössze annyit kell tenni, hogy a 6.27. táblázatban látható script „Rezil~~ 0*Önreg” utasításában (ez most a „#Residual covariances” sor fölött található) töröljük az Önreg faktornév előtti „0*” részt, mely a két faktor közötti 0 korrelációt fixálja. Ha ezt az elemzést is elvégezzük, akkor az illeszkedési mutatók tekintetében kiváló faktormodellt láthatunk (pl. RMSEA = 0,033, pClose = 1, CFI = 0,993, TLI = 0,990, SRMR = 0,038), de az Mmonit látens faktor modellbeli varianciája negatív (-0,092) lesz, ami miatt a modell nem elfogadható.

6.35. táblázat. A 6.27. táblázat R-scriptjének futtatása során kapott output 8. része (modifikációs indexek)

	lhs	op	rhs	mi	epc	sepc.lv	sepc.all	sepc.nox
42	Rezil	~~	Önreg	186.992	0.166	0.585	0.585	0.585
131	Önreg	==	PIK11R	114.998	0.703	0.351	0.360	0.360
117	Rezil	==	PIK02Ö	79.776	0.695	0.396	0.412	0.412
119	Rezil	==	PIK16Ö	63.727	0.603	0.343	0.349	0.349
118	Rezil	==	PIK05Ö	49.994	0.549	0.312	0.314	0.314
130	Önreg	==	PIK10R	44.025	0.478	0.239	0.222	0.222
245	PIK11R	~~	PIK05Ö	42.868	0.221	0.221	0.590	0.590
244	PIK11R	~~	PIK02Ö	37.674	0.205	0.205	0.406	0.406
99	AVhat	==	PIK13M	36.473	1.260	0.238	0.281	0.281
246	PIK11R	~~	PIK16Ö	32.188	0.195	0.195	0.322	0.322
132	Önreg	==	PIK14R	31.051	0.377	0.188	0.188	0.188
40	AVhat	~~	Rezil	28.635	-0.035	-0.321	-0.321	-0.321
111	Rezil	==	PIK13M	25.556	-0.357	-0.203	-0.240	-0.240
240	PIK10R	~~	PIK02Ö	22.972	0.179	0.179	0.423	0.423
189	PIK13M	~~	PIK07AV	22.814	0.120	0.120	0.269	0.269
101	AVhat	==	PIK10R	20.706	-1.203	-0.227	-0.211	-0.211
114	Rezil	==	PIK07AV	19.895	-0.288	-0.164	-0.208	-0.208
152	PIK03M	~~	PIK04AV	18.343	0.138	0.138	0.253	0.253
38	Mmonit	~~	Rezil	17.379	0.053	0.451	0.451	0.451
219	PIK07AV	~~	PIK10R	16.669	-0.124	-0.124	-0.314	-0.314
242	PIK10R	~~	PIK16Ö	15.546	0.140	0.140	0.277	0.277
41	AVhat	~~	Önreg	15.542	-0.016	-0.166	-0.166	-0.166
100	AVhat	==	PIK15M	14.895	0.942	0.178	0.190	0.190
247	PIK14R	~~	PIK02Ö	14.780	0.137	0.137	0.278	0.278
249	PIK14R	~~	PIK16Ö	14.437	0.131	0.131	0.222	0.222
122	Önreg	==	PIK06M	14.046	-0.256	-0.128	-0.133	-0.133
93	Mmonit	==	PIK05Ö	13.859	-1.316	-0.274	-0.276	-0.276
120	Önreg	==	PIK01M	12.489	-0.238	-0.119	-0.125	-0.125
91	Mmonit	==	PIK14R	12.352	1.310	0.273	0.273	0.273
194	PIK13M	~~	PIK14R	10.788	-0.103	-0.103	-0.205	-0.205
173	PIK06M	~~	PIK05Ö	10.710	-0.115	-0.115	-0.415	-0.415
39	Mmonit	~~	Önreg	10.599	-0.025	-0.237	-0.237	-0.237
207	PIK15M	~~	PIK16Ö	10.066	-0.103	-0.103	-0.193	-0.193
190	PIK13M	~~	PIK08AV	9.975	0.082	0.082	0.256	0.256
129	Önreg	==	PIK09AV	9.807	-0.186	-0.093	-0.108	-0.108
241	PIK10R	~~	PIK05Ö	9.552	0.117	0.117	0.371	0.371
200	PIK15M	~~	PIK08AV	9.252	0.088	0.088	0.297	0.297
104	AVhat	==	PIK02Ö	8.797	-0.694	-0.131	-0.137	-0.137
112	Rezil	==	PIK15M	8.781	-0.247	-0.141	-0.150	-0.150

```

121  Önreg =~ PIK03M  8.375 -0.185 -0.093 -0.097 -0.097
105  AVhat =~ PIK05Ö  7.872 -0.636 -0.120 -0.121 -0.121
115  Rezil =~ PIK08AV  7.393 -0.190 -0.108 -0.132 -0.132
157  PIK03M ~~ PIK11R  6.541 -0.086 -0.086 -0.161 -0.161
89   Mmonit =~ PIK10R  6.515  1.049  0.219  0.203  0.203
192  PIK13M ~~ PIK10R  6.328 -0.085 -0.085 -0.196 -0.196
37   Mmonit ~~ AVhat  6.302 -0.010 -0.265 -0.265 -0.265
102  AVhat =~ PIK11R  6.265 -0.590 -0.112 -0.114 -0.114
220  PIK07AV ~~ PIK11R  6.079 -0.068 -0.068 -0.144 -0.144
248  PIK14R ~~ PIK05Ö  6.049  0.086  0.086  0.235  0.235
87   Mmonit =~ PIK08AV  5.567 -0.736 -0.153 -0.187 -0.187
[1] "c:/_vargha/ropstat/aktualis/bifplot1.pdf"
attr(,"class")
[1] "knit_image_paths" "knit_asis"

```

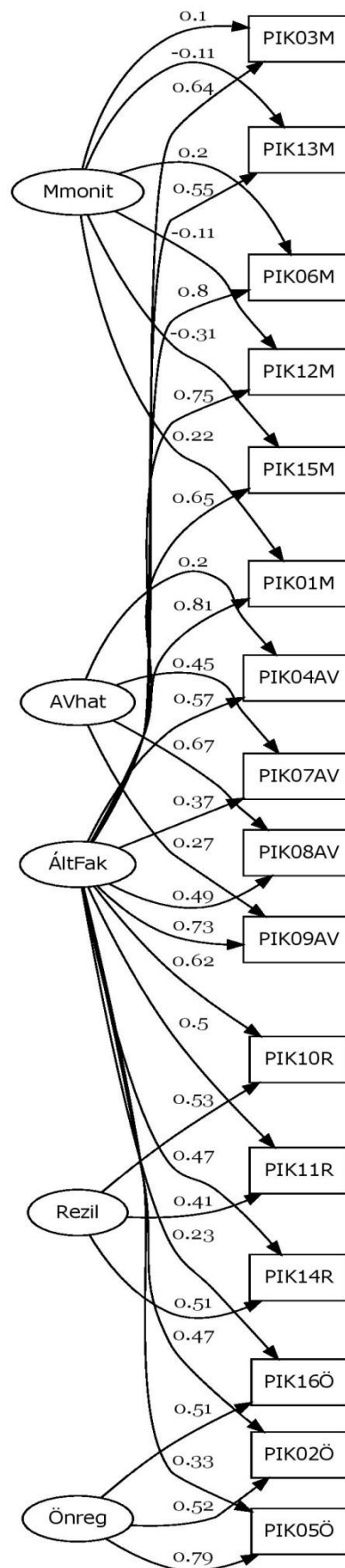
A szöveges eredménylista után érdemes megtekinteni a 6.27. táblázat scriptje által elkészített faktorábrát is (lásd 6.10. ábra). Ez ugyan lehetne kissé áttekinthetőbb is, olyan elrendezésű, mint például a 6.4. ábra (az általános faktort a tételek függőleges vonulatától nem balra, hanem jobbra helyezve el, hogy ne keresztezzék egymást a tételek specifikus faktorokhoz vezető, illetve az általános faktorhoz vezető nyilai), de ez a script csak erre alkalmas. A szabad mozgástér számunkra az ábra kinézetét illetően a script alján látható alábbi négy sor.

```

pl <- lavaanPlot(model = fit, graph_options = list(rankdir = "LR", fontsize = "10"),
node_options = list(shape = "box", fontname = "Helvetica"),
edge_options = list(color="black"), coefs=T, stand=T, covs=F, digits=2)
embed_plot_pdf(pl,'c:/_vargha/ropstat/aktualis/bifplot1.pdf')

```

1. Ha itt az első sorban LR helyett TD-t írunk, ezzel a faktorábra irányítását vízszintesről (LR = Left to Right = balról jobbra), függőlegesre (TD = Top to Down = felülről lefelé) változtatjuk.
2. A harmadik sorban coefs=T (hosszabban: coefs=TRUE) arra utasít, hogy a faktorsúlyok jelenjenek meg az ábrán. Sok item és sok faktor esetén ezek kissé áttekinthetetlené teszik az ábrát. Ha azt szeretnénk, hogy a faktorsúlyok ne jelenjenek meg az ábrán, a coefs=F (hosszabban: coefs=FALSE) paraméterezést kell alkalmaznunk.
3. Ugyancsak a 3. sorban covs=F arra utasít, hogy a faktormodell kovarianciái (a látens változók közötti vagy a modifikációs küszöbök alapján beépített kovarianciák) ne jelenjenek meg az ábrán. Ha azt akarjuk, hogy a kovarianciák értékei mégis megjelenjenek, akkor a covs=T (hosszabban: covs=TRUE) szöveget kell írunk. Azt azonban figyelembe kell venni, hogy a bifaktoros modell alapértelmezése szerint a látens változók nem korrelálnak, ezért kovarianciáik nem szerepelnek a modellben. Emiatt covs=T esetén a látens változókat összekötő nyilakon nem együtttható érték, hanem NA (not available) felirat lesz majd látható.
4. Ugyanebben a sorban stand=T (hosszabban: stand=TRUE) arra utasít, hogy a faktormodell regressziós együttthatói standardizált formában jelenjenek meg az ábrán. Persze ennek csak akkor van hatása, ha coefs=T és covs=T közül legalább az egyik szerepel a sorban. Ha a nem standardizált együttthatókat akarjuk megjelentetni, akkor a stand=F (hosszabban: stand=FALSE) paraméterezést kell alkalmaznunk.



6.10. ábra. A PIK16 kérdőív bifaktoros modelljének faktorábrája

5. Ugyanebben a sorban `digits=2` azt szabja meg, hogy az ábrán megjelenő számok hány tizedessel legyenek megjelenítve. Ezt a 2-es számot szabadon csökkenthetjük vagy növelhetjük.
6. Végül az utolsó sorban a `bifplot1.pdf` fájlnev módosításával tetszőleges nevű fájlba elmenthetjük a faktorábrát.

Mindezek a lehetőségek alkalmazhatók azoknak a faktormodelleknek a scriptjeire is, amelyeket a ROP-R más CFA elemzések során elkészít és elment (a `CFA.r`, `CFA2.r` nevű fájlokba).

III. RÉSZ

Klaszterező modulok

Az empirikus pszichológiai vizsgálatok célja, hogy feltárják a lelki élet törvényszerűségeit, a pszichológiai jellemzők közti összefüggéseket, okokat-okozatokat, csoportok közti különbségeket. A statisztika ebben úgy segít, hogy megmondja: hogyan kell a vizsgált pszichológiai populációkból megfelelő mintákat kiválasztani, érvényes módon mérni, adatokat feldolgozni, elemezni és alkalmas statisztikai módszerek segítségével helyes következtetéseket levonni. Bár ezekkel az elemzésekkel többnyire az emberekről szeretnénk többet megtudni, gyakran megrekedünk olyan mutatók (pl. átlag, korreláció stb.) vizsgálatának a szintjén, ami inkább a változókat, mintsem az egyéneket jellemzi.

Ha vizsgálatunk fókuszában változók közti kapcsolatok, összefüggésük mikéntje, rendszere áll, *változó-orientált* elemzésekről beszélünk, de ugyanez a helyzet akkor is, ha különböző csoportokat az átlag vagy más középérték szerint hasonlítunk össze (Vargha, 2019, 2020). Ezekben az elemzésekben az a közös, hogy statisztikai modelljeikben egy-egy konkrét személynek sehol se találunk helyet. A *személy-orientált* megközelítés szerint viszont a személyt jellemző adatokat, változóértékeket feldarabolatlan egységként kell tekinteni és kezelni. Ez a megközelítés az egyénnel kapcsolatos folyamatokra fókuszál és azt hangsúlyozza, hogy a kapott statisztikai eredményeknek informatívaknak kell lenniük az egyén fejlődésére, változására (Bergman és Lundh, 2015).

A személy-orientált (személyközpontú) többváltozós statisztika olyan eljárásokra fókuszál, amelyek esetében központi szerepet játszanak az egyének közti kvalitatív jellegű különbségek. Ezek hátterében típusmodellek állnak, amelyek jellemzően klasszifikációs módszerekkel tárhatók fel. A személyiség tipológiai megközelítése nem új, például Hippokratész-Galénosz vérmérsékleti tipológiája, az első olyan részletesen kidolgozott típustan, amely ma is jól ismert, két és félezer éves (vö. Bartha, 1980).

A legismertebb típusfeltáró eljárás a *klaszteranalízis* (*klaszterelemzés*), a személy-orientált pszichológiai kutatások kedvelt módszere. Ez matematikai módszerekkel keresi egy többváltozós minta olyan alcsoportjait, ún. *klasztereit*, amelyek egymástól határozottan különböznek, de amelyeken belül a vizsgált személyek erősen hasonlítanak egymásra. A típusfeltáráson kívül a klaszteranalízis használható az adatminta takarékos leírására is (adatredukció), melynek során egy nagyobb minta személyeit szakmailag értelmezhető és további elemzésekre alkalmas csoportokba soroljuk.

Ezen III. rész fejezeteiben bemutatjuk, hogyan lehet a ROP-R klaszterező moduljai segítségével összevonó és osztódó hierarchikus, *k*-középpontú nem hierarchikus, valamint modell-alapú klaszterelemzést végrehajtani. Mielőtt azonban belevágnánk a klaszteranalízis különböző típusainak a taglalásába, szólnunk kell közös alkalmazási feltételeikről.

Megbízható többváltozós elemzéshez legalább 300-400 személyre van szükség és jó, ha a személyek száma legalább tízszerese a változók számának (Vargha, 2019, 80. o.). Speciálisan a klaszteranalízis esetében jó, ha a mintaelemszám az 1000-et is meghaladja, hogy a ritkább, 5% alatti arányú típusokat képviselő klaszterek is feltárhatók legyenek. A mintában

ne legyenek szélsőséges értékű outlierek, mert a kilógó személyek rontják a feltárt klaszterstruktúra homogenitását.

Az input változók száma ne legyen túl nagy, ideális, ha 2-7 közötti. Törekedjünk arra, hogy ezek száma minél kisebb legyen, de azért fedjenek le egy releváns területet, amelyen belül a típusokat keressük. Kerüljük a redundanciát! Általában közepes szorosságú kapcsolatban lévő változókat érdemes az elemzésbe bevenni. Legyenek legalább intervallum-skálájúak, de megengedett dummy (kétértékű) változók jelenléte is. A változók pszichometriai reliabilitása legyen magas. Minél nagyobb ugyanis a változók mérési hibája, annál nehezebb egy populációban létező klaszterstruktúrát akár nagy minták segítségével is azonosítani (Vargha és Bergman, 2019). Ha a változók különböző mértékegységűek, szükséges a változókat azonos skálára hozni. Egy ilyen lehetőség a *z*-standardizálás, mellyel minden változónál szórás léptékű skálára térünk át. Olykor azonban érdemes elkerülni ezt a standardizálást, mert az eredeti skálákon jobban értelmezhetők szakmailag az eredmények (Moeller, 2025; Vargha, 2025). E cél elérésére más transzformációk is szóba jöhetnek. Például vannak, akik azt javasolják, hogy a közös léptékre hozás érdekében a változókat ne a szórásukkal osszuk le (mint pl. a *z*-standardizálás során), hanem az értéktartomány szélességével, vagyis a terjedelemmel, amit a POMS-transzformáció esetében is alkalmaznak (Milligan és Cooper, 1988; Schaffer és Green, 1996; Vargha és Grezsa, 2025).

Néha egyszerű lineáris transzformációval is közös skálára hozhatók az input változók. E célból érdemes például a kérdőívek különböző tételszámú skáláit a skálába tartozó tételek átlagolásával képezni, ami megőrzi a tételek közös értéktartományát. Ha pedig mondjuk egy *A* kérdőív [1–4] értéktartományú *X* és egy *B* kérdőív [1–7] értéktartományú *Y* skáláját akarjuk közös nevezőre hozni úgy, hogy *Y* skáláját választjuk közös mércének, ezt a következő módon tudjuk elérni. Először mindkét változóból kivonunk 1-et, amivel *X* skálája [0–3], *Y*-é pedig [0–6] lesz. Látjuk, hogy *Y*-é kétszer olyan széles ($6/3 = 2$), ezért *X*-et még beszorozzuk 2-vel, így végül *X* skálája is [0–6] lesz, ugyanolyan, mint *Y*-é. A változók ilyen átalakításai – a jelen esetben két lépésben – a *Transzformációk* menüpont segítségével hajthatók végre ROP-R-ben (és ROPstatban).

Ha nem standardizált, de közös skálájú változókkal hajtunk végre klaszteranalízist, a klaszterátlagok (centroidok) oszlopdiagramját érdemes Excelben az ún. *z_i*-standardizált átlagok segítségével is elkészíteni, amelyeket ROP-R úgy számít ki, hogy minden átlagból kivonja az *M_i* összátlagot (az összes változó teljes mintabeli átlagának átlagát), majd leoszt az *S_e* együttes szórással (a változók variancia-átlagának négyzetgyökével; lásd 7.4. alfejezet).

Fontos még, hogy a változók szórása/szóródása ne legyen nagyon kicsi, mert egy nagyon kis variabilitású változó értékei szorosan egyetlen centrum köré tömörülnek, így ilyen változó nem lehet valódi klaszterképző komponens. Az ilyen változókat hagyjuk ki az elemzésből! Az a jó, ha a változók eloszlása nem normális, különösen, ha az adatok több centrum köré is tömörülnek az eloszlásban (bimodális és multimodális eloszlások). Az azonban nem szerencsés, ha a normalitás csupán a változók erős ferdesége miatt sérül.

Számos többváltozós eljárás (pl. a faktoranalízis) alkalmazási feltétele az elemzett változók többdimenziós normalitása. A klaszteranalízisben ennek éppen a fordítottja igaz. Ha a klaszteranalízis változói többváltozós normális eloszlást követnek, a klaszteranalízis bizonyosan nem vezethet nekünk tetsző megoldásra (Vargha, 2022, 62. o.). Ha polinomiális regresszióban (lásd 2. fejezet) találunk markáns nem lineáris összefüggéseket az input változók között, ez esélyt adhat nem triviális klaszterstruktúra feltárására.

7. fejezet

Az összevonó hierarchikus klaszteranalízis (AHKA) modul

Ebben a fejezetben AHKA, a hagyományos összevonó (agglomeratív) hierarchikus klaszteranalízis módszerét és ROP-R-beli futtatásának módját ismertetjük. Az AHKA modulban a hierarchikus klaszteranalízis végrehajtására 6 lehetséges személytávolság és 8 választható klaszterösszevonási módszer áll rendelkezésre.

A 7.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 7.2. alfejezet az AHKA menüablak használatát mutatja be, a 7.3. alfejezet az eredménylista szerkezetét, a 7.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

7.1. Elméleti alapok

A személyeken végezhető legismertebb hagyományos klaszterelemzési módszer a *hierarchikus klaszteranalízis* (HKA). Ez valójában nem egyetlen klaszteranalízis, hanem olyan többlépéses klasszifikáció sorozat, amelynek minden lépésében vagy összevonunk két klasztert egy közös nagyobb klaszterbe (*összevonó, agglomeratív HKA*, röviden AHKA), vagy szétbontunk, felosztunk egy klasztert két kisebb klaszterre (*osztódó, felosztó HKA*, röviden OHKA). Ezek közül az AHKA az elterjedtebb, melynek első lépésében az egymáshoz legközelebbi két személyt vonjuk össze közös klaszterbe (1-elemű klaszternek tekintve őket), s minden további lépésben az egymáshoz legközelebbi két klasztert. Az összevonást addig folytatjuk, amíg vannak különböző klaszterek; a végén minden személy egyetlen nagy közös klaszterbe kerül. OHKA esetében a folyamat fordított, itt az első lépésben még egyetlen nagy klaszterben van mindenki, az utolsóban pedig mindenki egyedül egy 1-eleműben. Itt a szétbontás alapelve, hogy mindig a legheterogénebb klasztert bontjuk két alklaszterre úgy, hogy azok a lehető legjobban különbözzenek egymástól (részletesen lásd a 8. fejezetben). Mindezekből az is látható, hogy a klaszteranalízisben alapvető fontosságú a *hasonlóság* és a *különbözőség* fogalma mind a személyek, mind a klaszterek tekintetében. Az alábbiakban ezért ezekkel részletesebben is foglalkozunk.

7.1.1. Személyek távolsága

Klaszteranalízissel egy többváltozós minta homogén alcsoportjait keressük, vagyis olyan klasztereket, amelyekben belül a személyek hasonlítanak egymásra. A személyek klaszteranalízisének eredményét alapvetően befolyásolja a személyek közti távolság megválasztása, definiálása. Klaszteranalízisekben a személyek hasonlóságának mérése

számos távolságmérték ismeretes (Füstös et al., 2004, 169–173.; Takács, Makrai és Vargha, 2015). A leggyakrabban használt ilyen mutatókat a 7.1. táblázatban foglaltuk össze.

Az ED^{38} *euklideszi távolság* a szokásos síkbeli, illetve térbeli távolság a vizsgált változók ortogonális (egymással páronként derékszöget bezáró) terében. Ha a változóértékek különbségei közül büntetni akarjuk a legnagyobbakat (így csökkentve a legnagyobb különbségek hatását), akkor nem végezzük el a gyökvonást az euklideszi távolság képzésekor, amivel a SED^{39} *négyzetes euklideszi távolság*ot kapjuk. Ha ezt még leosztjuk a változók számával, az $ASED$ távolsághoz jutunk. SED és $ASED$ között csak egy konstans osztóban van különbség, így ezek minden klasszifikációs elemzésnél ugyanarra az eredményre vezetnek, de az $ASED$ távolság jobban értelmezhető, mert jelzi az egy változóra eső átlagos eltérést.

7.1. táblázat. A klaszteranalízisekben a személyek hasonlóságának mérésére használt leggyakoribb távolságmértékek

Távolság neve	Távolság definíciója
Euklideszi (ED)	Négyzetes eltérések összegének négyzetgyöke
Négyzetes euklideszi (SED)	Négyzetes eltérések összege
Átlagos négyzetes euklideszi (ASED)	Négyzetes eltérések átlaga
Minkowski, rögzített p értékkel	A változónkénti eltérések p -edik hatványait összegezzük, majd az összegből p -edik gyököt vonunk (ED távolság általánosítása tetszőleges p hatványra)
Mannhattan (city-block)	Abszolút eltérések összege (a Minkowski távolság $p = 1$ értékkel)
Canberra	A Manhattan távolság egy súlyozott változata
Csebisev (Maximum)	Abszolút eltérések maximuma
Adatsorok Pearson-féle r korrelációja	A lehetséges maximális pozitív együttjárástól való elmaradás: $d = 1 - r$

Ha a személyek hasonlóságának megítélésekor értékmintázataik hasonlósága a döntő, akkor a hasonlóság mérésére jó választás lehet a két személy adatsorának Pearson-korrelációja. Ha viszont a hasonlóságot a változónkénti értékek közelségével mérjük, akkor jobb az euklideszi távolság valamilyen változata. Személy-orientált kutatásokban a legkedveltebb személytávolság a SED , illetve az ezzel ekvivalens $ASED$.

7.1.2. Klaszterek távolsága

A klaszteranalízisben nem csak az fontos, hogy a feltárt klaszterek homogének legyenek, hanem az is, hogy jól elkülönüljenek egymástól. Ehhez pedig az szükséges, hogy a klaszterek kellő távolságra legyenek egymástól. *Két klaszter távolsága* olyan fogalom, amelyet matematikailag többféleképpen lehet definiálni, s amely ugyancsak döntő hatással van a feltárt klaszterstruktúrára, különösen AHKA esetén. A legismertebb klasztertávolságokat a 7.2. táblázatban foglaltuk össze.

³⁸ Euclidean Distance rövidítése

³⁹ Squared Euclidean Distance rövidítése

7.2. táblázat. Két klaszter távolságának legismertebb mértékei

Magyar elnevezés	Angol elnevezés	Alkalmazott klasztertávolság
1. Minimális távolság (legközelebbi szomszéd)	Nearest neighbor	A két klaszter egymáshoz legközelebbi elemének távolsága
2. Maximális távolság (legtávolabbi szomszéd)	Furthest neighbor	A két klaszter egymástól legtávolabbi elemének távolsága
3. Átlagos távolság	Between-groups distance	A két klaszter elemei közti távolságok átlaga
4. Centroid-távolság	Centroid distance	A két klaszter centroidjának (többdimenziós átlagának) a távolsága

7.1.3. Klaszterstruktúrák jóságának mérése

Minden klaszteranalízis használ valamilyen algoritmust arra, hogy optimalizáljon egy kritériumot a feltárt klaszterstruktúra jóságával, megfelelőségével kapcsolatban. A klaszterstruktúra jóságának megítélése fontos annak eldöntéséhez is, hogy a feltárt struktúra jobb-e, mint egy másik, amelyet például más távolságmértékkel vagy más módszerrel kapunk, vagy amelynél más a változók vagy a klaszterek száma. Az ilyen jellegű kérdések megválaszolására találták ki a különböző klaszter adekvációs mutatókat (angolul clustering quality coefficients), amelyeket a továbbiakban QC-ként rövidítünk. A QC-k tehát olyan mutatók, amelyek segítségével egy klasztermodell jósága megítélhető. A QC-ket részletesen ismerteti Vargha (2022, 4.4. alfejezet), ezért most csak a ROP-R-ben elérhetőekről szólunk.

Egy jó klasztermodelltől elvárjuk, hogy a klaszterek homogének legyenek, amit jól mér például a klaszterbeli személyek egymástól való átlagos ASED távolsága (vö. 7.1. táblázat), a HC⁴⁰ *klaszter homogenitási együttható*. Minél kisebb egy klaszter HC-értéke, annál homogénebb a klaszter. Ha a klaszterelemzést standardizált változókkal végezzük, 0,5-nél kisebb HC értékek jellemeznek egy igazán homogén klasztert. Ha nem standardizálunk, érdemes HC-t leosztani a változók teljes mintabeli varianciáinak átlagával (VarT) és ezen „standardizált” HC mutató (HCstan = HC/VarT) alapján ítéletet mondani a klaszter homogenitásáról. A teljes klaszterstruktúra homogenitását mérhetjük a HC, illetve HCstan értékek átlagával (HCátlag, illetve HCátlagS). Nem standardizált változók esetén a HCátlagS mutatót érdemes használni, melyre a HCátlagS = HCátlag/VarT összefüggés is érvényes. Jó struktúra esetén – saját tapasztalataim alapján – 0,50-et nem meghaladó HC, illetve HCstan értékek, valamint 1-nél érezhetően kisebb HCátlag, illetve HCátlagS szintek jellemzők.

Jó klaszterhomogenitási mutató EESS% is, a klaszterek által megmagyarázott varianciaarány, mely a varianciaanalízisből ismert eta-négyzet (η^2) hatásmérték többváltozós általánosítása (Vargha, Torma és Bergman, 2015; Vargha és Bergman, 2019; Vargha, 2022, 72. o.):

$$\text{EESS\%} = 100 \cdot (\text{SS}_{\text{klaszter}} - \text{SS}_{\text{total}}) / \text{SS}_{\text{total}}.$$

Itt SS_{klaszter}, az ún. *összhiba*, a klasztereken belüli – klasztercentroidtól való – négyzetes eltérések összegének az összege (a klaszterekre összegezve), míg SS_{total} a teljes mintában a minta centroidjától való négyzetes eltérések összege (a négyzetes eltéréseket minden esetben változónként kiszámítva, majd rájuk összegezve).

⁴⁰ Homogeneity Coefficient

Minél jobban megközelíti egy klaszterstruktúra EESS% értéke a 100-at, annál kisebb az összhiba, és annál jobban magyarázzák a struktúra klaszterei az elemzésbe bevont változók varianciáját, vagyis az egyének közti eltéréseket. Más szavakkal EESS% azt méri, hogy az egyének közti különbségekért milyen mértékben felelős az, hogy az egyének melyik klaszterbe tartoznak. Egy elfogadható klaszterstruktúra esetén EESS%-nak illik elérnie a 65, jó struktúra esetén pedig a 70-75%-ot. Az a jó, ha EESS% nagy, de a klaszterek száma viszonylag kevés. Ezt a két ellentétes szempontot kell az elemzés során összehangba hozni.

Egy klaszterstruktúra jóságának megítélésekor fontos arra is figyelni, hogy a klaszterek mennyire különböznek, mennyire szeparálhatók egymástól. A QC-k közül a Rousseeuw (1987) által bevezetett SC^{41} Silhouette együttható (vagy Silhouette index) közkedvelt szeparációs mutató, melyet a személyenként kiszámított

$$SC_i = (B - A) / \max(A, B)$$

értékek⁴² átlagolásával kapunk, ahol A az i -edik személy átlagos távolsága saját klaszterének többi személyétől, B pedig az i -edik személy átlagos távolsága a hozzá legközelebbi idegen klaszter személyeitől. Az eredeti Rousseeuw-féle SC (= SC_{orig}) ezeknek az SC_i értékeknek az átlaga a teljes mintára. Az SC_i értékek átlagát klaszterenként is érdemes kiszámítani ($SC(1)$, $SC(2)$, ..., $SC(k)$), ezek jelzik az egyes klaszterek szeparációját a többi klasztertől egy klaszterstruktúrában. Egy magas SC érték arra utal, hogy az adott klaszterstruktúrában a személyek általában közelebb vannak saját klaszterük egyedeihez, mint a legközelebbi idegen klaszter egyedeihez. Az SC értékek -1 és 1 közé esnek. Negatív értékek igen gyenge szeparációjú klaszterstruktúrát jeleznek. Jó struktúra esetén elvárt, hogy SC nagyobb legyen $0,5$ -nél, a 0 és $0,2$ közötti SC értékek pedig gyenge szeparációjú klaszterstruktúrára utalnak (Kaufman és Rousseeuw, 1990). SC általában könnyebben ér el kíváncsan magas értéket kevés számú klaszternél, ahol könnyebben különülnek el egymástól a klaszterek.

