

Minták eltérésének statisztikai elemzése

BBNPS0200

Készítette: Soltész-Várhelyi Klára

Analyzis of Variances elmélete
Független mintás ANOVA

Hol tartunk az anyagban?

Elmélet

A statisztika azért van, hogy megállapítsd elhíhated-e a mért mintában talált hatásokat a populációban is létezőnek.

Adatfeldolgozás

Adatok felkészítése a statisztikai elemzésekre, hibák, outlierok szűrése. A próbák végzése feltételekhez kötődik, megbeszéltük a négy leggyakrabban előforduló feltétel ellenőrzésének módjait (függetlenség, skála típusosság, normalitás és szóráshomogenitás)

Statisztikai próbák

Változók közötti kapcsolat vizsgálata

Van összefüggés a szubjektív jólét és az önismeret közt?

Parametrikus korreláció
Pearson korreláció

Nem parametrikus korreláció
Spearman, Kendall-féle tau

Minták közötti különbség vizsgálata

Különbözik-e férfiak és nők szorongásértéke? És a biztonságosan, elkerülően vagy ambivalensen kötődő személyek párkapcsolati elégedettsége? Magasabb-e tréninget követően az önismeret?

1 minta és konstans különbsége

Parametrikus
Egymintás t-próba

2 minta különbsége

Összefüggő minták

Parametrikus Összefüggő mintás
t-próba

Nem parametrikus Wilcoxon
előjeles rang teszt, McNemar teszt

Független minták

Parametrikus
Független mintás t-próba

Nem parametrikus
Mann-Whitney, Kolmogorov
Smirnov Z, Moses extreme
reaction, χ^2

Kettőnél több minta Különbsége

Összefüggő minták

Parametrikus
Repeated Measures ANOVA

Nem parametrikus
Friedman ANOVA

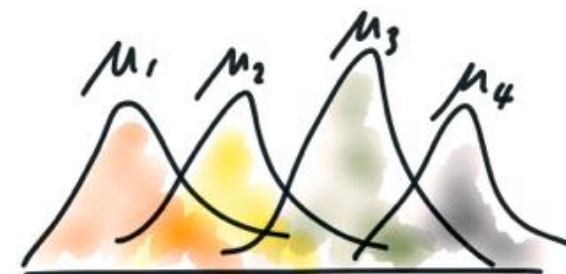
Független minták

Parametrikus
One-Way ANOVA

Nem parametrikus Kruskal-
Wallis teszt

Az ijesztő elméleti háttér

Sok ronda képlet, koncentrálj a koncepció megértésre!

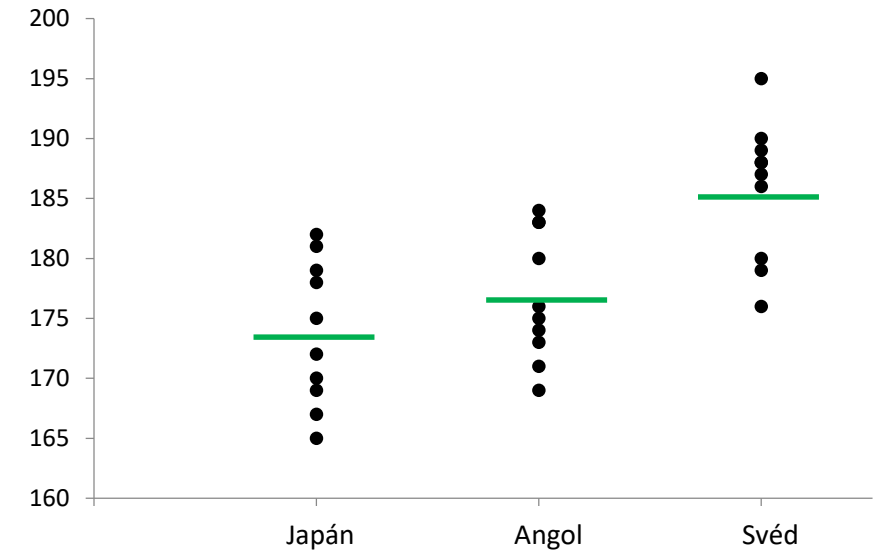


ANOVA

$$\mu_1 = \mu_2 = \mu_3 = \mu_4 ?$$

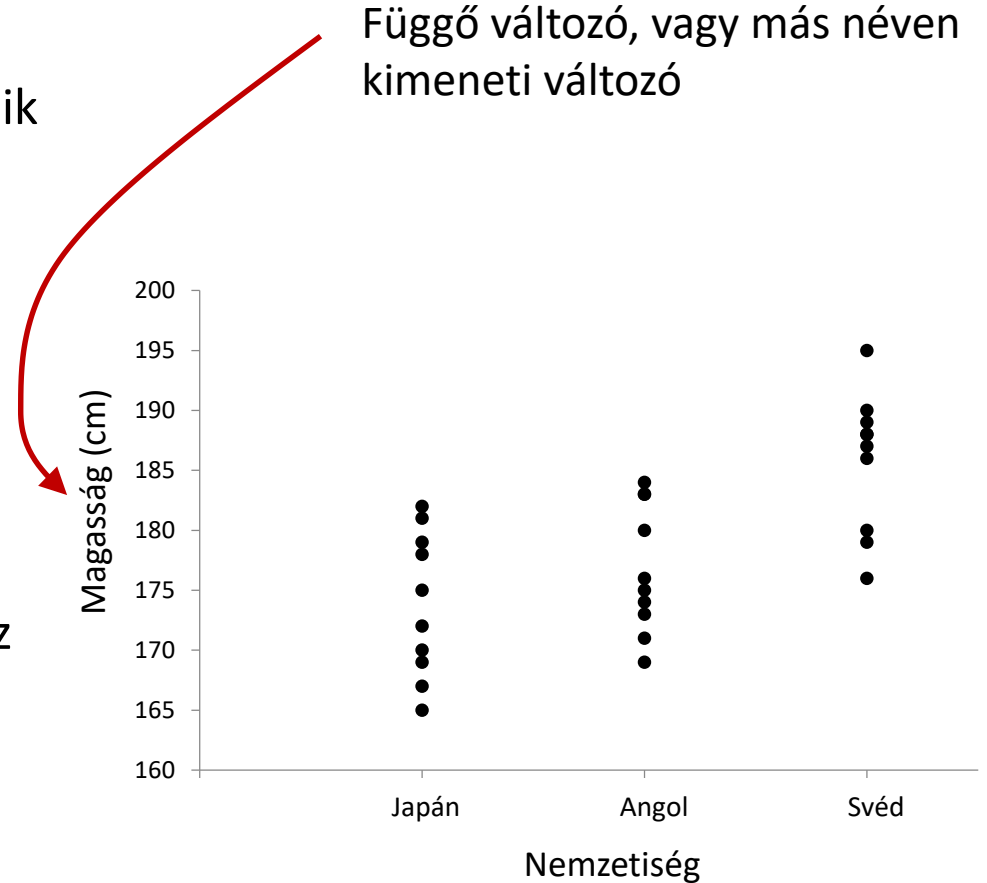
Miért nem sok t-próba

- Funkcióját tekintve
 - a t-próbához hasonlóan minták közötti különbséget mér,
 - de míg a t-próbával mindig csak két mintát lehetett összehasonlítani, az ANOVA több minta összehasonlítására is alkalmas.
 - Különbözik-e a magasság a japánok, angolok, svédek között?
- Miért nem végzünk sok t-próbát?
 - Mert minél több t-tesztet végzünk, annál jobban nő az elsőfajú hiba aránya.
 - Ez azért van, mert minden egyes teszteléskor 5%-ban lehetőséget adunk a tévedésre
 - Ezt nevezzük **familywise hibának** (FWER)– amikor ugyanazon az adaton sok azonos statisztikai próbát végzünk, elsőfajú hiba felhalmozódik
 - Öt minta esetén 10 párosítás lehetséges, tehát 10 t-próbát kellene végeznünk, hogy mindent mindennel összehasonlítsunk. Ekkor az $FWER = 1 - (.95)^{10} = 40\%$, azaz 40% valószínűsége van annak, hogy a kapott szignifikáns különbségeink közül legalább egy a véletlen műve



A változók

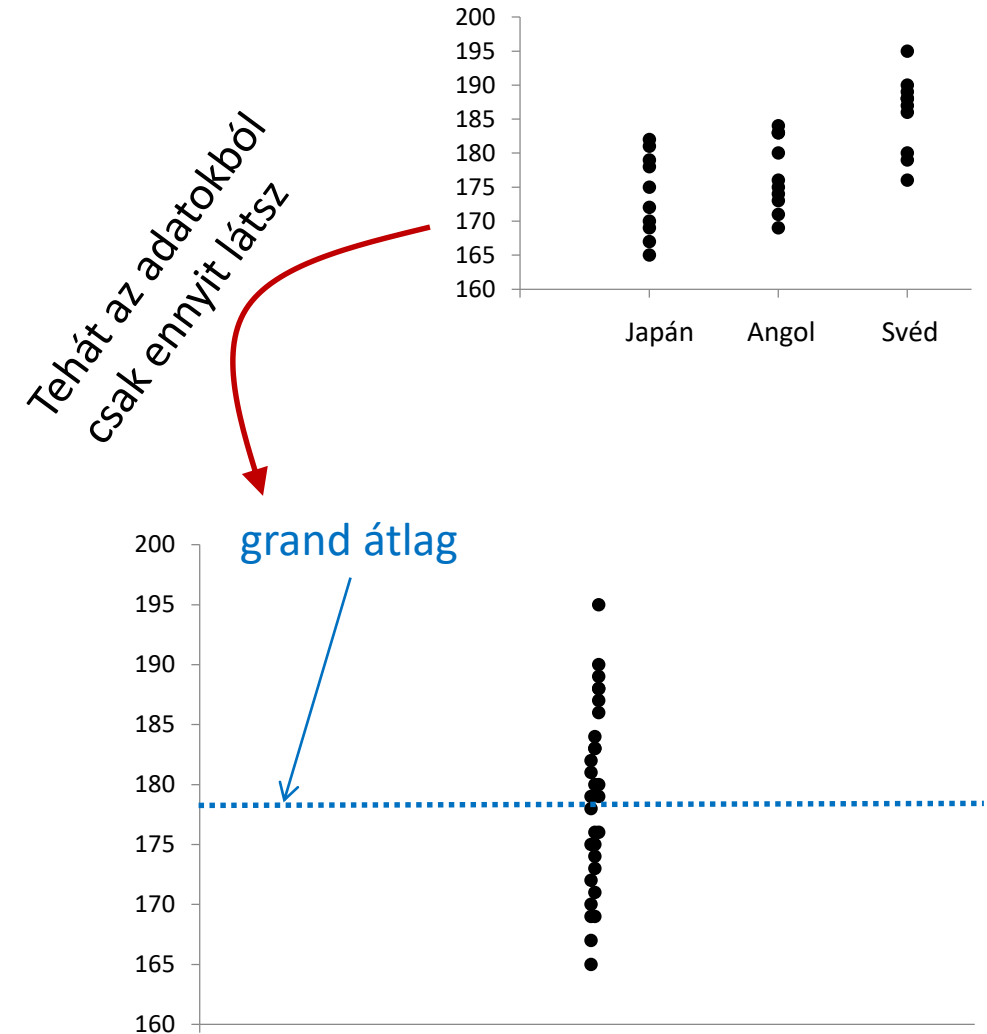
- Különbözik-e japánok, angolok, svédek magassága egymástól?
- **A nullhipotézis**
 - A három csoport nem különbözik egymástól
 - Tehát a három nemzet átlaga a populációban ugyanoda esik – magasság szempontjából egy populációból származnak
- **Függő változó**
 - Vagy más néven a **kimeneti változó** – a magasság
 - Kimenetinek nevezzük, mert ennek az értékét szeretnénk bejósolni
- **Független változó**
 - Vagy más néven a **prediktor változó**: a nemzetiség
 - Független mintás elemzésnél ez a csoportosító változó lesz
 - Összefüggő mintás elemzésnél egy kicsit bonyolultabb a helyzet, itt a prediktor szerepét a kondíciók töltik majd be
- Az elmélet a független mintás ANOVA mentén halad, de az itt megértett eljárás átültethető az összefüggő mintás elemzésre is



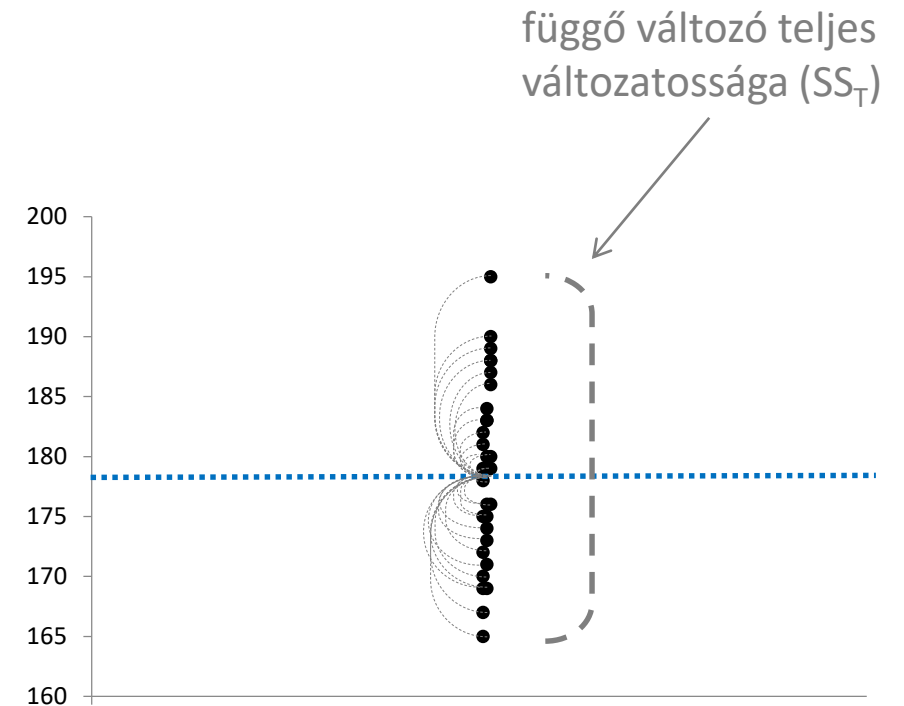
Független változó, vagy más néven a prediktor

A referencia modell

- Mivel itt már több változót hasonlítottunk össze, kell valami viszonyítási érték, egy referencia modell. Ettől fogjuk mérni a távolságokat a minták elhelyezkedésének vizsgálatakor.
- Mi legyen a referencia?
 - Mekkora magasságot jósolnánk akkor, ha nem tudnánk ki milyen nemzetiségű, csak megkapnánk egy csomó ember magasság adatát?
- **Referencia modell**
 - **Grand átlag** – a függő változó átlaga az összes csoportot egyben számolva
 - A legegyszerűbb modell, melyben még nem használjuk fel a független változó információját
 - Amihez majd hasonlítani tudjuk a modellünk (mennyivel másabb a predikció, ha figyelembe vesszük a prediktort)
- N_T = teljes elemszám
- $M_T = \frac{\sum x_i}{N_T}$ = grand átlag (mean of total) az összes ember magasságának átlaga (a nemzetiségtől függetlenül)



- Mekkora tévedünk ezzel a referencia modellel?
 - Azaz mekkora a különbség a grand átlaggal predikált érték és a ténylegesen mért értékek között
- **SS_T** azaz **Sum of Squares of Total**
 - Megadja, mekkora az összváltozatosság a kimeneti változóban (vedd észre, a képlete mennyire emlékeztet a variancia, illetve a szórás képletére!)
 - A mért értékek grand átlagtól való távolságának négyzetösszege
 - Itt még mindig nem számít, ki milyen nemzetiségű
- **SS_T** = $\sum (x_i - M_T)^2$



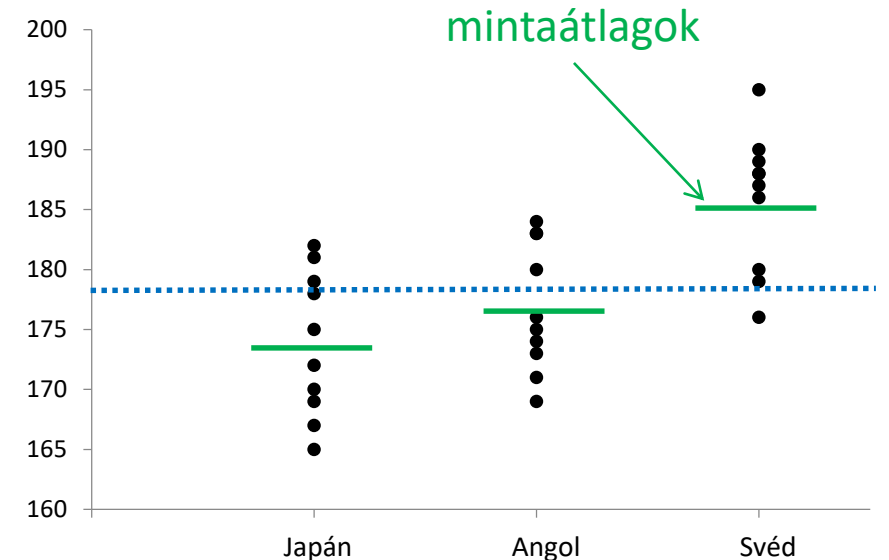
A predikált érték

- Mekkora magasságot jósolnánk akkor, ha tudnánk az emberek nemzetiségét is?
 - Ha már tudjuk a nemzetiséget, nyilván figyelembe vennénk a jóslatunkkor, és egy japán embernek a japán minta átlagát, a svédnek a svéd minta átlagot jósolnánk.
- **Predikált érték**
 - ANOVÁBAN a predikált értékek a mintaátlagok lesznek
 - a nemzetiséget, mint prediktort is figyelembe vesszük a becslés során, azaz a grand átlag helyett például a japánoknak a japán minta átlagát fogjuk predikálni, mert ezáltal – ha a prediktornak van hatása a függő változóra – pontosabb becslést tudunk tenni, mint amilyen a grand átlag lenne

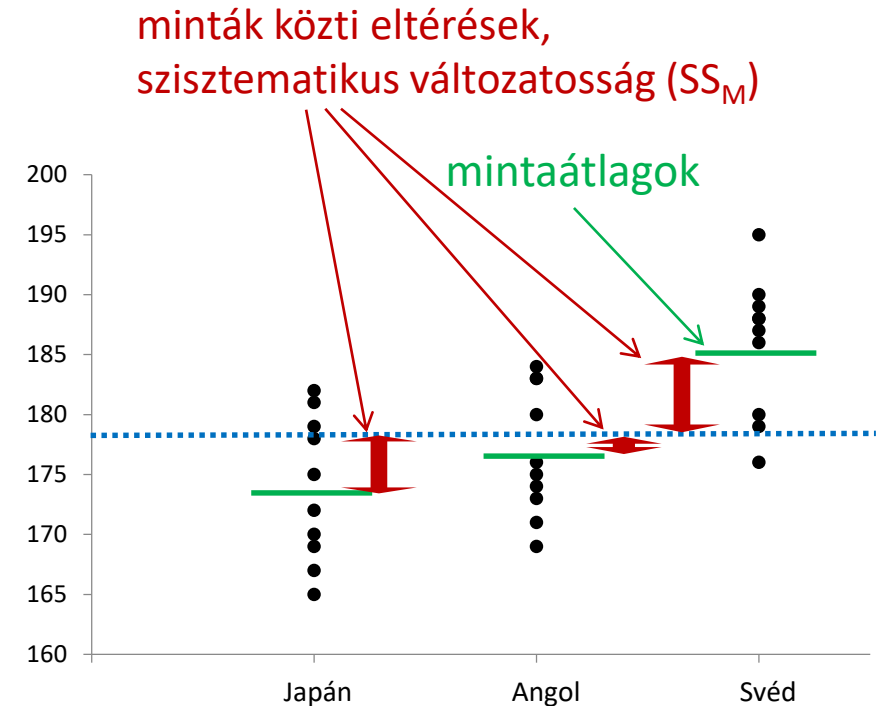
- $$M_{jap} = \frac{\sum x_{jap,i}}{N_{jap}}$$

- $$M_{ang} = \frac{\sum x_{ang,i}}{N_{ang}}$$

- $$M_{sv} = \frac{\sum x_{sv,i}}{N_{sv}}$$



- Az előbb azt mondtam, a független változót is figyelembe véve pontosabb becslést tudunk adni a függő változóra, mint amilyen a grand átlag volt. Mennyivel is predikálunk most mást? Annyival, amekkora különbség a mintaátlagok és a grand átlag között van.
- **SS_M** azaz **Sum of Squares of Model**
 - **A minták közötti eltérés**
 - Megadja, hogy mennyivel tudunk mást, többet mondani a nemzetiséget is figyelembe véve, mintha csak a grand átlagot használnánk a predikcióra, mennyiben tudunk a különböző nemzetiségekre különböző predikciót tenni
 - A modell által predikált értékek (mintaátlagok) és a grand átlag távolságának négyzetösszege
 - Minden egyed esetében kiszámoljuk a predikció (mintaátlag) és grandátlag közötti különbséget, mely azonban itt az egyes csoportokon belül azonos érték lesz
 - $SS_M = \sum (M_j - M_T)^2 * N_j$ tehát a japán minta átlaga és a grandátlag különbségének négyzete a japán minta elemszámával szorozva, plusz ugyanez az angol és svéd mintára is kiszámolva, és hozzáadva

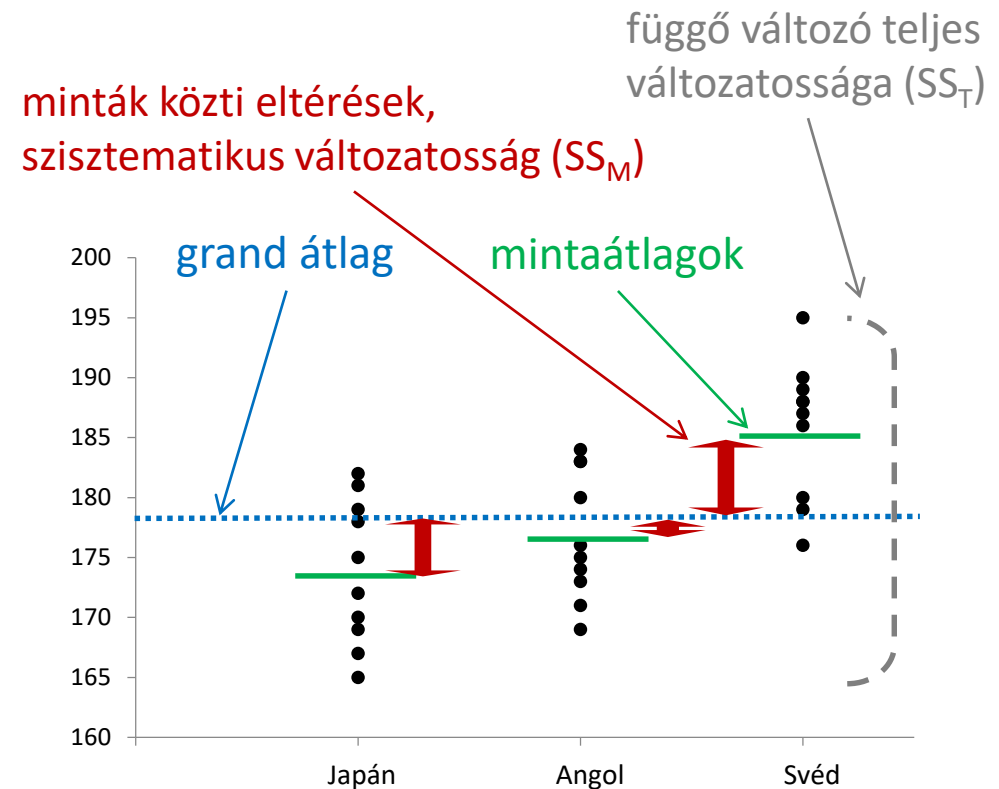


ANOVA effect-size mutatója

- **Effect size**

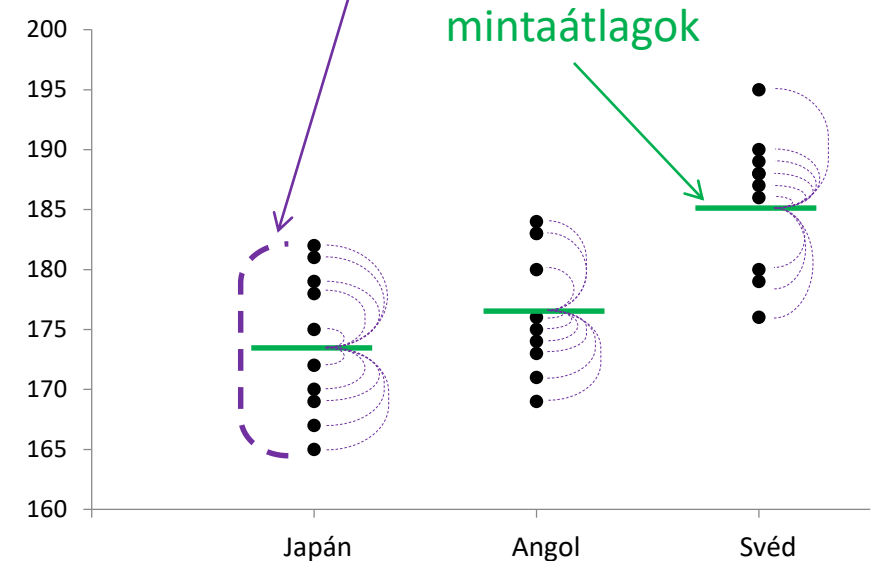
- Az effect-size a hatás nagyságát méri. Itt a képletből is könnyen látható, hogy mit jelent ez. Azt, hogy a függő változó változatosságának mekkora részét magyarázza a független változó, azaz hány százalékban magyarázza az emberek magasságában látható különbségeket az, hogy különböző nemzetiségűek.

- $R^2 = \frac{SS_M}{SS_T}$
- Az eddig tanult statisztikai próbákban r-értéket használtunk effect-size mutatóként. Ez ugyanaz a mutató, csak négyzetre emelve.
- Az négyzetes megjelenítés nagy előnye, hogy szemléletes a jelentése: az R^2 egy százalékban értendő érték, megadja, hogy a prediktor változó hány százalékát magyarázza a kimeneti változó változatosságának.



- Ez mind szép, a független változót figyelembe vevő modellünk pontosabb, mint a referencia modell, de azért ez semtökéletes.
 - Mekkora a tévedésünk, ha a nemzetiséget figyelembe véve jósoljuk meg a magasságot?
- **SS_R** azaz **Sum of Squares of Reziduals**
 - Mekkora az összeltérés a predikált és a ténylegesen mért értékek között – azaz vesszük a négyzetes eltérést minden személy értéke és a csoportjához tartozó mintaátlag között, és ezeket a négyzetes eltéréseket összegezzük
 - Megadja, hogy mekkorát tévedünk, amikor a magasságot a csoportátlagokkal predikáljuk
 - Azaz ez a teljes változatosság azon része, amit a modellel nem tudunk magyarázni (azt hogy vannak alacsonyabb és magasabb japánok, a nemzetiség nem magyarázza)
 - Nevezik a **minták belüli eltérések**nek, vagy nemszisztematikus varianciának is
 - $SS_R = \sum (x_{j,i} - M_j)^2$

mintán belüli eltérések,
nemszisztematikus változatosság (SS_R)



F-érték, a statisztikai érték

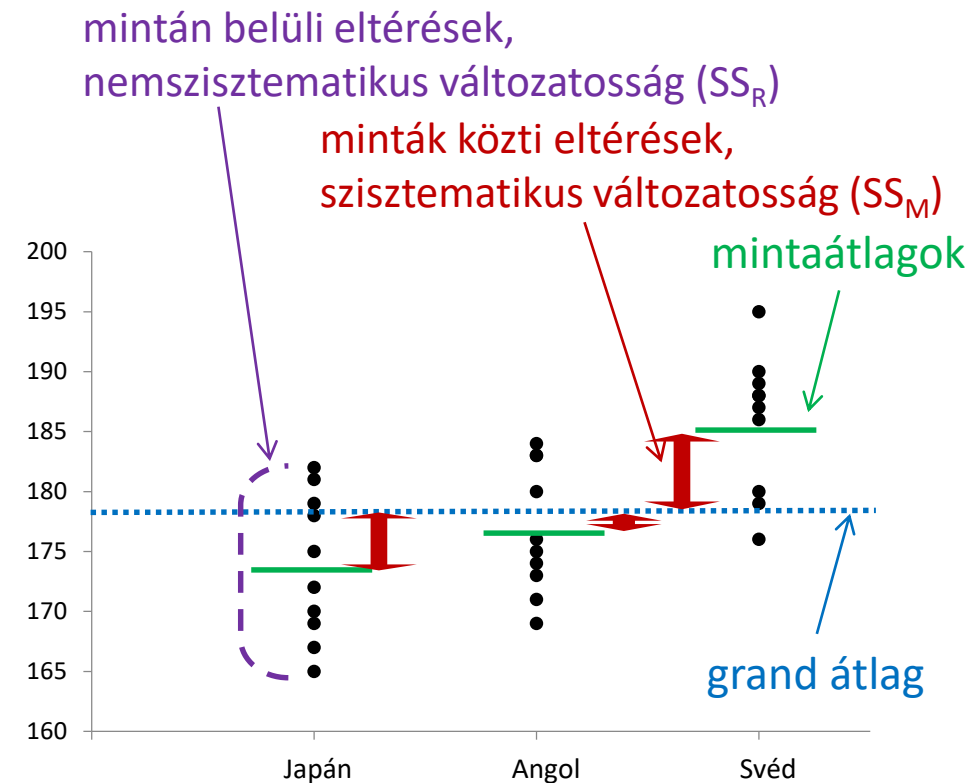
- **Statisztikai érték**

- Most jön pár számolás, mert a statisztikai értékhez ki kell számolnunk az átlagos hatást, és átlagos hibát. De nem történik más, minthogy átlagoljuk az előbb kiszámolt és megértett SS_M és SS_R értékeket.
- df_M = a modell szabadságfoka (a modell összetettségének mutatója – itt csoportok száma - 1)
- df_R = a hiba szabadságfoka (N - predikált paraméterek száma)
- $MS_M = SS_M / df_M$ = hatás
- $MS_R = SS_R / df_R$ = hiba