SC egyszerűsített formája (Sil.eh. vagy SC_{simple}) azt méri, hogy a mintabeli személyek átlagosan mennyivel vannak közelebb saját klaszter-centrumukhoz, mint a legközelebbi idegen centrumhoz (Vargha, Bergman és Takács, 2016; Vargha, 2022, 73. o.). Egy magas SC_{simple} érték arra utal, hogy az adott klaszterstruktúrában a személyek általában közelebb vannak saját klaszterük centrumához, mint a legközelebbi idegen klasztercentrumhoz. Negatív értékek esetén a személyek általában távolabb vannak saját klaszterük centrumától, mint a legközelebbi idegen klasztercentrumtól, ami nem szerencsés. Jó struktúra esetén elvárt, hogy SC_{simple} nagyobb legyen $0,5$ -nél. SC_{simple} általában $1-2$ tizeddel nagyobb szokott lenni, mint SC_{orig} .

A ROP-R klaszterező moduljai által elkészíthető Silhouette-ábrák SC_{orig} értékeit mutatják be különböző klaszterszámokra, de a ROP-R k -középpontú moduljának outputja egyaránt tartalmazza az SC_{orig} és az SC_{simple} értékeket a teljes mintára és klaszterenként külön is (lásd 9. fejezet). Ugyanakkor a ROPstat klaszterező moduljai csak SC_{simple} (Sil.eh.) értékét számítják ki, de a *Validálás* modul outputja SC_{orig} értékét is tartalmazza (a teljes mintára és klaszterenként külön-külön is).

Hasznos szeparációs mutató XBmod is, a módosított Xie-Beni index (Vargha, 2022, 74. o.), mely azt méri, hogy az egymáshoz legközelebbi két klasztercentrum távolsága relatíve mennyivel nagyobb, mint a saját klasztercentrumtól való átlagos távolság. XBmod kiértékelése és értelmezése a Silhouette együtthatóhoz hasonlóan végezhető el. Ha egy

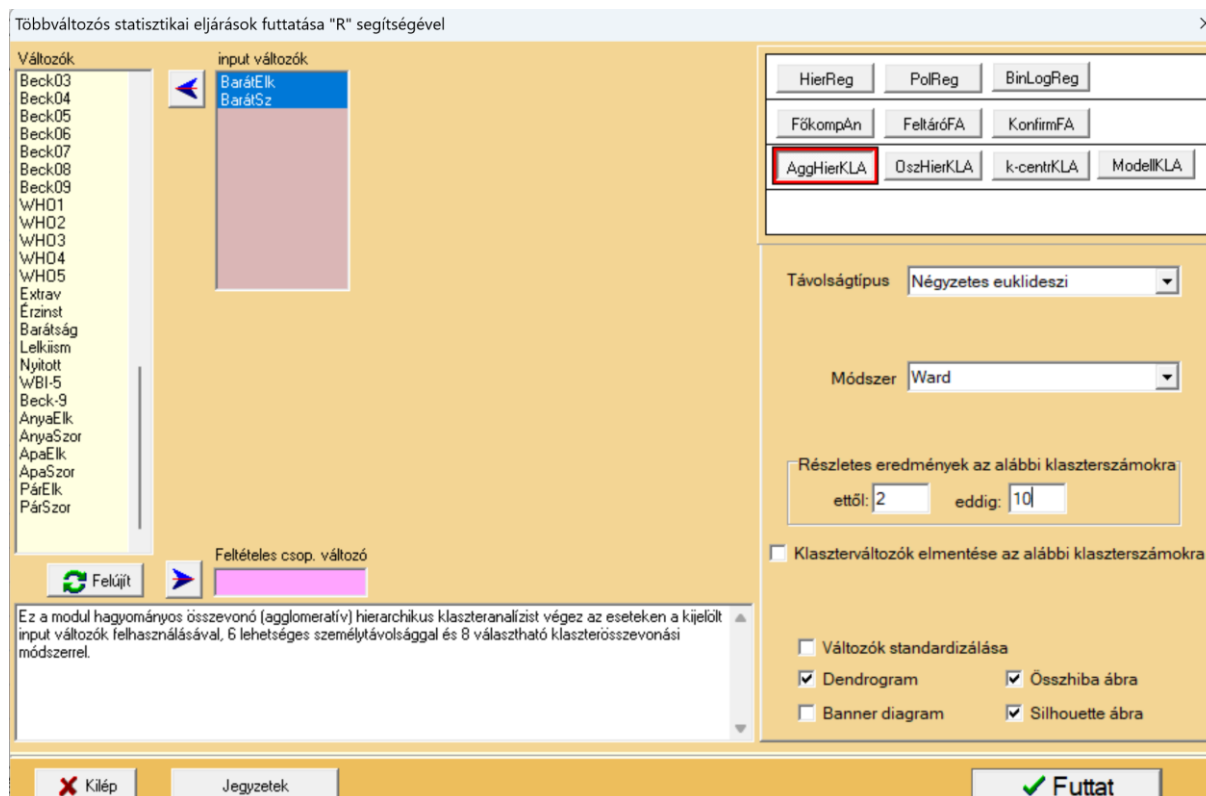
⁴¹ Silhouette Coefficient

⁴² Az i -edik személy SC_i értékét Silhouette-szélességnek (Silhouette width) nevezzük.

klaszterstruktúrában van egy-két olyan klaszter, amely igen nagy és heterogén (mondjuk $HCstan > 1,5$), továbbá van két olyan klaszter, amelyik igen közel van egymáshoz, akkor $XBmod$ értéke jóval a kívánatos 0,5-ös szint alá is mehet.

7.2. Az AHKA menüablak

AHKA elemzést ROP-R-ben az *összevonó (agglomeratív) hierarchikus klaszteranalízis* (röviden AHKA) modul segítségével végezhetünk el (lásd 7.1. ábra). Ez a modul AHKA-t végez a mintabeli eseteken a kijelölt változók felhasználásával. AHKA a *stats* (R Core Team, 2021), a *cluster* (Maechler et al., 2022), a *ggplot2* (Wickham, 2016) és a *factoextra* (Kassambara és Mundt, 2020) R-package-et használja fel elemzéseikhez. Az AHKA menüablak szolgál az elemzésbe bevonandó változók kiválasztására (input változók ablaka), a személytávolság (Távolságtípus) és az AHKA összevonási módszerének (Módszer) megválasztására, valamint egy sor opció kijelölésére (lásd 7.1. ábra), amelyekhez az alábbi alpontokban fűzünk eligazítást.



7.1. ábra. Az AHKA modul menüablaka ROP-R-ben

7.2.1. Input változók kiválasztása

Az AHKA-hoz szükséges változók kiválasztásához a Változók feliratú ablakból kell az elemzéshez szükséges változókat az „input változók” ablakba átküldeni a két ablak közti nyílra rákattintva.

7.2.2. Személytávolság megválasztása

Az AHKA modulban hat személytávolság áll rendelkezésre, a 7.1. táblázatban felsoroltak közül az utolsó sor (Pearson korreláció) kivételével mindegyik. ASSED csak azért hiányzik, mert a vele ekvivalens SED jelenléte feleslegessé teszi.

7.2.3. Klaszterösszevonás módszerének megválasztása

Ahogy már említettük, AHKA minden lépésében az egymáshoz legközelebbi két klasztert vonjuk össze közös klaszterbe. ROP-R AHKA moduljában nyolc klaszterösszevonási módszer áll rendelkezésre. Ezek közül az első négy a 7.2. táblázatban összefoglalt klasztertávolságon alapul, a maradék négy pedig az alábbi:

- *Medián*: ez a centroid módszer olyan variánsa, amely szimmetrikus eloszlású változók esetén hasonló, erősen ferde eloszlások esetén pedig gyakran jobb struktúrához vezet.
- *Ward*: ekkor minden lépésben azt a két klasztert egyesítjük, amelyekre vonatkozóan a klaszterstruktúra összhomogenitását mérő EESS% mutató a legkisebb mértékben csökken.
- *Flexibilis béta*: egy viszonylag bonyolult, de néha igen sikeres klasszifikációt eredményező összevonási módszer, melynél egy β (beta) paraméter értékének 0,1 és 1 között lehetséges beállításával különböző klasztertávolságok állíthatók be (a részletekről lásd Vargha, 2022, 86. o.). Például $\beta = 0,2$ a Ward-féle módszer, egy 1-hez közeli β pedig a minimális távolság esetén kapotthoz hasonló megoldásra vezet.
- *McQuitty*: az átlagos távolság módszerének egy olyan variánsa, amelynél a létrehozott új és a régi klaszterek távolságának meghatározásában egyenlővé tesszük az éppen összevont két klaszter esetleg eltérő elemszámának a hatását (Vargha, 2022, 85. o.).

Megjegyezzük, hogy a centroid-, a medián- és a Ward-féle módszer választása esetén ROP-R kötelezően mindig a SED személytávolságot állítja be.

Pszichológiai kutatásokban a fenti nyolc módszer mindegyike vezethet érdekes eredményre, de személy-orientált vizsgálatokban, típusok feltárására leginkább a Ward-féle módszert használják a leggyakrabban (Bergman, Magnusson és El-Khoury, 2003).

7.2.4. Segítő ábrák

Az AHKA eredményének értelmezését négy választható ábra segíti az AHKA modulban, amelyeket a menüablak jobb alsó paneljén lehet kijelölni:

- Dendrogram: az AHKA lépésenkénti összevonásait szemléltető fadiagram;
- Banner-ábra: a dendrogram jégcsap-ábrához hasonló változata;
- Összhiba ábra: a klasztereken belüli összhiba, vagyis SSklaszter lejtődiagramja 1 és 8 klaszterszám között;
- Silhouette-ábra: a Silhouette együttható értékének ábrázolása 1 és 8 klaszterszám között.

7.2.5. Egyéb lehetőségek

Az AHKA menüablakában megadható a klaszterszámok egy övezete, amelyen belül minden megoldásra megkapjuk a klaszterstruktúra legfontosabb QC adekvációs mutatóit (EESS%, XBmod, HCátlag stb.) a klaszterstatisztikákat (átlag, szórás, minimum, maximum), valamint a standardizált átlagok mintázatát. Külön kérésre a klaszterszámok egy megadott övezetére a klaszterkódot személyenként megadó klaszterváltozók elmenthetők és a menüablakban állítható be az is, hogy standardizáljuk-e az input változókat (alapértelmezés szerint igen).

7.2.6. Az elemzés után létrejövő fájlok

Egy AHKA elemzés végrehajtása után a „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzéshez elkészített ideiglenes adatfájlt (tmpdat.txt), a kért klaszterszámokhoz tartozó klaszterváltozókkal kiegészített ideiglenes adatfájlt (tmpdat2.txt), a futtatott R-scriptet (AHCA.r), valamint a kért diagramokat jpg vagy pdf fájlban (pl. Dendr1.jpg vagy Banner1.pdf). Ha feltételes csoportosító változót is kijelölünk, akkor minden feltételes csoport elemzése során elkészülnek a kért diagramok, a diagramokat tartalmazó fájlok nevében megjelenő számok ezen csoportok sorszámát jelzik.

7.3. Mit tartalmaz az AHKA eredménylistája?

Az AHKA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- Minden klasztermegoldásra:
 - o Klaszterstruktúra adekvációs mutatói (EESS% stb., vö. 7.1.3. alpont)
 - o Klaszterstatisztikák (elemszám, átlag, szórás, minimum, maximum)
 - o Nem standardizált átlagok (nem standardizált centroidok)
 - o Standardizált átlagok (standardizált centroidok)
 - o Standardizált átlagok mintázata
- Különböző klasztermegoldások adekvációs mutatóinak összefoglaló táblázata.

Ha kérjük a változók standardizálását, akkor az átlagok standardizálása változónként a saját átlaggal és szórással, ha pedig nem kérjük, akkor a változók összátlagával (változóátlagok átlaga) és együttes szórásával (változóvarianciák átlagának gyöke) lesz végrehajtva.

7.4. Az AHKA szemléltetése valódi adatokon

A KÖT2016 kötődéskutatási adatokon (lásd B2.2. alpont) a baráttal kapcsolatos elkerülés (BarátElk) és szorongás (BarátSz) skálájával mint input változóval szeretnénk klaszterelemzést végezni. Fraley és munkatársai (2011) modelljében biztonságos, jó kötődéssel rendelkeznek azok, akik az elkerülés és a szorongás tekintetében egyaránt alacsony szinten vannak, míg a félelemteli, elkerülő kötődésűek mindkét dimenzión magas értékűek. A magas elkerülés – alacsony szorongás kombináció az elutasító–elkerülő, a magas szorongás – alacsony elkerülés kombináció pedig az elárasztott–megszállott típusra jellemző (vö. Jantek és Vargha, 2016, 1. ábra). Kíváncsiak vagyunk, sikerül-e ezt a négy típust, mint klasztert AHKA-val, a baráttal kapcsolatos kötődés tekintetében mintánkban feltárni.

Mielőtt belevágnánk a klaszterelemzésekbe, érdemes kiszűrni az outliereket és normalitásvizsgálatot végezni. Az outliereket a ROP-R FKA moduljával, a BarátElk és BarátSz skála összesen 9 tételén végzett elemzés során azonosítottuk az „Extrém esetek (outlierek) azonosítása” opció bejelölésével (a GMD általánosított Mahalanobis távolság alapján; vö. 4.2. alfejezet) és az Outli nevű outlier változó elmentésével. Ezek közül az esetek közül a 12 legnagyobb GMD értékűt a GMD távolsággal és a barátra vonatkozó 9 tétel értékével együtt a 7.3. táblázatban foglaltuk össze.

7.3. táblázat. A Dmin érték feletti GMD távolságok, a távolsághoz tartozó esetsorszám (Eset), valamint a skála 9 tételének az értéke az adott személynél (Elk=Elkerülés, Sz=Szorongás)

GMD	Eset	Elk1	Elk2	Elk3	Elk4	Elk5	Elk6	Sz7	Sz8	Sz9
90,43	229	4	4	2	5	7	1	7	1	4
76,68	111	1	5	5	6	7	1	1	5	5
67,95	223	2	3	3	2	1	1	1	2	2
62,71	322	1	3	2	3	2	2	2	2	7
59,74	321	2	2	6	2	6	2	1	1	1
54,88	196	7	7	7	1	7	7	1	1	1
51,06	323	1	2	2	2	6	6	6	4	1
48,17	297	1	1	1	1	6	5	1	1	5
46,53	125	1	2	2	1	1	1	4	1	6
46,22	213	3	3	3	2	2	1	1	6	1
45,88	76	2	1	1	1	7	1	1	1	1
45,04	153	1	3	3	3	4	4	1	6	1

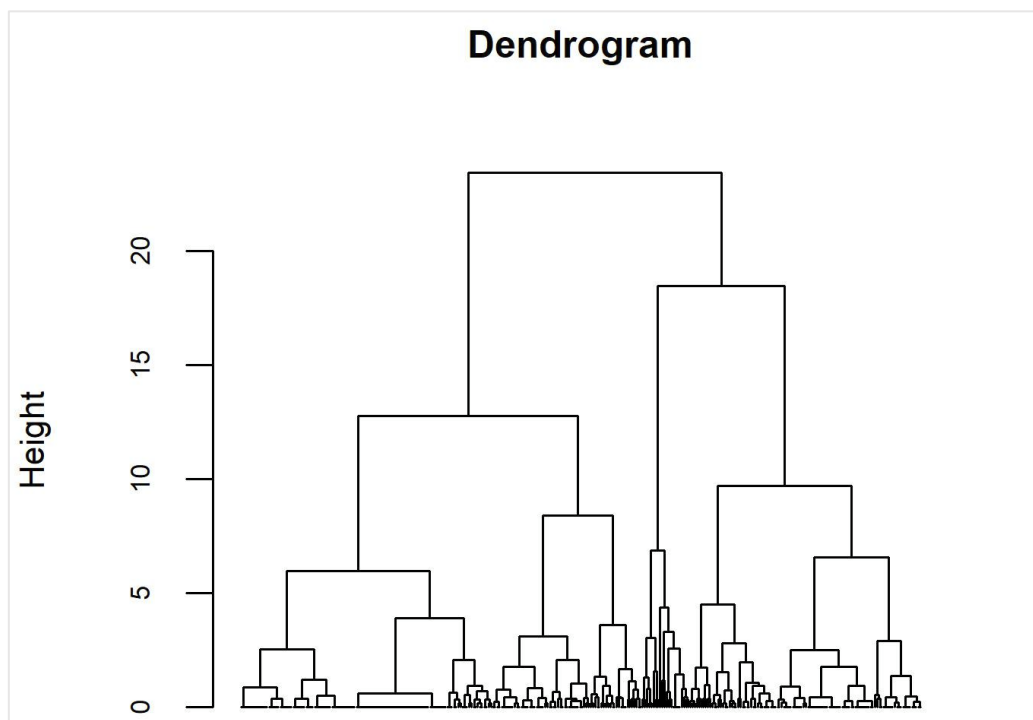
E 12 eset közül az első kettőnél (229-es és 111-es), valamint az 5–8. (321, 196, 323, 297 sorszámú) esetében fordult elő többszörös inkonzisztencia a tételek értékei között, így ezt a 8 esetet érdemes kihagyni a klaszterelemzések során. Például a 229-es esetnél Elk5 és Elk6, valamint Sz7 és Sz8 értéke között van teljes inkonzisztencia. A harmadik (223-as) eset extrém ugyan a nagyon alacsony értékszintje miatt, de az értékek között nincs inkonzisztencia. A negyedik (322-es) esetében is látunk egy következetesen alacsony értékszintet, amiből az utolsó tétel (Sz9) kiugró 7-es értéke meglepő, de ha az ilyen egy értéknél előforduló inkonzisztencia alapján már kihagynánk egy esetet a mintából, az már túlzott mértékű beavatkozás lenne a kapott adatmintába. Mindezen megfontolások alapján tehát az említett 8 esetet szeretnénk outliernek kijelölni a GMD változónál. Ezt úgy tudjuk megtenni, hogy a Változók deklarációi ablakban az Outli változónál – övezetes csoportkijelöléssel – a 48 alatti értékekkel definiáljuk a normál övezetet (0 alsó és 48 felső határral), de az adatállományban módosítjuk a 223-as és a 322-es sorszámú személy Outli értékét egy 48 alatti értékre (pl. 40-re), hogy ezek ne legyen outlierként számontartva. Ezek után ez a 8 outlier személy oly módon hagyható el a további elemzésekből, hogy az Outli változót feltételes csoportosító változóként jelöljük majd ki. Ennek nyomán a 336 fős mintából – más adathiányzásokat is figyelembe véve – 314 személynél volt érvényes értéke a BarátElk és a BarátSz változónak, amelyekkel klaszterelemzéseket végzünk.

A normalitás vizsgálatára kiszámítottuk – a ROPstat szoftver segítségével – a ferdeségi és a csúcsossági együtthatót, s ez a BarátSz skála esetében jelezte a normalitás súlyos sérülését (ferdeség = 1,70, csúcsosság = 2,87, mindkét esetben $p < 0,001$). A többdimenziós normalitás sérülését a ROP-R polinomiális regresszioelemzése jelezte (vö. 2. fejezet), függő változó = BarátSz, független változó = BarátElk szerepkiosztással, 5-ös maximális hatvány beállításával. Ez esetben ugyanis a független változónak mind a 2. hatványa ($p = 0,046$), mind az 5. hatványa ($p = 0,034$) szignifikánsan megemelte a megmagyarázott varianciát, a nemlineáris hatások pedig ellentmondanak a többdimenziós normalitásnak. Az egy- és a többdimenziós normalitás ily módon detektált sérülése esélyt adhat egy nem triviális klaszterstruktúra feltárására.

Ezután AHKA-t futtattunk az alapértelmezés szerinti Ward-módszerrel, a változókat nem standardizálva (mindkét változó csak az 1-7 intervallumon vehet fel értékeket, így azonos értékskáláról beszélhetünk), a 2–10 klaszterszámokra kérve részletes eredményeket (lásd 7.1. ábra). Az ábrák közül dendrogramot, valamint összhiba- és Silhouette-ábrát kértünk. Elsőként a Dendr1.jpg fájlban (a „c:_vargha\ropstat\aktualis” mappában) elhelyezett dendrogramot érdemes szemügyre venni (lásd 7.2. ábra).

A 7.2. ábra dendrogramjának vízszintes tengelyén az elemzett minta személyei foglalnak helyet. Egy-egy függőleges vonal alatti hierarchikus struktúra egy-egy klasztert képvisel, amelynek magassága a Height (= összevonási szint) elnevezésű függőleges tengelyen azt méri, hogy a klaszter létrejöttékor annak két alkotó komponense, amelyek összevonásából a klaszter létrejött, milyen távol volt egymástól. Így minél magasabb szinten vonódik össze két klaszter, annál heterogénebb klaszter jön létre.

Ha a dendrogramot a vízszintes tengellyel párhuzamosan valamilyen szinten egy egyenessel elmetsszük, annyi klasztert kapunk, ahány függőleges vonalat átmetszünk. Például egy 18 fölötti szinten meghúzott vonal 2, a 15-nél meghúzott 3, a 10-es szinten meghúzott pedig 4 klasztert jelöl ki. A dendrogram alapján oly módon dönthetünk egy optimális klaszterszámról, hogy megnézzük, milyen szinten metszhetjük el vízszintes vonallal a dendrogramot úgy, hogy a metszési szint viszonylag alacsony legyen és a metszéspontok száma se legyen túl nagy. Jelen esetben a 10-es szinten történő metszéssel kapott 4 klaszter elfogadható megoldásnak tűnik, ahogy a 7,5 körüli metszéssel kapott 6 klaszter is.



7.2. ábra. Az AHKA eredményének ábrázolása dendrogram segítségével a ROP-R AHKA moduljában (Ward-módszer, nem standardizált változók)

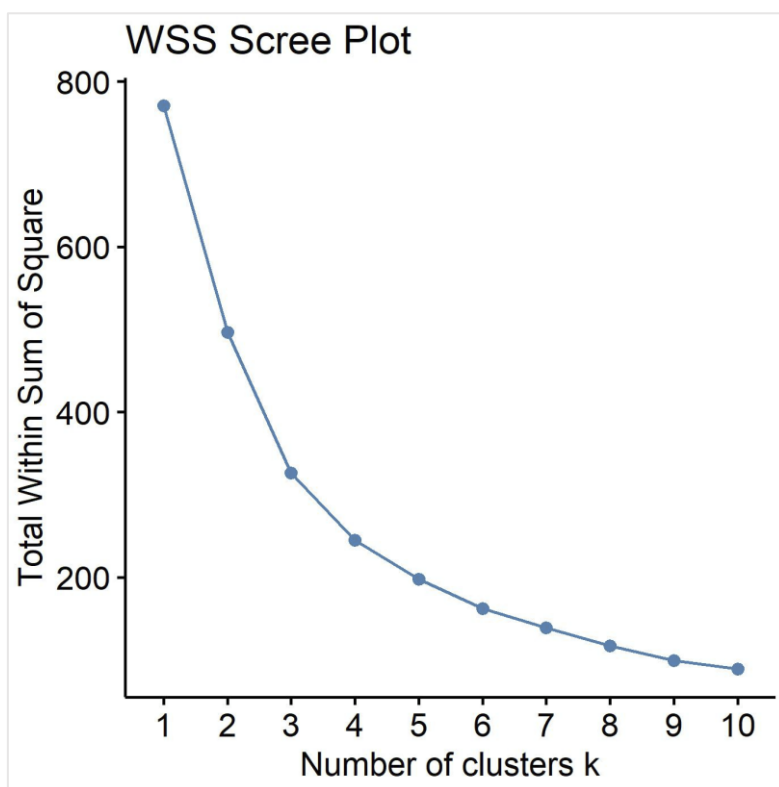
Az összhiba lejtődiagramját (WSS Scree Plot) ROP-R a szokásos „aktualis” mappában, egy WSSplot1.jpg nevű fájlban helyezi el (lásd 7.3. ábra). Ez a jelen esetben nem jelez olyan értéket a klaszterszámok tengelyén (Number of clusters k), ahol a görbe hirtelen

lecsökkenne, s attól kezdve ellaposodna, ezért ezzel az ábrával nem jutunk előbbre az optimális klaszterszám meghatározásában.

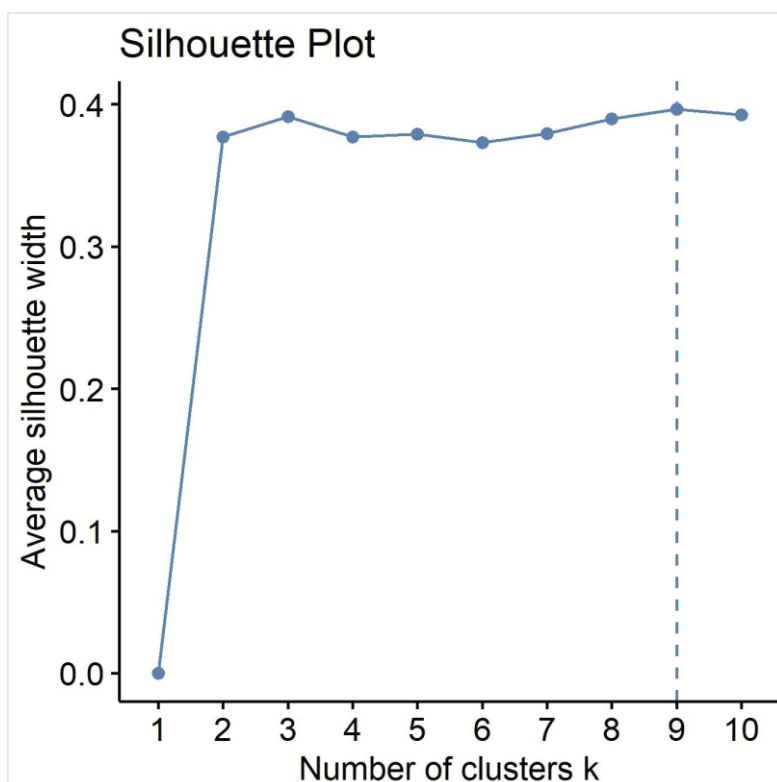
Szintén segíthet az optimális klaszterszám meghatározásában a Silhouette1.jpg fájlban elhelyezett SOrig (átlagos Silhouette-szélesség/Average silhouette width) grafikon, röviden Silhouette-ábra (lásd 7.4. ábra). A Silhouette-ábra esetében az első nagy kiugrásnál várjuk az optimális klaszterszámot, legalább 0,50-es értékszinttel (Rousseeuw, 1987).

Esetünkben a 7.4. ábrán az első markáns csúcs $k = 3$ értéknél látható, ez után a görbe nagyjából ugyanazon a 0,40 körüli nagyságsszinten mozog. A maximális SOrig értéket a $k = 9$ értéknél találjuk, de markáns kiugrás itt sincs. Egyértelmű támpontot tehát sem az összhiba lejtődiagramja, sem a Silhouette-ábra nem nyújt nekünk.

Nézzük meg ezután az AHKA eredménylistájának alján a különböző klasztermegoldások jellemzőit összefoglaló táblázatot, amelyet ROP-R a kért 2–10 klaszterszámokra készít el (lásd 7.4. táblázat). Minthogy az input változókat nem standardizáltuk, HCátlag helyett ennek mindig a standardizált változatát, a HCátlagS mutatót érdemes az értelmezésekben alapul venni, a HCstan értékek minimumával (HCminS) és maximumával (HCmaxS) együtt.



7.3. ábra. A klasztereken belüli összhiba lejtődiagramja a ROP-R AHKA moduljában (Ward-módszer, nem standardizált változók)



7.4. ábra. A Silhouette-ábra a ROP-R AHKA moduljában
(Ward-módszer, nem standardizált változók)

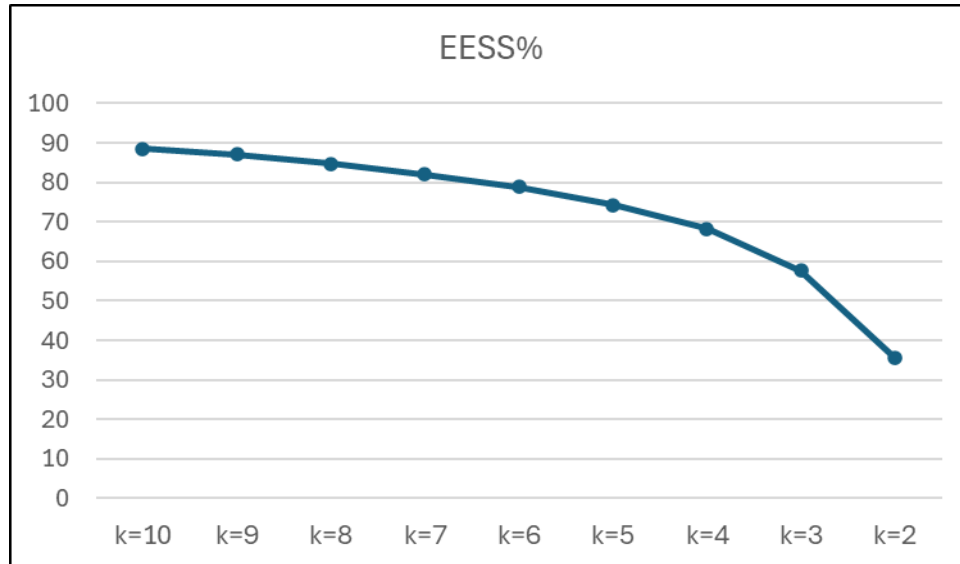
7.4. táblázat. Különböző klasztermegoldások jellemzői AHKA-ban
(Ward-módszer, nem standardizált változók)

k	EESS%	XBmod	HCátlag	HCátlagS	HCminS	HCmaxS
10	88,43	0,542	0,299	0,243	0,073	1,307
9	87,10	0,490	0,331	0,269	0,073	1,307
8	84,78	0,641	0,388	0,315	0,110	1,307
7	81,97	0,620	0,457	0,372	0,195	1,307
6	78,91	0,555	0,532	0,432	0,195	2,000
5	74,33	0,665	0,645	0,524	0,239	2,000
4	68,22	0,586	0,794	0,645	0,239	2,000
3	57,65	0,642	1,053	0,856	0,760	2,000
2	35,58	0,562	1,592	1,293	0,760	2,058

A 7.4. táblázat alapján az alábbi következtetéseket vonhatjuk le.

1. A szeparációt mérő XBmod $k = 5$ esetén a legnagyobb (XBmod = 0,665), de minden esetben közelíti ($k = 9$) vagy meghaladja a 0,50-es szintet, ezért az optimális klaszterszámról most a homogenitás mutatók alapján döntünk.
2. EEES% oszlopát nézegetve azt látjuk, hogy értéke – 10-től lefele haladva – a 4-es klaszterszámig (azaz a $k = 4$ értékig) elfogadható, mert végig 70% felett van ($k > 4$ esetén), illetve nem esik nagyon ez alá ($k = 4$ esetén 68,22).
3. Radikális esés csak ez után figyelhető meg (EEES% értéke $k = 3$ esetén már 60% alá csökken; lásd 7.5. ábra), így a $k < 4$ klaszterszámú megoldások negligálhatók.

4. A HCátlagS oszlopának elemzése során azt látjuk, hogy HCátlagS értéke az 5-ös klaszterszámig a kívánatos 0,50-es szint alatt vagy kicsivel fölötte van. Ez után hirtelen nagyot ugrik (kb. 0,12-t), $k = 3$ -ra pedig még nagyobb (több mint 0,20-at).
5. A legheterogénebb klaszter homogenitását jelző HCmaxS oszlopát megnézve azt láthatjuk, hogy 10 és 2 között minden megoldásban előfordul legalább egy nagyon heterogén klaszter, ahol HCmaxS értéke határozottan 1 feletti (minimum 1,3). Nagy ugrás (1,307-ről 2 fölé) $k = 7$ után, 7 alatti k értékekre figyelhető meg.