- **A statisztikai érték a hatás és hiba aránya**

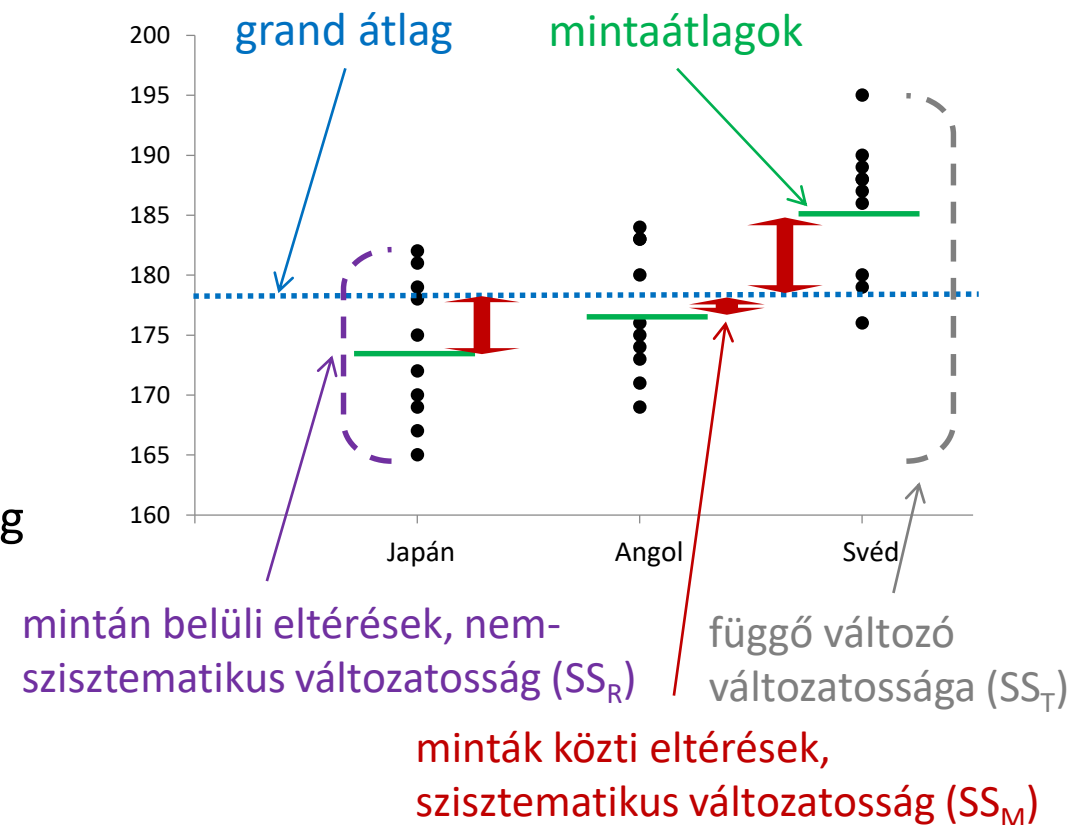
$$F = \frac{MS_M}{MS_R} = \frac{\frac{SS_M}{df_M}}{\frac{SS_R}{df_R}}$$

- Az F -értékhez a szabadságfokok ismeretében hozzárendelhető egy p -érték is



Összegezve

- Összegezve az ANOVÁban mindig van egy független, azaz **prediktor**, és egy függő, azaz **kimeneti változó**. A prediktor segítségével szeretnénk bejósolni a kimeneti változó értékét. Azaz szeretnénk jóslatot tenni arra ki milyen magas a nemzetiségét is figyelembe véve. Ezek a jóslatok, azaz **predikált értékek** a **mintaátlagok** lesznek.
- Látjuk, hogy a kimeneti változónak van valamekkora változatossága, hiszen vannak a alacsonyabb, és magasabb emberek is. Ezt nevezzük **teljes változatosságnak**.
- Ezt a változatosságot **részben meg tudjuk magyarázni** a prediktorral, mivel a különböző nemzetiségű személyekből álló **minták között különbség** van a magasságban (a japánok általában alacsonyabbak, a svédek általában magasabbak).
- Részben pedig **nem tudjuk megmagyarázni**, hiszen a **mintán belül** is vannak **különbségek** (azaz vannak alacsonyabb és magasabb japánok is)
- Az ANOVA effect-size mutatója megadja, a teljes változatosság mekkora részét tudjuk megmagyarázni.
- Az ANOVA statisztikai értéke pedig a átlagosan megmagyarázott és meg nem magyarázott változatosság aránya.



minta	érték
-------	-------

1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5



Elemsszámok:

N_{Total} a teljes vizsgálat elemszáma

$$N_T = 15$$

N_{Minta} a minták elemszámai

$$N_{\text{Minta1}} = 5 \quad N_{\text{Minta2}} = 4 \quad N_{\text{Minta3}} = 6$$

Szabadságfokok:

df_{Total} a teljes mintához tartozó szabadságfok

$$df_T = N_T - 1 = 14$$

df_{Minta} a mintákhoz tartozó szabadságfok

$$df_{\text{Minta1}} = N_{\text{Minta1}} - 1 = 4 \quad df_{\text{Minta2}} = 3 \quad df_{\text{Minta3}} = 5$$

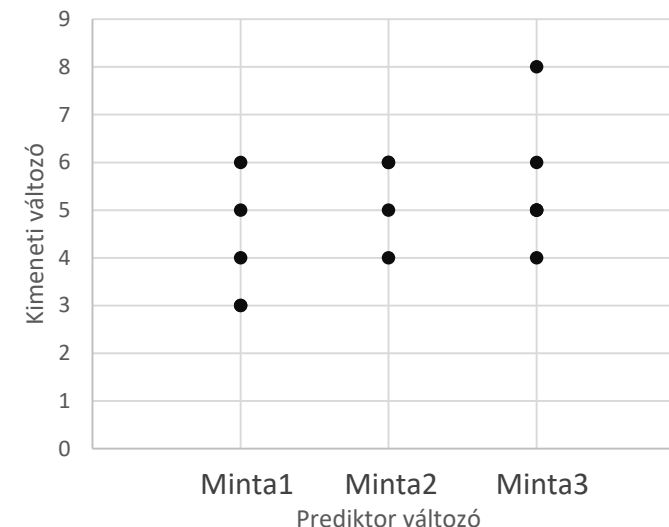
df_{Residual} a hibához (mintákon belüli változatosság) tartozó szabadságfok

$$df_R = df_{\text{Minta1}} + df_{\text{Minta2}} + df_{\text{Minta3}} = 4 + 3 + 5 = 12$$

df_{Model} a modelhez (minták közötti változatosság) tartozó szabadságfok. Hányszor tudok a minták közül szabadon választani? Mintaszám – 1 alkalommal

$$df_M = 3 - 1 = 2$$

Példa 3 mintával



Minta	N	df	M
Minta1	5	4	
Minta2	4	3	
Minta3	6	5	
Total	15	14	
ANOVA	df	SS	MS
Model	2		
Residual	12		
Total	14		

minta	érték
-------	-------

1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5



Átlagok:

Grand átlag (M_{Total}) az összes érték átlaga

A grandátlag lesz a referencia

(X-szel az egyedek vannak jelölve)

$$M_T = \frac{\sum X}{N} = \frac{3+4+5+3+6+5+6+4+6+4+5+8+6+5+5}{15} = 5$$

Minták átlagai (a mintához tartozó egyedek átlagai)

A mintaátlagok lesznek a különböző mintákhoz predikált értékek

(X_{Minta1} -gyel az első mintához tartozó egyedek vannak jelölve)

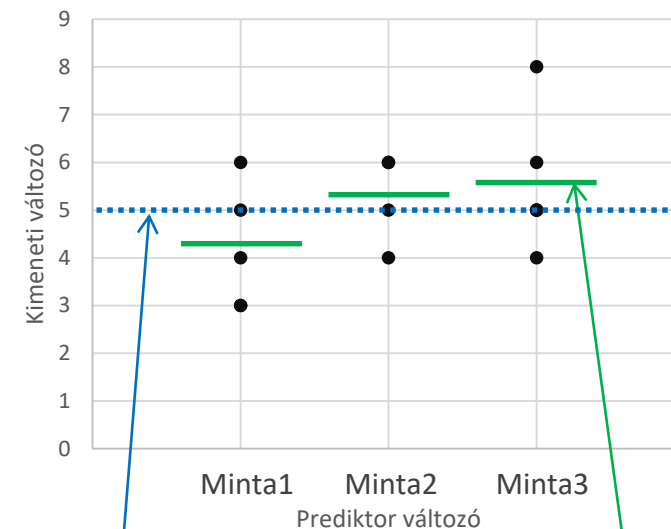
$$M_{\text{Minta1}} = \frac{\sum X_{\text{MINTA1}}}{N_{\text{MINTA1}}} = \frac{3+4+5+3+6}{5} = 4.2$$

$$M_{\text{Minta2}} = \frac{\sum X_{\text{MINTA2}}}{N_{\text{MINTA2}}} = \frac{5+6+4+6}{4} = 5.25$$

$$M_{\text{Minta3}} = \frac{\sum X_{\text{MINTA3}}}{N_{\text{MINTA3}}} = \frac{4+5+8+6+5+5}{6} = 5.5$$

Minta	N	df	M
Minta1	5	4	4.2
Minta2	4	3	5.25
Minta3	6	5	5.5
Total	15	14	5
ANOVA	df	SS	MS
Model	2		
Residual	12		
Total	14		

Példa 3 mintával



grand átlag

mintaátlagok

minta	érték
-------	-------

1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5



Total Sum of Squares

SS_{Total} – összes mért érték (függetlenül attól melyik mintába tartozik) és a grandátlag távolságának négyzetösszege

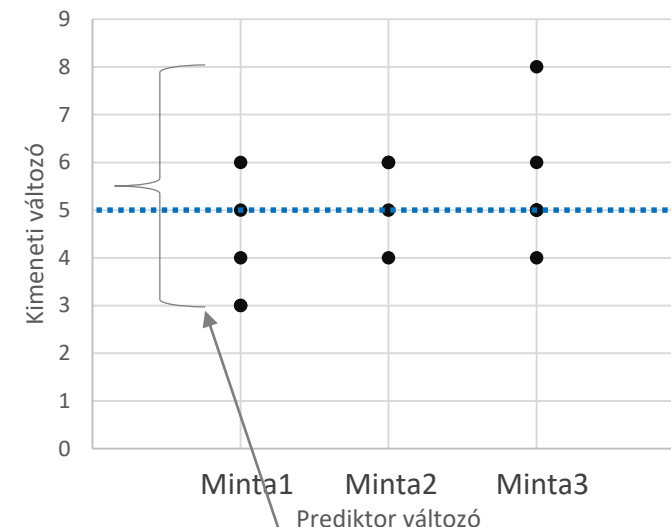
Ez a függő változó teljes változatossága

Számolása során minden személy értékéből kivonjuk a grandátlagot, a kapott különbséget négyzetre emeljük és ezeket összeadjuk.

$$SS_T = \sum (x_i - M_T)^2 = (3 - 5)^2 + (4 - 5)^2 + (5 - 5)^2 + (3 - 5)^2 + (6 - 5)^2 + (5 - 5)^2 + (6 - 5)^2 + (4 - 5)^2 + (6 - 5)^2 + (4 - 5)^2 + (5 - 5)^2 + (8 - 5)^2 + (6 - 5)^2 + (5 - 5)^2 + (5 - 5)^2 = 24$$

minta	N	df	M
Minta1	5	4	4.2
Minta2	4	3	5.25
Minta3	6	5	5.5
Total	15	14	5
ANOVA	df	SS	MS
Model	2		
Residual	12		
Total	14	24	

Példa 3 mintával



Függő változó teljes varianciája

minta	érték
1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5



Model Sum of Squares

SS_{Model} – minden mért érték helyén a modell által predikált érték (mintaátlag) és a grandátlag távolságának négyzetösszege

Megadja, mennyivel tér el a modell által predikált érték a grandátlagtól.

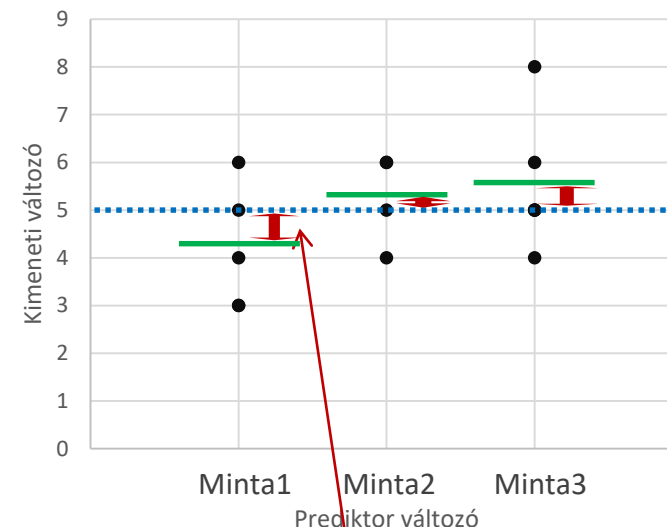
ANOVÁban kicsit más a megfogalmazás, itt azt mondjuk, ez a minták közötti eltérések összege. De e kettő ugyanazt jelenti – ha nem lenne a minták között eltérés, akkor minden mintaátlag a grandátlagnál lenne, tehát a predikált érték (mintaátlag) és grandátlag távolsága 0 lenne.

Minta	N	df	M
Minta1	5	4	4.2
Minta2	4	3	5.25
Minta3	6	5	5.5
Total	15	14	5
ANOVA	df	SS	MS
Model	2	4.95	
Residual	12		
Total	14	24	

Számolása: egy mintán belül minden egyednél a mintaátlagot predikáljuk, tehát minden minta esetén a mintaelemszám-szor kell a mintaátlag és a grandátlag közötti különbség négyzetét összeadni

$$SS_M = \sum N_{MINTA,j} * (M_{MINTA,j} - M_T)^2 = 5*(4.2 - 5)^2 + 4*(5.25 - 5)^2 + 6*(5.5 - 5)^2 = 4.95$$

Példa 3 mintával



minták közti eltérések,
szisztematikus variancia

minta	érték
-------	-------

1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5



Residual Sum of Squares

SS_{Residual} – mintákon belül a mért értékek és a mintaátlagok távolságának négyzetösszege
Egy mintán belül az egyedekhez a mintaátlagot predikáljuk, és ezzel valamilyen mértékben tévedünk. ANOVÁ-ban a hiba a mintán belüli eltérésekből adódik.

Számolása: minden egyed és a hozzá tartozó mintaátlag különbségének négyzetösszege

$$SS_{\text{RMinta1}} = (3 - 4.2)^2 + (4 - 4.2)^2 + (5 - 4.2)^2 + (3 - 4.2)^2 + (6 - 4.2)^2 = 6.8$$

$$SS_{\text{RMinta2}} = 2.75$$

$$SS_{\text{RMinta3}} = 9.5$$

Majd ezeket a kapott hibákat összeadjuk:

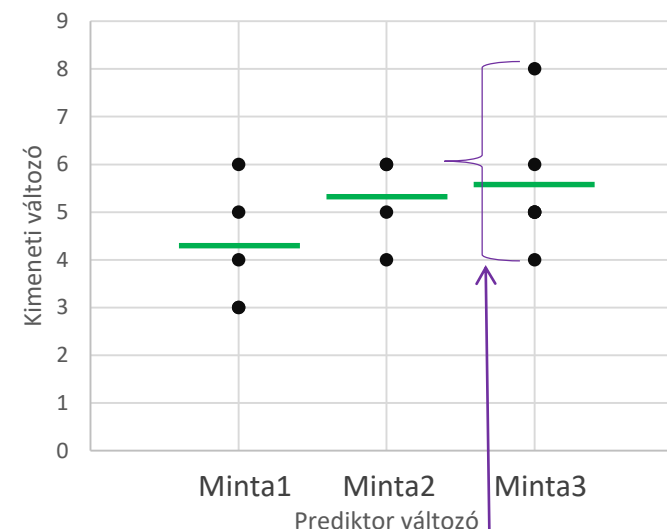
$$SS_R = SS_{\text{RMINTA1}} + SS_{\text{RMINTA2}} + SS_{\text{RMINTA3}} = 19.05$$

Ellenőrzésként: a teljes változatosság felbontható a modell által megmagyarázott és meg nem magyarázott részre, így ez utóbbi kettő összege ki kell hogy adja az előbbit.

$$SS_M + SS_R = SS_T \checkmark$$

Minta	N	df	M
Minta1	5	4	4.2
Minta2	4	3	5.25
Minta3	6	5	5.5
Total	15	14	5
ANOVA	df	SS	MS
Model	2	4.95	
Residual	12	19.05	
Total	14	24	

Példa 3 mintával



mintán belüli eltérések,
nemszisztematikus variancia

Példa 3 mintával



Mean Squares értékek

Ezt követően kiszámoljuk a modell és hibához tartozó varianciát. ANOVÁban ezt szisztematikus és nemszisztematikus varianciának nevezzük

MS_{Model} - a modellhez tartozó variancia (szisztematikus variancia)

$$MS_M = \frac{SS_M}{df_M} = \frac{4.95}{2} = 2.475$$

$MS_{\text{Residuals}}$ – a hibához tartozó variancia (nemszisztematikus variancia)

$$MS_R = \frac{SS_R}{df_R} = \frac{19.05}{12} = 1.588$$

Ki lehetne számolni az SST-hez tartozó teljes varianciát is, ennek értéke megegyezik a függő változó varianciájával. ANOVÁ-ban a számoláshoz ezt nem használjuk, ezért a táblázatokban is üresen szokás hagyni

MS_{Total} – a függő változó teljes varianciája

$$MS_T = \frac{SS_T}{df_T} = \frac{24}{14} = 1.714$$

minta	érték
1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5

Minta	N	df	M
Minta1	5	4	4.2
Minta2	4	3	5.25
Minta3	6	5	5.5
Total	15	14	5
ANOVA	df	SS	MS
Model	2	4.95	2.475
Residual	12	19.05	1.588
Total	14	24	1.714

minta	érték
1	3
1	4
1	5
1	3
1	6
2	5
2	6
2	4
2	6
3	4
3	5
3	8
3	6
3	5
3	5



R² (effect-size)

Az R² a megmagyarázott és teljes változatosság arányát adja meg

SPSS One-Way ANOVÁ-nál nem számolja ki :(

$$R^2 = \frac{SS_M}{SS_T} = \frac{4.95}{24} = \mathbf{0.206}$$

F érték (próba statisztikai értéke)

Az F érték a hatás és a hiba arányából adódik

$$F = \frac{MS_M}{MS_R} = \frac{2.475}{1.588} = \mathbf{1.559}$$

p érték

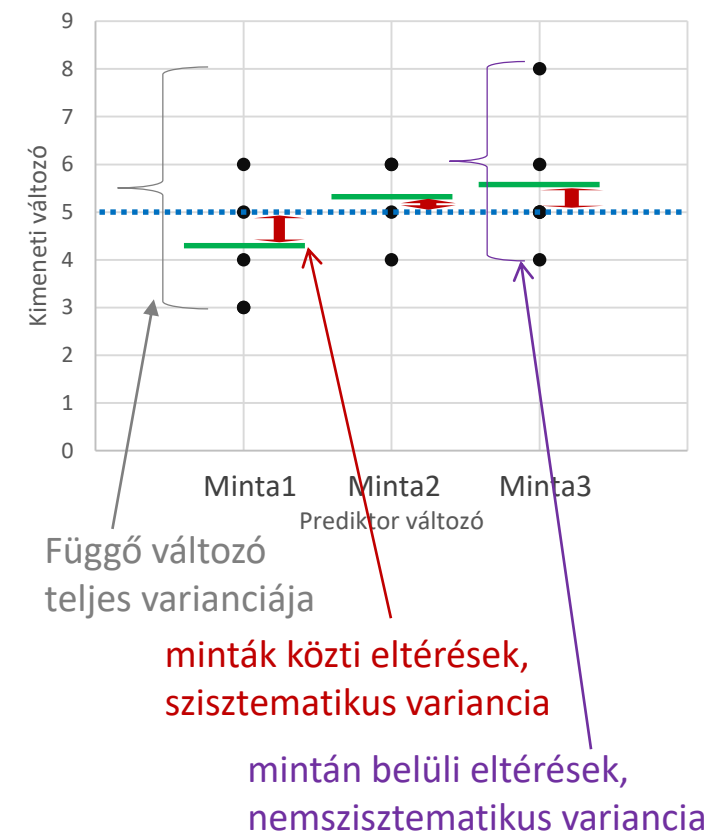
A p értéket az F-érték a df_M és df_R segítségével az F-eloszlás táblából lehet kikeresni

$$\mathbf{p = .250}$$

Minta	N	df	M
Minta1	5	4	4.2
Minta2	4	3	5.25
Minta3	6	5	5.5
Total	15	14	5

ANOVA	df	SS	MS	F	Sig.
Model	2	4.95	2.475	1.559	.250
Residual	12	19.05	1.588		
Total	14	24	1.714		

Példa 3 mintával

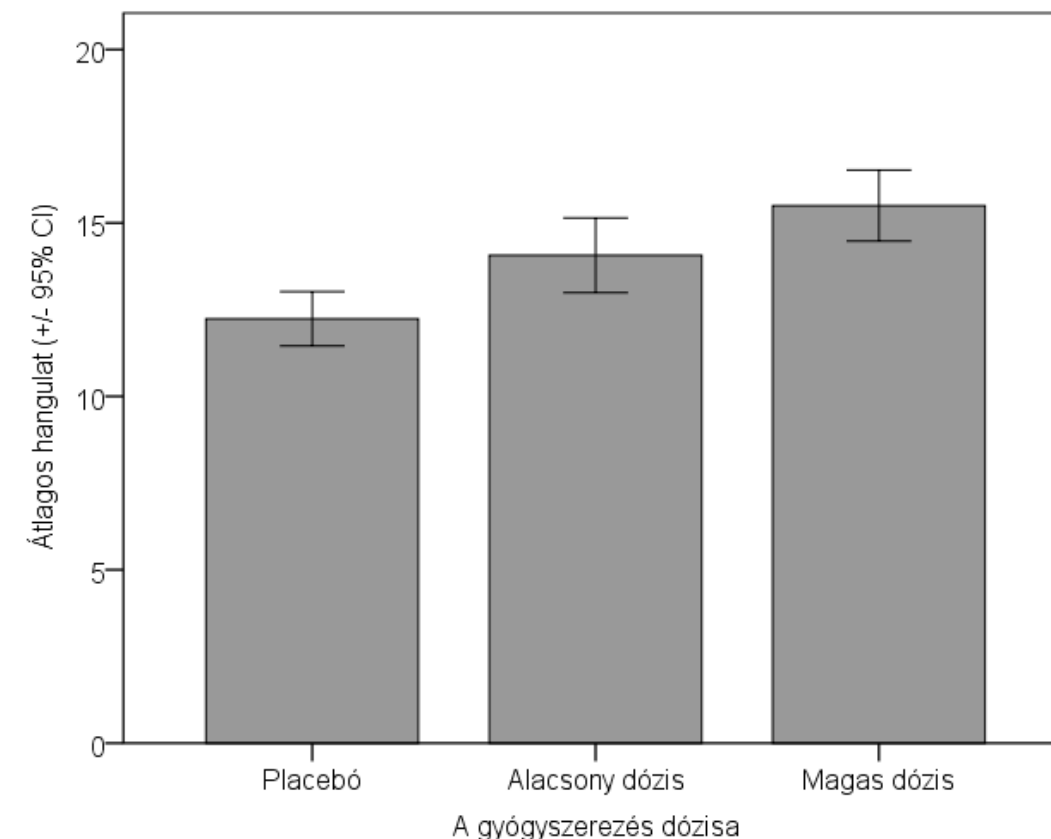


	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	4,950	2	2,475	1,559	,250
Within Groups	19,050	12	1,588		
Total	24,000	14			

One-Way ANOVA azaz a független mintás egyszempontos ANOVA

Lássuk, a ijesztő elmélet után mennyire egyszerű is az egész a gyakorlatban!

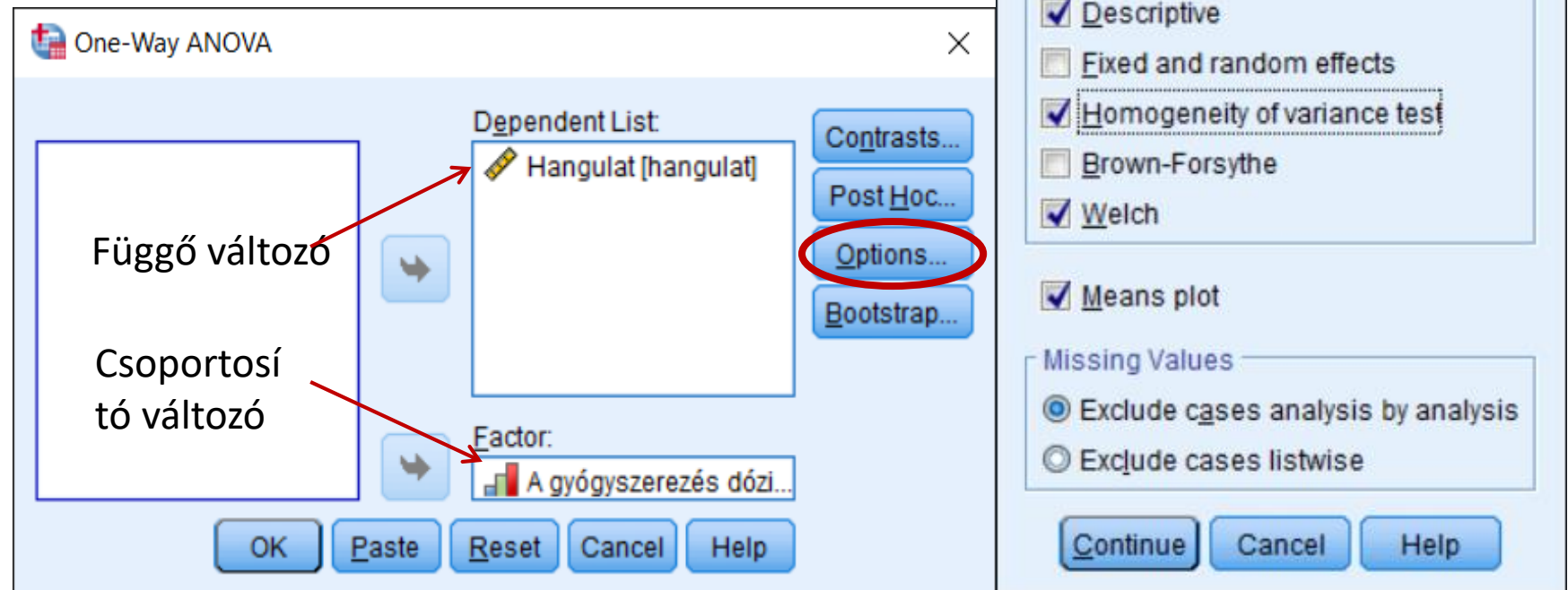
- `statgyakGY2_04_anova1_gyogyszerezes.sav`
- (A példa Andy Field ötletén alapul, a hatások tendenciája hasonló, de az itt szereplő adatbázis a könyvben szereplőtől eltérő értékeket tartalmaz.)
- Egy új hangulatjavító gyógyszer hatásosságát szeretnénk vizsgálni. A kérdés az, hogy van-e hatása a gyógyszernek, azaz van-e különbség a három csoport hangulatértéke között.
- Adatok:
 - Dózis:
 - A minta depressziós személyek adatait tartalmazza. Három csoportot hozunk létre:
 - Placebó: úgy tudják, gyógyszert kapnak, pedig valójában csak vitaminokat. Ezzel a csoporttal szeretnénk a placebo-hatást ellenőrizni.
 - Alacsony dózis: a hangulatjavító gyógyszerből alacsony dózist kapnak.
 - Magas dózis: a gyógyszerből magas dózist kapnak.
 - Hangulat: A személyek hangulatértékét tartalmazza egy hét gyógyszerzedést követően.



- A kitöltők függetlenek egymástól
- Csoportok függetlenek
- Függő változó legalább intervallum típusú
- Függő változó normális eloszlású
 - A függő változó normál eloszlása **csoportonként ellenőrizendő**
 - Analyze / Descriptive / Explore menün belül található a Plots-ok között
- Szóráshomogenitás
 - Ellenőrizhető az eddig tanult módon az Analyze / Descriptive / Explore menün belül is, de a feltétel ellenőrzésére használt Levene-teszt be van építve az ANOVÁba, így ott is elegendő kikérni.
 - Ha a csoportok elemszáma megegyezik, elég robosztus rá a próba
 - De ha a nagyobb csoport nagyobb varianciával rendelkezik, az ANOVA szigorodik, ha a kisebb csoport rendelkezik nagyobb varianciával, az ANOVA megengedőbb lesz

One-Way ANOVA az SPSS-ben

- Analyze / Compare Means / One-Way ANOVA
- Descriptive: átlagok és szórások miatt fontos
- Homogeneity of variance – a Levene féle szóráshomogenitás teszt itt is kikérhető
- Welch – ha a szóráshomogenitás nem teljesül, korrigálni kell a próba eredményeit. Ezt teszi a Welch-teszt. Csak akkor használjuk, ha nem teljesül a szóráshomogenitás. Most tudjuk, hogy teljesül a feltétel, és csak azért jelöltük be, hogy megnézzük, hogy néz ki a táblázata.
- Means plot: nem túl jó grafikon az eredmények értelmezéséhez. A dolgozatba ne ezt tedd be, de az hatások megértéséhez ez is jó.