7.5. ábra. EESS% lejtődiagramja (AHKA, Ward-módszer, nem standardizált változók)

Mindezek alapján arra a következtetésre jutunk, hogy a KÖT2016 mintát nem lehet AHKA segítségével kevés számú homogén alcsoportra bontani. Viszont létezik részleges megoldás, egy-két heterogén, valamint több elfogadhatóan homogén klaszterrel. Ezt a megoldást a 7.2. és a 7.5. ábra, valamint a 7.4. táblázat alapján $k = 5$ körül kell keresnünk. Mivel változóink azonos skálán mérnek (elvi minimum = 1, elvi maximum = 7), a beállítások során nem kértük a változók standardizálását. A kapott klaszterstruktúra értelmezésekor valahogy meg kell ítélni, hogy a klaszterekben a változók átlagai mikor tekinthetők alacsonynak, mikor magasnak. Bár változónként nem standardizálunk, a klaszterátlagok nagyságának megítéléséhez célszerű kiszámítani a változók M_i összátlagát és S_e együttes szórását (az átlagvariancia négyzetgyökeként) és ezekkel standardizálni a klaszterátlagokat minden változóra, az $(M_{ij} - M_i)/S_e$ lineáris transzformáció segítségével. Ezt a standardizálást a hagyományos z -standardizálástól való megkülönböztetés érdekében a továbbiakban z_i -standardizálásnak nevezzük.

Az AHKA elemzés eredménylistáján megtalálhatók ezen átlagok és mintázataik táblázata mindazon klaszterszámokra, amelyekre részletes eredményeket kérünk. A $k = 6, 5, 4$ értékekre most az ily módon standardizált átlagok mintázatát a 7.5., 7.6. és 7.7. táblázat tartalmazza. Ezekben A és M az egyes klaszterátlagoknak a változók összátlagához viszonyított alacsony, illetve magas szintjét jelzi (minél több + kíséri, annál jobban), a sima pont (.) átlagos szintet, a zárójelbe írt A vagy M – (A), illetve (M) – összátlaghoz viszonyítva enyhén alacsony, illetve enyhén magas szintet jelez.

7.5. táblázat. A z_i -standardizált átlagok mintázata a 6-klasztteres AHKA megoldásban

Klaszter	BarátElk	BarátSz	KLgyak	HC	HCstan
KL1	M+	A	67	0,51	0,41
KL2	M+	.	40	0,56	0,46
KL3	(A)	A	116	0,29	0,24
KL4	M+++	M++++	22	2,46	2,00
KL5	.	.	45	0,24	0,20
KL6	(A)	M+	24	0,48	0,39

7.6. táblázat. A z_I -standardizált átlagok mintázata az 5-klasztteres AHKA megoldásban

Klaszter	BarátElk	BarátSz	KLgyak	HC	HCstan
KL1	M+	A	67	0,51	0,41
KL2	M+	.	40	0,56	0,46
KL3	(A)	A	116	0,29	0,24
KL4	M+++	M++++	22	2,46	2,00
KL5	.	.	69	0,84	0,68

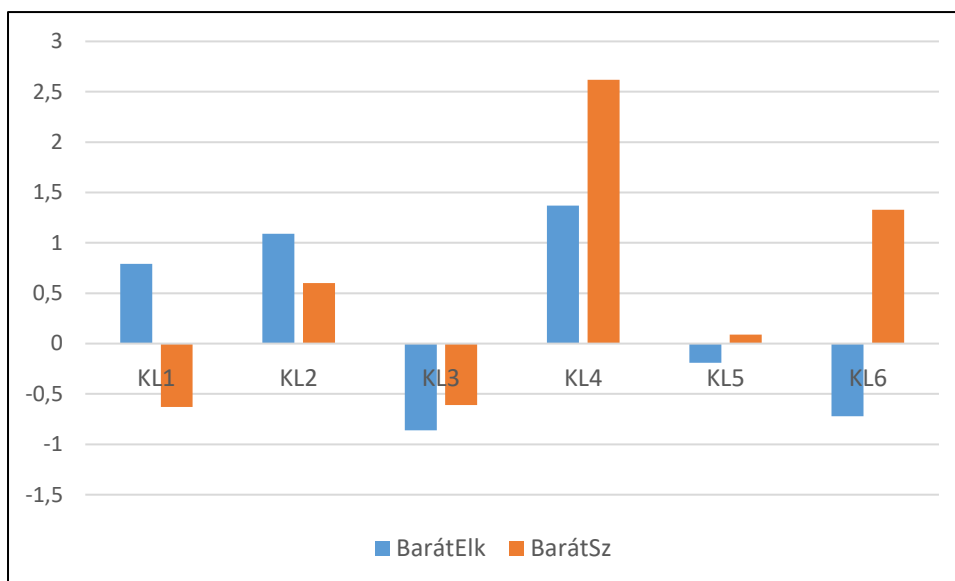
7.7. táblázat. A z_I -standardizált átlagok mintázata a 4-klasztteres AHKA megoldásban

Klaszter	BarátElk	BarátSz	KLgyak	HC	HCstan
KL1	M+	.	107	0,97	0,79
KL2	(A)	A	116	0,29	0,24
KL3	M+++	M++++	22	2,46	2,00
KL4	.	.	69	0,84	0,68

A 7.5. és a 7.6. táblázatot összevetve megállapíthatjuk, hogy a 6-klasztteres megoldásban az előnyösen homogén 45 fős KL5 (HCstan = 0,20) és a 24 fős KL6 (HCstan = 0,39) klaszter egyesítésével jött létre a 5-klasztteres megoldás 69 fős KL5 klasztere, melynek homogenitása már érezhetően gyengébb (HCstan = 0,68). A homogenitás gyengülésénél nagyobb baj, hogy két olyan klasztert vonunk össze, amelyek szakmailag nem tekinthetők közös típusnak. A 6-klasztteres megoldásban a 45 fős KL5 a mindkét változó tekintetében átlagos személyeket tömöríti egy erősen homogén (HCstan = 0,20) klaszterben, a csaknem fele akkora KL6 viszont egy érdekes nemtriviális, [(A), M+] mintázatú (BarátElk tekintetében az átlagosnál enyhén alacsonyabb, BarátSz tekintetében pedig az átlagosnál hangsúlyosan magasabb szintű) típust képvisel. Ezt a KL6-ot beolvastva a nála sokkal nagyobb [., .] mintázatú, vagyis mindkét változóban átlagos szintű KL5 klaszterbe, fontos szint veszítünk el a típusfeltárás során. Ennek alapján az AHKA-ban már meg is állhatunk a 6-klasztteres megoldásnál. Ha az egyszerűbb struktúrát erőltetve mégis továbbmennénk, a 7.6. és a 7.7. táblázatbeli mintázatokat összehasonlítva hasonló típusvesztést tapasztalhatunk, amikor az 5-klasztteres megoldásban az [M+, A] mintázatú KL1, valamint az [M+, .] mintázatú KL2 klasztert egyesítjük a 4-klasztteres megoldás [M+, .] mintázatú KL1 klaszterében.

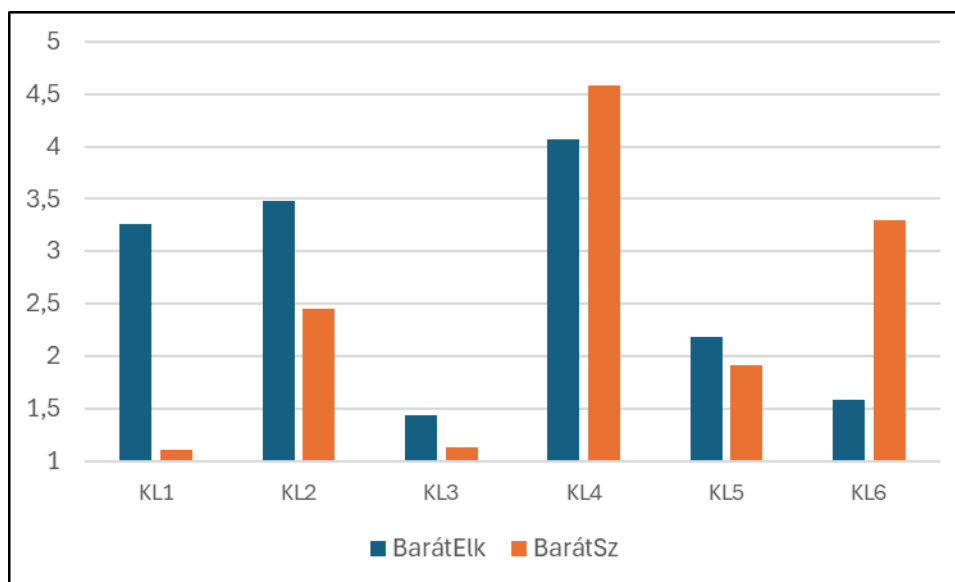
Mindezeket figyelembe véve AHKA-ban a 6-klasztteres megoldásnál célszerű megállni. E struktúra szemléltetésére két ábrát érdemes elkészíteni (pl. Excelben). A klaszterátlagok összátlaghoz viszonyított helyzetét a fenti módon leírt és 7.5. táblázat készítésekor is felhasznált z_I -standardizált átlagok ábrázolásával szemléltetjük (lásd 7.6. ábra), de informatív a nyers klaszterátlagok ábrázolása is, saját eredeti skálájukon (lásd 7.7. ábra). Az ábrákhoz szükséges átlagtáblázatok megtalálhatók az AHKA eredménylistáján a $k = 6$

klaszterszámhoz tartozó részben (futtatás előtt az AHKA feladatablakában jeleztük, hogy a $k = 6, 5, 4$ értékekre kérünk részletes eredményeket).



7.6. ábra. A 6-klaszteres AHKA-megoldás z_i -standardizált átlagainak oszlopdiagramja (Ward-módszer, nem standardizált változók)

Mind a 7.6., mind a 7.7 ábra jól mutatja, hogy ebben a 6-klaszteres megoldásban két nem szokványos mintázatú klaszter fedezhető fel. A 67 fős KL1-nek az összatlagnál érezhetően kisebb és a minimumhoz nagyon közeli baráti szorongás szintje magas baráti elkerülés szinttel jár együtt, míg a 24 fős KL6 klaszter mintázata ezzel éppen ellentétes.



7.7. ábra. A 6-klaszteres AHKA-megoldás nem standardizált átlagainak oszlopdiagramja (Ward-módszer, nem standardizált változók)

Ezzel a 6-klaszteres megoldással a Fraley és munkatársai (2011) által leírt mind a négy típust sikerült azonosítani (AA=KL3, AM=KL6, MA=KL1, MM=KL4&KL2). Az eltérés

mindössze annyi, hogy kaptunk egy Fraley és munkatársainak modelljéből hiányzó, minden tekintetben átlagos kötődésű, nem túlságosan nagy, de igen homogén típust (KL5, $n = 45$), továbbá a barát viszonylatban bizonytalanul kötődők – magas elkerüléssel és magas szorongással jellemezhető – típusát két klaszter képviseli: KL2 ($n = 40$) a homogénebb és a kötődési zavarok (különösen a baráttal szembeni szorongás) tekintetében enyhébb szintet képviseli, a sokkal heterogénebb, de kisebb KL4 ($n = 22$) pedig súlyosabb szintet, egy átlagosnál intenzívebb baráti elkerüléssel és szorongással.

A 7.6. és a 7.7 ábra között az a fő különbség, hogy előbbiről könnyebben leolvashatók változónként az összátlagtól való eltérések irányai, utóbbiról pedig a klaszterátlagok nagyság szintjei a változók közös értékskáláján. Előbbiről elsősorban a változók összátlagához viszonyított relatív elhelyezkedések, utóbbiról pedig a változók közös értékskáláján elfoglalt abszolút értékszintek azonosíthatók.

Egy-egy klaszter centroidjának értelmezéséhez érdemes többnyire mindkét ábrát szemügyre venni. Például a KL2 klaszter centroidjának mintázata a 7.6 ábra szerint a baráti elkerülés tekintetében kiemelkedő, kb. 1 szórással emelkedik az összátlag fölé, emellett a baráttal szembeni szorongás tekintetében is érezhetően átlag fölötti. A 7.7 ábra viszont arról is tájékoztat, hogy ezek a szintek a változók közös értékskáláján (1–7) érezhetően az elméleti skálaközép (4) alá esnek, különösen a baráttal szembeni szorongás tekintetében, ami annak köszönhető, hogy mintánkban a barát domén viszonylatában mind az elkerülési tendencia, mind a szorongás általában alacsony szintű, jóval közelebb van az elvi minimumhoz (1), mint az elvi maximumhoz (7).

8. fejezet

Az osztódó hierarchikus klaszteranalízis (OHKA) modul

Ebben a fejezetben OHKA, az *osztódó hierarchikus klaszteranalízis* módszerét és ROP-R-beli futtatásának módját ismertetjük. Az OHKA modul segítségével DIANA osztódó hierarchikus klaszteranalízis (Kaufman és Rousseeuw, 1990) végezhető a mintabeli személyeken a kijelölt kvantitatív változók felhasználásával, 6 lehetséges személytávolsággal.

A 8.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 8.2. alfejezet az OHKA menüablak használatát mutatja be, a 8.3. alfejezet az eredménylista szerkezetéről tájékoztat, a 8.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

8.1. Elméleti alapok

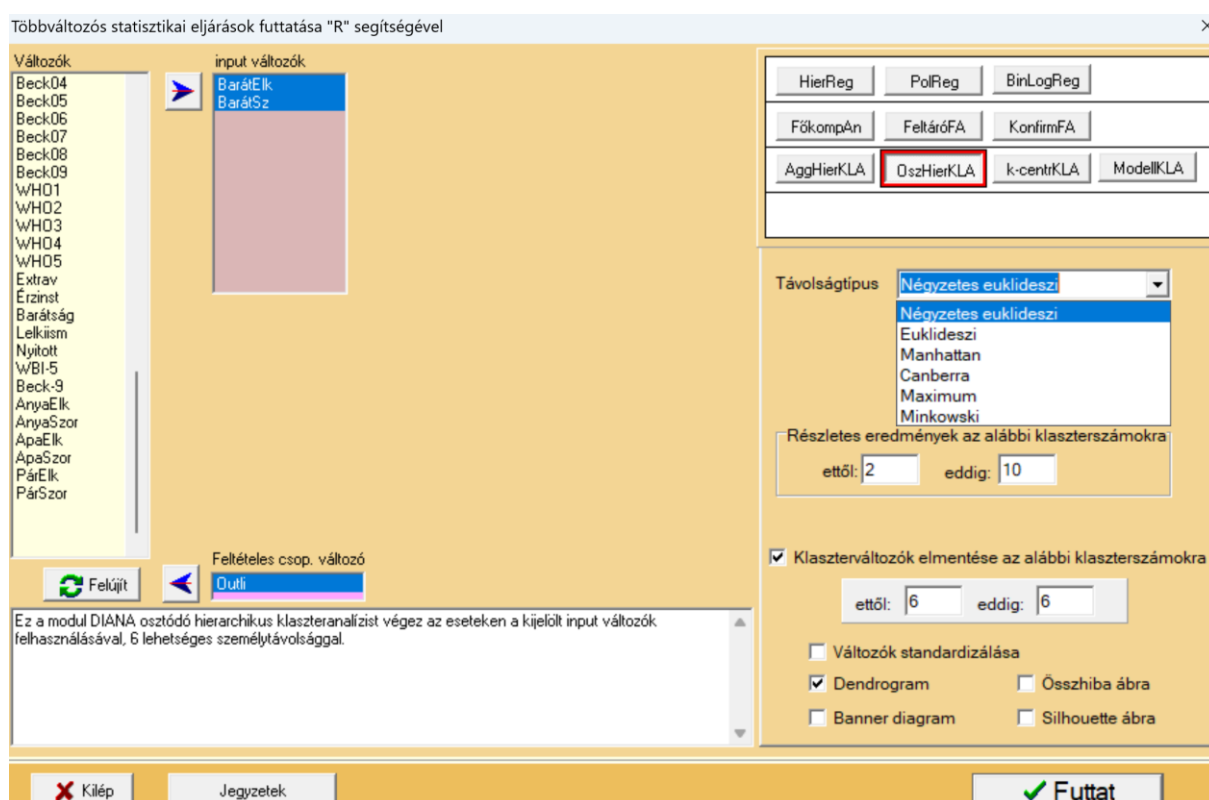
A legismertebb OHKA módszer a DIANA (DIvisive ANAlysis Clustering) névre keresztelt algoritmuson alapuló (Kaufman és Rousseeuw, 1990), melynek menete az alábbi:

- Első lépésben a minta összes egyede egy klaszterben van és ezt a következőképpen bontjuk két alklaszterre. Kiválasztjuk azt a személyt (egyedet), amelynek a többitől való átlagos távolsága a legnagyobb, s ezt egy különálló egyelemű klaszternek tekintjük. Ezután egyenként addig csatolunk ehhez az új klaszterhez újabb egyedeket a régi klaszterből, ameddig azok közelebb esnek az új klaszterhez, mint a maradék egyedek által alkotott régi klaszterhez.
- A következő lépésekben, amikor már több klaszterünk van, minden lépésben kiválasztjuk a klaszterek közül a legheterogénebbet (pl. amelynél a legnagyobb a *klaszter átmérője / cluster diameter*, az egymástól legtávolabbi két egyed távolsága), majd ugyanúgy bontjuk azt két alklaszterre, ahogy az első lépésben az összes egyedet tartalmazó nagy klasztert.
- Mindezt addig folytatjuk, amíg végül minden egyed egy-egy 1-elemű klaszterbe kerül.

Más algoritmusoknál a legheterogénebb klasztert máshogy azonosítják (pl. lehet ez a legnagyobb HC-értékű), vagy máshogy bontják fel két alklaszterre (pl. a k -közép klaszteranalízis algoritmusával; vö. 9. fejezet). OHKA segítségével általában a nagyobb klaszterek azonosíthatók jól, míg AHKA-t alkalmasabbnak tartják a kisebb méretű klaszterek azonosítására. AHKA több szoftverben elérhető és jóval ismertebb, mint OHKA. Például az SPSS, jamovi, JASP statisztikai szoftverek egyike sem tartalmaz az osztódó hierarchikus klaszteranalízis végrehajtására alkalmas modult.

8.2. Az OHKA menüablak

OHKA elemzést ROP-R-ben az osztódó hierarchikus klaszteranalízis (röviden OHKA) modul segítségével végezhetünk el. Ez a modul OHKA-t végez a mintabeli eseteken a kiválasztott változók felhasználásával. OHKA a *cluster* (Maechler et al., 2022), a *ggplot2* (Wickham, 2016) és a *factoextra* (Kassambara és Mundt, 2020) R-package-et használja fel elemzéseikhez. Az OHKA menüablak szolgál az elemzésbe bevonandó változók kiválasztására (input változók ablaka), a személytávolság (Távolságtípus) megválasztására, valamint egy sor opció kijelölésére. Az OHKA menüablak (lásd 8.1. ábra) csak annyiban különbözik AHKA menüablakától, hogy nincs benne választható módszer, ez ugyanis a rögzített DIANA.



8.1. ábra. Az OHKA modul menüablaka ROP-R-ben

Egy OHKA elemzés végrehajtása után a „c:_vargha\ropstat\aktualis” mappában megtaláljuk az elemzéshez elkészített ideiglenes adatfájlt (tmpdat.txt), a kért klaszterszámokhoz tartozó, klaszterváltozókkal kiegészített ideiglenes adatfájlt (tmpdat2.txt), a futtatott R-scriptet (DHCA.r), valamint a kért diagramokat jpg vagy pdf fájlban (pl. Dendr_1.jpg vagy WSSplot_1.pdf). Ha feltételes csoportosító változót is kijelölünk, akkor minden feltételes csoport elemzése során elkészülnek a kért diagramok, a diagramokat tartalmazó fájlok nevében megjelenő számok ezen csoportok sorszámát jelzik.

8.3. Mit tartalmaz az OHKA eredménylistája?

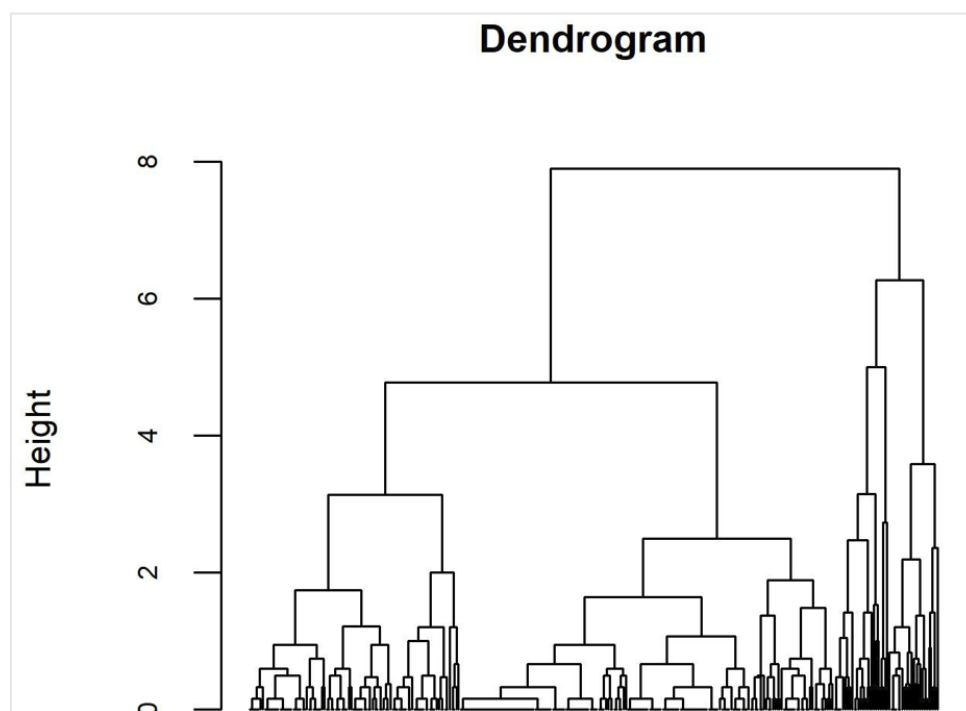
Az OHKA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- Minden klasztermegoldásra:
 - o Adekvációs mutatók (EESS%, HCátlag vagy HCátlagS, XBmod)
 - o Klaszterstatisztikák (elemszám, átlag, szórás, minimum, maximum)
 - o Nem standardizált átlagok (nem standardizált centroidok)
 - o Standardizált átlagok (standardizált centroidok)
 - o Standardizált átlagok mintázata
- Különböző klasztermegoldások adekvációs mutatóinak összefoglaló táblázata.

Ha kérjük a változók standardizálását, akkor az átlagok standardizálása változónként a saját átlaggal és szórással (z -standardizálás), ha pedig nem kérjük, akkor a változók összátlagával és együttes szórásával lesz végrehajtva (z_T -standardizálás).

8.4. Az OHKA szemléltetése valódi adatokon

Az OHKA elemzés szemléltetésére a B2.2. alponthban ismertetett kötődéskutatás KÖT2016 mintájának BarátElk és a BarátSz input változójával, az elemzést SED személytávolsággal és a változókat nem standardizálva végeztük el, ugyanúgy a 2–10 klaszterszámokra kérve részletes eredményeket, mint a 7.4. alfejezetben az AHKA elemzés során. Az ábrák közül most csak dendrogramot (lásd 8.2. ábra) kértünk, mert az összhiba és a Silhouette-ábra OHKA-ban mindig ugyanúgy néz ki, mint AHKA-ban (vö. 7.3. és 7.4. ábra).



8.2. ábra. Az OHKA eredményének ábrázolása dendrogram segítségével a ROP-R OHKA moduljában (nem standardizált változók, SED személytávolság)

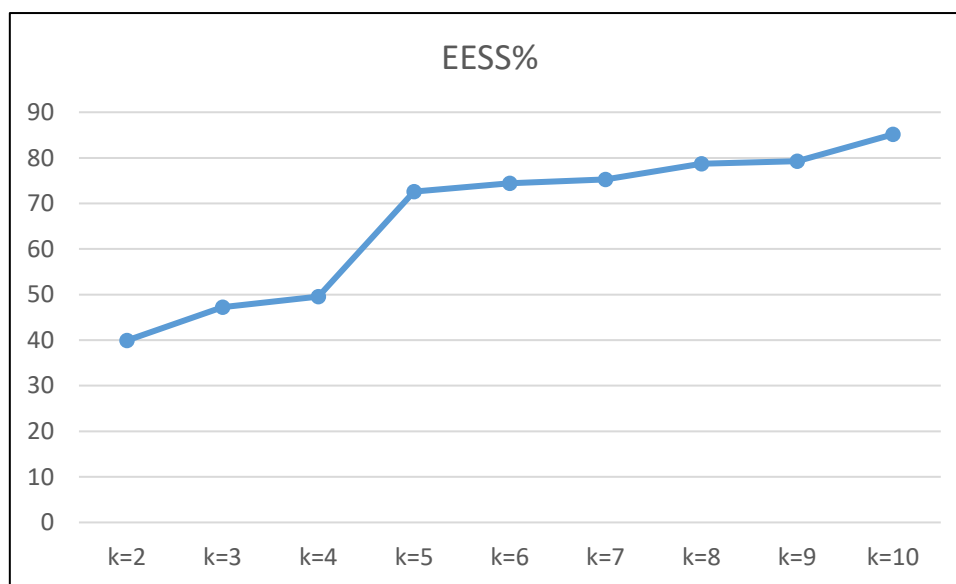
A 8.2. ábrán látható dendrogram kissé más képet mutat, mint AHKA esetén (vö. 7.2. ábra). Itt is oly módon dönthetünk egy optimális klaszterszámról, hogy megnézzük, milyen

szinten metszhetjük el vízszintes vonallal a dendrogramot úgy, hogy a metszési szint viszonylag alacsony legyen és a metszéspontok száma se legyen túl nagy. Jelen esetben egy elfogadhatóan alacsony (2,5-ös) szinten történő metszés már 8 klaszteres megoldást mutat.

Nézzük meg most is a különböző klasztermegoldások jellemzőit összefoglaló táblázatot az OHKA eredménylistájának alján, amelyet ROP-R a kért 2–10 klaszterszámokra készít el (lásd 8.1. táblázat). A klaszterszámok a 7.4. táblázattal ellentétben itt most növekedők, mert OHKA-ban a klaszterszám az osztódó struktúrának megfelelően lépésenként 1-gyel nő.

8.1. táblázat. Különböző klasztermegoldások jellemzői OHKA-ban (SED távolság, nem standardizált változók)

k	EESS%	XBmod	HCátlag	HCátlagS	HCminS	HCmaxS
2	40,23	0,808	1,479	1,202	1,031	2,149
3	47,50	0,724	1,302	1,058	1,031	1,341
4	49,79	0,724	1,248	1,014	0,705	1,254
5	72,59	0,765	0,688	0,559	0,474	1,254
6	74,43	0,781	0,644	0,523	0,474	1,254
7	75,29	0,764	0,624	0,507	0,438	1,254
8	78,68	0,579	0,542	0,440	0,310	1,254
9	79,24	0,590	0,523	0,425	0	0,779
10	85,17	0,707	0,378	0,307	0	0,779



8.3. ábra. Az EESS% mutató lejtődiagramja OHKA-ban

A 8.1. táblázat adatai azt mutatják, hogy a klaszterstruktúra homogenitása a lépésenkénti összevonások során – 2-től felfelé haladva – az 5-ös klaszterszámnál ér el először elfogadható szintet (EESS% itt kerül először, mégpedig igen nagy ugrással a 70%-os szint fölé, HCátlagS pedig meredek eséssel 1 alá, a 0,56-os szintre). A szeparációt mérő XBmod minden esetben a kívánatos 0,50-es szint felett van. Ha elkészítjük Excelben az EESS% lejtődiagramot (lásd 8.3. ábra), ez vizuálisan is megerősíti, hogy $k = 5$ alatt EESS% értékszintje elfogadhatatlanul

alacsony, ahogy az AHKA esetében is volt, $k = 5$ fölött viszont már sehol se látunk hirtelen ugrást, vagyis most is az 5- és a 6-klasztteres megoldás mintázatát érdemes alaposabban szemügyre venni. A fő kérdés az, hogy érdemes-e az 5-klasztteresről a finomabb 6-klasztteresre áttérnünk. Ennek eldöntésére megnéztük az OHKA outputján az átlagokkal és az együttes szórással standardizált átlagok mintázatának táblázatát az 5- és a 6-klasztteres részben (lásd 8.2. és 8.3. táblázat).

8.2. táblázat. A z_i -standardizált átlagok mintázata az OHKA 5-klasztteres megoldásában

Klaszter	BarátElk	BarátSz	KLgyak	HC	HCstan
KL1	M+	(A)	96	0,65	0,53
KL2	(A)	(A)	170	0,58	0,47
KL3	M+++	M+	21	0,87	0,71
KL4	.	M+++	23	1,30	1,06
KL5	M++++	M++++	4	1,54	1,25

8.3. táblázat. A z_i -standardizált átlagok mintázata az OHKA 6-klasztteres megoldásában

Klaszter	BarátElk	BarátSz	KLgyak	HC	HCstan
KL1	M+	(A)	96	0,65	0,53
KL2	(A)	(A)	170	0,58	0,47
KL3	M+++	M+	21	0,87	0,71
KL4	.	M++	18	0,63	0,51
KL5	M++++	M++++	4	1,54	1,25
KL6	(M)	M++++	5	0,96	0,78

A 8.2. és a 8.3. táblázatot összevetve megállapíthatjuk, hogy a 6-klasztteres megoldáshoz az 5-klasztteresből úgy jutunk, hogy ez utóbbinak kétfelé bontjuk a 23 fős KL4 klaszterét (egy 18 és egy 5 fős klaszterre, jelük a 6-klasztteres megoldásban KL4 és KL6). Mivel már maga a bontandó KL4 klaszter is igen kicsi, szakmailag nincs értelme az 5-klasztteres struktúrát tovább bontani. Ugyanezen okból szakmailag az 5-klasztteres megoldásban is csak a két nagyobb klaszter (KL1 és KL2) értelmezhető, amelyek a baráti elkerülés szintje tekintetében ellentétesek [(A) vs. M+], de a szorongás tekintetében egyformán enyhén átlag alattiak.

Ha az OHKA-ban kapott legjobbnak tűnő 5-klasztteres megoldást (lásd 8.2. táblázat) összevetjük az AHKA legjobbnak tűnő 6-klasztteres megoldásával (lásd 7.5. táblázat), a két legnagyobb klaszter (OHKA-ban KL1 és KL2, AHKA-ban KL1 és KL3) mintázata páronként nagy hasonlóságot mutat, kellően alacsony (0,55-nél kisebb) HCstan értékekkel, azaz megfelelő homogenitással. Ez egybevág azzal a korábbi megállapításunkkal, hogy OHKA elsősorban a nagyobb méretű klaszterek azonosításában sikeres, AHKA viszont képes olyan kisebb méretű, de szakmailag értelmezhető homogén klasztereket feltárni (lásd a KL2 és a KL5 klasztert a 7.5. táblázatban), amelyek OHKA-ban nem jelennek meg.