One-Way ANOVA outputja

- Leíró statisztikák
 - A szokásos leíró statisztikákat láthatjuk, elemszám, átlag, szórás igazán érdekes számunkra.
 - Látjuk, hogy az elemszámok hasonlóak, 15-18 fő van a csoportokban, de a szórások is hasonlóak, 1,5 és 2 között mozog mind a három csoport.
 - Az átlagok alapján a legrosszabb hangulata a placebót kapó csoportnak van (12.24 pont), ezt követi az alacsony (14.07) és a magas dózis (15.5 pont). Ha a csoportok közötti különbség szignifikáns, akkor pont a hipotézisünknek megfelelően alakulnak az átlagok, azaz pongyolán minél több gyógyszert kapott valaki, annál jobb hangulatú volt.

Descriptives

Hangulat	95% Confidence Interval for Mean							
	N	Mean	Std. Deviation	Std. Error	Lower Bound	Upper Bound	Minimum	Maximum
Placebó	17	12,24	1,522	,369	11,45	13,02	10	16
Alacsony dózis	15	14,07	1,944	,502	12,99	15,14	11	18
Magas dózis	18	15,50	2,065	,487	14,47	16,53	12	19
Total	50	13,96	2,285	,323	13,31	14,61	10	19

- Test of Homogeneity of Variances
 - Ha nem ellenőriztük előzetesen a szóráshomogenitást, akkor keressük meg a Levene-teszt tábláját.
 - Látjuk, hogy ez nem szignifikáns, azaz, nincs szignifikáns különbség a csoportok szórásában. Ha nincs különbség, akkor a szórások hasonlóknak tekinthetők, azaz homogének, a szóráshomogenitás feltétele teljesül. (A gondolat menetet átismételheted az előző féléves anyag feltétel-tesztelés diásorában.)
 - $F(2, 47) = 1.467 \quad p = .241$
 - Ha a szóráshomogenitás teljesül, akkor nézhetjük az ANOVA táblázatát, viszont ha sérül, akkor az ANOVA értékeit korrigálni kell, ezt teszi meg a Welch statisztika, amit a Robust Tests of Equality of Means táblában találunk.

Test of Homogeneity of Variances

Hangulat			
Levene Statistic	df1	df2	Sig.
1,467	2	47	,241

ANOVA

Hangulat

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	93,428	2	46,714	13,512	,000
Within Groups	162,492	47	3,457		
Total	255,920	49			

Robust Tests of Equality of Means

Hangulat

	Statistic ^a	df1	df2	Sig.
Welch	14,563	2	30,247	,000

a. Asymptotically F distributed.



- ANOVA tábla
 - A szóráshomogenitás teljesült, tehát nézhetjük az ANOVA táblát.
- Nézzük az elméleti rész matekja hogyan jelenik meg a táblázatban:
- **Between Groups:** a kísérleti manipulációnk hatása, a csoportok közötti különbség (hatás)
 - **Sum of Squares** (SS_M) – a mintaátlagok közötti különbség / megmagyarázott variancia
 - **Mean Square** (MS_M) – átlagos különbség a mintaátlagok között
- **Within group:** a csoportokon belüli eltérések (zaj)
 - **Sum of Squares** (SS_R) – a pontok mintaátlagtól való távolsága / meg nem magyarázott variancia
 - **Mean Square** (MS_R) – átlagosan zaj a mintákon belül
- **Total:** függő változó teljes változatossága

ANOVA

Hangulat

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	93,428	2	46,714	13,512	,000
Within Groups	162,492	47	3,457		
Total	255,920	49			

- Az ANOVA eredménye:
 - **df**
 - Szabadságfokok. A F-értéknek két szabadságfoka van, egy a hatásához tartozik, számolása csoportok száma - 1 (itt 2 az értéke), egy pedig a zajhoz, számolása elemszám – hatás szabadságfoka (itt az értéke 47)
 - **F-érték és szignifikancia**
 - Az F-érték egy statisztikai érték, tehát a hatás és zaj hányadosából adódik. Ehhez tudjuk a szabadságfokok (df) segítségével hozzárendelni a szignifikancia-értéket, melyből megtudjuk, hogy a csoportok között látható különbség szignifikáns-e, azaz hihető-e, általánosítható-e a populációra.
 - Az ANOVA szignifikáns, a csoportok között szignifikáns különbség van.
- Ha a szóráshomogenitás sérült volna, akkor a Welch-táblát néznénk, melyben szintén megtaláljuk a két szabadságfokot, az F- és szignifikancia-értéket.

ANOVA

Hangulat

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	93,428	2	46,714	13,512	,000
Within Groups	162,492	47	3,457		
Total	255,920	49			

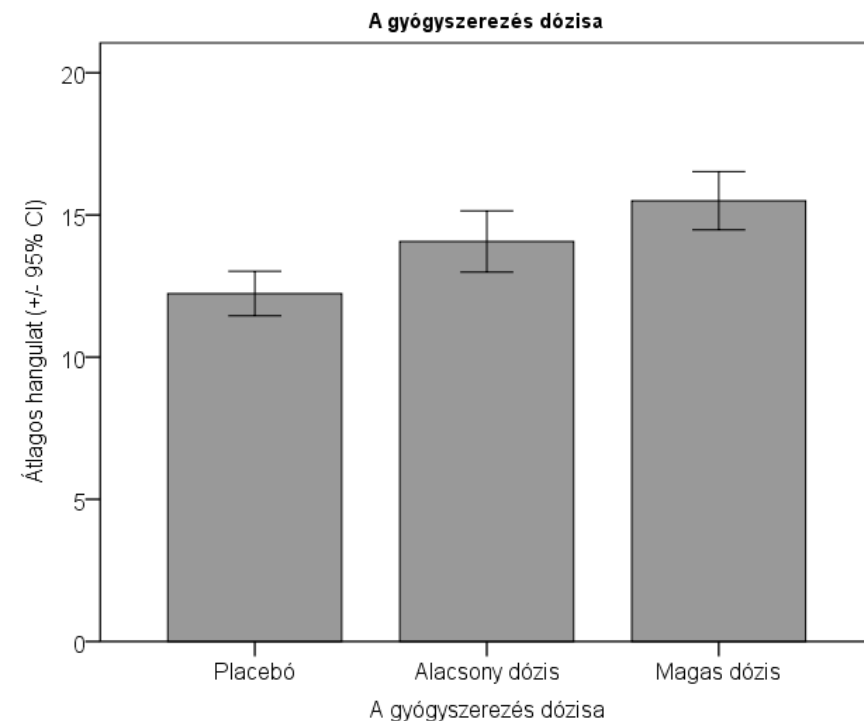
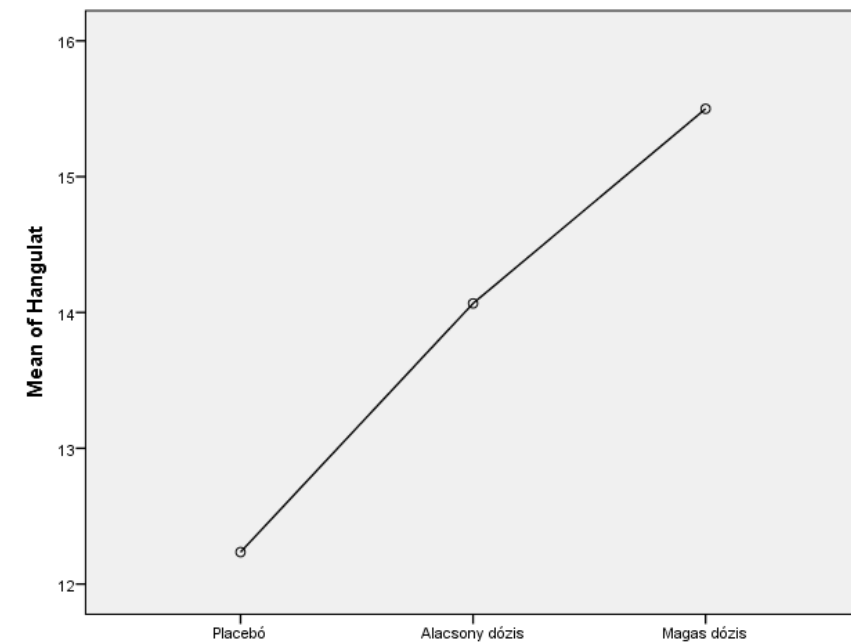
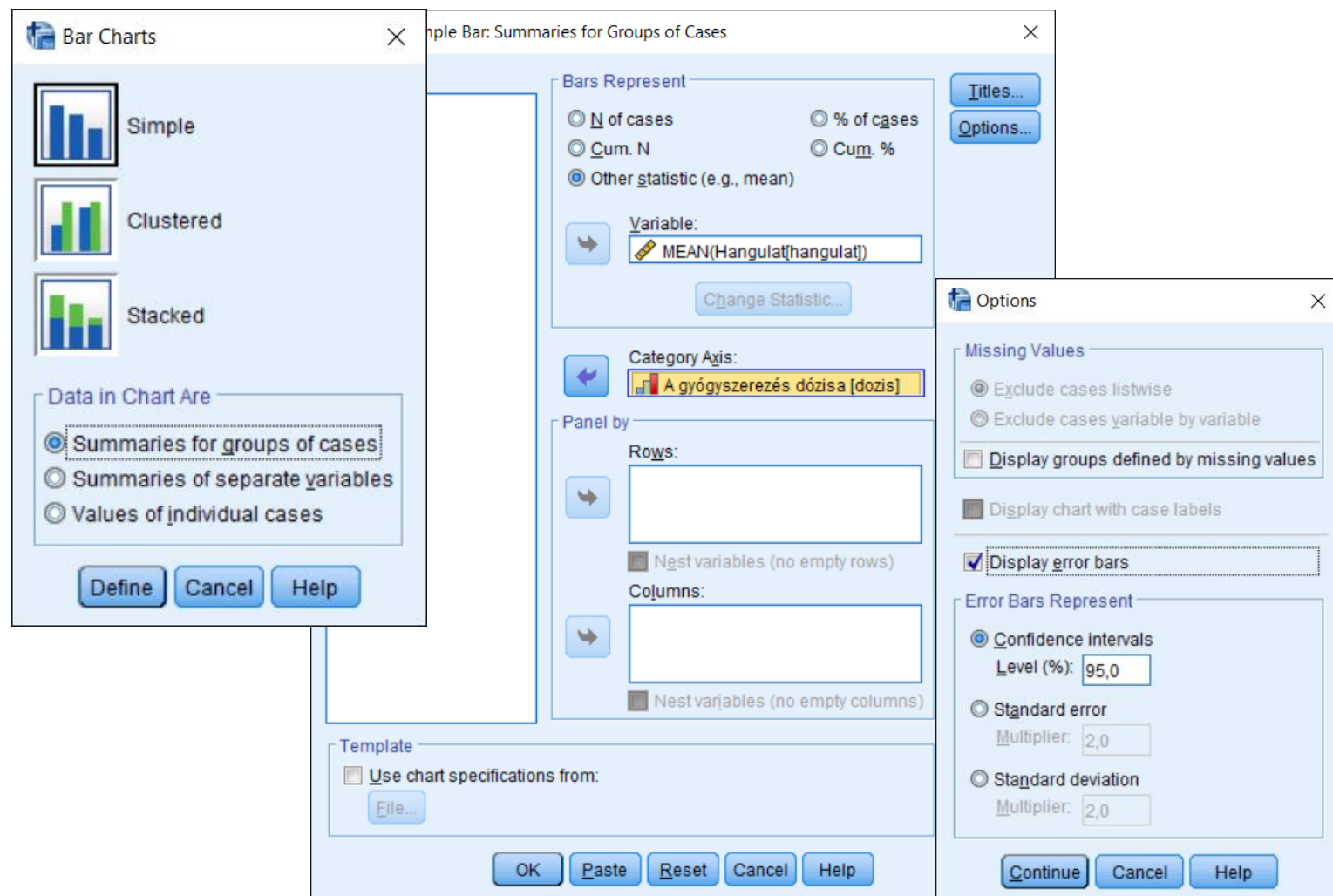
Robust Tests of Equality of Means

Hangulat

	Statistic ^a	df1	df2	Sig.
Welch	14,563	2	30,247	,000

a. Asymptotically F distributed.

- Végül kapunk egy grafikont, mely az eredmények értelmezésénél elegendő, de a dolgozatba oszlopdiagrammot tegyél be hibásávv!
- Graphs / Legacy Dialogs / Bar



- Az ANOVA publikálásakor is a korrelációs együttthatót használjuk effect-size mutatóként, illetve elterjedtebb annak négyzetre emelt változata (r^2), ráadásul a neve is más, ANOVÁnál az elnevezése η^2 (eta²). Jelentése azonban megegyezik a korábban tanultakkal, a hatás nagyságát adja meg, értékéből megtudhatod, hogy a modell hány százalékát magyarázza a teljes varianciának.
- Hátránya, hogy túlbecsüli a hatást, mert azt mutatja meg, mekkora a megmagyarázott variancia aránya a mintában, nem pedig a populációban. Nagy elemszámnál ez a bias csökken.
- Számolását tekintve a modell által megmagyarázott variancia (SS_M) és az adatokban lévő teljes variancia (SS_T) aránya.

$$\bullet \eta^2 = \frac{SS_M}{SS_T} = \frac{93.428}{255.920} = .365$$

- Azaz az, hogy ki mennyi gyógyszert kapott a hangulatban lévő változatosságot 36.5%-ban magyarázza.

ANOVA

Hangulat

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	93,428	2	46,714	13,512	,000
Within Groups	162,492	47	3,457		
Total	255,920	49			

- Létezik egy kevésbé torzított effect-size mutató is, az ω^2 . Ez figyelembe veszi az elemszámot is.
- Pontosabb effect-size mutató, de feltétele, hogy a csoportok elemszáma hasonló legyen.

$$\omega^2 = \frac{SS_M - df_M * MS_R}{SS_T + MS_R} = \frac{93.428 - 2 * 3.457}{255.920 + 3.457} = .334$$

- Azaz az, hogy ki mennyi gyógyszert kapott a hangulatban lévő változatosságot 33.4%-ban magyarázza. Látjuk, hogy az értéke egy kicsit kisebb, mint az η^2 volt, ez a korrekció miatt mindig így van.