9. fejezet

A k -középpontú klaszteranalízis (KKA) modul

Ebben a fejezetben a k -középpontú vagy k -centrumú klaszteranalízis (KKA) módszerét és ROP-R-beli futtatásának mikéntjét ismertetjük. KKA ugyanúgy számos klaszteranalízis gyűjtőneve, ahogy HKA. A k -középpont jelző arra utal, hogy KKA-ban mindig egy előre megadott számú (k) klasztert hozunk létre, másrészt arra, hogy a klaszterek középpontjai, centrumai kiemelt fontosságú többdimenziós pontok, a klaszterbeli személyek reprezentatív értékmintázatai. Centrumként szóba jöhet a *többdimenziós átlag (centroid)*, a *geometriai medián*, valamint a *medoid* is. Minden KKA elemzés lényege, hogy először készítünk egy kezdeti felosztást k centrummal, majd egy többlépéses iterációs folyamat, az ún. *relokáció* segítségével addig javítjuk a klaszterstruktúrát a személyek egyik klaszterből a másikba átrakásával, amíg el nem érünk egy lehetséges maximumot a klaszterek összhomogenitása tekintetében. KKA-ban a személytávolság rögzített, a k -közép módszereknél egységesen SED (illetve a vele ekvivalens ASSED), a k -medián és a k -medoid módszernél pedig ED.

A 9.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük. A 9.2. alfejezet a KKA menüablak használatát mutatja be, a 9.3. alfejezet az eredménylista szerkezetét szemlélteti, a 9.4. alfejezet pedig a modullal végrehajtható statisztikai elemzésekre mutat példákat valódi adatok segítségével.

9.1. Elméleti alapok

Ebben az alfejezetben a KKA különböző módszereit ismertetjük.

9.1.1. A klasztercentrum kiválasztása

A KKA elemzések elsőként abban különböznek egymástól, hogy mi a választott klasztercentrum típusa. Attól függően, hogy ez a centrum a centroid, a geometriai medián vagy a medoid, beszélünk k -közép, k -medián vagy k -medoid módszerről. Egy p -dimenziós minta geometriai mediánjának azt a p -dimenziós pontot nevezzük, amelyiknek az átlagos euklideszi távolsága a mintabeli elemektől a legkisebb. A medoid pedig a többdimenziós mintának az az eleme, amelyik a legközelebb van (euklideszi távolsággal) a mintabeli többi egyedhez. A medoiddal ellentétben a centroid és a geometriai medián olyan többdimenziós pont, amelyik nem biztos, hogy fellep a minta valamelyik egyedénél. A k -medián és a k -medoid módszert főleg erősen ferde eloszlású, szélsőséges adatokat is tartalmazó minták esetén szokták az alapértelmezett k -közép módszer helyett javasolni.

9.1.2. A k -közép módszer

Minden k -közép elemzés minden iterációs lépésében minden személyt megpróbálunk úgy átrakni egy másik klaszterbe, hogy a klaszterstruktúra javuljon. Ha nem lehet, nem rakjuk át.

Az induló klasztercentrumokat úgy határozzuk meg, hogy véletlenszerűen kiválasztunk a mintából k számú személyt, s ezeket tesszük meg induló centrumoknak. Alkalmazott algoritmus alapján három k -közép típust szoktak megkülönböztetni:

- Hartigan-Wong
- MacQueen
- Lloyd (más néven Forgy).

Ezek részletes ismertetését lásd Vargha (2022, 136–137. o.). Röviden összefoglalva, a Hartigan-Wong algoritmus a klasztereken belüli összhibát próbálja meg minimalizálni (s egyben az EESS% megmagyarázott varianciaarányt maximalizálni), a másik kettő pedig a relokációk iterációs lépéseiben akkor tesz át egy személyt egy másik klaszterbe, ha az a másik klaszter centrumához közelebb van, mint saját klaszterének centrumához. A Lloyd algoritmus csak annyiban különbözik a MacQueen-félétől, hogy itt a klasztercentroidok újra számolása nem történik meg minden áthelyezésnél, csak egy-egy teljes iterációs lépéssorozat végén.

9.1.3. A k -medián módszer

A k -medián módszer algoritmus ugyanaz, mint a MacQueen-féle k -közép módszeré, annyi különbséggel, hogy klasztercentrumként nem a centroidot, hanem a geometriai mediánt, személytávolságként pedig nem SED-et, hanem a sima ED euklideszi távolságot használjuk.

9.1.4. A k -medoid módszer

A k -medoid módszer algoritmus ugyanaz, mint a Hartigan-Wong-féle k -közép módszeré. A különbség mindössze annyi, hogy a centroidok szerepét a medoidok veszik át, és nem a klasztereken belüli összhibát, hanem a saját klasztermedoidtól való személytávolságok összegét próbáljuk minimalizálni. Az alkalmazott személytávolság ROP-R-ben itt is a sima euklideszi ED (és nem SED, mint minden k -közép módszernél).

9.2. A KKA menüablak

A k -centrumú klaszteranalízis ROP-R-beli modulja (röviden KKA) többnyire ugyanazokat az R-package-eket (*stats*, *cluster*, *factoextra*, *ggplot2*) használja fel elemzéseikhez, mint AHKA és OHKA, de ha k -medián elemzést futtatunk, akkor még a Gmedian package-re is szüksége van (Cardot, 2022). Megjegyezzük, hogy a KKA elemzések eredménye – a kezdeti random besorolások miatt – a legtöbb esetben kismértékű random ingadozást mutathat, ezért érdemes ezeket többször futtatni, majd az eredmények közül a legjobbat kiválasztani.

A KKA modul lehetséges beállításait a menüablak 9.1. ábrán látható része mutatja be. Ezek közül csak az *Optimális klaszterszámhoz ábrakészítés* panel használata igényel részletesebb magyarázatot. Ezek az ábrák a menüablak jobb alsó paneljén kérhetők (lásd 9.1. ábra), ahol egy elemzéshez mindig egyetlen ábra választható az alábbi négy közül:

- Silhouette
- EESS%
- Átlagos heterogenitás
- $f(K)$ torzulás.

Megjegyezzük, hogy ezen ábrák kérésekor a maximális klaszterszám 5 és 20 közötti értékre állítható be.

The screenshot shows the 'Módszer' (Method) section with three radio buttons: 'k-közép elemzés' (selected), 'k-medoid elemzés', and 'k-medián elemzés'. To the right is the 'Algoritmus' (Algorithm) section with three radio buttons: 'Hartigan-Wong' (selected), 'MacQueen', and 'Lloyd/Forgy'.

Below these is a section for 'Klaszterek száma' (Number of clusters) with a text box containing '3'. Underneath are three unchecked checkboxes: 'Végző klaszterváltozó mentése', 'Végző klaszterek ábrázolása', and 'Változók standardizálása'.

Next is 'Iterációk maximális száma' (Maximum number of iterations) with a text box containing '25' and a range '(1 - 200)'.

Below that is a checked checkbox 'Optimális klaszterszámhoz ábrakészítés' (Create plots for optimal cluster number). Underneath is a section for 'Maximális klaszterszám' (Maximum cluster number) with a text box containing '10' and a range '(5 - 20)'.

Finally, there is a 'Klaszterszám meghatározása' (Determine cluster number) section with four radio buttons: 'Silhouette' (selected), 'EESS%', 'Átlagos heterogenitás', and 'f(K) torzulás'.

9.1. ábra. A KKA modul lehetséges beállításai ROP-R-ben

A Silhouette- és az EESS%-ábra értelmezése nem igényel magyarázatot. Az *átlagos heterogenitás* a klasztereken belüli összes páronkénti személytávolság átlaga euklideszi távolsággal, egyfajta HCátlag logikájú mutató, az *f(K) mutató* pedig egy random egyenletes eloszláson végzett KKA eredményéhez viszonyítja a kapott struktúrát (vö. Pham, Dimov és Nguyen, 2005). Olyan – lehetőleg alacsony – k érték tekinthető optimálisnak, amelyre a Silhouette-együttható (SCorig), illetve EESS% értéke nagy, az átlagos heterogenitás és $f(K)$ értéke pedig kicsi, utóbbi lehetőleg 0,85-nél is kisebb. Megjegyezzük, hogy a SCorig torzít a kis k értékek felé, pusztán azon okból, hogy kevés klaszter jobban el tud különülni egymástól, mint sok. Jó tudni, hogy ezek az ábrák mind k -közép elemzésekkel készülnek, és nem függenek attól, hogy a menüablakban milyen KKA módszert állítunk be.

Egy KKA elemzés végrehajtása után a szokásos „aktualis” mappában találjuk az elemzéshez elkészített ideiglenes adatfájlt (tmpdat.txt), a kért klaszterszámhoz tartozó, klaszterváltozóval kiegészített ideiglenes adatfájlt (tmpdat2.txt), a futtatott R-scriptet (KCA.r), valamint az optimális klaszterszám meghatározásához kért diagramokat jpg vagy pdf fájlban. ROP-R a választott ábrát a KKA futtatása után az ábra sorszámának megfelelően optk1_1.jpg (Silhouette), optk2_1.jpg (EESS%), optk3_1.jpg (Átlagos heterogenitás) vagy optk4_1.jpg (f(K) torzulás) néven menti képfájlba.

Ha feltételes csoportosító változót is kijelölünk, akkor minden feltételes csoport elemzése során elkészülnek a kért diagramok, a diagramokat tartalmazó fájlok nevében megjelenő számok ezen csoportok sorszámát jelzik. A „Végző klaszterek ábrázolása” opciót bejelölve ROP-R a k -közép, k -medoid, illetve k -medián módszer esetén egy kmeansplot.jpg, kmedoidsplot.jpg, illetve kmediansplot.jpg nevű fájlban menti el a kapott klaszterstruktúra

ábráját. Ha az elemzésre kiválasztott változók száma 2-nél nagyobb, ROP-R a klasztereket az elemzésbe bevont változók FKA elemzéséből kapott első két főkomponens terében ábrázolja.

9.3. Mit tartalmaz a KKA eredménylistája?

A KKA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- A megadott klaszterszámhoz tartozó megoldásra:
 - o Adekvációs mutatók (EESS%, HCátlag/HCátlagS, XBmod, SCorig, SCsimple)
 - o Klaszterstatisztikák (elemszám, átlag, szórás, minimum, maximum)
 - o Nem standardizált centrumok (átlagok, medoidok vagy geometriai mediánok)
 - o Standardizált centrumok (átlagok, medoidok vagy geometriai mediánok)
 - o Standardizált centrumok mintázata.

Ha kérjük a változók standardizálását, akkor a centrumok standardizálása változónként külön-külön a saját átlaggal és szórással (z -standardizálás), ha pedig nem kérjük, akkor egységesen a változók átlagával és együttes szórásával lesz végrehajtva (z_i -standardizálás).

9.4. A KKA szemléltetése valódi adatokon

Ebben az alfejezetben egyrészt a KÖT2016 kötődéskutatási adatokon (lásd B2.2. alpont) a barátai kapcsolatos elkerülés (BarátElk) és szorongás (BarátSz) skálájával mint input változóval végzünk klaszterelemzéseket: a 9.4.1. alpontban k -közép, a 9.4.2. alpontban pedig k -medián és k -medoid elemzést. A 9.4.3. alpontban a PIK16 kérdőív négy skálájával (Mmonit, AVhat, Rezil és Önreg; vö. B2.3. alpont) végzünk k -közép elemzést.

9.4.1. k -közép elemzés két kötődésváltozóval

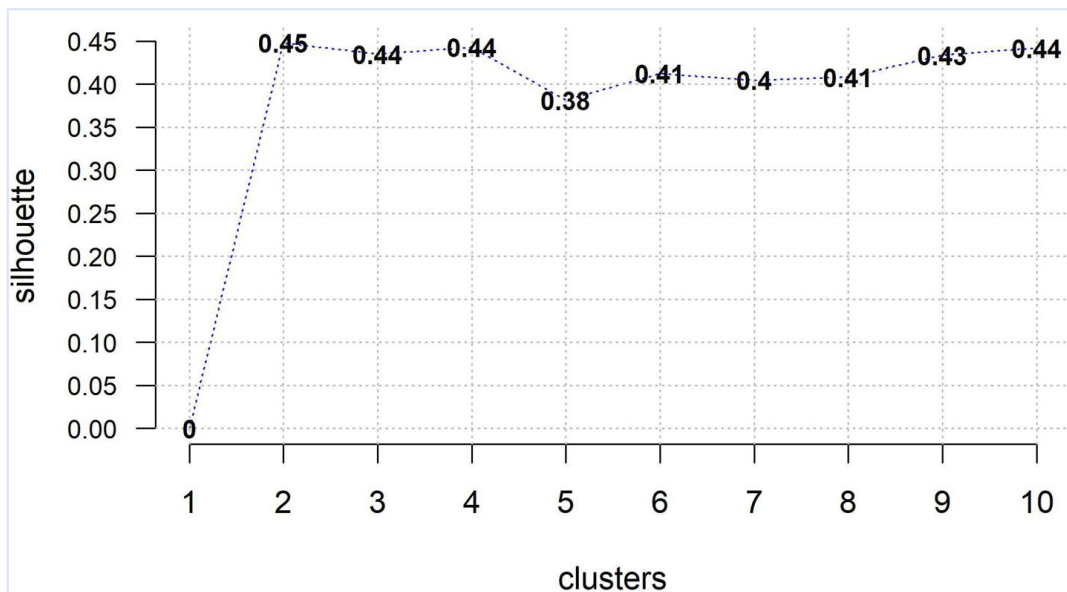
A k -közép elemzést először az alapértelmezés szerinti Hartigan-Wong algoritlussal hajtottuk végre, az AHKA és OHKA elemzés tapasztalatait is felhasználva a 4–6 klaszterszámokra (vö. 7.4. és 8.4. alfejezet). Mivel BarátElk és BarátSz értékei egyaránt az [1–7] tartományban változhatnak (vö. B.4. táblázat), a változókat nem standardizáljuk, bár BarátSz átlaga (1,8) észlelhetően alacsonyabb, mint a BarátElk (2,4) skáláé. Ilyen esetben érdemes lehet az elemzéseket standardizált változókkal is elvégezni és a kapott eredményeket egymással összevetni.

A klaszterelemzést a $k = 6$ klaszterszámmal kezdtük, a 8 outlierrel csökkentett 314 fős KÖT2016 mintán (vö. 7.4. alfejezet). KKA-t négyszer futtattuk, rendre más optimumkereső ábrát beállítva, amelyeket a futtatások után a szokásos „aktualis” mappában találtunk meg (optk1_1.jpg, ..., optk4_1.jpg néven, lásd 9.2–9.5. ábrák). Mivel a k -közép elemzés véletlen kiválasztással kezdődik, különböző futások gyakran különböző végeredményre vezetnek. Ez okból mind a négy futás klaszterváltozóját elmentettük, s közülük a legjobbat tartottuk meg.

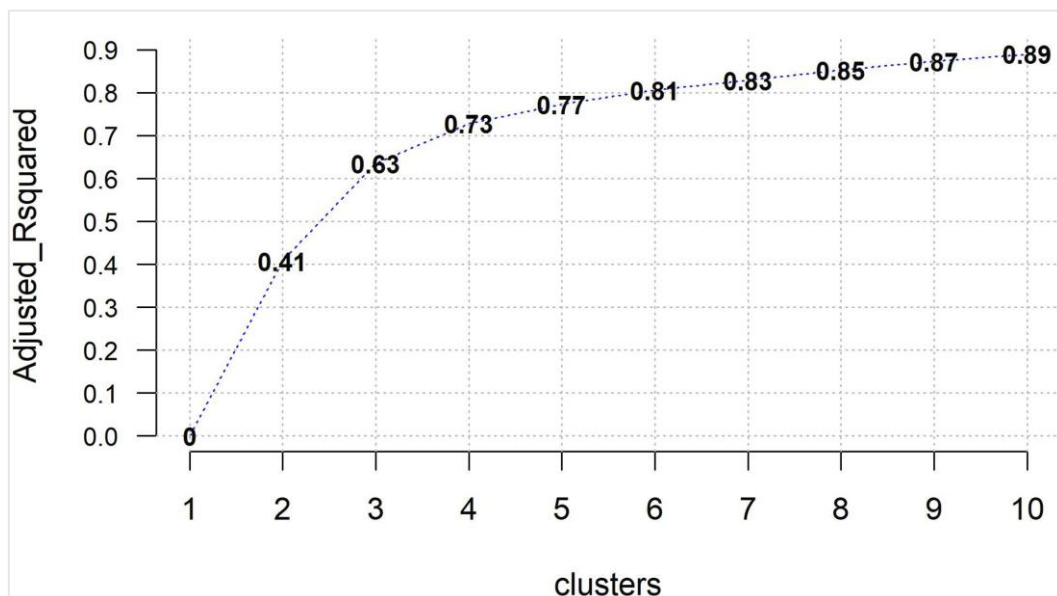
A kapott ábrákat megismerve az alábbi megállapításokat tehetjük.

1. A Silhouette-ábrán az SCorig értékek végig a 0,38–0,45 szűk sávban mozognak, így itt egyetlen markáns kiugrást sem találunk (lásd 9.2. ábra).

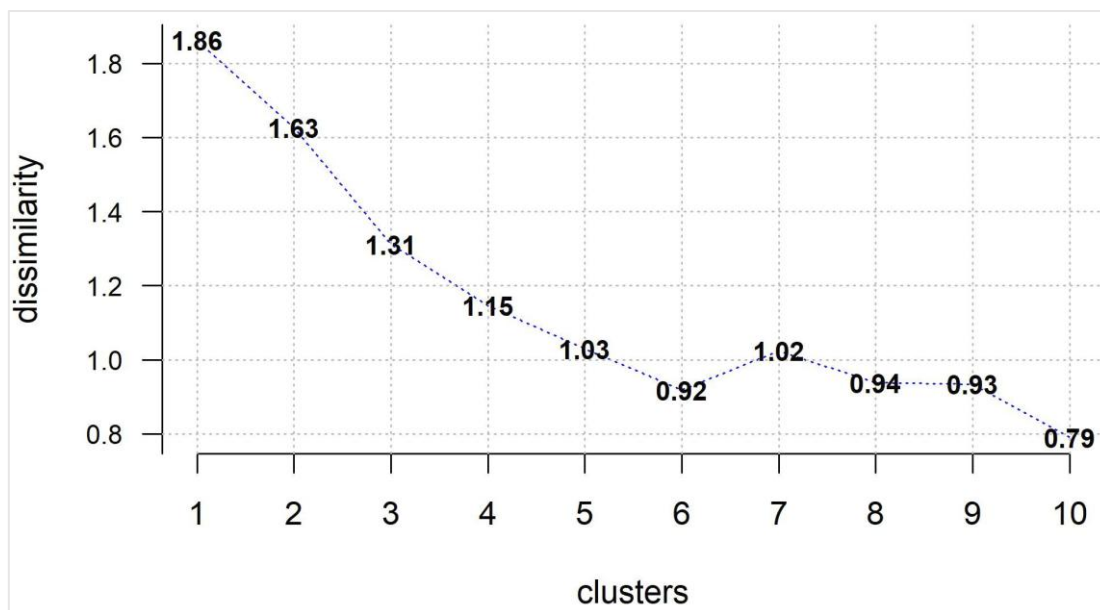
2. A 9.3. ábra azt mutatja, hogy az EESS% (100-zal beszorozva az EESS arányt) már $k = 4$ esetén meghaladja a 70-es szintet, de nagy ugrások híján még nem tudjuk, hogy 4 és 6 között melyik klaszterszám vezet a legjobb eredményre.
3. Az átlagos heterogenitás 9.4. ábráján (az ábrán *dissimilarity*) $k = 6$ értéknél van egy határozott helyi minimum az átlagos klaszteren belüli különbözöségre, ami megerősíti az AHKA-ban is optimálisnak tűnő 6-os klaszterszámot.
4. Végül az $f(K)$ -torzulás ábrája (lásd 9.5. ábra) $k = 3$ mellett teszi le a voksot, de ezt is bizonytalanul, mert $k = 3$ esetén az $f(K)$ -érték messze nem esik a kívánatos 0,84-es szint alá. Emellett $k = 3$ esetén EESS% értéke olyan alacsony (63%, vö. 9.3. ábra), hogy ez a megoldás semmiképpen nem elfogadható.



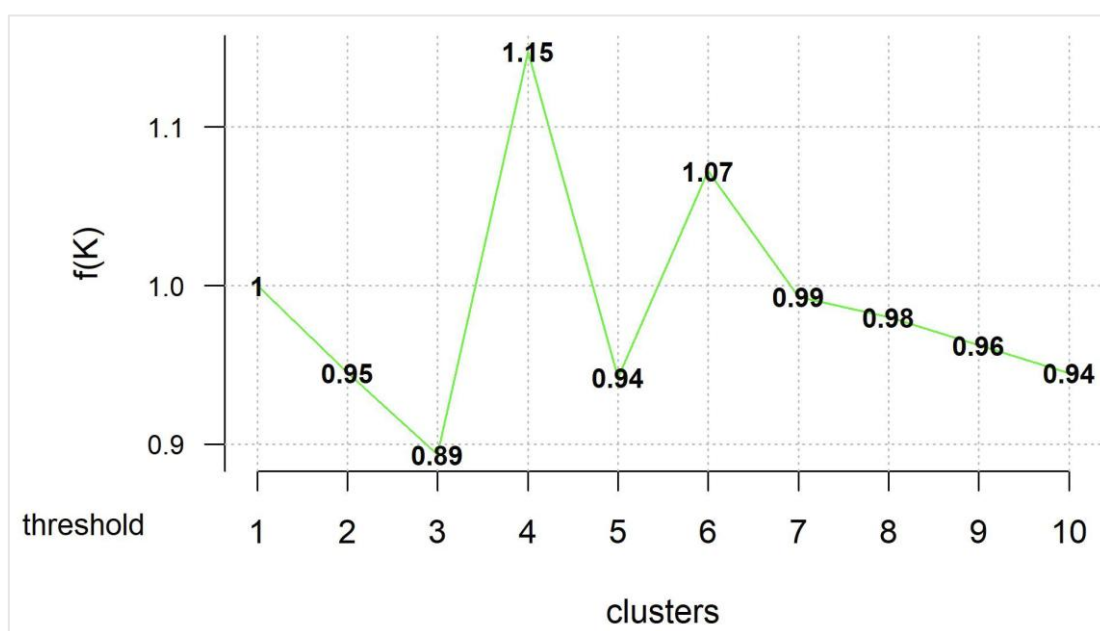
9.2. ábra. A Silhouette értékek függése a klaszterszámtól a k -közép módszer esetén



9.3. ábra. Az EESS-arány függése a klaszterszámtól a k -közép módszer esetén



9.4. ábra. Az átlagos heterogenitás függése a klaszterszámtól a k -közép módszer esetén



9.5. ábra. Az $f(K)$ -torzulás függése a klaszterszámtól a k -közép módszer esetén

Bár az ábrák nem jeleznek egyértelmű értéket az optimális klaszterszámról, indikációik alapján az optimális megoldást most is 4 és 6 közötti k értékekre kell keresnünk. Az elemzéseket tehát ezekre a k értékekre végeztük el (minden esetben 4 futásból kiválasztva a legmagasabb EESS% értékűt, mint legjobbat). Elsőként a klaszterstruktúrát jellemző főbb adekvációs értékeket hasonlítottuk össze (lásd 9.1. táblázat). A szokásos QC-k közül az összhomogenitást legjobban mérő EESS%, a szeparációra is érzékeny QC-k közül XBmod és SCorig szerepelt, a klaszterek homogenitását pedig standardizált HC értékük átlagával (HCátlagS), minimumával (HCminS) és maximumával (HCmaxS) mértük. Ezen kívül az eredménylista homogenitás százalékok táblázatából kigyűjtöttük a „HCstan < 0,5” adatot is,

mely azt jelzi, hogy a 0,5-nél kisebb HCstan értékű (tehát jó homogenitású) klaszterekbe a teljes minta összesen hány százaléka esik (lásd 9.1. táblázat).

9.1. táblázat. Különböző k -közép klasztermegoldások főbb jellemzői

klaszter-szám	EESS%	XBmod	SCorig	HCátlagS	HCminS	HCmaxS	HCstan < 0,5
4	73,1	0,754	0,595	0,546	0,22	1,77	40,5%
5	77,7	0,502	0,518	0,455	0,27	1,52	58,3%
6	81,5	0,621	0,546	0,380	0,11	1,27	86,3%

A 9.1. táblázat alapján a következő megállapításokat tehetjük.

1. A táblázat SCorig értékei nem egyeznek meg a 9.2. ábra Silhouette értékeivel. A ROP-R KKA modulja SCorig-ot a *cluster* package *silhouette* függvénye segítségével számítja ki, míg a 9.2. ábrát a *ClusterR* package *jpeg* függvénye segítségével. Mindenesetre közös bennük, hogy értékük $k = 4$ esetén a legnagyobb és $k = 5$ esetén a legkisebb.
2. Az összhomogenitás már $k = 4$ esetén is elfogadható (EESS% = 73,1), de $k > 4$ esetén több tekintetben különösen tetszetős (pl. HCátlagS < 0,50).
3. A $k = 6$ megoldás mutatói minden tekintetben jobbak, mint a $k = 5$ megoldásé. A homogenitás mutatók (EESS%, HC mutatók) esetében ez persze nem meglepő, de ugyanezt látjuk a szeparációs mutatók (XBmod és SCorig), valamint az $f(K)$ -torzulás esetén is (lásd 9.5. ábra), továbbá a 6-klaszteres struktúrát támogatja az is, hogy ebben az esetben a minta 86,3%-a 0,5-nél kisebb HCstan értékű, tehát erősen homogén klaszterbe tömörül (lásd a 9.1. táblázat utolsó oszlopát), míg $k = 5$ esetén ez az érték csupán 58,3%.

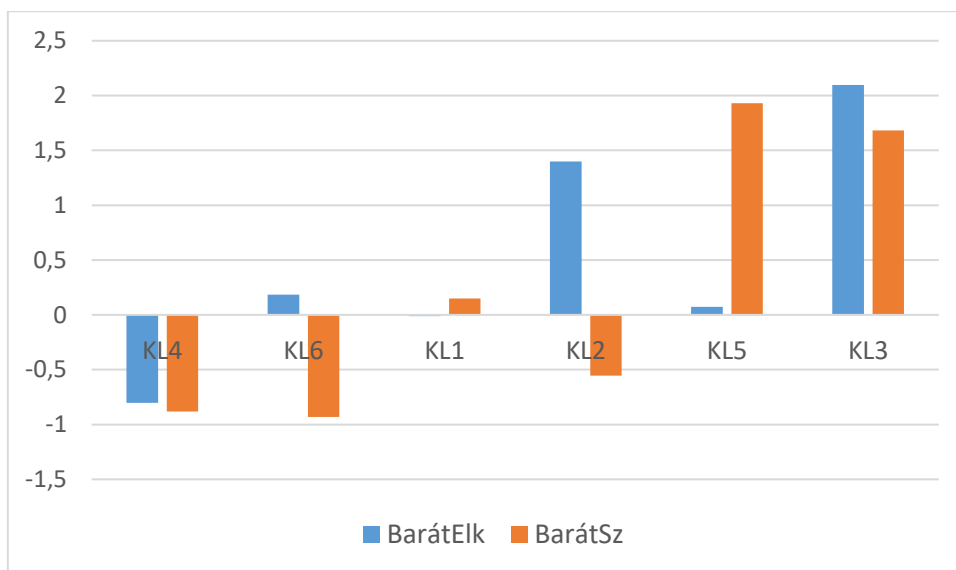
Mindezek alapján a 6-klaszteres k -közép megoldást fogadjuk el optimálisként. Az ehhez tartozó z_r -standardizált átlagok mintázata a 9.2. táblázatban látható, melyet a könnyebb áttekinthetőség érdekében az alábbiak szerint formáltunk át.

- Minden klaszterre kiszámítottuk a két változó z_r -standardizált átlagának átlagát (z_r -átlag; lásd a táblázat utolsó oszlopát). Negatív z_r -átlagok (KL4 és KL6) a jó kötődés, pozitív z_r -átlagok (KL2, KL5 és KL3) pedig a rossz kötődés jelei.
- A klasztereket a z_r -átlag értéke szerint növekvő rendbe raktuk. Eszerint a táblázat első sora a legjobb, utolsó sora pedig a legrosszabb kötődésű klaszter adatait tartalmazza.

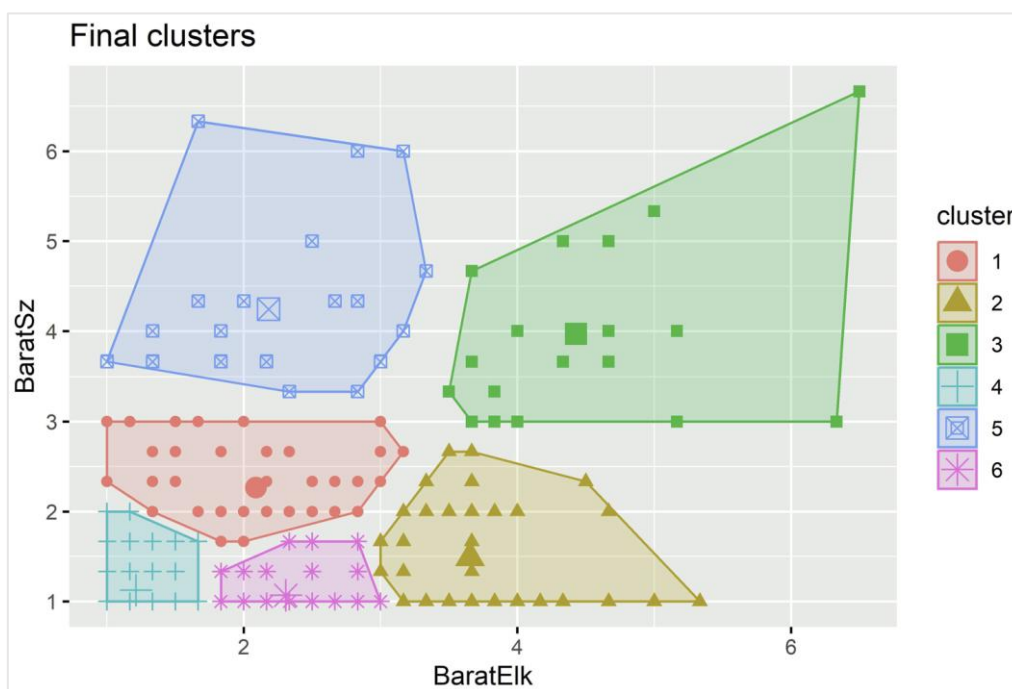
9.2. táblázat. A z_r -standardizált átlagok mintázata a 6-klaszteres k -közép megoldásban (M=Magas, A=Alacsony, a + jelek száma a szélsőségesség mértékét jelzi)

Klaszter	BarátElk	BarátSz	KLgyak	HCstan	SCorig	SCsimple	z_r -átlag
KL4	A	A	80	0,11	0,76	0,83	-0,84
KL6	.	A	68	0,13	0,66	0,74	-0,37
KL1	.	.	60	0,40	0,33	0,62	0,07
KL2	M+	(A)	63	0,45	0,45	0,72	0,42
KL5	.	M+++	22	1,03	0,35	0,63	1,00
KL3	M++++	M+++	21	1,27	0,46	0,75	1,89
			átlag:	0,380	0,546	0,727	

A klaszterek könnyebb értelmezéséhez két ábrát is készítettünk. Az elsőt a z_i -standardizált klaszterátlagok segítségével Excelben (lásd 9.6. ábra), a másodikat pedig ROP-R KKA moduljával, a KKA menüablakban bejelölve a „Végső klaszterek ábrázolása” opciót (lásd 9.7. ábra). Ez utóbbi a k -közép elemzés során a szokásos „aktualis” alkönyvtárban jön létre kmeans6plot1.jpg néven (értelemszerűen jpg formátumú képfájlként). Ha valamely rögzített klaszterszámmal (pl. $k = 6$) egymás után több elemzést is elvégzünk (pl. variálva a k -közép elemzés algoritmusát) és mindegyikhez szeretnénk ilyen ábrázolást kapni, akkor célszerű a létrehozott jpg fájlt minden elemzés után átnevezni, mert ha ROP-R talál az alapértelmezett névvel már ilyen fájlt, azt automatikusan felülírja.



9.6. ábra. A 6-klaszteres k -közép megoldás Exceles ábrázolása (a klasztereket a kötődés jósága szerint növekvő sorrendbe raktuk)



9.7. ábra. A 6-klaszteres k -közép megoldás ROP-R-es ábrázolása

A 9.2. táblázat és a fenti ábrák alapján a következő megállapításokat tehetjük.

1. A leghomogénebb ($HC_{stan} = 0,11$) és legbővebb ($n = 80$) klaszter a biztonságosan kötődők típusaként azonosítható $[A, A]$ – vagy más jelöléssel $[Elk-Sz-]$ ⁴³ típus (KL4). Ez a 9.6. ábrán bal oldalt, a 9.7. ábrán a bal alsó sarokban, a 9.2. táblázatban pedig a legfelső sorban látható. Jellemzője még az erősen negatív ($-0,8$ alatti) z_I -átlag és a többi klasztertől való erős szeparáció (SC_{orig} és SC_{simple} itt a legnagyobb).
2. Egy másik markáns mintázatú klaszter (KL3) a félelemteli-elkerülő kötődők típusát képviseli $[M+++]$ vagy $[Elk+Sz+]$ jelöléssel. Ez a 9.6. ábrán a jobb oldalon, a 9.7. ábrán a jobb felső régióban, a 9.2. táblázatban pedig a legalsó sorban található. Jellemzője az erősen pozitív ($1,8$ feletti) z_I -átlag, a kis elemszám ($n = 21$) és az erős heterogenitás ($HC_{stan} = 1,27$). Ez is jól elkülönül a többi klasztertől (lásd 9.7. ábra), amit a $0,50$ -hez közeli SC_{orig} és a $0,75$ -öt elérő magas SC_{simple} érték is jelez.
3. Karakteres a KL2 klaszter által képviselt $[M+, (A)]$, azaz $[Elk+Sz-]$ elutasító-elkerülő típus is, mely az Elkerülés tekintetében hangsúlyosan magas, a Szorongás tekintetében pedig átlag alatti (vö. Jantek és Vargha, 2016). Ez a homogenitás ($HC_{stan} = 0,45$) és a többi klasztertől való elkülönülés ($SC_{orig} = 0,45$, $SC_{simple} = 0,72$) tekintetében egyaránt jó konstrukció, mely összességében átlag alatti kötődést jelez (z_I -átlag = $0,42$).
4. Ennek ellentéte az elárasztott-megszállott típust képviselő KL5 klaszter $[., M+++]$ vagy $[Elk0Sz+]$ jelöléssel, ahol magas szorongáshoz átlagos szintű elkerülés társul. Ez hasonlít KL3-hoz abban, hogy eléggé heterogén ($HC_{stan} = 1,03$), mérete kicsi ($n = 22$) és gyenge kötődést jelez (z_I -átlag = $1,00$).
5. KL4-hez némileg hasonlít a KL6 klaszter, ahol az alacsony szorongás átlagos szintű elkerüléssel társul, ily módon ez is az átlagosnál egyértelműen jobb kötődést jelez (z_I -átlag = $-0,37$). Ez az $[Elk0Sz-]$ jelölésű klaszter a 9.6. és a 9.7. ábrán közvetlenül K4 mellett található, szintén jó homogén ($HC_{stan} = 0,11$) és méretes ($n = 68$) klaszter, mely jól elkülönül a többi klasztertől ($SC_{orig} = 0,66$).
6. Végül van még egy mindenben átlagoshoz típus (KL1), $[Elk0Sz0]$ jelöléssel, mely méretes ($n = 60$) és KL2-höz hasonlóan jobb homogenitású ($HC_{stan} = 0,40$).

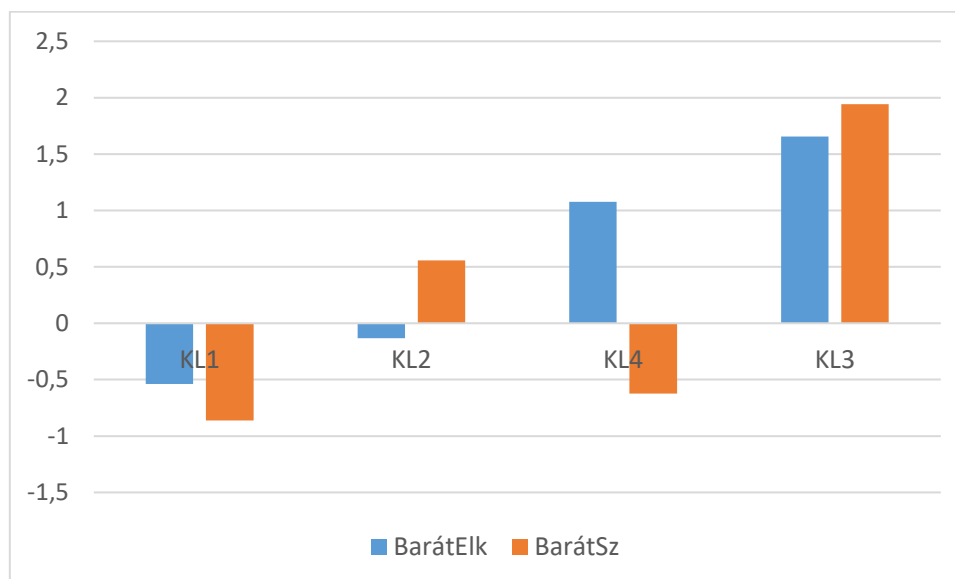
Ebben az igen jó statisztikai mutatókkal rendelkező 6-klasztteres k -közép megoldásban megtaláljuk a Fraley és Shaver (2000) nevéhez fűződő kötődéselméletnek mind a 4 típusát, ami jó. Új elemként azonban megjelenik a biztonságosan kötődő típus egyik variánsaként az átlagos elkerülés – alacsony szorongás típusa (KL6), valamint a mindenben átlagoshoz típus (KL1), ami felvetheti a Fraley-Shaver kötődéselmélet finomításának igényét. Ennek eldöntésére azonban mintánk nem elég nagy, egy stabilabbnak tekinthető klaszterstruktúra feltárásához legalább 500-600 fős mintára lenne szükség.

Annak eldöntésére, hogy a kapott 6-klasztteres megoldásban melyek a legstabilabb típusok, érdemes megnézni a 4-klasztteres megoldást és összevetni a 6-klasztteresével. E célból elkészítettük a 4-klasztteres megoldás z_I -standardizált átlagainak mintázatát a 9.2. táblázatnak megfelelő formátumban (lásd 9.3. táblázat), valamint ennek Exceles ábrázolását (lásd 9.8. ábra).

⁴³ Elk=Elkerülés, Sz=Szorongás, $[Elk-Sz-]$ = alacsony elkerülés, alacsony szorongás

9.3. táblázat. A z_i -standardizált átlagok mintázata a 4-klaszteres k -közép megoldásban (M=Magas, A=Alacsony, a + jelek száma a szélsőségeség mértékét jelzi)

Klaszter	BarátElk	BarátSz	KLgyak	HCstan	SCorig	SCsimple	z_i -átlag
KL1	(A)	A	127	0,22	0,80	0,84	-0,70
KL2	.	(M)	57	0,64	0,35	0,60	0,21
KL4	M+	(A)	100	0,53	0,54	0,72	0,23
KL3	M+++	M+++	30	1,77	0,39	0,72	1,80
			Átlag:	0,546	0,595	0,746	



9.8. ábra. A 4-klaszteres k -közép megoldás Exceles ábrázolása (a klasztereket a kötődés jósága szerint növekvő sorrendbe raktuk)

A 9.3. táblázat és a 9.8. ábra alapján megállapítható, hogy a 4-klaszteres k -közép megoldás lényegében megerősíti a Fraley-Shaver-féle kötődésméletet, csak kissé gyengébb statisztikai mutatókkal, mint amelyek a 6-klaszteres megoldást jellemzik. Az itt feltárt négy típus klasztere mind megfeleltethető a 6-klaszteres megoldás egy-egy klaszterének, ami jelzi stabilitásukat. A 6-klaszteres megoldás két új típusának megerősítéséhez egy nagyobb mintájú, új vizsgálat szükséges.

9.4.2. k -medián és k -medoid elemzés két kötődésváltozóval

Figyelembe véve a BarátSz változó igen magas ferdeségét (lásd B.4. táblázat) felmerülhet, hogy a k -medoid vagy a k -medián KKA módszer jobban strukturált, homogénebb megoldásra vezet. Ezért KKA elemzést végeztünk ugyanazon a 314 fős KÖT2016 mintán, mint a 9.4.1. alpontban, csak nem a k -közép, hanem a k -medián és a k -medoid módszerrel, $k = 4$ klaszterszámmal. Mindkét módszer esetén háromszor futtattuk le ugyanazt az elemzést (mindig elmentve a kapott klaszterváltozót), és az eredmények közül a legjobbat választottuk ki végső értékelésre. Megjegyezzük, hogy a k -medoid elemzés eredménye mindhárom esetben ugyanaz volt és a k -mediáné is csak csekély variabilitást mutatott.

A globális áttekintéshez először elkészítettük a különböző klasztermegoldások legfontosabb jellemzőit tartalmazó összefoglaló táblázatot (lásd 9.4. táblázat), az összehasonlíthatóság kedvéért feltüntetve a 4-klaszteres k -közép elemzés megfelelő adatait is.

9.4. táblázat. Különböző 4-klaszteres megoldások jellemzői KKA-ban

klaszter-szám	EESS%	XBmod	SCorig	HCátlagS	HCminS	HCmaxS	HCstan < 0,5
k -közép	73,1	0,754	0,595	0,546	0,22	1,77	40,5%
k -medián	72,7	0,749	0,442	0,554	0,22	1,82	70,7%
k -medoid	68,0	0,295	0,344	0,649	0,17	1,90	61,5%

A 9.4. táblázat alapján az alábbi észrevételeket tehetjük.

- A k -medián módszer QC értékei alig térnek el a k -közép módszerétől, viszont a 0,5-nél homogénebb klaszterbe tartozó személyek aránya (70,7%) jelentősen nagyobb, mint a 4-klaszteres k -közép megoldás esetén (40,5%).
- A k -medoid módszer QC értékei érezhetően gyengébbek, mint a k -medián és a k -közép megoldásáé, viszont a 0,5-nél homogénebb klaszterbe tartozó személyek aránya (61,5%) itt is érezhetően nagyobb, mint a 4-klaszteres k -közép megoldásé (40,5%).
- A fentiek alapján megállapíthatjuk, hogy a k -medián és a k -medoid módszer segíthet homogénebb klasztereket feltárni, mint a k -közép elemzés.

A három módszer eredményeinek tartalmi összehasonlításához meg kell néznünk a z_i -standardizált átlagok mintázatát is a különböző megoldásokban. A z_i -standardizált geometriai mediánok és medoidok mintázatát a 9.5. és a 9.6. táblázat tartalmazza, melyeket a k -közép módszerrel való összehasonlíthatóság érdekében ugyanolyan formátumban készítettünk el, mint a 9.3. táblázat esetében.

9.5. táblázat. A z_i -standardizált geometriai mediánok mintázata a 4-klaszteres k -medián megoldásban (M=Magas, A=Alacsony, a + jelek száma a szélsőségesség mértékét jelzi)

Klaszter	BarátElk	BarátSz	KLgyak	HCstan	SCorig	SCsimple	z_i -átlag
KL4	(A)	A	125	0,22	0,597	0,818	-0,72
KL1	M	A	97	0,49	0,409	0,751	0,15
KL3	.	(M)	58	0,64	0,267	0,61	0,16
KL2	M+++	M++	34	1,82	0,266	0,696	1,61
			Átlag:	0,554	0,442	0,746	

9.6. táblázat. A z_i -standardizált medoidok mintázata a 4-klaszteres k -medoid megoldásban (M=Magas, A=Alacsony, a + jelek száma a szélsőségesség mértékét jelzi)

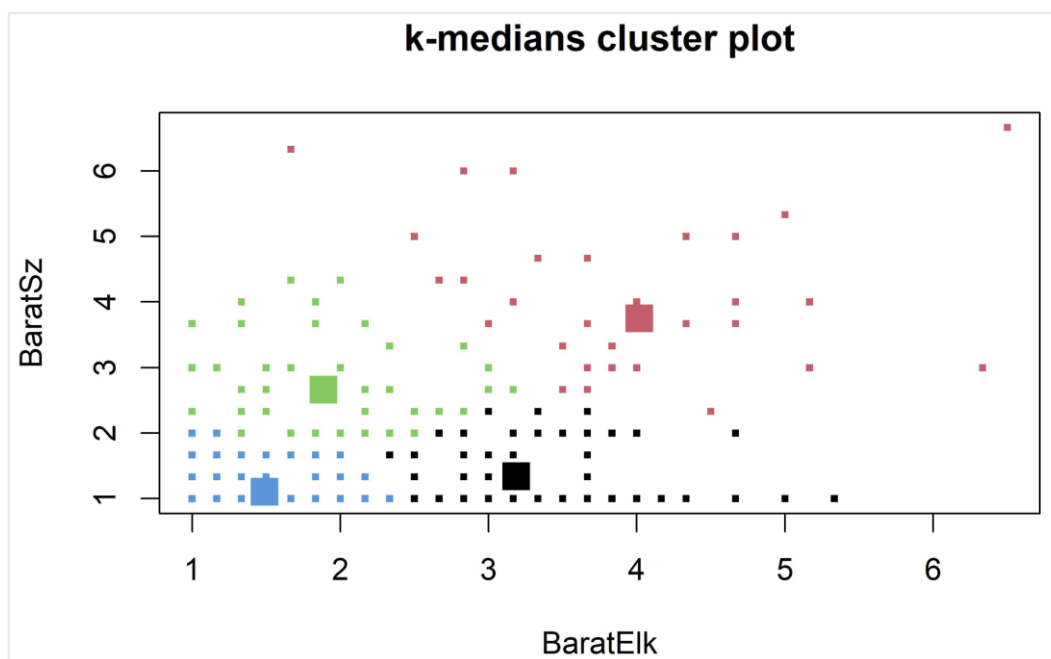
Klaszter	BarátElk	BarátSz	KLgyak	HCstan	SCorig	SCsimple	z_i -átlag
KL2	A	A	72	0,17	0,63	0,80	-0,99
KL4	.	A	121	0,38	0,23	0,63	-0,32
KL1	M++	A	65	0,59	0,38	0,72	0,36
KL3	M	M++	56	1,90	0,19	0,64	1,11
			Átlag:	0,649	0,344	0,689	

A k -medián elemzés eredményét tartalmazó 9.6. táblázatot összevetve a k -közép elemzés 4-klaszteres eredményét tartalmazó 9.3. táblázattal azt láthatjuk, hogy mind a négy klaszter jól megfeleltethető egymásnak, az egyezés teljes mértékű. Ezt erősíti a Jaccard index (0,940) és a korrigált Rand index (0,956) magas értéke, valamint a cellaegyezés arányok igen magas szintje (lásd 9.7. táblázat) a két megoldás összehasonlításában, amiket a ROPstat Exacon moduljával lehet kiszámítani. Ugyanakkor a k -medoid elemzés eredményeképpen kapott z_i -standardizált átlagok mintázata jelentősen eltér a k -közép elemzés 4-klaszteres eredményétől.

9.7. táblázat. Cellaegyezés arányok a 4-klaszteres k -közép és k -medián megoldás összehasonlításában

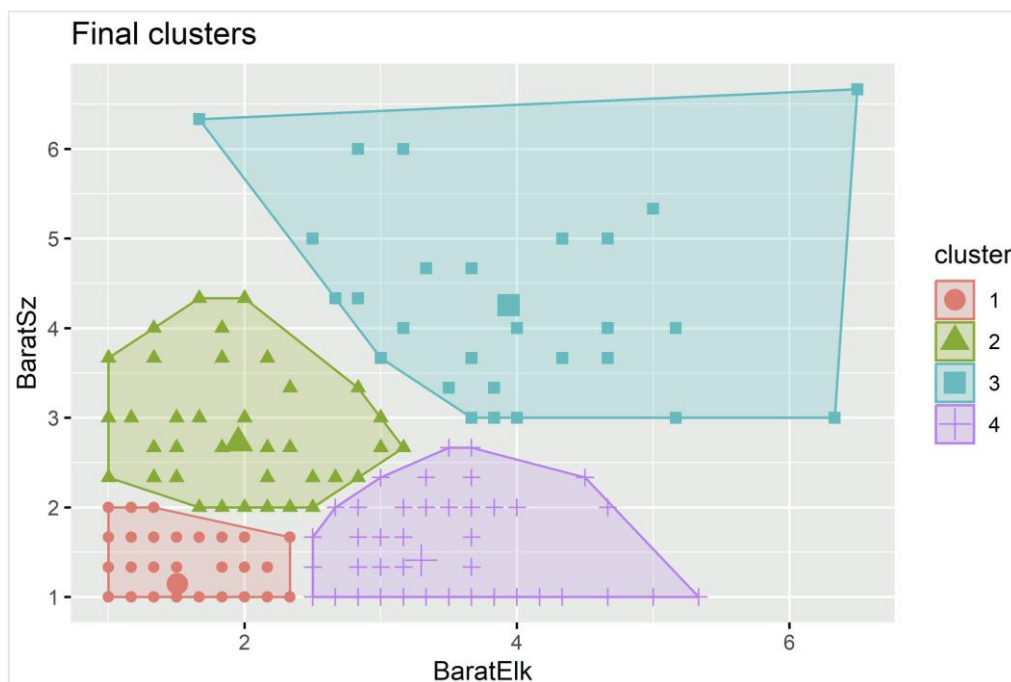
k -közép	k -medián megoldás				
megoldás	KL1	KL2	KL3	KL4	Összesen
KL1	-	-	-	0,99	127
KL2	-	-	0,99	-	57
KL3	-	0,94	-	-	30
KL4	0,97	-	-	-	100
Összesen	97	34	58	125	314

Mindezek alapján a k -közép és a k -medián elemzés közös eredményét fogadjuk el, mégpedig utóbbi képviselésével, melyben a kellően homogén klaszterbe tartozó személyek aránya érezhetően nagyobb. Erre a ROP-R a 9.9. ábrán látható pontdiagramot készítette el. Itt a jobb felső sarokból kiinduló széles tartományban szóródó lila pontok nyilván a legheterogénebb KL2 klaszter (lásd 9.5. táblázat utolsó sora) elemei, a bal alsó sarok köré tömörülő világoskék pontok a leghomogénebb KL4-é, a zöld pontok nyilván KL3-é, végül a BarátElk tengely mentén fölfelé hosszan elnyúló sötétkék pontok pedig KL1-é.



9.9. ábra. A 4-klaszteres k -medián megoldás ábrája

Persze a 4-klaszteres k -közép megoldás ROP-R-es ábrája (lásd 9.10. ábra) segítségével is könnyen azonosíthatjuk a klasztereket. Erről is leolvasható a bal alsó sarokponthoz közeli három klaszter elfogadható homogenitása, ahol a klasztercentrumok jó reprezentánsai a klaszter által képviselt típusnak, ami már kevésbé igaz az erősen heterogén [M+++, M+++] jelzetű klaszterre (a 9.3. táblázatban KL3, a 9.5. táblázatban KL2). Itt a klasztercentrum helyett jobban képviseli a klasztert maga a szélsőséges mintázat, mely lehet olykor mérsékeltebb, máskor viszont szélsőségesebb.

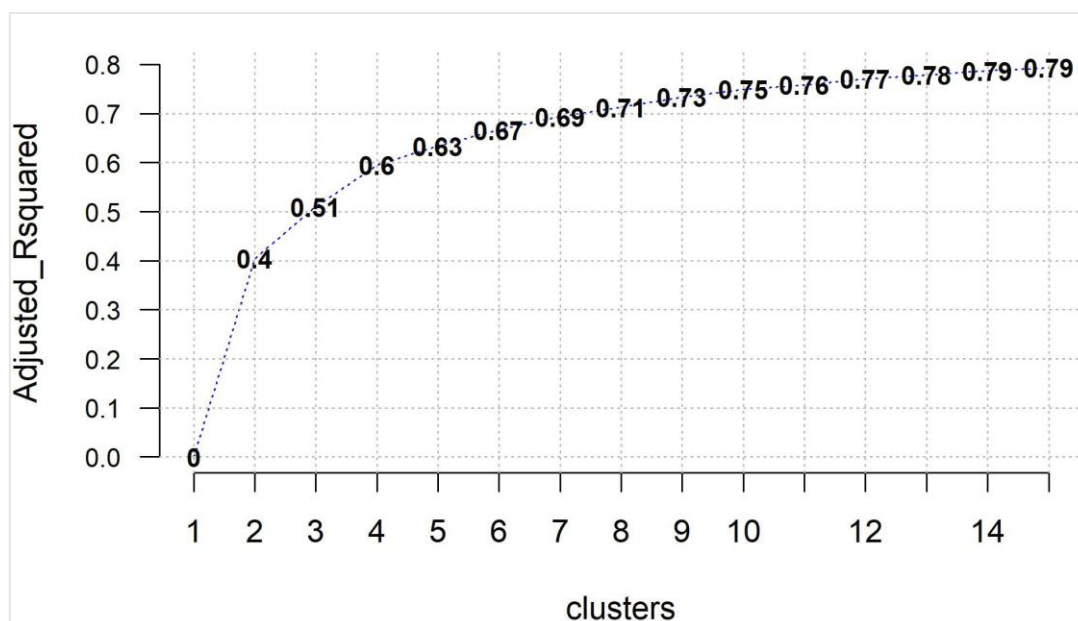


9.10. ábra. A 4-klaszteres k -közép megoldás ROP-R-es ábrázolása

9.4.3. k -közép elemzés a PIK16 kérdőív négy skálájával

A PIK16 kérdőívet (vö. B.6. táblázat) a pszichológiai immunrendszer állapotának és összetevőinek feltérképezésére szerkesztették (Oláh 2005). Négy skálája (Mmonit, AVhat, Rezil és Önreg, vö. B2.3. alpont) a pszichológiai immunrendszer négy fő immunkompetencia pillérét méri. Most az 1003 fős Btérkép2022 minta alapján (vö. B2.3. alpont) e négy skála segítségével szeretnénk megnézni, hogy vannak-e a pszichológiai immunrendszernek olyan mintázata, amelyek típusként értelmezhetők. E probléma tisztázására k -közép elemzéseket végzünk. Az 5.4. alfejezetben ugyanezen a mintán, ugyanezekkel a változókkal kapcsolatban nem találtunk kihagyandó eseteket, ezért most is a teljes, 1003 fős mintán dolgoztunk.

Az első kérdés, hogy milyen k -érték lenne optimális a k -közép elemzéshez. Ehhez először a KKA modul alapértelmezéseként 3-klaszteres elemzéseket futtatunk, a változókat nem standardizálva, egymás után kérve három optimumkeresésre alkalmas ábrát (Silhouette, EESS% és $f(K)$ torzulás), 15-re állítva az ábrakészítés paneljén a maximális klaszterszámot. Ezek közül se a Silhouette-, se az $f(K)$ torzulás ábrájának nem vettük hasznát, mert mindkettő $k = 2$ esetén jelzett optimális értéket. Az EESS arány ábrájáról azt olvashatjuk le, hogy az EESS% $k = 8$ esetén lép először a 70% fölé, de a fokozatos kis emelkedések miatt nem tudjuk, hogy 7 és 10 között melyik klaszterszám vezet a legjobb eredményre. (lásd 9.11. ábra). Az elemzéseket ezért a 7–10 közötti k -értékekre végeztük el, nem kérve az input változók standardizálását.



9.11. ábra. Az EESS-arány függése a klaszterszámtól a PIK16 kérdőív négy skálájával végzett k -közép módszer esetén

A 7–10 közötti klaszterszámok eredményeinek főbb jellemzői a 9.8. táblázatban láthatók, melynek alapján a következő döntést hozhatjuk. Már a 7-klasztteres megoldás is elfogadható, mert EESS% gyakorlatilag eléri a 70%-os szintet, HCátlagS jelentősen 1 alatt van úgy, hogy még a legheterogénebb klaszter HC-szintje ($HC_{maxS} = 0,77$) is határozottan 1 alatt van. Emellett itt a legmagasabb a klaszterek szeparációját mérő XBmod mutató értéke is. Ha pedig megnöveljük a klaszterszámot, a homogenitás sosem nő olyan mértékben, ami indokolná, hogy egy egy nagyobb klaszterszámú, bonyolultabb klasztermodellt fogadjunk el.

9.8. táblázat. Különböző megoldások jellemzői a k -közép klaszterelemzésben 7 és 10 között

klaszter-szám	EESS%	XBmod	SCorig	HCátlagS	HCminS	HCmaxS	HCstan < 0,5
7	69,65	0,469	0,376	0,611	0,45	0,77	26,2%
8	71,91	0,257	0,347	0,566	0,35	0,77	33,2%
9	73,75	0,279	0,351	0,530	0,33	0,78	47,2%
10	75,20	0,418	0,371	0,501	0,29	0,75	52,1%

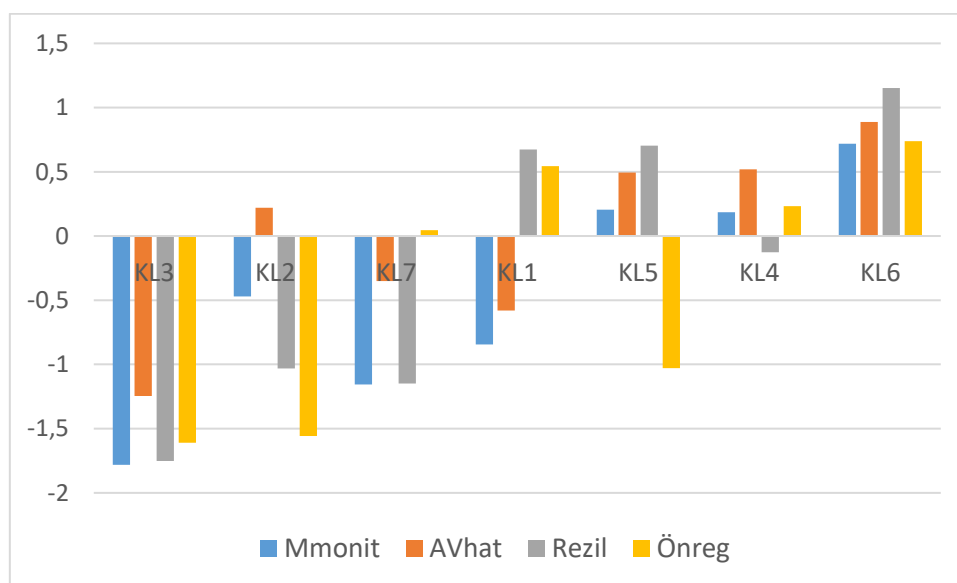
A 7-klasztteres megoldáshoz tartozó z_i -standardizált átlagok mintázatának táblázatát a 9.9. táblázat tartalmazza. A táblázat utolsó oszlopában megadjuk a standardizált skálaátlagok átlagát (z_i -átlag) is, mely a klaszterek általános immunkompetencia szintjét jelzi. Ez a szint annál alacsonyabb, minél negatívabb a z_i -átlag értéke és annál magasabb, minél pozitívabb. A 9.8. és a 9.9. táblázat alapján az alábbi következtetések vonhatók le.

- Bár minden klaszter HCstan értéke kisebb 1-nél, hiányzanak az igazán homogén (pl. a 0,30-nál kisebb HCstan értékű klaszterek), a 0,5-nél kisebb HCstan értékű klaszterekbe tartozó személyek aránya is csak 26,2%.

- A negatív z_t -átlagú, vagyis a gyengén működő pszichológiai immunrendszert képviselő klaszterek a legheterogénebbek ($HCstan > 0,7$), mindegyikük nagyobb $HCstan$ értékű, mint bármelyik pozitív z_t -átlagú klaszter.
- Úgy látszik, hogy a nem megfelelő pszichológiai működéssel jellemezhető személyek klasztere mindig heterogén, a pszichológiailag jól működőké pedig homogén, vagyis egyszerű szavakkal megfogalmazva: a boldogtalansághoz több út is vezet, a boldogsághoz azonban csak egy.
- Szerencsésnek mondható körülmény, hogy a két legjobb immunkompetencia szintű klaszter (KL4 és KL6) egyben a két legnagyobb klaszter is (184, ill. 263 fővel).
- Egészen természetes, hogy a leggyengébb immunkompetenciájú KL3 klaszter mindenben alacsony szintű (legalább A+, de két esetben A+++ minősítésű), a legjobb KL6 klaszter pedig mindenben magas (M vagy M+ minősítésű), de találunk nem triviális mintázatú klasztereket is (pl. KL1, KL5 és KL7).