ANOVA

Hangulat

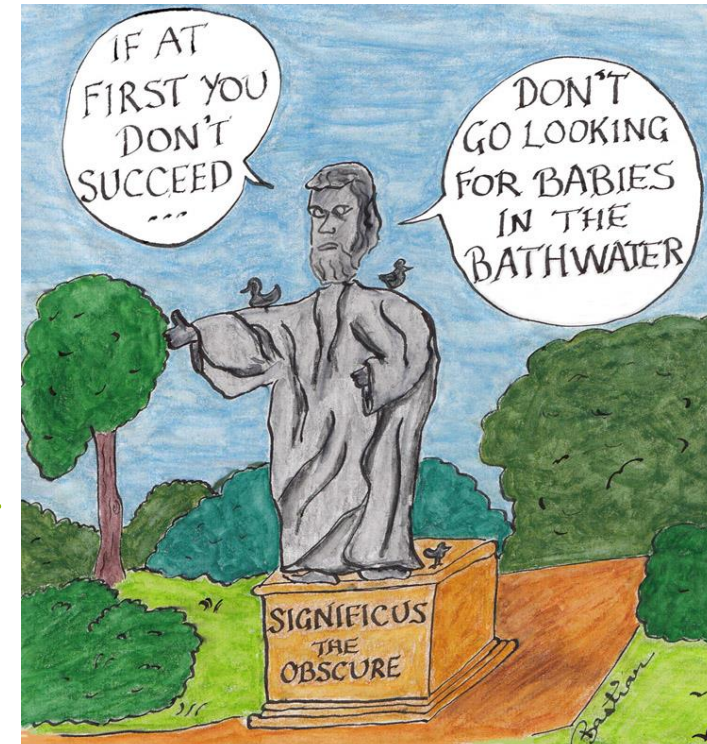
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	93,428	2	46,714	13,512	,000
Within Groups	162,492	47	3,457		
Total	255,920	49			

Eredmények közzlése

- F-érték közzlése:
 - $F([df_M], [df_R]) = [F] \ p = [Sig]$ +effect-size mutató
 - $F(2, 47) = 13.512 \ p < .001 \ \omega^2 = .334$
- Szövegesen többféleképp megfogalmazható:
 - Független-mintás ANOVA vizsgálatot használva azt találtuk, a medikáció szignifikáns hatással van a hangulatra $F(2, 47) = 13.512 \ p < .001 \ \omega^2 = .334$.
 - Független-mintás ANOVA vizsgálatot használva a csoportok között szignifikáns különbséget találtunk $F(2, 47) = 13.512 \ p < .001 \ \omega^2 = .334$. A medikáció 33.4%-ban magyarázza a hangulat varianciáját.

Hangulat	ANOVA				
	Sum of Squares	df_M df	Mean Square	F	Sig.
Between Groups	93,428	2	46,714	13,512	,000
Within Groups	162,492	47	3,457		
Total	255,920	49			

One-Way ANOVA utóvizsgálatai azaz a Post hoc és kontraszt vizsgálatok



AN EARLY STATISTICAL PHILOSOPHER, SIGNIFICUS THE OBSCURE WAS KNOWN FOR MANGLED APHORISMS AND AN INTENSE HATRED OF POST HOC ANALYSES.

Melyik minták különböznek egymástól?

- Az ANOVA eredménye az F-érték egy úgynevezett omnibusz (átfogó) statisztika:
 - Megtudjuk belőle, hogy a kísérleti manipulációnk sikeres volt-e, hogy a minták között **van-e különbség**
 - Nem tudjuk meg, **melyik minták között van különbség**
 - Több lehetséges változat is van: lehet, hogy $A \neq B \neq C$ vagy $A = B \neq C$ vagy $A \neq B = C$ vagy $B = C \neq A$
 - F próba csak azt mondja meg, a minták nem egyformák, azt nem mondja meg, melyek különböznek egymástól
- **FONTOS!!**
 - **Utóvizsgálatot CSAK akkor szabad végezni, ha az ANOVA szignifikáns, ha az ANOVA alapján nincs szignifikáns különbség, akkor TILOS az utóvizsgálatokban „valami publikálhatót keresgélni”**
 - Tehát, ha az előzőleg elvégzett ANOVA szerint szignifikáns különbség van a csoportok között, akkor válassz valamilyen utóvizsgálatot, hogy megtudd, mely csoportok különböznek egymástól.
 - Ha az ANOVA alapján nincs szignifikáns különbség, akkor azt kell publikálni, hogy nincs, és kész.

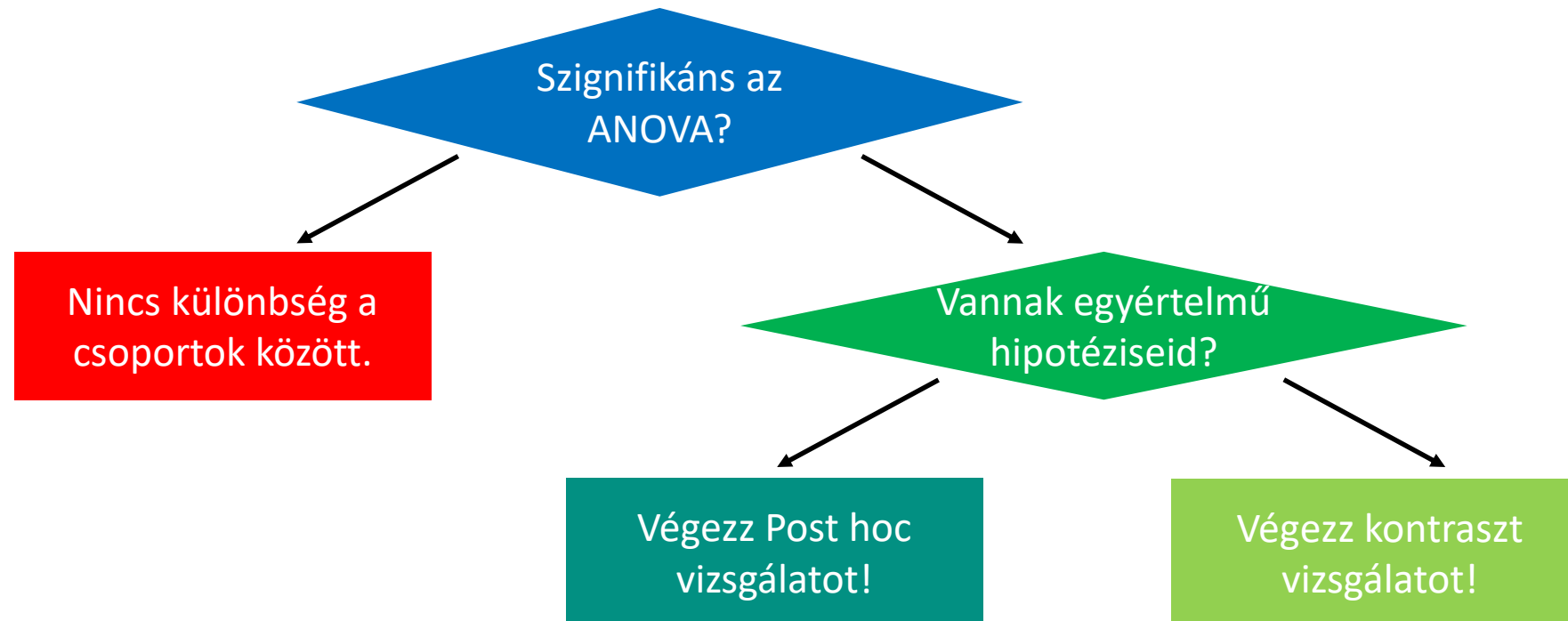


IF YOU TORTURE THE
DATA LONG ENOUGH, IT
WILL CONFESS.

- Két megközelítés létezik a minták összehasonlítására az elsőfajú hiba esélyének növelése nélkül. A kérdéseinknek megfelelően választjuk ki melyikre van inkább szükségünk
- **Post hoc**
 - „**Mindent mindennel, de szigorúan**”
 - Minden lehetséges párosításnál megnézzük, van-e különbség (tehát vizsgáljuk a Placebo vs. Alacsony, Placebo vs. Magas és Alacsony vs. Magas különbségét), de a szignifikancia szintet szigorúbbra vesszük, így kerüljük ki a familywise hibát
 - Úgy állítjuk be az összehasonlításoknál a szignifikancia szintet, hogy az összegzett, familywise hiba határa legyen 5%.
 - Mivel szigorúbbak a használt szignifikancia szintek, ezért kevesebb különbség lesz szignifikáns
 - Feltáró vizsgálat esetén
- **Kontrasztok**
 - „**Csak a hipotéziseinkben szereplő összehasonlítások, de azok érzékenyen**”
 - Ha kész hipotéziseink vannak, melyik minták vagy mintaegységek között várunk különbséget, megtehetjük, hogy csak a hipotéziseknek megfelelő különbségeket ellenőrizzük.
 - Nem hasonlítunk össze mindent mindennel, viszont amit igen, ott érzékeny a próba
 - Nem csak egyedi mintákat, de minták egységét is össze lehet hasonlítani
 - Különbség az orthogonális és nem orthogonális kontrasztok között (később)

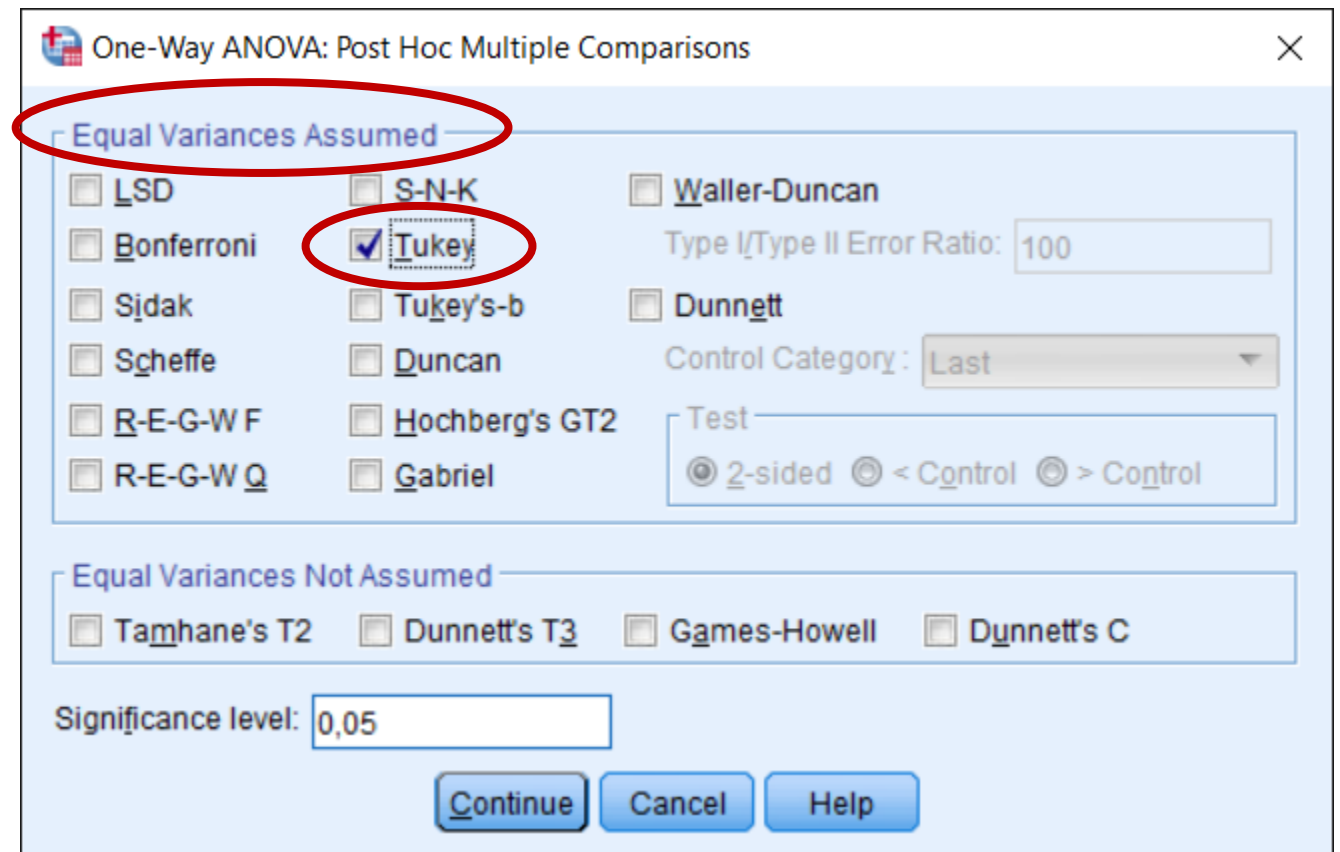
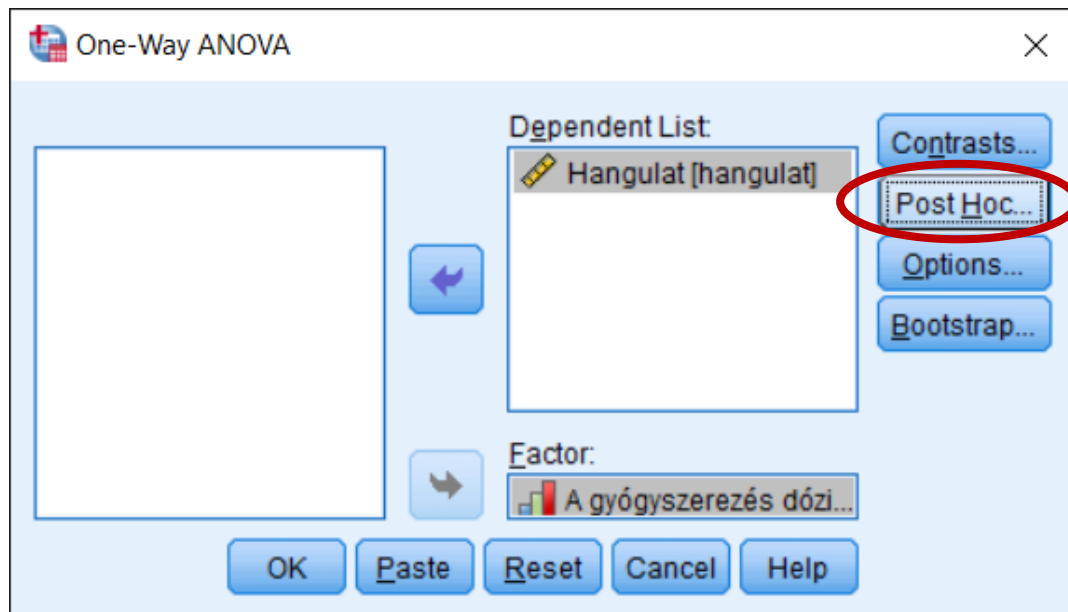
Kontraszt VAGY post hoc

- A kontraszt és post hoc vizsgálatok egymás alternatívái, de együtt soha nem végezhetők, mert az a familywise hiba megjelenéséhez vezet!
- Nem helyes kutatói hozzáállás, ha az érzékenyebb kontraszt vizsgálattal ellenőrizzük a minket érdeklő összefüggéseket, majd elmegyünk még halászni egy kicsit a zavarosban a Post hoc segítségével.
- Ha egyértelmű hipotéziseid vannak, válaszd a kontraszt vizsgálatot, mely szűken, érzékenyen vizsgál. Ha feltáró elemzést végzel, válaszd a Post hoc-ot, mely átfogóbb viszont szigorúbb elemzés.