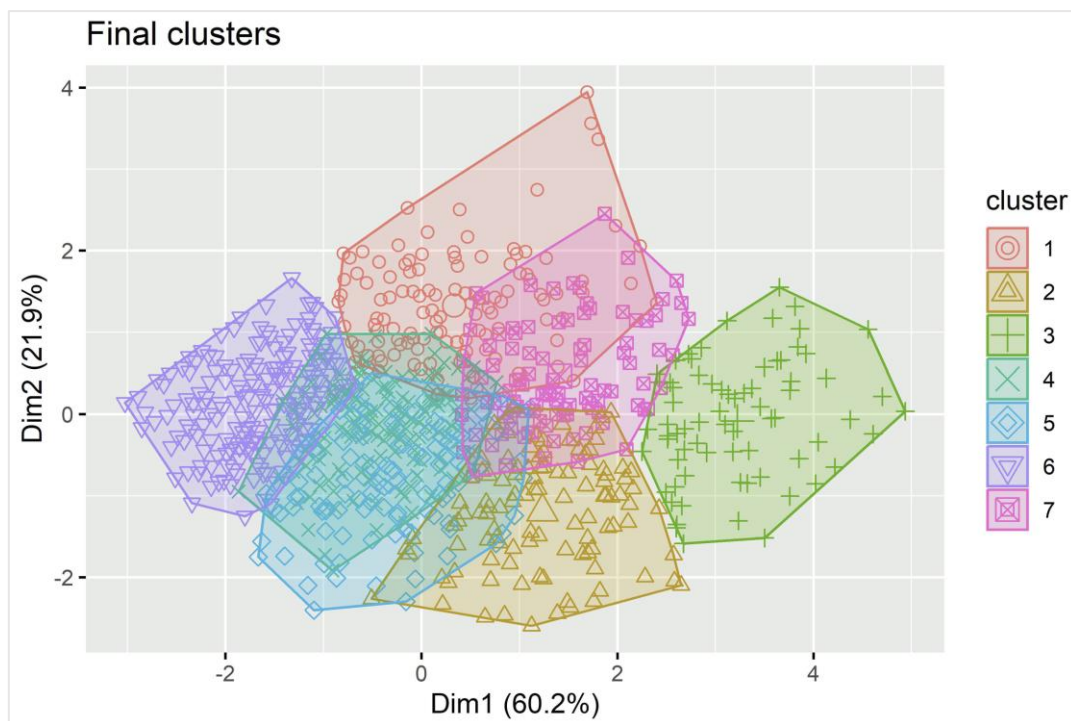
9.9. táblázat. A z_t -standardizált átlagok mintázata a 7-klasztères k -közép megoldásban (M=Magas, A=Alacsony, a + jelek száma a szélsőségesség mértékét jelzi)

Klaszt.	Mmonit	AVhat	Rezil	Önreg	gyak	SCorig	SCsimple	HCstan	z_t -átlag
KL1	A	(A)	M	(M)	130	0,284	0,535	0,75	-0,05
KL2	(A)	.	A+	A++	116	0,307	0,513	0,76	-0,71
KL3	A+++	A+	A+++	A++	75	0,456	0,636	0,77	-1,60
KL4	.	(M)	.	.	184	0,331	0,482	0,52	0,20
KL5	.	(M)	M	A+	127	0,305	0,511	0,63	0,09
KL6	M	M	M+	M	263	0,518	0,640	0,45	0,87
KL7	A+	.	A+	.	108	0,323	0,497	0,71	-0,65
					átlag	0,376	0,551	0,611	



9.12. ábra. A 7-klasztères k -közép megoldás z_t -standardizált átlagainak oszlopdiagramja

A klaszterek mintázata jól azonosítható a klaszterstruktúra ábráján (lásd 9.12. ábra). A ROP-R KKA menüablakában a 7-klasztáros elemzés kijelölésekor kértük még a kapott klasztermegoldás ábrázolását is (lásd 9.13. ábra), mely kissé máshogy néz ki, mint amikor két változóval végeztünk k -közép elemzést (vö. 9.7. és 9.10. ábra).



9.13. ábra. A 7-klasztáros k -közép megoldás klasztereinek ábrázolása a két első főkomponens terében

Kettőnél több input változó esetén ROP-R a változókon végzett FKA elemzés első két főkomponensének a terében ábrázolja a klasztereket, így az ábra helyes értelmezéséhez érdemes nekünk is elvégeznünk ezt az elemzést ROP-R FKA moduljának a segítségével (lásd 4. fejezet), két főkomponens megtartását kérve és forgatás nélkül. A jelen esetben ez a 9.10. táblázatban látható eredményre vezetett.

9.10. táblázat. A PIK16 kérdőív négy skáláján végzett FKA első két főkomponense

Változó	Fkomp1	Fkomp2
Mmonit	0,873	-0,302
AVhat	0,785	-0,491
Rezil	0,799	0,277
Önreg	0,628	0,682

A 9.10. táblázatból látható, hogy az első főkomponensen (Fkomp1) minden változó erősen súlyozódik, tehát ez a pszichológiai immunrendszer általános szintjét méri, ami a PIK16 teszt összpontszámának feleltethető meg. A második főkomponens (Fkomp2) pedig a négy skála közül az önreguláció (Önreg) vs. alkotó-végrehajtó hatékonyság (AVhat) dimenzióját képviseli, melyben az immunrendszer két összetevőjének egyedi vonása (belső vezérlés vs. külső viselkedés) van szembeállítva.

A 9.10. táblázat információját is felhasználva az alábbi következtetések vonhatók le a 9.13. ábra alapján (a 9.12. ábrát is figyelembe véve).

- A 9.13. ábrán jobbról balra, tehát a 9.12. ábrával pont ellentétes irányban haladva találjuk az immunrendszer működése szempontjából növekvő szintű klasztereket. Ez az ellentétes irány amiatt van, mert a főkomponensanalízisben a főkomponensek irányítása (előjele) nem mindig ugyanolyan, de ez a lényegen nem változtat.
- A jobb szélső (3-as indexű, zöld plusz jelzetű) felel meg a mindenben alacsony (a 9.12. ábrán KL3), a bal szélső (6-os indexű, lila fordított háromszög) pedig a mindenben magas (a 9.12. ábrán KL6) klaszternek.
- A legfelső (1-es indexű) klaszter felel meg az Önreg-ben és Rezilben magas, a másik két skálán pedig alacsony klaszternek (a 9.12. ábrán ez KL1).
- Az ezzel átellenben alul elhelyezkedő 2-es indexű barnás színű klaszter (a 9.12. ábrán KL2) AVhat tekintetében átlagos, minden más skála szerint viszont – Önreg esetében kirívóan – alacsony szintű.

Az egyes klaszterek jelentésének finomabb szakmai kibontásához persze további vizsgálatokra van szükség, amelyek megvalósítása könyvünk kereteit meghaladja.

10. fejezet

A modell-alapú klaszteranalízis (MKA) modul

ROP-R MKA moduljával *modell-alapú klaszteranalízis* (MKA) végezhető, melynek során a program beállított klaszterszámok és modelltípusok mindegyikére (14 típus közül bármely részegyüttes, akár mind a 14 típus kiválasztható) a program megkeresi a legjobb maximum likelihood illesztést.

Az MKA elemzés gondolatmenetében az a fő feltételezés, hogy mintánk adatai egy többdimenziós keverékeloszlást követnek, ahol a komponensek mindegyike többdimenziós normális eloszlású, de más-más centrummal és esetleg más varianciákkal és kovarianciákkal. Ebben a keretben MKA célja, hogy azonosítsa az összekevert többdimenziós eloszlások számát (optimális klaszterszám) és megadja az eloszlások jellemzőit (Fraley és Raftery, 2002; Gergely és Vargha, 2021; Vargha, 2022, 7. fejezet).

A 10.1. alfejezetben az érintett statisztikai módszerekkel kapcsolatos elméleti alapokat ismertetjük, a 10.2. alfejezet az MKA menüablak használatát mutatja be, a 10.3. alfejezet az eredménylista szerkezetét, a 10.4. alfejezet pedig a modullal végrehajtható statisztikai elemzéseket szemlélteti valódi adatok segítségével.

10.1. Elméleti alapok

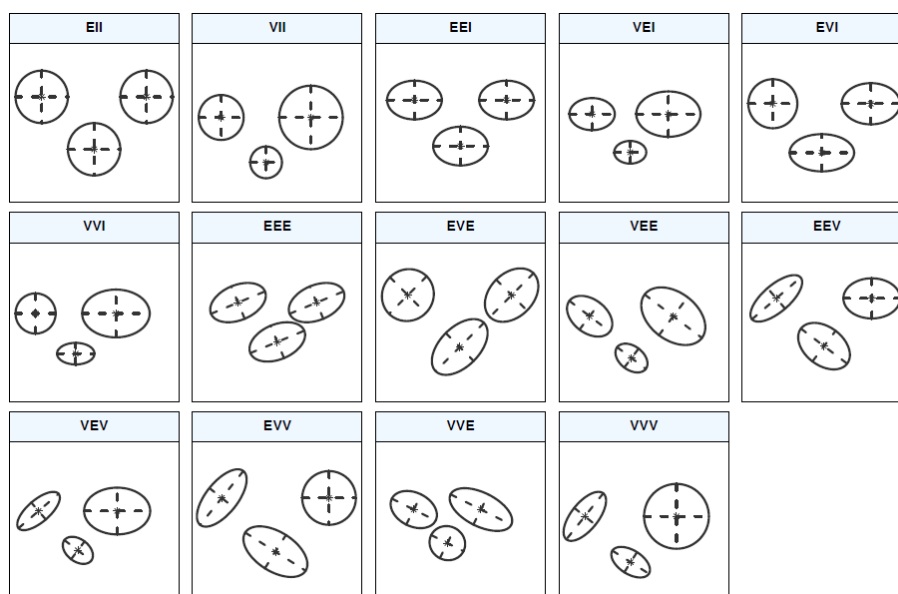
MKA a két legnépszerűbb hagyományos klaszterezési módszer, a HKA és a KKA eléggé körülményes modellválasztási eljárásának (lásd 7–9. fejezet) automatizálására jött létre. E módszereknél a klaszterek számának meghatározására szolgáló standard eljárás az, hogy kiválasztunk néhány releváns QC klaszter adekvációs mutatót (pl. ROP-R-ben EESS%, XBmod, SCorig, HCátlag), és több lehetséges klaszterstruktúrát ezek szerint kiértékelve döntünk arról, hogy melyik tekinthető a legjobbnak (Vargha, Bergman és Takács 2016).

A legnagyobb probléma ezzel az eljárással kapcsolatban a kutató szubjektivitása, mert különböző klaszterezési módszerek és különböző QC-k alkalmazásával a kutatók teljesen eltérő megoldásokra juthatnak, így e vizsgálatok megismételhetősége kérdéses lehet. MKA egy automatikusabb módszert alkalmaz az optimális klaszterszám kiválasztására. Elsődlegesen olyan többdimenziós eloszláskomponenseket tárunk fel, amelyeknek egy alkalmas keverékére mintánk elfogadhatóan illeszkedik. Ha feltártunk egy ilyen struktúrát, az eloszláskomponensek segítségével egy klaszterstruktúrát is definiálhatunk. Minden eloszlás egy klasztert képvisel és minden személyt abba a klaszterbe sorolunk, amelyhez tartozó eloszláskomponens a személy értékmintázatát a legnagyobb valószínűséggel produkálja, vagyis amelynél maximális az eloszlás többváltozós sűrűségfüggvényének a személy értékmintázatához tartozó értéke. A legjobb klaszterstruktúra megtalálása abból áll, hogy különböző számú és típusú komponensekből álló keverékmódelleket összehasonlítunk és kiértékelünk.

Ha a két input változónak megfelelő kétdimenziós térben három kétváltozós normális eloszlást összekeverünk, a klaszterstruktúra típusa 14-féle lehet. A keverékarány által meghatározott relatív klaszterméret/klasztérnagyság (volume) szerint egyforma (E) vagy változó (V). A főtengelyeik hossza által meghatározott alakjuk szerint szintén egyforma (E) vagy változó (V). Végül az eloszlások két főtengelyének orientációja szerint a tengelyekkel – és ekkor természetesen egymással is – megegyező irányú (I), a tengelyekkel ferde szöget bezáró, de egymással egyező irányultságú (E), illetve eltérő/változó irányultságú (V). Ezeket a lehetséges variánsokat a 10.1. táblázat, illetve a 10.1. ábra foglalja össze szokásosan használt hárombetűs kódjukkal (Scrucca, Fop, Murphy és Raftery, 2016).

10.1. táblázat. Egy többdimenziós normális keverékeloszlás 14 különböző lehetséges struktúrája három komponens és két elemzett változó esetén

Index	Modell neve	Eloszlások típusa	Klaszterek mérete	Eloszlások alakja	Eloszlások orientációja
1.	EII	Gömb	Egyforma	Egyforma	–
2.	VII	Gömb	Változó	Egyforma	–
3.	EEI	Ellipszoid	Egyforma	Egyforma	Koordinátatengelyek
4.	VEI	Ellipszoid	Változó	Egyforma	Koordinátatengelyek
5.	EVI	Ellipszoid	Egyforma	Változó	Koordinátatengelyek
6.	VVI	Ellipszoid	Változó	Változó	Koordinátatengelyek
7.	EEE	Ellipszoid	Egyforma	Egyforma	Egyforma
8.	EVE	Ellipszoid	Egyforma	Változó	Egyforma
9.	VEE	Ellipszoid	Változó	Egyforma	Egyforma
10.	VVE	Ellipszoid	Változó	Változó	Egyforma
11.	EEV	Ellipszoid	Egyforma	Egyforma	Változó
12.	VEV	Ellipszoid	Változó	Egyforma	Változó
13.	EVV	Ellipszoid	Egyforma	Változó	Változó
14.	VVV	Ellipszoid	Változó	Változó	Változó



10.1. ábra. Egy többdimenziós normális keverékeloszlás 14 különböző lehetséges struktúrája három komponens és két változó esetén

Az eloszlások típusa gömb (szferikus / spherical), ha a többdimenziós eloszlás összetevői egyforma varianciájúak és korrelálatlanok (pl. EII és VII esetén), egyébként ellipszoid. Különböző varianciák és páronkénti korrelálatlanság esetén az orientáció a koordinátatengelyekkel azonos irányú. Ha a páronkénti korrelálatlanság nem áll fenn, akkor az orientáció egyforma vagy változó attól függően, hogy a változók kovariancia-mátrixa minden klaszterben ugyanolyan, vagy pedig eltérő.

Ez a 14 keveréktípus természetesen kettőnél több elemzett változó és háromnál több klaszter esetén is a struktúra fontos jellemzője akkor is, ha egy-egy típus olykor még további altípusokra bontható. Egy MKA elemzés során tehát egyaránt feltárandó a klaszterek (eloszláskomponensek) optimális száma, valamint a klaszterstruktúra típusa a 10.1. táblázat, illetve 10.1. ábra szerinti besorolással.

MKA tehát lényegében abból áll, hogy mind a 14 típusra és egy megadott klaszterszám tartomány (pl. $2 \leq k \leq 9$) minden értékére megkeressük azt a keverékmodellt, amely a lehető legjobban illeszkedik a vizsgált mintára. A modellbecslés az ML módszerrel történik, mely szerint annál a modellenél a legjobb az illeszkedés, amely szerint az adott minta bekövetkezése a legnagyobb valószínűségű, vagyis a legnagyobb log-likelihood (LL) értékű. A minősítő kritérium tehát az ML becslés LL értéke. Azonosítva a legjobban illeszkedő keverékmodellt minden típus x klaszterszám kombinációra, ki kell választani közülük a legjobbat. Ezt azonban nem az LL nyers értéke, hanem a komponensek számát is figyelembe vevő BIC együttható alapján döntjük el (Vargha, 2022, 160. o.). Fraley és Raftery (2002) szerint minél nagyobb egy adott modellhez tartozó BIC-érték, annál nagyobb az evidenciája ennek a modellnek az adott klaszterszámmal.

Megjegyezzük, hogy a klasztermodellek kiértékelésére itt kiszámított ML alapú BIC mutató hasonló logikájú, mint a CFA-modellek kiértékelése során használt BIC kritérium (lásd 6.1. táblázat), azzal a különbséggel, hogy az itteni BIC megőrzi az LL mennyiség negatív előjelét, míg CFA esetén ez át van fordítva pozitívrá. Ez okból a CFA elemzések esetén a kisebb BIC-érték jelez jobb illeszkedést, míg az MKA elemzések esetén a nagyobb.

Miután kiválasztottuk a legnagyobb BIC-értékű keverékmodellt, annak k számú eloszláskomponense segítségével a minta személyei k klaszterbe sorolhatók. Minden személyt abba a klaszterbe teszünk, amelyhez tartozó eloszláskomponens a személy értékmintázatát a legnagyobb valószínűséggel produkálja. Ezen ún. *klaszterbesorolási valószínűségek* szerint egyébként ún. *fuzzy klaszterezés* is megvalósítható, melynél nem soroljuk a személyeket fix klaszterekbe, hanem csak azt mondjuk meg, hogy melyik klaszterbe milyen eséllyel tartoznak. Az MKA ROP-R-beli futtatása során egyaránt lekérhetők a klaszterbe tartozást megadó kódok és a klaszterbesorolási valószínűségek.

Biernacki, Celeux és Govaert (2000) szerint a BIC kritérium alkalmas a komponens-eloszlások feltárására, de nem optimális egy olyan klaszterbesorolásra, ahol feltételezzük, hogy minden személy egyetlen klaszterbe tartozik. Emiatt, ha a fentebb felvázolt keverékfelbontási módszert valódi klaszteranalízisre akarjuk használni, akkor a BIC kritériumon kissé módosítani kell. E probléma orvoslására Biernacki, Celeux és Govaert (2000) egy ICL⁴⁴ javított likelihood kritériumot javasolt BIC helyett, mely BIC olyan módosított változata, amelynek képletében a klaszterek átfedését egy entrópia tag „bünteti”. Ez az entrópia egy olyan bizonytalansági mérőszám, amely akkor magas, ha több személy esetén nagy a bizonytalanság afelől, hogy a személy melyik klaszterbe tartozik. Az ICL-re építő módszer nem javasol más eljárást a komponenseloszlások becslésére, mint a BIC-re

⁴⁴ Integrated Complete-data Likelihood

építő, csak a végső modellválasztás során BIC helyett az ICL kritérium értéke alapján döntünk arról, hogy melyik modellt tekintjük optimális struktúrának.

Gergely és Vargha (2021) vizsgálatai szerint az MKA automatikus alkalmazása gyakran vezethet az optimálisnál gyengébb eredményre. A BIC- vagy az ICL-ábrán a szabálytalan, hirtelen ugrásokkal jellemezhető vagy torzó grafikonok ritkán azonosítják a legjobb megoldást.

10.2. Az MKA menüablak

MKA elemzést ROP-R-ben a modell-alapú klaszteranalízis (röviden MKA) modul segítségével végezhetünk a mintabeli eseteken a kijelölt változók felhasználásával. ROP-R MKA modulja az *mclust* (Scrucca et al., 2016), a *ggplot2* (Wickham, 2016) és a *factoextra* (Kassambara és Mundt, 2020) R-package-et használja fel elemzéseikhez.

HierReg PolReg BinLogReg

Főkompan FeltáráFA KonfirmFA

AggHierKLÁ ÖszHierKLÁ k-centrKLÁ **ModellKLÁ**

☐ Változók standardizálása

Klaszterszám limitek
ettől: 2 eddig: 9 (2 - 25)

Modelltípus (több is választható)

☒ Mind ☒ EII ☒ VII ☒ EEI
☒ VEI ☒ EVI ☒ VI ☒ EEE
☒ VEE ☒ EVE ☒ VVE ☒ EEV
☐ Semmi ☒ VEV ☒ EVV ☒ VVV

Táblázatkészítés
☒ BIC-értékek ☒ ICL-értékek ☐ Besorolási p-értékek

Ábramentés jpg fájlba
☒ BIC ábra ☐ ICL ábra ☐ Legjobb klaszterstruktúra

Mentés (adatfájlhoz illesztés)
☐ Klaszterkód értékek ☐ Bizonytalanság értékek

☒ Futtat

10.2. ábra. Az MKA modul menüablakának jobb oldala ROP-R-ben

Az MKA menüablak (lásd a beállítások szempontjából fontos jobb oldalát a 10.2. ábrán) szolgál az elemzésbe bevonandó változók kiválasztásán túl

- a klaszterszám limitek beállítására
- az elemzendő modell típusok kiválasztására
- speciális táblázatok (BIC, ICL, besorolási p -értékek) kérésére
- speciális ábrák (BIC, ICL, legjobb klaszterstruktúra) fájlban való elmentésére
- speciális ábrák (klasszifikációs ábra, bizonytalanság értékek, sűrűségábra) fájlban való elmentésére legfeljebb 3 input változó esetén
- klaszterkód változó és bizonytalanság értékek elmentésére.

A fentebb említett néhány speciális ábra jelentése az alábbi:

1) *Klasszifikációs ábra*: ez az adatok pontdiagramja, amelyen a különböző klaszterekbe tartozó személyek eltérő színekkel és formákkal vannak jelölve.

2) *Bizonytalanság értékek* ábrája: ez az adatok olyan pontdiagramja, amelyen a különböző klaszterekbe tartozó személyek eltérő színű kitöltött körökkel vannak jelölve úgy, hogy a nagyobb besorolási bizonytalanságot nagyobb kör képviseli. Egy személy *besorolási bizonytalanságát* úgy kapjuk, hogy 1-ből kivonjuk a személy p_{max} legnagyobb besorolási valószínűségét (Gormley, Murphy és Raftery, 2023).

3) *Sűrűségábra*: a többdimenziós adatok klaszterenkénti sűrűségeloszlása.

Egy MKA elemzés végrehajtása után a szokásos „aktualis” mappában találjuk az elemzéshez elkészített ideiglenes adatfájlt (tmpdat.txt), az elemzés által optimálisnak tekintett modellhez tartozó klaszterváltozóval kiegészített ideiglenes adatfájlt (tmpdat2.txt), a futtatott R-scriptet (MBCA.r), valamint a kért ábrákat jpg fájlban (BIC1.jpg, ICL1.jpg, mbplot1.jpg, Class1.jpg, Uncert1.jpg, Density1.jpg). Ha kérjük a besorolási p -értékek táblázatát, akkor ezt – terjedelme miatt – nem az eredménylistán, hanem egy külön – pval1.txt nevű – fájlban találjuk meg, szintén az „aktualis” mappában.

Ha feltételes csoportosító változót is kijelölünk, akkor minden feltételes csoport elemzése során elkészülnek a kért ábrák, az ezeket tartalmazó fájlok nevében megjelenő számok ezen csoportok sorszámát jelzik (1, 2, ...).

10.3. Mit tartalmaz az MKA eredménylistája?

Az MKA eredménylistája az alábbi elemeket tartalmazza.

- Alapstatisztikák a kiválasztott változókra
- A kijelölt modell tartományon belül a program által legjobbnak tartott megoldás (modell típus és klaszterszám)
- A program által azonosított legjobb megoldásra:
 - o Adekvációs mutatók (EESS%, HCátlag vagy HCátlagS, XBmod)
 - o Klaszterstatisztikák (elemszám, átlag, szórás, minimum, maximum)
 - o Nem standardizált átlagok (nem standardizált centroidok)
 - o Standardizált átlagok (standardizált centroidok)
 - o Klasztercentroidok páronkénti ASED távolságai
 - o Standardizált átlagok mintázata
- BIC értékek táblázata (ha kérjük)
- ICL értékek táblázata (ha kérjük).

Ha kérjük a változók standardizálását, akkor az átlagok standardizálása változónként a saját átlaggal és szórással (z -standardizálás), ha pedig nem kérjük, akkor a változók összátlagával és együttes szórásával lesz végrehajtva (z_T -standardizálás).

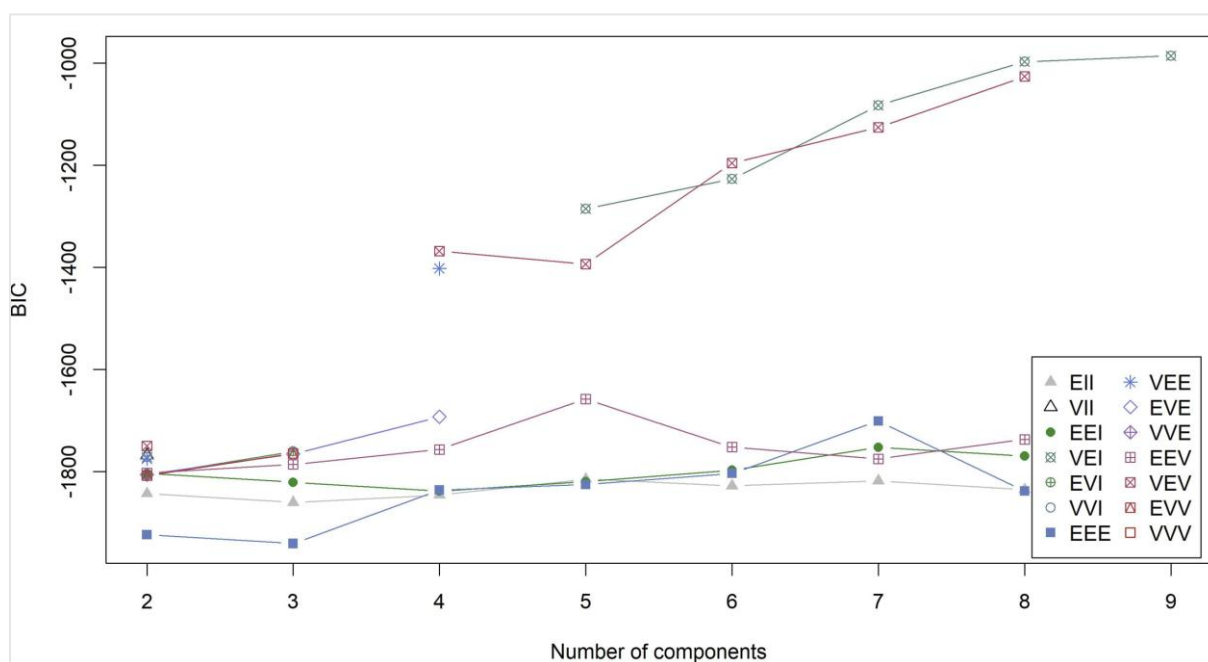
10.4. Az MKA szemléltetése valódi adatokon

Ebben az alfejezetben egyrészt a KÖT2016 kötődéskutatási adatokon (lásd B2.2. alpont) a baráttal kapcsolatos elkerülés (BarátElk) és szorongás (BarátSz) skálával végzünk MKA elemzést a 10.4.1. alpontban (a 8 outlierrel csökkentett KÖT2016 mintán, vö. 7.4. alfejezet), majd a PIK16 kérdőív négy skálájával (vö. B2.3. alpont) a 10.4.2. alpontban.

10.4.1. MKA elemzés két kötődésváltozóval

Az MKA elemzést először az alapértelmezés szerinti 2–9 klaszterszámokra hajtottuk végre, a változókat nem standardizálva, az összes modelltípusra, BIC- és ICL-ábrát is kérve. ROP-R eredménylistáján az első fontos információ, hogy a legjobb modell 2–9 között az összes típust figyelembe véve: VEI, $k = 9$ klaszterszámmal. Jelöljük ezt a megoldást röviden VEI9-cel.

Ehhez már érdemes megnéznünk a BIC1.jpg fájlba mentett BIC-ábrát is (lásd 10.3. ábra). És valóban, a 10.3. ábrán $k = 6$ fölött már minden klaszterszámmal a VEI modell típus BIC-értéke a legnagyobb, a maximum pedig $k = 9$ értéknél látható. Lehet, hogy ha szélesebb (pl. 2–12) klaszterszám övezetet jelölnénk ki, magasabb BIC-csúcsot is láthatnánk. Kipróbáltuk, nem kaptunk. Az ábrán nem látható mind a 14 típus esetén grafikonon, illetve a látható grafikonok se húzódnak több esetben végig a teljes 2–9 klaszterszám tartományon. Ez azért van, mert előfordulhat, hogy egy modellbecslésnél nem konvergál a megoldás, s ilyenkor nincs érvényes BIC-érték sem. Ez az outputon a BIC-értékek – itt nem közölt – táblázatából is kiolvasható.



10.3. ábra. Az ECR-RS kérdőív BarátElk és BarátSz mutatóján elvégzett MKA elemzés BIC-értékeinek grafikonja (összes modell $k = 2$ és 9 között) ROP-R-ben

Megnéztük egyébként az ICL-ábrát is, mely csak minimális mértékben különbözött a BIC-ábrától és ugyanazt a VEI9 modellt jelezte optimális megoldásnak. Nézzük meg tehát alaposabban ezt! Az output VEI9-et leíró részei csaknem ugyanolyanok, mint a KKA elemzések során. Elsőként a globális adekvációs mutatók láthatók: EESS% = 71,25, HCátlagS = 0,589, XBmod = -0,417. Ezek EESS% és HCátlagS tekintetében a homogenitás közepes szintjét jelzik, XBmod határozottan negatív értéke pedig arra utal, hogy a klaszterek háttérében álló komponens eloszlások erősen átfedők. Nézzük meg ezért a z_I -standardizált átlagok mintázatának táblázatát is (a klasztereket a BarátElk és BarátSz standardizált átlagának átlaga, a z_I -átlag szerint sorba rendezve, lásd 10.2. táblázat), valamint a z_I -standardizált klasztercentroidok páronkénti ASED távolságait tartalmazó táblázatot (lásd 10.3. táblázat). Ez utóbbi kvantifikált formában jelzi, hogy a kapott klaszterstruktúrában mely klaszterek esnek egymáshoz igen közel. 0,15 alatti távolságértékek esetén igen közeli, 0,50 alatti távolságértékek esetén közeli klasztercentrumokról beszélhetünk.

10.2. táblázat. A z_I -standardizált átlagok mintázata a VEI9 megoldásban

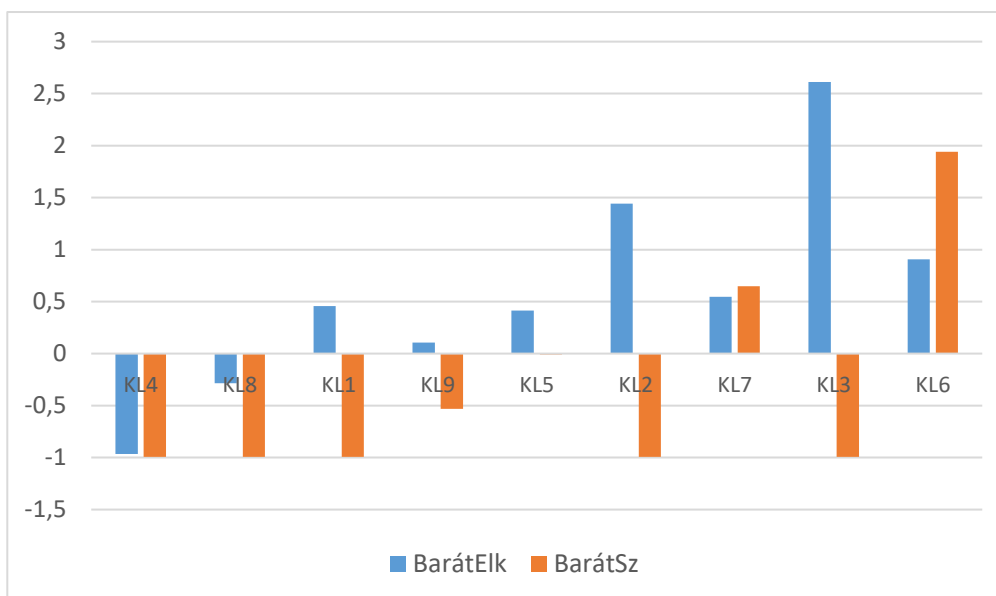
Klaszter	BarátElk	BarátSz	KLgyak	HCstanS	z_I -átlag
KL4	A	A	42	0	-0,98
KL8	.	A	47	0,05	-0,64
KL1	(M)	A	31	0,03	-0,27
KL9	.	(A)	47	0,56	-0,21
KL5	.	.	55	0,73	0,20
KL2	M++	A	24	0,11	0,23
KL7	(M)	(M)	26	1,38	0,60
KL3	M++++	A	4	0,06	0,81
KL6	M	M+++	38	2,00	1,42

10.3. táblázat. A klasztercentroidok páronkénti d ASED távolságai a VEI9 megoldásban

Klaszter	KL1	KL2	KL3	KL4	KL5	KL6	KL7	KL8	KL9
KL1	0	0,49*	2,33	1,01	0,48*	4,40	1,35	0,28*	0,17*
KL2	0,49*	0	0,69	2,90	1,01	4,45	1,75	1,49	1,00
KL3	2,33	0,69	0	6,40	2,90	5,76	3,49	4,20	3,25
KL4	1,01	2,90	6,40	0	1,43	6,05	2,48	0,23*	0,68
KL5	0,48*	1,01	2,90	1,43	0	2,02	0,23*	0,73	0,18*
KL6	4,40	4,45	5,76	6,05	2,02	0	0,90	5,01	3,38
KL7	1,35	1,75	3,49	2,48	0,23*	0,90	0	1,69	0,79
KL8	0,28*	1,49	4,20	0,23*	0,73	5,01	1,69	0	0,18*
KL9	0,17*	1,00	3,25	0,68	0,18*	3,38	0,79	0,18*	0

Jelölés: *: $d < 0,50$; **: $d < 0,15$

A 10.2. táblázat és a klaszterstruktúra értelmezését elősegítendő, elkészítettük még a klaszterek ábráját is a z_I -standardizált átlagok alapján (lásd 10.4. ábra), a klasztereket itt is a z_I -átlag szerint sorba rendezve.