- Minden lehetséges párosításnál megnézzük, van-e különbség, de a szignifikancia szintnél valahogy figyelembe vesszük a familywise hiba jelenségét
- **Szóráshomogenitás** esetén, több Post hoc teszt közül is lehet választani:
 - **LSD** - legelső megoldás volt. Kiszámolja, hogy mekkora a legkisebb különbség két csoport között, mely már szignifikáns lenne – ehhez hasonlítja, mekkora a valódi különbség – ne használd, mert **túlzottan megengedő**
 - **Bonferroni korrekció**: legnépszerűbb megoldás. Elosztjuk a szignifikancia szint értékét a próbák számával. – ne használd, mert túl szigorú, **elsőfajú hibát megoldja, de megnöveli a másodfajú hiba arányát.**
 - **Tukey teszt** – LSD-hez hasonló megoldás, de megbízhatóbb. **Szigorú! Egyenlő elemszám esetén jó!**
 - **Hochberg's GT2 és Gabriel's test** – **nem egyenlő elemszámot kezel.**
 - **Scheffe** – legrobosztusabb az ANOVA feltételeinek megszegésére. Legszigorúbb.
- **Szóráshomogenitás megszegése** esetén:
 - **Tamhane's T2** – nagyon szigorú
 - **Dunnnett's tesztek** - szigorúak – **T3** kis elemszám, **C** nagy elemszám esetén ajánlott
 - **Games-Howell** – túlzottan megengedő

- Kérjük ki újra az előző ANOVÁt! Analyze / Compare Means / One-Way ANOVA
 - A vizsgálatunkban most teljesül a szóráshomogenitás és hasonlóak az elemszámok, ezért Tukey tesztet választunk



- Az előbb azt írtam, hogy a Post hocnál szigorúbb szignifikancia-szintet állítunk be. A tanult eljárások ezt a korrekciót elvégzik, így nekünk kényelmesen továbbra is az $\alpha = .05$ -ös határt kell nézni.

Itt látod, milyen Post hoc-ot választottál

Dependent Variable: Hangulat
Tukey HSD

Multiple Comparisons

(I) A gyógyszerelés dózisa		(J) A gyógyszerelés dózisa	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Placebó	VS	Alacsony dózis	-1,831 [*]	,659	,021	-3,43	-,24
		Magas dózis	-3,265 [*]	,629	,000	-4,79	-1,74
Alacsony dózis	VS	Placebó	1,831 [*]	,659	,021	,24	3,43
		Magas dózis	-1,433	,650	,081	-3,01	,14
Magas dózis	VS	Placebó	3,265 [*]	,629	,000	1,74	4,79
		Alacsony dózis	1,433	,650	,081	-,14	3,01

*. The mean difference is significant at the 0.05 level.

Mean difference – két csoport átlaga közötti különbség + Standard error

A szignifikanciát mindig 2-tailed értelmezzük (hiszen nincs előre meghatározott hipotézisünk)

Konfidencia intervallum a különbségre. A populáció- különbség várhatóan e két érték között helyezkedik el

- Post hoc közlése
 - Az átlagokat és szórásokat a leíró statisztikában találod.
 - Tukey HSD post hoc vizsgálat eredményei alapján a placebo csoport ($M = 12.24$ $SD = 1.522$) és az alacsony dózisú medikációt kapó csoport ($M = 14.07$ $SD = 1.944$) szignifikánsan különbözik egymástól ($p = .021$) az átlagos hangulatértékben. A medikációt kapó csoport hangulata magasabb. Szintén szignifikáns különbséget találtunk a placebo és a magas dózisú ($M = 15.50$ $SD = 2.065$) csoport között ($p < .001$). Nem találtunk szignifikáns különbséget az alacsony és magas dózisú gyógyszert kapó csoportok között ($p = .081$)

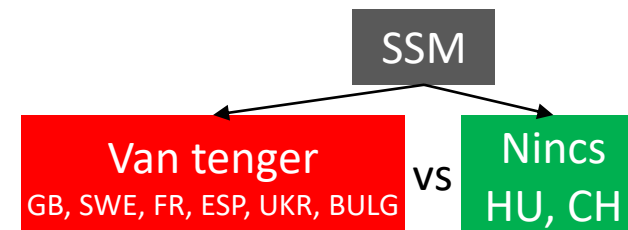
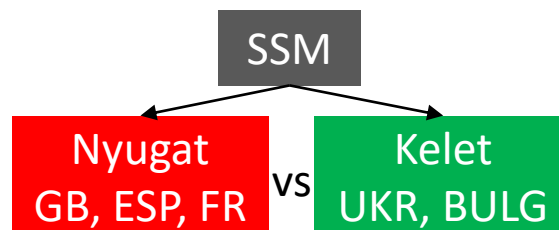
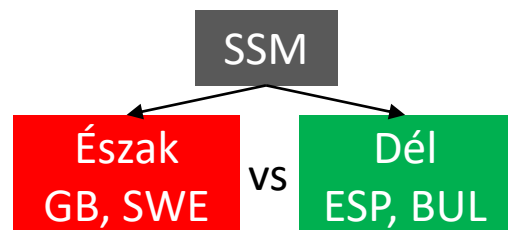
Descriptives						Multiple Comparisons					
Dependent Variable: Hangulat						Tukey HSD					
Hangulat						(I) A gyógyszerelés dózisa	(J) A gyógyszerelés dózisa	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval Lower Bound Upper Bound
	N	Mean	Std. Deviation	Std. Error	L						
Placebó	17	12,24	1,522	,369		Placebó	Alacsony dózis	-1,831*	,659	,021	-3,43 -,24
							Magas dózis	-3,265*	,629	,000	-4,79 -1,74
Alacsony dózis	15	14,07	1,944	,502		Alacsony dózis	Placebó	1,831*	,659	,021	,24 3,43
							Magas dózis	-1,433	,650	,081	-3,01 ,14
Magas dózis	18	15,50	2,065	,487		Magas dózis	Placebó	3,265*	,629	,000	1,74 4,79
							Alacsony dózis	1,433	,650	,081	-,14 3,01
Total	50	13,96	2,285	,323							

*. The mean difference is significant at the 0.05 level.

- Az utóvizsgálatok másik lehetséges útja a kontrasztok végzése
 - A kontrasztok a megmagyarázott varianciát „darabolják”
 - Egyedi minták helyett minták együttese hasonlítható így össze
 - Egy kontrasztban mindig a variancia két részét hasonlítjuk össze
 - Rengeteg csoportegüttest (kontrasztot) létre lehet hozni. Példák:
 - A) 3 csoport: placebo, alacsony dózisú gyógyszer, magas dózisú gyógyszer



- B) 8 csoport: Magyaró., N.Brit., Svédo., Franciao, Spanyolo., Ukrajna, Cseho., Bulgária



- Különbség van az úgynevezett ortogonális és nem ortogonális kontrasztok között.

- Ortogonális kontrasztok**

- A varianciát „darabolják” egyre kisebb részekre, mintha egy tortát szeletelnénk.
- Ha a csoportokat már két részre osztottam, akkor az egy kontrasztba került csoportokat szabad további kontrasztokba bontani, de a külön oldalra került csoportok már NEM hasonlíthatók össze!

- 1. kontraszt: Kontroll vs Medikáció

- ✓ 2. kontraszt: Alacsony vs Magas dózis

- ✗ ~~3. kontraszt: Placebo vs Alacsony dózis~~

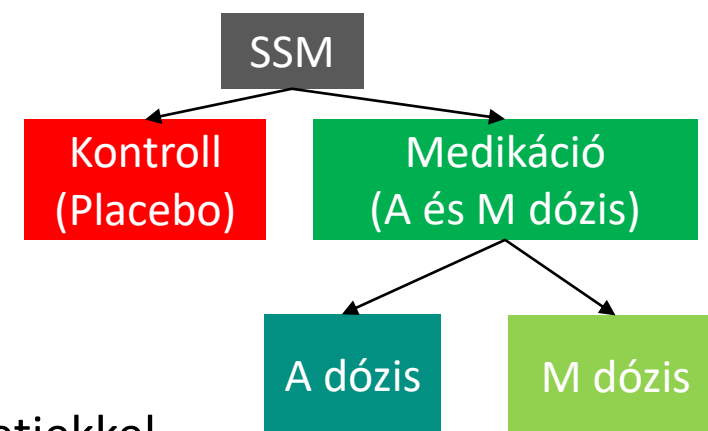
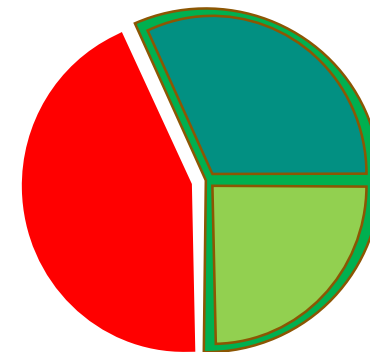
Ez ortogonális az elsőre

Ez NEM ortogonális az elsőre!

- Így mindig csoportszám-1 kontraszt hozható létre
- Érdemes először a kontrol csoportot (ha van) összehasonlítani a kísérletiekkel
- Nem minden változópárra nézve lesz információm,
- Az orthogonális kontraszt előnye, hogy NEM kell a familywise hibára kontrollálni

- Nem ortogonális kontrasztok**

- Ekkor ugyanazokat a csoportokat újra és újra felhasználhatom, újracsoportosíthatom a különböző kontrasztokban (pl. az előző dia országos példája).
- Mivel az összehasonlítások átfedésben vannak egymással, így a familywise hiba újra megjelenik. A kurzuson a nem ortogonális kontrasztokat kerüljük, ha lehet.



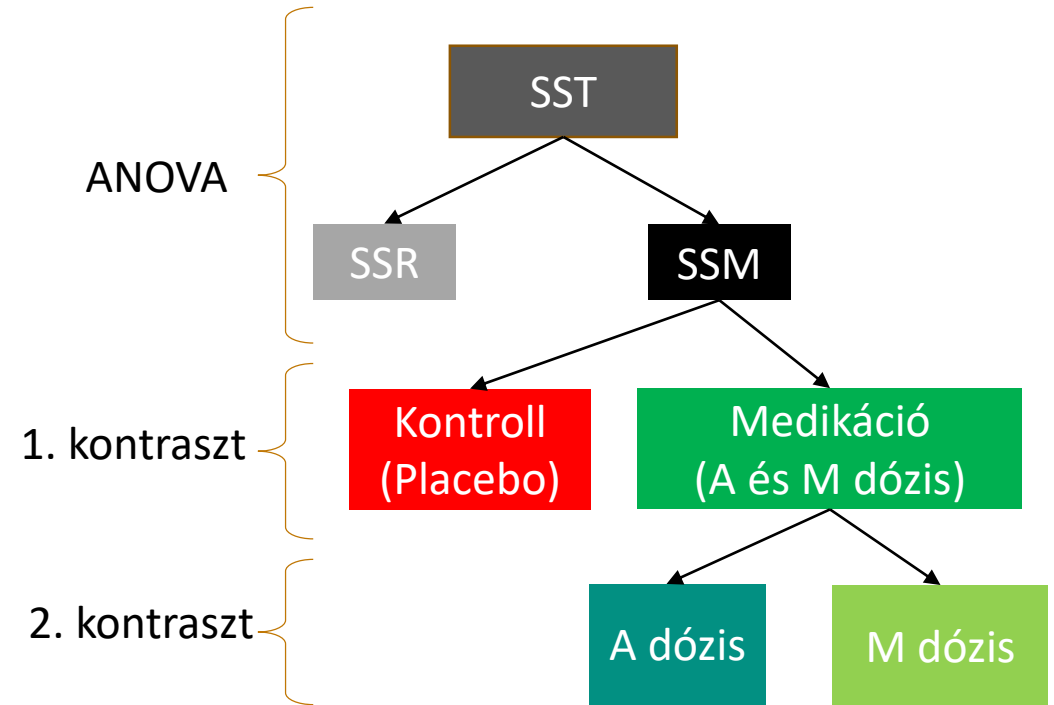
Az ANOVA a függő változó varianciáját (SS_T) a csoportosító változóval megmagyarázható (hatás, SS_M) és meg nem magyarázható (hiba, SS_R) részre bontja

Az első kontraszt a függő változóból megmagyarázott varianciát a kontroll és medikáció által megmagyarázott részre bontja

A második kontraszt a medikáció által megmagyarázott varianciát alacsony és magas dózissra bontja.

- Szabályok az ortogonális kontrasztok kiosztására

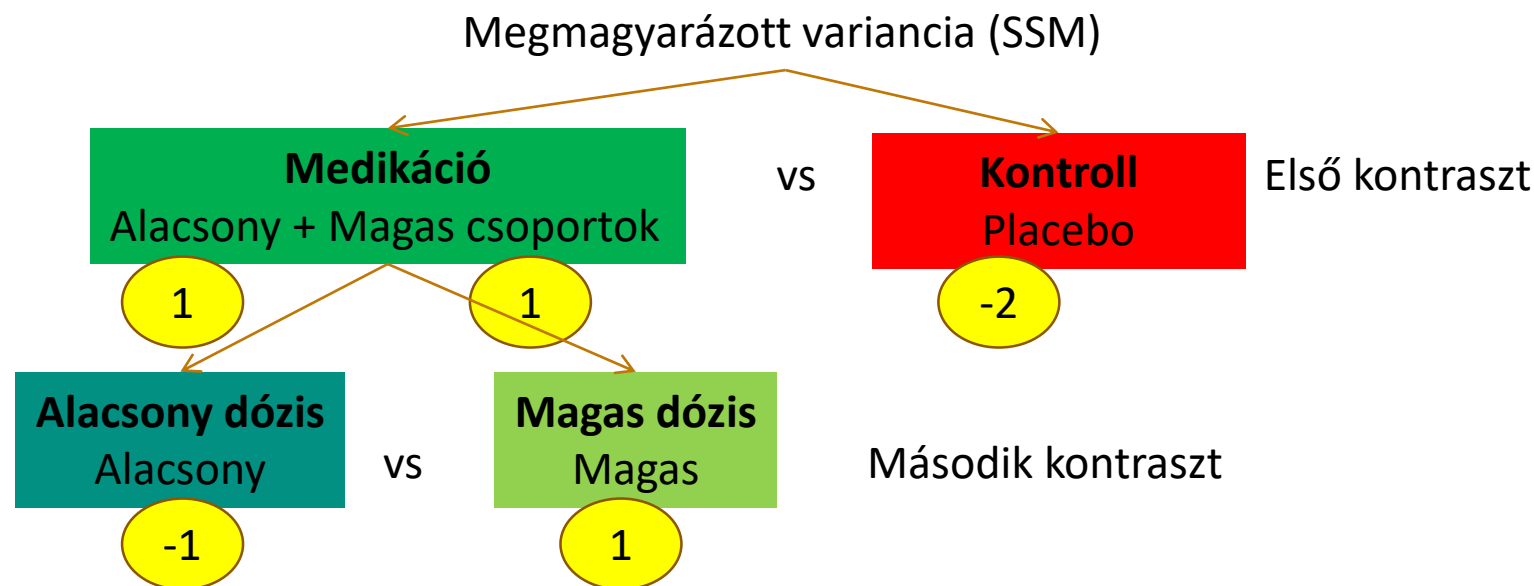
- Minden kontrasztban mindig két varianciadarabot hasonlítunk össze (azaz a kontrasztnak mindig két oldala van)
- Az egyik oldal csoportjaihoz pozitív, a másik oldal csoportjaihoz negatív számokat rendelünk
- A pozitív és negatív számok összege meg kell egyezzen egymással
- Ha egy csoport nem vesz részt a kontrasztban, a hozzá rendelt érték 0



- Gyakoroljunk! Rajzoljuk meg a farajzot, rendeljünk a csoportokhoz értékeket!
- Három csoportunk van: placebo, alacsony és magas dózis. Nézzünk kontrasztot először a placebo és nem placebo között (kontrol vs. Kísérleti csoportok) majd egy második kontrasztot az alacsony és magas dózis között!

csoport	Első kontraszt	Második kontraszt
1-placebó		
2-alacsony		
3-magas		

- Gyakoroljunk! Rajzoljuk meg a farajzot, rendeljünk a csoportokhoz értékeket!
- Három csoportunk van: placebo, alacsony és magas dózis. Nézzünk kontrasztot először a placebo és nem placebo között (kontrol vs. Kísérleti csoportok) majd egy második kontrasztot az alacsony és magas dózis között!