10.4. ábra. Az ECR-RS kérdőív BarátElk és BarátSz mutatóján elvégzett MKA elemzés 9-klasztteres VEI9 megoldásában a z_i -standardizált átlagok oszlopdiagramja

A 10.2. táblázat és a 10.4. ábra alapján az alábbi következtetéseket vonhatjuk le.

- VEI9-ben három olyan típus is van, amelyik rokonítható az ugyanezen a mintán, ugyanezen változókkal végrehajtott AHCA és k -közép megoldásokkal.
 - o A KL4 klaszter által képviselt biztonságos kötődés [Elk–Sz–] típusa. Különösen figyelemre méltó, hogy itt egy jelentős nagyságú ($n = 42$), gyakorlatilag maximálisan homogén klaszter ($HCstanS = 0$) típusához jutottunk.
 - o A KL9, KL5, KL8 klaszterek által képviselt típus fő jellemzője a baráti elkerülés átlagos szintje, amelyet átlagos (KL5) vagy az átlagosnál kissé alacsonyabb (KL9 és KL8) baráti szorongás szint kísér. Ezek egymáshoz való közelségét a $d(KL9, KL5)$ és a $d(KL9, KL8)$ alacsony (0,20-nál kisebb) távolságok is jelzik (vö. 10.3. táblázat). Nyilván ezen közelségek az okai XBmod határozottan negatív szintjének. Kiemelendő a KL8 által képviselt [Elk0Sz+] variáns, mely jelentős nagyságú ($n = 47$) és szintén nagyon homogén ($HCstanS = 0,05$).
 - o A KL2 klaszter a korábbi elemzésekben már feltárt [M+, (A)] típus egyik variánsa, az összátlagtól kicsit erőteljesebben eltérő baráti elkerülés (M++) és baráti szorongás (A) szinttel.
- VEI9 fontos tulajdonsága, hogy klasztereinek többsége (a 9-ből 5) rendkívül homogén, 0 és 0,12 közötti $HCstan$ -értékű, amelyek együtt a 314 fős minta közel felét (47%-át) képviselik. Kiemelkedik ezek közül a maximális homogenitású 42 fős KL4 klaszter ($HCstan = 0$).
- Persze elképzelhető, hogy ezek a homogén klaszterek részben annak köszönhetik létüket, hogy a struktúra 9-klasztteres, ami nyilván kedvez a kisebb homogén klaszterek létrejöttének. Ezért elvégeztük a 9-klasztteres k -közép elemzést is. Ennek az összhomogenitást maximalizáló KK9 megoldása érthetően sokkal jobb globális

adekvációs mutatókkal rendelkezik, mint VEI9 (EESS% = 88,10, XBmod = 0,576, HCátlagS = 0,249), a klasztereire vonatkozó z_I -standardizált átlagok mintázatát a 10.4. táblázatban helyeztük el (itt is a z_I -átlag szerint rendezve a sorokat).

10.4. táblázat. A z_I -standardizált átlagok mintázata a k -közép elemzés 9-klasztteres megoldásában

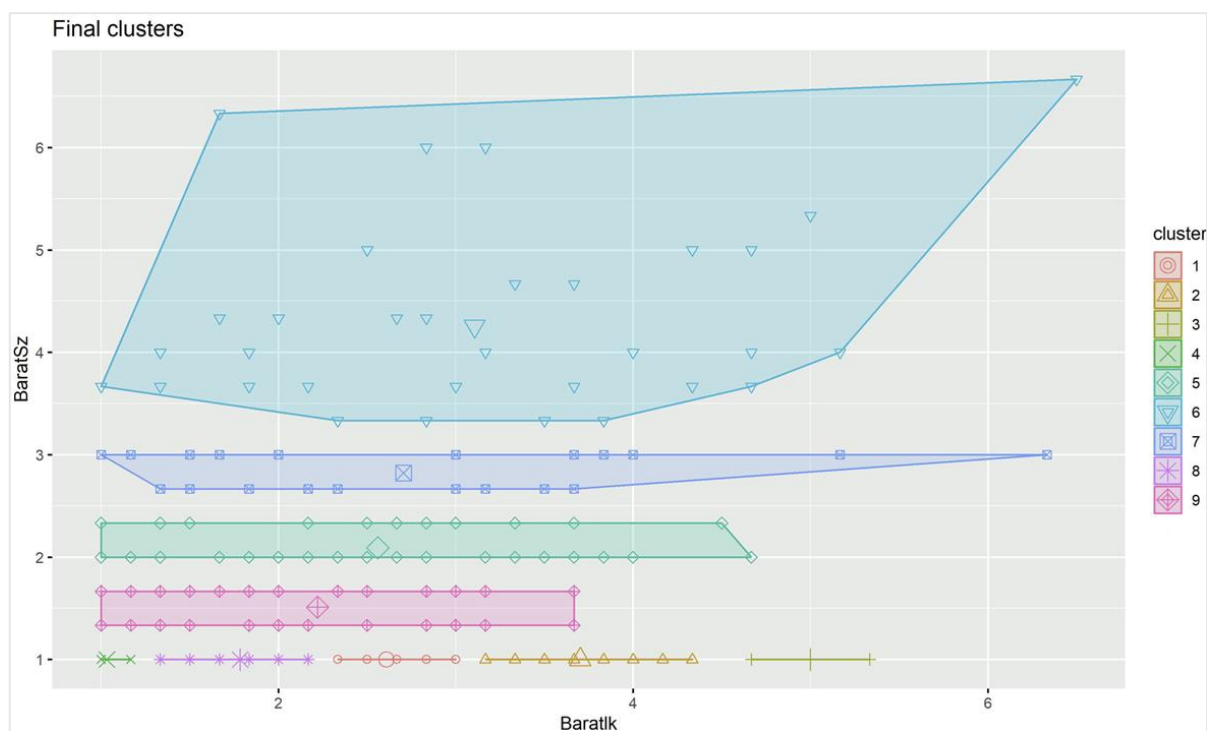
Klaszter	BarátElk	BarátSz	KLgyak	HCstan	z_I -átlag
KL9	A	A	62	0,05	-0,91
KL3	.	A	50	0,06	-0,54
KL2	.	.	42	0,24	-0,18
KL1	M	A	51	0,13	-0,06
KL5	.	M+	18	0,30	0,44
KL6	M+++	A	24	0,32	0,52
KL7	M	.	41	0,31	0,59
KL8	M	M++++	10	0,92	1,61
KL4	M++++	M++++	15	1,30	2,12

A VEI9-re vonatkozó 10.2. és a k -közép elemzés KK9 megoldás klasztereit összefoglaló 10.4. táblázatot összevetve az alábbi következtetéseket vonhatjuk le.

- VEI9-ben két olyan klaszter van, amelyik igen jól megfeleltethető a KK9 struktúra egy-egy klaszterének: a biztonságos kötődés típusát képviselő – sárgával kiemelt – KL4 (KK9-ben KL9), valamint a zöld színnel kiemelt sorban található és szintén a biztonságos kötődés egyik típusvariánsát képviselő, de már az átlagos szint körüli baráti elkerüléssel KL8 (KK9-ben KL3). Ezek közös jellemzője, hogy rendkívül homogének (HCstan < 0,10).
- Jó megfelelést találunk még a szintén zölddel kiemelt és mindenben átlag körüli szintet képviselő VEI9-beli KL5 és a KK9-beli KL2 klaszter mintázata között, amelyek homogenitása közepes szintű.
- Elfogadható egyezést találunk a világosbarnával kiemelt sorú VEI9-beli KL2 és a KK9-beli KL6 klaszter mintázata között. Ezek közös nemtriviális mintázata egy igen magas (VEI9-ben M++, KK9-ben M++) szintű baráti elkerülés és egy alacsony szintű (A+) baráti szorongás, erős homogenitású klaszterekkel képviselve (HCstan < 0,35).
- Bár KK9-cel összehasonlítva VEI9-ben többen vannak igen nagy homogenitású, 0,10 alatti HCstan értékű klaszterben (VEI9: 39,5%, KK9: 35,7%) KK9 összességében jóval homogénebb klaszterstruktúra, mint VEI9. Ezt nemcsak a magasabb EESS% érték jelzi (VEI9: 71,3%, KK9: 88,1%), hanem az is, hogy KK9-ben a személyek 91,7%-a 0,50 alatti HCstan értékű klaszterben van, míg VEI9-ben ez az arány csak mindössze 47,1%.

Mindezek alapján talán helytálló az a végkövetkeztetés, hogy olyan struktúrák esetén, ahol a teljes populáció nem bontható fel teljeskörűen homogén alcsoportokra, vagyis ahol a klaszteranalízis eredménye eleve csak részleges lehet, az MKA módszer talán sikerebben képes igen erősen homogén típusokat azonosítani. Ezt a konklúziót erősíti a legjobb VEI9 klaszterstruktúra ábrája (lásd 10.5. ábra), melyen a leghomogénebb klaszterek (KL4, KL8, KL1, KL2, KL3) egyenes szakaszokkal vannak ábrázolva (ebben a sorrendben), közvetlenül a

BarátElk tengelyen. Ez azt is jelzi egyben, hogy ezekben a klaszterekben a személyek szorongásértéke minimális variabilitást mutat, amiben különböznek: az az elkerülés dimenzió. A BarátSz változó egyébként egy kivétellel minden klaszterben nagyon kis variabilitású. Ez a kivétel a KL6, mely minden tekintetben nagyon heterogén és abban különbözik a többitől, hogy egyesíti az összes kb. 3,5-nél nagyobb értékű személyt. A klaszterek különbözőségének logikája tehát az alábbiak szerint fogalmazható meg.



10.5. ábra. Az ECR-RS kérdőív BarátElk és BarátSz mutatóján elvégzett MKA elemzés legjobb VEI9 megoldásának klaszterábrája

- Van öt rendkívül homogén klaszter (KL4, KL8, KL1, KL2, KL3), amelyek BarátSz tekintetében mind a minimális szinten vannak, s ami megkülönbözteti őket, az a BarátElk változó szintje.
- Van egyetlen rendkívül heterogén klaszter (KL6), amelyik egyesíti az összes olyan személyt, aki egy meglehetősen magas szintnél (kb. 3,5-nél) nagyobb BarátSz értékkel rendelkezik.
- A többi klaszter (KL9, KL5, KL7) BarátSz tekintetében kis variabilitású, de egymáshoz viszonyítva növekvő szintű. Közös jellemzőjük még, hogy BarátElk tekintetében viszont viszonylag nagy variabilitásúak (lásd 10.5. ábra).

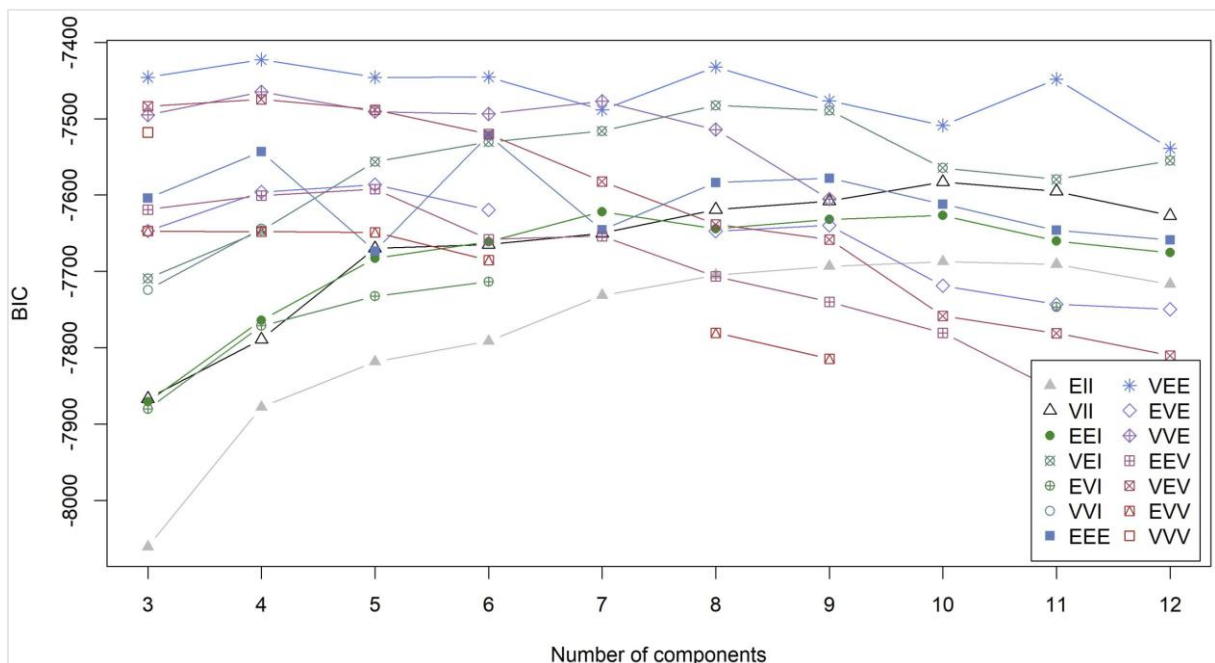
A fentebb megfogalmazott következtetések érvényességét némileg csorbítja mintánk kis elemszáma ($n = 314$). Megbízható típusfeltáráshoz sokkal nagyobb mintára van szükség és annál nagyobbra, minél több a lehetséges típus. Esetünkben 9 stabil típus feltáráshoz legalább 1000 fős mintára lenne szükség.

Végül megemlíjtük, hogy különböző klasztermegoldások összehasonlításának más útja is van, mint amit ebben az alponthan követtünk. A ROPstat keretű ROP-R szempontjából jó tudni, hogy a ROPstat *Mintázatfeltáró elemzések* statisztikai menüpontjában mind a

Keresztábrák celláinak elemzése (Exacon), mind a *Centroidok összehasonlítása* modul hasznos információkkal szolgál klaszterstruktúrák összehasonlítására, továbbá ugyanitt, a *Validálás* modulban a ROP-R-ben kapottnál lényegesen több adekvációs mutatóval értékelhető ki egy klasztermegoldás. Például elérhető itt a MORI index is, amely jelzi, hogy egy klasztermegoldás mennyivel markánsabb, eredetibb, mint egy véletlen adathalmazon feltárt optimális klaszterstruktúra. Az ezekkel kapcsolatos részleteket illetően lásd Vargha (2022, 8. fejezet).

10.4.2. MKA elemzés a PIK16 kérdőív négy skálájával

Az MKA elemzést szemléltető 2. példánk a 9.4.3. alponthoz ismertetett vizsgálat, melyben most MKA-val nézzük meg a PIK16 kérdőív négy skálája segítségével (Mmonit, AVhat, Rezil és Önreg, vö. B2.3. alpont és B.6. táblázat), hogy vannak-e a pszichológiai immunrendszernek olyan mintázatai, amelyek típusként értelmezhetők. Az elemzést ebben az alponthoz is az 1003 fős Btérkép2022 mintán végezzük el (vö. B2.3., illetve 9.4.3. alpont).



10.6. ábra. A PIK16 kérdőív 4 skáláján elvégzett MKA elemzés BIC-értékeinek grafikonja ROP-R-ben (összes modell $k = 3$ és 12 között)

Tekintettel arra, hogy a k -közép elemzés során a 9.4.3. alponthoz egy 9-klasztteres megoldást találtunk a legjobbnak ugyanezekre a változókra, az elemzést először a 3–12 klaszterszámokra hajtottuk végre az összes modell típusra a változókat nem standardizálva, BIC- és ICL-ábrát is kérve. ROP- R eredménylistája azt jelezte, hogy a legjobb modell 3–12 között az összes típust figyelembe véve: VEE, $k = 4$ klaszterszámmal (VEE4). Ugyanez jól látható a BIC-ábrán is (lásd 10.6. ábra), hiszen a BIC-értékek maximuma a VEE görbén található, $k = 4$ értéknél.

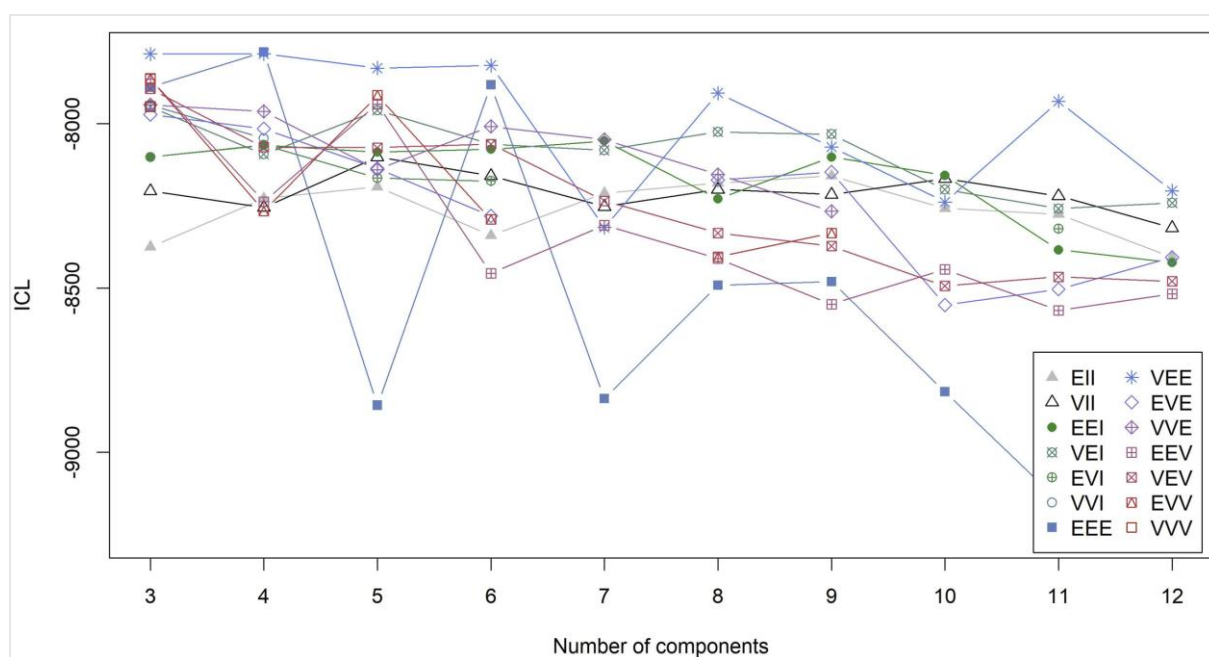
A VEE4 megoldás meglehetősen gyenge globális adekvációs mutatókkal rendelkezik (EESS% = 44,9, HCátlagS = 1,104, XBmod = 0,066), további jellemzői a standardizált átlagok mintázatának táblázatából olvasható ki (lásd 10.5. táblázat). A táblázat szerint az

MKA-val feltárt struktúrában mindössze egyetlen kellően homogén ($HC_{stan} = 0,14$) klaszter található, amelyik triviális típusú (minden változó szerint magas szintű) és viszonylag kis méretű ($n = 64$, a teljes minta 7%-át se éri el).

10.5. táblázat. A z_i -standardizált átlagok mintázata az MKA elemzés VEE4 megoldásában

Klaszter	Mmonit	AVhat	Rezil	Önreg	KLgyak	HCstan
KL1	M+	M+	M++	M+	64	0,14
KL2	A	.	A	(A)	426	1,77
KL3	.	(M)	M	.	465	0,68
KL4	(M)	(M)	M	A+++	48	0,65

Mint ahogy a VEE4 megoldás nehezen fogadható el optimálisnak, kell keresnünk egy másikat helyette. Most megtekintjük ezért az ICL-ábrát is (lásd 10.7. ábra). A BIC-ábrán a 2. legmagasabb BIC-értékű egyértelműen a VEE8 modellé, míg az ICL-ábrán ez nem olyan egyértelmű. Itt van egy tömörülés a 3 és 4 klaszterszámnál, de mint már láttuk, kevés klaszter nem elég a 4 változó sokféle lehetséges mintázatának lefedésére, így ismét a VEE görbe jöhet szóba, a VEE6 vagy a VEE5 modellel.



10.7. ábra. A PIK16 kérdőív 4 skáláján elvégzett MKA elemzés ICL-értékeinek grafikonja ROP-R-ben (összes modell $k = 3$ és 12 között)

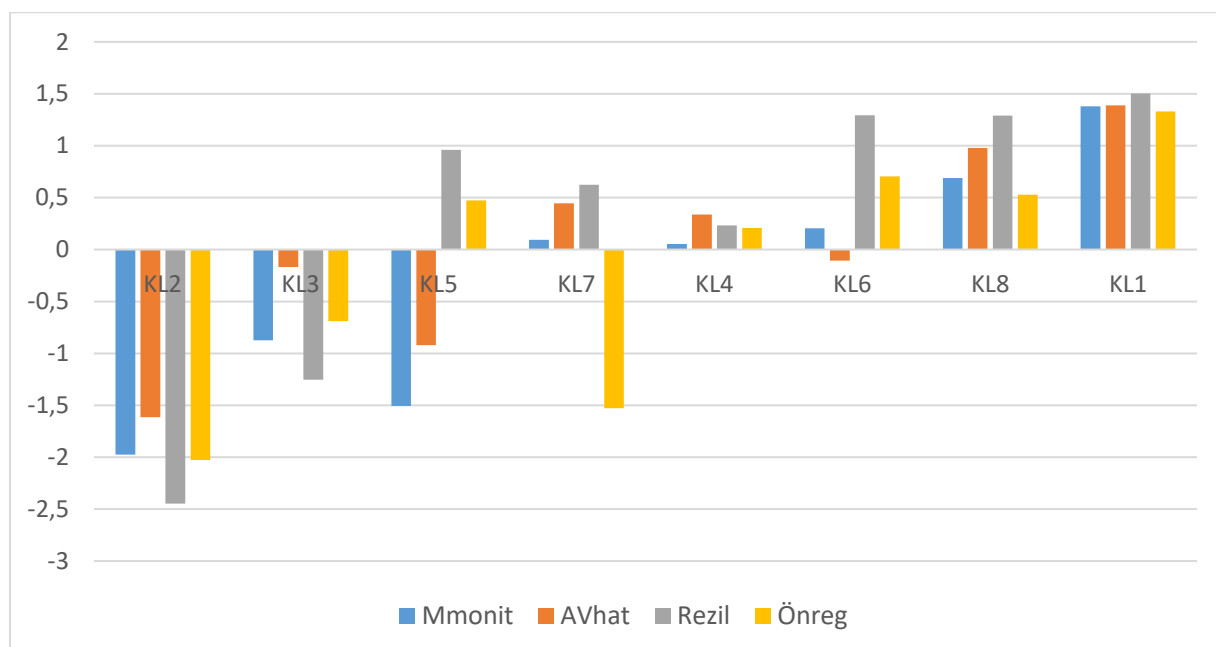
Ezek alapján a VEE8, VEE6, VEE5 modellekkel is elvégeztük az MKA-t, a VEE modell típussal, a megfelelő klaszterszámot minden esetben alkalmas klaszterszám limitekkel beállítva. A futtatások során kapott globális adekvációs mutatók értékeit a 10.6. táblázatban foglaltuk össze (az összehasonlítás kedvéért a VEE4-re vonatkozó eredményekkel együtt). A 10.6. táblázat alapján nem tudunk olyan megoldást kiválasztani, amely megfelelően klaszterezné a teljes mintát. Ezek az adatok arra utalnak, hogy legjobb esetben is csak részleges megoldásra számíthatunk. És mivel VEE8 esetén van legnagyobb esélye néhány

igazán homogén klaszter feltárására ($HC_{minS} = 0,05$), ezt a modellt vizsgáljuk meg a továbbiakban részletesen.

10.6. táblázat. Négy VEE típushoz tartozó MKA klasztermegoldás jellemzői

Modell	EESS%	XBmod	HCátlagS	HCminS	HCmaxS
VEE4	44,90	0,066	1,104	0,14	1,77
VEE5	47,96	0,088	1,043	0,17	1,71
VEE6	52,51	0,113	0,953	0,17	1,52
VEE8	58,34	-0,201	0,838	0,05	1,44

VEE8 kiértékeléséhez elkészítettük szokásos módon a z_t -standardizált átlagok oszlopdiagramját (lásd 10.8. ábra) a z_t -standardizált átlagok mintázatának táblázatát (lásd 10.7. táblázat), valamint a z_t -standardizált klasztercentroidok páronkénti ASED távolságait tartalmazó táblázatot (lásd 10.8. táblázat).



10.8. ábra. A PIK16 kérdőív 4 skáláján elvégzett MKA elemzés 9-klaszteres VEE8 megoldásában a z_t -standardizált átlagok oszlopdiagramja

10.7. táblázat. A z_r -standardizált átlagok mintázata az MKA elemzés VEE8 megoldásában

Klaszter	Mmonit	AVhat	Rezil	Önreg	KLgyak	HCstan	z_r -átlag
KL2	A+++	A++	A++++	A+++	17	0,33	-2,02
KL3	A	.	A+	A	291	1,44	-0,75
KL5	A++	A	M	(M)	41	1,14	-0,25
KL7	.	(M)	(M)	A++	83	0,83	-0,09
KL4	.	.	.	.	325	0,67	0,21
KL6	.	.	M+	M	64	0,32	0,52
KL8	M	M	M+	(M)	164	0,37	0,87
KL1	M+	M+	M++	M+	18	0,05	1,40

10.8. táblázat. A klasztercentroidok páronkénti d ASED távolságai a VEE8 megoldásban

Klaszter	KL1	KL2	KL3	KL4	KL5	KL6	KL7	KL8
KL1	0	11,78	4,79	1,43	3,67	1,01	2,87	0,33*
KL2	11,78	0	1,63	5,03	4,64	7,12	4,55	8,58
KL3	4,79	1,63	0	1,03	1,80	2,40	1,39	2,93
KL4	1,43	5,03	1,03	0	1,155	0,40*	0,80	0,51
KL5	3,67	4,64	1,80	1,16	0	0,94	2,14	2,13
KL6	1,01	7,12	2,40	0,40*	0,94	0	1,44	0,36*
KL7	2,87	4,55	1,39	0,80	2,14	1,44	0	1,33
KL8	0,33*	8,58	2,93	0,51	2,13	0,36*	1,33	0

Jelölés: *: $d < 0,50$; **: $d < 0,15$

Mindezek alapján az alábbi következtetéseket vonhatjuk le.

- A klaszterek többsége (KL2, KL3, KL4, KL8, KL1) triviális mintázatú (mindenben hasonló szintű), de vannak egyéni mintázatúak is (KL5, KL7 és KL6), ahol egy vagy két klaszter kilóg a közös szintből.
- Érdekes, hogy a két szélsőségesebb klaszter egyaránt igen homogén. A 17 fős és mindenben alacsony KL2 (z_r -átlag = -2,02) esetén HCstan = 0,33, míg a 18 fős és mindenben magas KL1 (z_r -átlag = 1,40) esetén HCstan = 0,05.
- Ez a két szélsőséges klaszter megjelenik a 7-klaszteres k -közép (röviden KK7) megoldásban is, csak határozottan heterogénebb formációban (lásd KL1 és KL7 sorát a 9.9. táblázatban). A KK7-beli KL7 klaszter mindenben magas mintázata egyébként erősen hasonlít a VEE8-beli KL8 mintázatára is.
- Még három klaszterpár erős hasonlóságát érdemes megemlíteni. A triviális típusok közül a VEE8-beli KL2 és a KK7-beli KL3, a nem triviális típusok közül pedig a VEE8-beli KL4 és KL5, illetve a KK7-beli KL5 és KL7 hasonlóságát,
- A pszichológiai immunkompetencia szempontjából legjobb két klaszter, KL1 és KL8 centruma esik egymáshoz a legközelebb ($d = 0,33$; vö. 10.8. táblázat), és mivel mindkettő elfogadhatóan homogén (HCstan < 0,40; vö. 10.7. táblázat) és hasonló mintázatú (lásd 10.8. ábra), akár össze is lehetne őket vonni egyazon klaszterbe.

Végző konklúzióként itt is helytálló az a megállapítás, hogy az MKA módszer sikerebben tud néhány erősen homogén típust azonosítani, mint a k -közép klaszteranalízis, ha csak részleges megoldás létezik a teljes populáció klasszifikációjában.

Záró gondolatok

A jelen könyv fő célja, hogy a ROP-R tíz többváltozós statisztikai moduljának megfelelő tíz fejezetével segítséget nyújtson a szoftver értő és kompetens használatához. Ezt a célt három módon próbálta elérni: 1) a szoftverhasználat pontos leírásával, 2) az érintett statisztikai módszerek és modellek elméleti hátterének és főbb fogalmainak ismertetésével, végül 3) valódi pszichológiai adatelemzések szemléltető mintapéldáival.

A ROP-R ingyenes szoftver, mint az R (R Core Team, 2021), a JASP (JASP Team, 2022) és a jamovi (The jamovi project, 2021; Şahin és Aybek, 2019). Bárki térítésmentesen használhatja, ráadásul használatát nem nehezítik meg alkalmanként váratlanul felbukkanó reklámok vagy más oda nem illő tartalmak. Nem kíván olyan átfogó statisztikai szoftver lenni, mint a JASP és a jamovi, mégis vannak figyelemre méltó előnyei, amelyek az alábbiak szerint fogalmazhatók meg.

- A ROPstat Magyarországon széles körben ismert. A programot rendszeresen használók (Magyarországon több ezer pszichológus és más szakmabeli kolléga) örömmel vehetik, hogy a ROP-R-rel futtathatnak standard többváltozós statisztikai elemzéseket a ROPstatéval teljesen megegyező keretben (a ROP-R-ben ugyanolyan a változók kezelése, a fájlkezelés, az adatok szerkesztése, a transzformációk stb.). A ROP-R-ben mindaz végrehajtható, ami a ROPstatban, kivéve annak statisztikai elemzéseit.
- A ROP-R szoftver a ROPstatra specifikus *.msw adatfájlokat is mindenféle konverzió nélkül be tudja olvasni, de automatikusan be tud olvasni adatokat egy standard Excel fájlból (ha az adattáblázat az aktív fülön található), valamint SPSS (sav és por) fájlokat.
- A ROP-R az angolon kívül magyar nyelvre is beállítható, amikor is nemcsak a szoftver kommunikációs nyelve, hanem az eredménylista is magyar nyelvű.
- A ROP-R modulok számos olyan fontos statisztikai elemzés (pl. polinomiális regresszió, k -medoid és k -medián nemhierarchikus klaszterelemzés vagy modell-alapú klaszteranalízis) végrehajtását teszik lehetővé, amelyek más felhasználóbarát szoftverekben jelenleg nem elérhetők.

A ROP-R-rel kapcsolatban kis nehézség a hozzá kötődő R szoftver és a Bevezetés B1.2. alpontjában felsorolt R-package-ek installálása, de ezt szerencsére csak egyszer kell megtenni. A ROP-R-nek szintén hiányossága, hogy csak Windows operációs rendszerben használható, és jóval szűkebb kínálatú, mint a JASP és a jamovi, amelyek tartalmukban SPSS és SAS szintű komplex szoftverek. Mindamelllett ROP-R a ROPstattal együtt lefedi a pszichológia BA és MA szakjának statisztikai anyagát és igen hasznos a doktori képzésben is. A ROP-R ROPstattal összekötött honlapja⁴⁵ nem olyan kimunkált, mint a legtöbb profi szoftveré, de megítéléséhez azt is figyelembe kell venni, hogy a ROP-R-t teljesen magánerebből hoztuk létre, ehhez sem pályázati, sem egyetemi, sem üzleti jellegű forrás vagy támogatás nem állt rendelkezésre.

⁴⁵ www.ropstat.com

A ROP-R megalkotásakor az a cél vezérelte a szerzőket (a jelen könyv szerzőjét és Bánsági Péter matematikus mérnököt), hogy a pszichológiai kutatások számára legfontosabb többváltozós statisztikai eljárások modern R-package-ek alkalmazásával a ROPstat kényelmes keretében és menürendszerében futtathatók legyenek standard beállítások, paraméterezések mellett. Ugyanakkor arra is figyeltünk, hogy eleget tehessünk a tapasztaltabb felhasználók olykor a standard szintet meghaladó igényeinek is. Ezt szolgálja ki a ROP-R azon lehetősége, hogy minden futás során fájlba menti az általa használt R-scriptet, amit a felhasználó – az adott package ismeretében – kedvére finomíthat, módosíthat, bonyolíthat. Például a konfirmatív faktorelemzésben a bifaktoros faktormodellek vizsgálata már eleve csak így módon lehetséges (lásd pl. a 6.4.5. alpont szemléltető mintapéldáját vagy a Vargha, 2024 cikkében található példákat).

Mindezekkel együtt a profi fizetős szoftverek rendelkeznek olyan tulajdonságokkal, amelyekkel ROP-R és az ingyen szoftverek nem. Például az Mplus (Muthén és Muthén, 1998-2017) grafikus lehetőségei jóval nagyobbak, mint a ROP-R-é. Több dologban a jamovi és a JASP lehetőségei is tágabbak (lásd pl. a mediációs elemzések kapcsán Vargha, 2023a, a konfirmatív faktoranalízis kapcsán pedig Vargha, 2024).

Végül jó hírként megemlíjtük, hogy a ROPstat – standard statisztikai menüje mellett – a ROP-R többváltozós menüjét is megjeleníti, ha a ROPstat észleli, hogy a számítógépen a ROP-R is szabályosan telepítve van (vagyis hogy a ROP-R.exe fájl ott van ugyanabban a „c:_vargha\ropstat” mappában, ahol a ropstat.exe). Ez azt jelenti, hogy gyakorlatilag a ROPstaton belül végrehajtható pl. egy speciális KKA vagy MKA elemzés, ennek elmenthető a klaszterváltozója, majd ezt a megoldást a ROPstat szoftver „Mintázatfeltáró elemzések/Validálás” menüpontjában a QC mutatók széles választéka és a MORI együtthatók segítségével egyben ki is lehet értékelni.

Irodalom

- Aggarwal, C. C. (2017). *Outlier Analysis, 2nd Edition*. Springer.
- Akaike, H. (1974). A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19(6), 716–723. <https://doi.org/10.1109/TAC.1974.1100705>
- Akinwande, M. O., Dikko, H. G., & Samson, A. (2015). Variance inflation factor: as a condition for the inclusion of suppressor variable (s) in regression analysis. *Open Journal of Statistics*, 5(07), 754. <https://doi.org/10.4236/ojs.2015.57075>
- Alshaya, D. S. (2022). Genetic and epigenetic factors associated with depression: An updated overview. *Saudi Journal of Biological Sciences*, 29(8), 103311. <https://doi.org/10.1016/j.sjbs.2022.103311>
- Babakus, E., Ferguson, C. E., & Jöreskog, K. G. (1987). The sensitivity of confirmatory maximum likelihood factor analysis to violations of measurement scale and distributional assumptions. *Journal of Marketing Research*, 24(2), 222–228. <https://doi.org/10.1177/002224378702400209>
- Baron, R. M., & Kenny, D. A. (1986). The moderator–mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations. *Journal of Personality and Social Psychology*, 51(6), 1173–1182. <https://doi.org/10.1037/0022-3514.51.6.1173>
- Bartha, L. (1980) (szerk.). *Pszichológiai alapfogalmak kis enciklopédiája*. Budapest: Tankönyvkiadó.
- Bech, P. (1996). *The Bech, Hamilton and Zung scales for mood disorders: screening and listening (2nd ed.)*. Springer.
- Bech, P. (2012). *The Bech, Hamilton and Zung scales for mood disorders: screening and listening: a twenty years update with reference to DSM-IV and ICD-10*. Springer Science & Business Media.
- Beck, A. T., & Beck, R. W. (1972). Screening depressed patients in family practice: A rapid technic. *Postgraduate Medicine*, 52(6), 81–85. <https://doi.org/10.1080/00325481.1972.11713319>
- Bergman, L. R., & Lundh, L. G. (2015). Introduction: The person-oriented approach: Roots and roads to the future. *Journal for Person-Oriented Research*, 1(1–2), 1–6. <https://doi.org/10.17505/jpor.2015.01>
- Bergman, L. R., Magnusson, D., & El-Khoury, B. M. (2003). *Studying individual development in an interindividual context. A Person-oriented approach*. Mahwa, New Jersey, London: Lawrence-Erlbaum Associates.
- Biesanz, J. C., Falk, C. F., & Savalei, V. (2010). Assessing mediational models: Testing and interval estimation for indirect effects. *Multivariate Behavioral Research*, 45(4), 661–701. <https://doi.org/10.1080/00273171.2010.498292>
- Biernacki, C., Celeux, G., & Govaert, G. (2000). Assessing a mixture model for clustering with the integrated completed likelihood. *IEEE transactions on pattern analysis and machine intelligence*, 22(7), 719–725. <https://doi.org/10.1109/34.865189>
- de Jong, P. F. (1999). Hierarchical regression analysis in structural equation modeling. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(2), 198–211. <https://doi.org/10.1080/10705519909540128>

- Caprara, G. V., Alessandri, G., Eisenberg, N., Kupfer, A., Steca, P., Caprara, M. G., Yamaguchi, S., Fukuzawa, A., & Abela, J. (2012). The Positivity Scale. *Psychological Assessment*, 24(3), 701–712. <https://doi.org/10.1037/a0026681>
- Cardot, H. (2022). *Gmedian: Geometric Median, k-Medians Clustering and Robust Median PCA. R package version 1.2.7*. <https://CRAN.R-project.org/package=Gmedian>
- Cattell, R. B. (1966). The scree test for the number of factors. *Multivariate Behavioral Research*, 1(2), 245–276. https://doi.org/10.1207/s15327906mbr0102_10
- Chen, F. F. (2007). Sensitivity of Goodness of Fit Indexes to Lack of Measurement Invariance. *Structural Equation Modeling: A Multidisciplinary Journal*, 14(3), 464–504. <https://doi.org/10.1080/10705510701301834>
- Cheng, C., Spiegelman, D., & Li, F. (2021). Estimating the natural indirect effect and the mediation proportion via the product method. *BMC Medical Research Methodology*, 21(1), 1–20. <https://doi.org/10.1186/s12874-021-01425-4>
- Curran, P. J., West, S. G., & Finch, J. F. (1996). The robustness of test statistics to nonnormality and specification error in confirmatory factor analysis. *Psychological Methods*, 1(1), 16–29. <https://doi.org/10.1037/1082-989X.1.1.16>
- Darvay, S. (szerk.), (1997). *Módszertani összeállítás a 0–10 éves korú gyermekek növekedésének és fejlődésének vizsgálatához. 52. sz. Módszertani levél*. Budapest: Anonymus Kiadó.
- DiStefano, C., Liu, J., Jiang, N., & Shi, D. (2018). Examination of the weighted root mean square residual: Evidence for trustworthiness? *Structural Equation Modeling*, 25(3), 453–466. <https://doi.org/10.1080/10705511.2017.1390394>
- Dolan, C. V. (1994). Factor analysis of variables with 2, 3, 5 and 7 response categories: A comparison of categorical variable estimators using simulated data. *British Journal of Mathematical and Statistical Psychology*, 47(2), 309–326. <https://doi.org/10.1111/j.2044-8317.1994.tb01039.x>
- Dunn, K. J., & McCray, G. (2020). The place of the bifactor model in confirmatory factor analysis investigations into construct dimensionality in language testing. *Frontiers in Psychology*, 11, 1357. <https://doi.org/10.3389/fpsyg.2020.01357>
- Fraley, R. C., Heffernan, M. E., Vicary, A. M., & Brumbaugh, C. C. (2011). The Experiences in Close Relationships-Relationship Structures questionnaire: A method for assessing attachment orientations across relationships. *Psychological Assessment*, 23(3), 615–625. <https://doi.org/10.1037/a0022898>
- Fraley, C., & Raftery, A. E. (2002). Model-based clustering, discriminant analysis and density estimation. *Journal of the American Statistical Association*, 97(458), 611–631. <https://doi.org/10.1198/016214502760047131>
- Fraley, R. C., & Shaver, P. R. (2000). Adult romantic attachment: Theoretical developments, emerging controversies, and unanswered questions. *Review of General Psychology*, 75(5), 132–154. <https://doi.org/10.1037/1089-2680.4.2.132>
- Füstös, L., Kovács, E., Meszéna, Gy. & Simonné Mosolygó, N. (2004). *Alakfelismerés (Sokváltozós matematikai módszerek)*. Budapest: Új Mandátum.
- Gao, C., Shi, D., & Maydeu-Olivares, A. (2019). Estimating the maximum likelihood root mean square error of approximation (RMSEA) with non-normal data: A Monte-Carlo study. *Structural Equation Modeling: A Multidisciplinary Journal*, 27(2), 192–201. <https://doi.org/10.1080/10705511.2019.1637741>

- Gergely, B., & Vargha, A. (2021). How to Use Model-Based Cluster Analysis Efficiently in Person-Oriented Research. *Journal for Person-Oriented Research*, 7(1), 22–35. <https://doi.org/10.17505/jpor.2021.23449>
- Gormley, I. C., Murphy, T. B., & Raftery, A. E. (2023). Model-Based Clustering. *Annual Review of Statistics and Its Application*, 10, 573–595. <https://doi.org/10.1146/annurev-statistics-033121-115326>
- Hayes, A. F., and Scharkow, M. (2013). The relative trustworthiness of inferential tests of the indirect effect in statistical mediation analysis: does method really matter? *Psychological Science*, 24(10), 1918–1927. <https://doi.org/10.1177/0956797613480187>
- Holtmann, J., Koch, T., Lochner, K., & Eid, M. (2016). A Comparison of ML, WLSMV, and Bayesian methods for multilevel structural equation models in small samples: A simulation study. *Multivariate Behavioral Research*, 51(5), 661–680. <https://doi.org/10.1080/00273171.2016.1208074>
- Hu, L., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling*, 6(1), 1–55. <https://doi.org/10.1080/10705519909540118>
- Hutchinson, S. R., & Olmos, A. (1998). Behavior of descriptive fit indexes in confirmatory factor analysis using ordered categorical data. *Structural Equation Modeling: A Multidisciplinary Journal*, 5(4), 344–364. <https://doi.org/10.1080/10705519809540111>
- Hunyadi, L. & Vita, L. (2008). *Statisztika közgazdászoknak*. KSH, Budapest.
- Jantek, G., & Vargha, A. (2016). A felnőtt kötődés korszerű mérési lehetősége: A közvetlen kapcsolatok élményei – kapcsolati struktúrák (ECR-RS) kötődési kérdőív magyar adaptációja párkapcsolatban élő felnőtt személyeknél. *Magyar Pszichológiai Szemle*, 71(3), 447–470. <http://dx.doi.org/10.1556/0016.2016.71.3.3>
- JASP Team (2022). *JASP (Version 0.16.2)* [Computer software]. <https://jasp-stats.org/>
- John, O. P., Donahue, E. M., & Kentle, R. L. (1991). *The Big Five Inventory - Versions 4a and 54*. Berkeley: University of California, Institute of Personality and Social Research.
- John, O. P., Naumann, L. P., & Soto, C. J. (2008). *Paradigm shift to the integrative Big Five trait taxonomy: History, measurement, and conceptual issues*. In O. P. John, R. W. Robins & L. A. Pervin, *Handbook of personality: Theory and research* (pp. 114–158). Guilford Press.
- Johnson, D. R., & Creech, J. C. (1983). Ordinal measures in multiple indicator models: A simulation study of categorization error. *American Sociological Review*, 48(3), 398–407. <https://doi.org/10.2307/2095231>
- Kassambara, A., & Mundt, F. (2020). *factoextra: Extract and Visualize the Results of Multivariate Data Analyses*. R package version 1.0.7. <https://CRAN.R-project.org/package=factoextra>
- Kaufman, L., & Rousseeuw, P. J. (1990). *Finding Groups in Data: an Introduction to Cluster Analysis*. New York: John Wiley & Sons.
- Kendall, K. M., Van Assche, E., Andlauer, T. F. M., Choi, K. W., Luykx, J. J., Schulte, E. C., & Lu, Y. (2021). The genetic basis of major depression. *Psychological Medicine*, 51(13), 2217–2230. <http://dx.doi.org/10.1017/S0033291721000441>
- Kelley, K. (2007): *Methods for the Behavioral, Educational, and Social Sciences (MBESS)* [Computer software and manual]. <https://CRAN.R-project.org/package=MBESS>
- Kenny, D. A. (2020). Measuring Model Fit. <http://davidakenny.net/cm/fit.htm>, letöltve 2023.06.14.

- Kim, E. S., & Yoon, M. (2011). Testing measurement invariance: A comparison of multiple-group categorical CFA and IRT. *Structural Equation Modeling*, 18(2), 212–228. <https://doi.org/10.1080/10705511.2011.557337>
- Kline, R. B. (2010). *Principles and practice of structural equation modeling*, 3rd ed. New York, NY: Guilford Publications.
- Koltai, J. A. (2014). Túl a regresszió – Újfajta modellek felhasználási lehetőségei a társadalomtudományokban. *Socio.hu Társadalomtudományi Szemle*, 4(2), 52–57. <https://socio.hu/index.php/so/article/view/440>
- Kopp, M., Skrabski, Á., & Czakó, L. (1990). Összehasonlító mentálhigiénés vizsgálatokhoz ajánlott módszertan. *Végeken*, 1(2), 4–24.
- Kyriazos, T., & Poga-Kyriazou, M. (2023). Applied Psychometrics: Estimator Considerations in Commonly Encountered Conditions in CFA, SEM, and EFA Practice. *Psychology*, 14(5), 799–828. <https://doi.org/10.4236/psych.2023.145043>
- Lee, S. Y., Poon, W. Y., & Bentler, P. M. (1995). A two-stage estimation of structural equation models with continuous and polytomous variables. *British Journal of Mathematical and Statistical Psychology*, 48(2), 339–358. <https://doi.org/10.1111/j.2044-8317.1995.tb01067.x>
- Li, Y. L. (2014). *Confirmatory factor analysis with continuous and ordinal data: An empirical study of stress level*. Doctoral thesis. Uppsala University.
- Li, C. H. (2016). The performance of ML, DWLS, and ULS estimation with robust corrections in structural equation models with ordinal variables. *Psychological Methods*, 21(3), 369–387. <https://doi.org/10.1037/met0000093>
- Li, C. H. (2021). Statistical estimation of structural equation models with a mixture of continuous and categorical observed variables. *Behavior Research Methods*, 53(5), 2191–2213. <https://doi.org/10.3758/s13428-021-01547-z>
- Lim, S., & Jahng, S. (2019). Determining the number of factors using parallel analysis and its recent variants. *Psychological Methods*, 24(4), 1–16. <https://doi.org/10.1037/met0000230>
- Lishinski, A. (2021). *lavaanPlot: Path Diagrams for 'Lavaan' Models via 'DiagrammeR'*. R package version 0.6.2. <https://CRAN.R-project.org/package=lavaanPlot>
- MacCallum, R. C., Roznowski, M., & Necowitz, L. B. (1992). Model modifications in covariance structure analysis: the problem of capitalization on chance. *Psychological Bulletin*, 111(3), 490–504. <https://doi.org/10.1037/0033-2909.111.3.490>
- MacKinnon, D. P. (2012). *Introduction to statistical mediation analysis*. New York: Routledge.
- MacKinnon D. P., Warsi G, & Dwyer J. H. (1995). A simulation study of mediated effect measures. *Multivariate Behavioral Research*, 30(1), 41–62. https://doi.org/10.1207/s15327906mbr3001_3
- Maechler, M., Rousseeuw, P., Struyf, A., Hubert, M., & Hornik, K. (2022). *cluster: Cluster Analysis Basics and Extensions*. R package version 2.1.4. <https://CRAN.R-project.org/package=cluster>
- Marzjarani, M. (2015). Sample size and outliers, leverage, and influential points, and Cooks distance formula. *International Journal of Arts and Commerce*, 4(2), 83–86. <https://api.semanticscholar.org/CorpusID:55026567>
- Mbachu, H. I., Nduka, E. C., & Nja, M. E. (2012). Designing a pseudo R-Squared goodness-of-fit measure in generalized linear models. *Journal of Mathematics Research*, 4(2), 148–154. <https://doi.org/10.5539/jmr.v4n2p148>

- Milligan, G. W. & Cooper, M. C. (1988). A study of standardization of variables in cluster analysis. *Journal of Classification*, 5(2), 181–204. <https://doi.org/10.1007/BF01897163>
- Mîndrilă, D. (2010). Maximum likelihood (ml) and diagonally weighted least squares (dwls) estimation procedures: a comparison of estimation bias with ordinal and multivariate nonnormal data. *International Journal of Digital Society*, 1(1), 60–66. <https://www.researchgate.net/publication/291006787>
- Moeller, J. (2025). Why and when you should avoid using z-scores in graphs displaying profile or group differences. *Journal for Person-Oriented Research*, 11(2), 58–78. <https://doi.org/10.17505/jpor.2025.28091>
- Mouselimis, L. (2022): *ClusterR: Gaussian Mixture Models, K-Means, Mini-Batch-Kmeans, K-Medoids and Affinity Propagation Clustering*. R package version 1.2.6. <https://CRAN.R-project.org/package=ClusterR>
- Muthén, B. O. (1994). Multilevel Covariance Structure Analysis. *Sociological Methods & Research*, 22(3), 376–398. <https://doi.org/10.1177/0049124194022003006>
- Muthén, B., du Toit, S. H. C. & Spisic, D. (1997). Robust inference using weighted least squares and quadratic estimating equations in latent variable modeling with categorical and continuous outcomes. *Unpublished technical report*. https://www.statmodel.com/papers_date.shtml
- Muthén, L. K., & Muthén, B. O. (1998-2017). *Mplus User's Guide. Eighth Edition*. Los Angeles, CA: Muthén & Muthén.
- Nagelkerke, N. J. (1991). A note on a general definition of the coefficient of determination. *Biometrika*, 78(3), 691–692. <https://doi.org/10.1093/biomet/78.3.691>
- Navruz, B. (2016). *The behaviors of robust weighted least squares estimation techniques for categorical/ordinal data in multilevel CFA models*. PhD Dissertation. Texas A&M University.
- Nevo, D., Liao, X., & Spiegelman, D. (2017). Estimation and inference for the mediation proportion. *The International Journal of Biostatistics*, 13(2), 20170006. <https://doi.org/10.1515/ijb-2017-0006>
- O'Brien, R. M. (2007). A caution regarding rules of thumb for variance inflation factors. *Quality & Quantity*, 41, 673–690. <https://doi.org/10.1007/s11135-006-9018-6>
- Oláh, A. (2005). *Érzelmek, megküzdés és optimális élmény*. Budapest: Trefort Kiadó.
- Oláh A. (2021). *A Globális Jól-lét Modell kidolgozása és empirikus validitásának igazolása a személyiségtényezők figyelembe vételével – A K1A: 116965. számú NKFI kutatás zárójelentése*. <https://www.otka-palyazat.hu/download.php?type=zarobeszamolo&projektid=116965>
- Oláh, A., Vargha, A., & Caprara, G. V. (2025). Hungarian validation of the Positivity Scale. *Mentálhigiéné és Pszichoszomatika*, 26(3), 150–163. <https://doi.org/10.1556/0406.2025.00088>
- Osborne, J. W. (2014). *Best practices in exploratory factor analysis*. CreateSpace Independent Scotts Valley: CreateSpace Independent Publishing.
- Pham, D. T., Dimov, S. S., & Nguyen, C. D. (2005). Selection of K in K-means clustering. *Proceedings of the Institution of Mechanical Engineers, Part C: Journal of Mechanical Engineering Science*, 219(1), 103–119. <https://doi.org/10.1243/095440605X8298>
- R Core Team (2021). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing, Vienna, Austria. <https://www.R-project.org/>

- Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48(2), 1–36. <https://doi.org/10.18637/jss.v048.i02>
- Rousseeuw, P. J. (1987). Silhouettes: a graphical aid to the interpretation and validation of cluster analysis. *Journal of Computational and Applied Mathematics*, 20(1), 53–65. [https://doi.org/10.1016/0377-0427\(87\)90125-7](https://doi.org/10.1016/0377-0427(87)90125-7)
- Rózsa, S., Kö, N., & Oláh, A. (2006). *Strukturált személyiség-kérdőívek*. In: Rózsa, S., Nagybányai Nagy, O., & Oláh, A (szerk.): *A pszichológiai mérés alapjai*. Budapest: Bölcsész Konzorcium. 223–255.
- Rózsa, S., Tárnok, Z., & Nagy, P. (2020). *A gyermekpszichiátriában alkalmazott kérdőívek, interjúk és tünetbecslő skálák*. Budapest: EFOP-2.2.0-16.2016.00002 Gyermek és ifjúságpszichiátriai, addiktológiai és mentálhigiénés ellátórendszer infrastrukturális feltételeinek fejlesztése projekt.
- Şahin, M., & Aybek, E. (2019). Jamovi: an easy to use statistical software for the social scientists. *International Journal of Assessment Tools in Education*, 6(4), 670–692. <https://doi.org/10.21449/ijate.661803>
- Saunders, C. T., & Blume, J. D. (2018). A classical regression framework for mediation analysis: fitting one model to estimate mediation effects. *Biostatistics*, 19(4), 514–528. <https://doi.org/10.1093/biostatistics/kxx054>
- Savalei, V., & Rhemtulla, M. (2013). The performance of robust test statistics with categorical data. *British Journal of Mathematical and Statistical Psychology*, 66(2), 201–223. <https://doi.org/10.1111/j.2044-8317.2012.02049.x>
- Schaffer, C. M., & Green, P. E. (1996). An empirical comparison of variable standardization methods in cluster analysis. *Multivariate Behavioral Research*, 31(2), 149–167. https://doi.org/10.1207/s15327906mbr3102_1
- Schermelleh-Engel, K., Moosbrugger, H., & Müller, H. (2003). Evaluating the fit of structural equation models: Tests of significance and descriptive goodness-of-fit measures. *Methods of Psychological Research Online*, 8(2), 23–74. <http://www.mpr-online.de>
- Schivinski, B. (2013). Implementing second-order CFA model for the factorial validity of brand equity. *PhD Interdisciplinary Journal*, 3, 53–59.
- Schreiber, J. B., Nora, A., Stage, F. K., Barlow, E. A., & King, J. (2006). Reporting Structural Equation Modeling and Confirmatory Factor Analysis Results: A Review, *The Journal of Educational Research*, 99 (6), 323–338. <https://doi.org/10.3200/JOER.99.6.323-338>
- Schwarz, G. E. (1978). Estimating the dimension of a model, *Annals of Statistics*, 6(2), 461–464. <https://doi.org/10.1214/aos/1176344136>
- Scrucca L., Fop, M., Murphy, T. B., & Raftery, A. E. (2016). mclust 5: Clustering, Classification and Density Estimation Using Gaussian Finite Mixture Models. *The R Journal*, 8(1), 289–317. <https://doi.org/10.32614/RJ-2016-021>
- Shi, D., & Maydeu-Olivares, A. (2020). The effect of estimation methods on SEM fit indices. *Educational and psychological measurement*, 80(3), 421–445. <https://doi.org/10.1177/0013164419885164>
- Shrout, P. E., & Bolger, N. (2002). Mediation in experimental and nonexperimental studies: new procedures and recommendations. *Psychological Methods*, 7(4), 422–455. <https://doi.org/10.1037/1082-989X.7.4.422>
- Susánszky, É., Konkoly Thege, B., Stauder, A., & Kopp, M. (2006). A WHO Jól-lét Kérdőív rövidített (WBI-5) magyar változatának validálása a Hungarostudy 2002 országos

- lakossági egészségfelmérés alapján. *Mentálhigiéné és Pszichoszomatika*, 7(3), 247–255.
<https://doi.org/10.1556/Mental.7.2006.3.8>
- T. Kárász, J., Nagybányai Nagy, O., Széll, K., & Takács, S. (2022). Cronbach-alfa: vele vagy nélküle?. *Magyar Pszichológiai Szemle*, 77(1), 81–98.
<https://doi.org/10.1556/0016.2022.00004>
- Tabachnick, B. G., & Fidell, L. S. (2013). *Using multivariate statistics* (6th ed.). Pearson.
- Takács, Sz., Makrai, B., & Vargha, A. (2015). Klasszifikációs módszerek mutatói. *Psychologia Hungarica Caroliensis*, 3(1), 67–88.
<https://doi.org/10.12663/PsyHung.3.2015.1.5>
- The jamovi project (2021). *jamovi* (Version 1.6) [Computer Software]. Letöltés:
<https://www.jamovi.org>
- Tjur, T. (2009). Coefficients of determination in logistic regression models—A new proposal: The coefficient of discrimination. *The American Statistician*, 63(4), 366–372.
<https://doi.org/10.1198/tast.2009.08210>
- Tucker, L. R., & Lewis, C. (1973). A reliability coefficient for maximum likelihood factor analysis. *Psychometrika*, 38(1), 1–10. <https://doi.org/10.1007/BF02291170>
- Vandenberg, R. J., & Lance, C. E. (2000). A Review and Synthesis of the Measurement Invariance Literature: Suggestions, Practices, and Recommendations for Organizational Research. *Organizational Research Methods*, 3(1), 4–70.
<https://doi.org/10.1177/109442810031002>
- Vargha, A. (2016). A ROPstat statisztikai programcsomag. *Statisztikai Szemle*, 94(11-12), 1165–1192. <https://doi.org/10.20311/stat2016.11-12.hu1165>
- Vargha, A. (2019). *Többváltozós statisztika dióhéjban: változó-orientált módszerek*. Budapest: Pólya Kiadó.
- Vargha, A. (2020). *Normális vagy? És ha nem? Statisztikai módszerek nem normális eloszlású változókkal pszichológiai kutatásokban*. Budapest: Pólya Kiadó.
- Vargha, A. (2022). *Személy-orientált többváltozós statisztika: Klasszifikációs módszerek*. Budapest: Pólya Kiadó.
- Vargha, A. (2023a). Mediációs elemzések pszichológiai kutatásokban. *Alkalmazott Pszichológia*, 25(2), 93–128. <https://doi.org/10.17627/ALKPSZICH.2023.2.93>
- Vargha, A. (2023b). Klaszterelemzések személy-orientált pszichológiai kutatásokban a ROP-R szoftver segítségével. *Alkalmazott Pszichológia*, 25(4), 115–149.
<https://doi.org/10.17627/ALKPSZICH.2023.4.115>
- Vargha, A. (2024). Konfirmatív faktoranalízis ROP-R-rel. *Alkalmazott Pszichológia*, 26(1), 149–191. <https://doi.org/10.17627/ALKPSZICH.2024.1.149>
- Vargha, A. (2025). Comments on Julia Moeller’s paper from a person-oriented perspective. *Journal for Person-Oriented Research*, 11(2), 79–84.
<https://doi.org/10.17505/jpor.2025.28093>
- Vargha, A. & Bánsági, P. (2022). ROP-R: a free multivariate statistical software that runs R packages in a ROPstat framework. *Hungarian Statistical Review*, 5(2), 3–29.
<https://doi.org/10.35618/HSR2022.02.en003>. Letölthető: <https://www.ksh.hu/hungarian-statistical-review#/year/2022?c=h#02>
- Vargha, A., Bánsági, P., & Jantek, G. (2024). Statisztikai elemzések a ROP-R szoftver segítségével és szemléltetésük egy kötődéskutatás adataival. *Mentálhigiéné és Pszichoszomatika*, 25(1), 36–55. <https://doi.org/10.1556/0406.2024.00028>

- Vargha, A. & Bergman, L. R. (2019). MORI coefficients as indicators of a “real” cluster structure. *Hungarian Statistical Review*, 2(1), 3–23. <https://doi.org/10.35618/hsr2019.01.en003>
- Vargha, A., Bergman, L. R., & Takács, S. (2016). Performing cluster analysis within a person-oriented context: Some methods for evaluating the quality of cluster solutions. *Journal for Person-Oriented Research*, 2(1–2), 78–86. <https://doi.org/10.17505/jpor.2016.08>
- Vargha, A., & Grezsa, F. (2024). Exploring types of parent attachment via the clustering modules of a new free statistical software, ROP-R. *Journal for Person-Oriented Research*, 10(1), 1–15. <https://doi.org/10.17505/jpor.2024.26255>
- Vargha, A., Jantek, G., Grezsa, F., Mirnics, Z., & Vass, Z. (2016). A szülői kötődés főbb típusai és kapcsolatuk függőségi változókkal 15-16 éves tanulóknál. In: Spannraft, M., Korpics, M., & Németh, L. (szerk.) (2016). *A család és a közösség szolgálatában. Tanulmányok Komlósi Piroska tiszteletére*. Károli Gáspár Református Egyetem, L'Harmattan Kiadó, Budapest, 269–289.
- Vargha, A. & Postáné Török, R. (2026). A new method to interpret cluster analysis results in the presence of heterogeneous clusters. *Journal for Person-Oriented Research*, 12(1), 1–12. <https://doi.org/10.17505/jpor.2026.29047>
- Vargha, A., Zábó, V., Török, R., & Oláh, A. (2020). A jóllét és a mentális egészség mérése: a Mentális Egészség Teszt. *Mentálhigiéné és Pszichoszomatika*, 21(3), 281–322. <https://doi.org/10.1556/0406.21.2020.014>
- Vigil-Colet, A., Navarro-González, D., & Morales-Vives, F. (2020). To reverse or to not reverse Likert-type items: That is the question. *Psicothema*. 32(1), 108–114. <https://doi.org/10.7334/psicothema2019.286>
- Wickham, H. (2016). *ggplot2: Elegant Graphics for Data Analysis*. Springer Verlag.