Három csoport van, tehát két kontrasztot tudunk kiosztani.
 Első kontrasztban alacsony+ magas csoportok értéke együtt 2, ezzel szembe kerül a placebo -2-je
 Második kontrasztban magas 1, ezzel szemben az alacsony -1. A placebo itt 0, mert nem vesz részt a kontrasztban

csoport	Első kontraszt	Második kontraszt
1-placebó	-2	0
2-alacsony	1	-1
3-magas	1	1

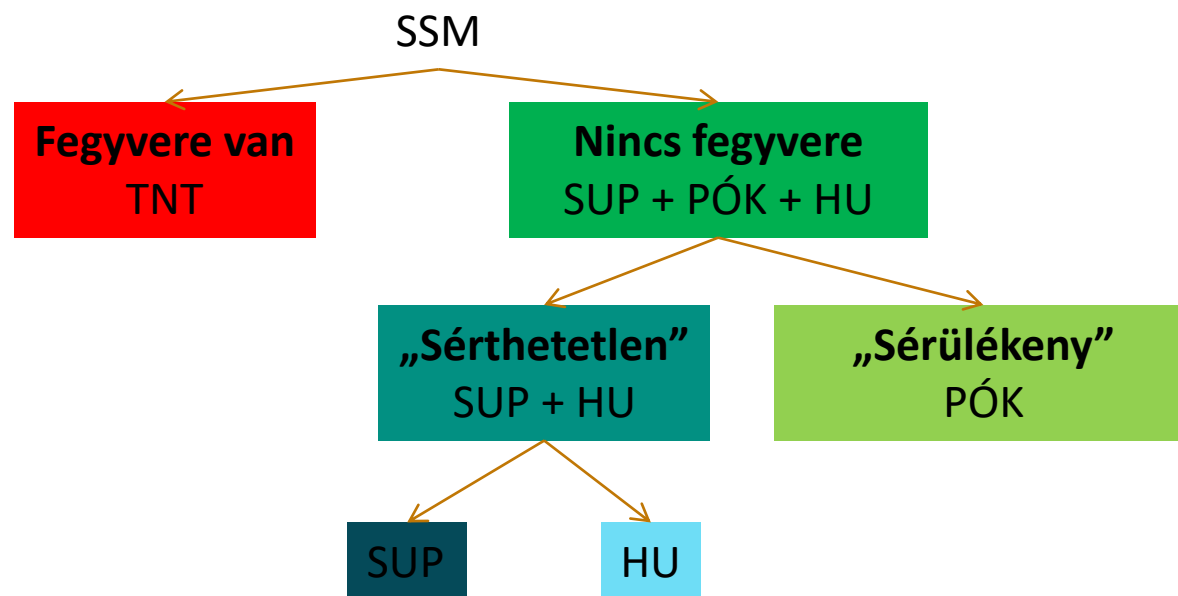
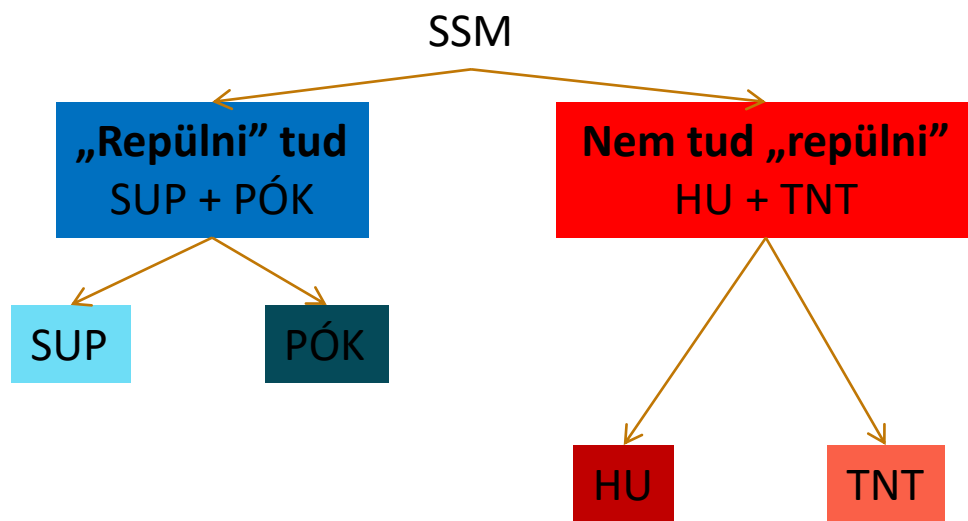
- Gyakoroljunk!
- Négy csoportunk van: Kontrol (tudja, hogy nem kap gyógyszert), Placebo (azt hiszi gyógyszert kap), Alacsony dózis, Magas dózis.

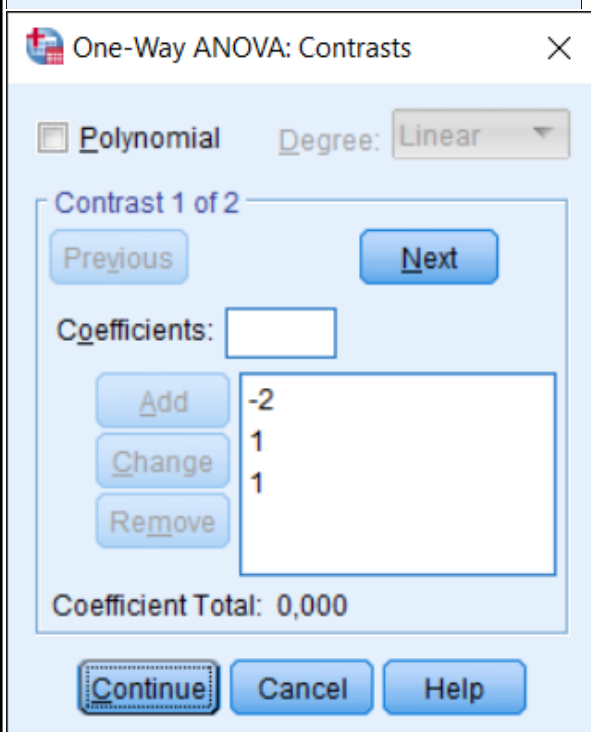
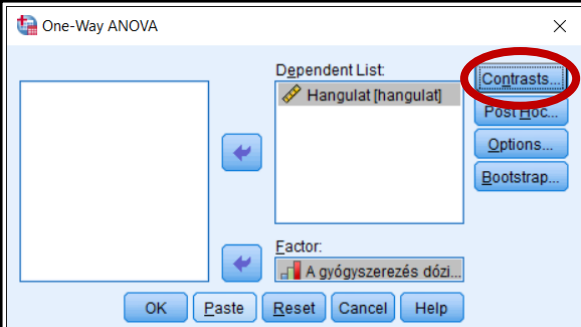
csoport	Első kontraszt	Második kontraszt	Harmadik kontraszt
1-placebó			
2-alacsony			
3-magas			
4-kontrol			

- Gyakoroljunk!
- Halloween alatt jelmezbe öltöző gyerekeket vizsgálunk abból a szempontból, hogy mennyi kisebb-nagyobb baleset éri őket. 4 csoport van, azok a gyerekek, akik Supermannek öltöznek, akik Pókembernek, akik Hulknak, és akik Tininindzsa teknőcnek.



- Kontraszt-kiosztásra több jó megoldás is lehet, fontos azonban, hogy a hipotézis szempontjából értelmes csoportosításokat csináljunk.
 - Például kézenfekvő kiosztás lenne a piros-kék és zöld jelmezek szétválasztása, de mégsem értelmes, hiszen a balesetek szempontjából valószínűleg nem releváns a jelmez színe.
 - De például működhet a következő két kiosztás bármelyike. Az elsőt akkor használnám, ha azt feltételezem, a sérülések leginkább abból adódnak, hogy figyelmetlenül ugrabugrálnak, és azok a gyerekek, akik olyan hősnek öltöznek be, akik tudnak repülni, megpróbálnak majd asztalokról leugrálni. A második kiosztást akkor használnám, ha azt gondolom, a sérülések leginkább a gyerekek egymás közötti játékos harcaiból adódnak.





csoport	Első kontraszt	Második kontraszt
1-placebó	-2	0
2-alacsony	1	-1
3-magas	1	1

- A kontraszt vizsgálathoz kérjük ki újra az ANOVA-t!
- **Kontrasztok:**
 - Vegyük elő az előző táblázatot, ahol kitöltöttük a kontrasztoknál a változókhoz tartozó értékeket
 - **Contrast 1 of 2:** itt látjuk, melyik kontrasztot szerkesztjük, nekünk most két kontrasztunk lesz, ha bevittük őket, a Previous és Next gombokkal lépegethetünk közöttük.
 - **Coefficients**-hez beírjuk az első kontraszt első értékét, azaz a -2-t, majd az **Add**-re kattintunk, bevisszük a második csoporthoz tartozó értéket, azaz az 1-et, rákattintunk az Add-re, végül bevisszük a harmadik csoport 1-esét is, és rákattintunk az Add-re
 - Ha végeztünk az első kontraszttal, **Next**-tel átlépünk a másodikra, és bevisszük a második kontraszt értékeit is
 - Fontos, a kontrasztokat mindig olyan sorrendben kell megadni, ahogy a változókat kódoltuk (most: először a placebót, mert annak 1 a kódja, stb.)

- Infótábla arról, milyen kontrasztokat kértél ki, ez csak egy emlékeztető a beállításokról.
- A kontraszt vizsgálat érzékeny a **szóráshomogenitásra**, ezért a kontraszt tábla is két részből áll. (Azt, hogy a szóráshomogenitás teljesül-e, korábban vizsgáltuk, és a Test of homogeneity of Variance táblában tudod ellenőrizni.) Ha a szóráshomogenitás teljesül (a Levene teszt nem szignifikáns), akkor az Assume equal variances részt kell nézni, ha nem teljesül (a Levene teszt szignifikáns volt), akkor a Does not assume equal variances részt.
- **Value of Contrast:** a csoportok súlyozott átlaga közötti különbség + SE, például az első kontraszt 5.1-es értéke azt jelenti, hogy a mediációt kapó személyek (az alacsony és magas dózis együtt) átlagosan 5.1 ponttal magasabb hangulatértéket értek el, mint a placebót kapó személyek.
- A **t-próba** megadja, szignifikáns különbség van-e a kontraszt két oldala között. Ez 2-tailed szignifikancia. Ha egyértelmű hipotézisünk volt az irányra is, akkor használhatunk 1-tailed szignifikanciát, amit úgy kapunk, hogy a Sig. Értékét elosztjuk kettővel

Contrast Coefficients

Contrast	A gyógyszerelés dózisa		
	Placebó	Alacsony dózis	Magas dózis
1	-2	1	1
2	0	-1	1

Contrast Tests

		Contrast	Value of Contrast	Std. Error	t	df	Sig. (2-tailed)
Hangulat	Assume equal variances	1	5,10	1,112	4,584	47	,000
		2	1,43	,650	2,205	47	,032
	Does not assume equal variances	1	5,10	1,017	5,012	40,491	,000
		2	1,43	,699	2,050	30,498	,049

- **Értelmezés:**

- A szóráshomogenitás tudjuk, hogy teljesült, tehát a felső blokkot nézzük majd.
- Első kontraszt
 - A medikációt kapó személyek átlagosan 5.1 ponttal magasabb hangulatértéket értek el, mint a placebo csoport.
 - A hipotézisünk az volt, hogy a medikációt kapó csoportok jobb hangulatot érnek el, mint a placebo csoport, ezért használhatunk 1-tailed szignifikanciát, de már a 2-tailed is $p < .001$ szinten szignifikáns.
- Második kontraszt
 - A magas dózist kapó személyek átlagosan 1.43 ponttal magasabb hangulatot érnek el, mint az alacsony dózist kapó személyek.
 - A hipotézis az volt, hogy a magasabb dózisú csoport jobb hangulatot ér el, mivel egyirányú a megfogalmazás, használhatunk 1-tailed szignifikanciát: $p = .032/2 = .016$ – beigazolódott.

Contrast Tests

		Contrast	Value of Contrast	Std. Error	t	df	Sig. (2-tailed)
Hangulat	Assume equal variances	1	5,10	1,112	4,584	47	,000
		2	1,43	,650	2,205	47	,032
	Does not assume equal variances	1	5,10	1,017	5,012	40,491	,000
		2	1,43	,699	2,050	30,498	,049

- Kontrasztok közlése
 - A kontrasztvizsgálat szerint szignifikánsan magasabb hangulatot találtunk a medikációt kapó csoportokban, mint a placebo csoportban $t(47) = 4.584$ $p < .001$ (1-tailed) $r = .556$, a különbség átlaga 5.10 ($SE = 1.112$).
 Szignifikánsan magasabb hangulatértéket találtunk a magas dózist kapó csoportban az alacsony dózis kapó csoporthoz viszonyítva $t(47) = -2.205$ $p = .016$ (1-tailed) $r = .306$, a különbség átlaga 1.43 ($SE = 0.650$)

Contrast Tests

		Contrast	Value of Contrast	Std. Error	t	df	Sig. (2-tailed)
Hangulat	Assume equal variances	1	5,10	1,112	4,584	47	,000
		2	1,43	,650	2,205	47	,032
	Does not assume equal variances	1	5,10	1,017	5,012	40,491	,000
		2	1,43	,699	2,050	30,498	,049