

Filozofická fakulta TU v Trnave



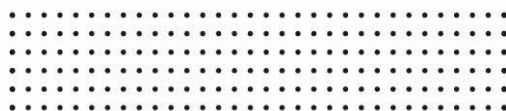
1998

:Michal Kohút

Sprievodca základmi analýzy psychologických dát v programe jamovi



Filozofická fakulta TU v Trnave



:Michal Kohút

Sprievodca základmi analýzy psychologických dát v programe jamovi



Michal Kohút

Sprievodca základmi analýzy psychologických dát v programe jamovi

2. aktualizované vydanie

(Vysokoškolská učebnica)

Tnava

2026

Recenzenti:

doc. Ing. Mgr. Róbert Hanák, PhD.

PhDr. Eva Ballová Mikušková, PhD.

© Michal Kohút, 2026

© Filozofická fakulta TU v Trnave, 2026

Ilustrácia na obálke: © Ján Kurinec, 2026

ISBN 978-80-568-0962-4

Publikácia neprešla jazykovou korektúrou.

Obsah

Úvod.....	7
1. Kvantitatívny výskum v psychológii.....	8
2. Základy práce v programe jamovi.....	10
2.1 Dátový súbor	11
2.2 Vytvorenie dátového súboru	12
2.2.1 Vytvorenie dátového súboru pri metóde papier-pero.....	13
2.2.2 Vytvorenie dátového súboru pri online zbere dát	13
2.3 Vloženie dátového súboru do jamovi a nastavenie premenných	14
2.4 Transformácie hodnôt	17
2.5 Matematické operácie.....	21
2.6 Filtrovanie dát	23
3. Základná deskriptívna štatistika	26
3.1 Miery polohy	26
3.2 Miery variability.....	29
3.3 Frekvenčná analýza	32
3.4 Deskriptívna analýza v programe jamovi.....	33
4. Vnútna konzistencia.....	38
5. Inferenčná štatistika.....	42
5.1 Testovanie signifikancie nulovej hypotézy	42
5.2 Alternatívne hypotézy	44
5.3 Bivariačná a multivariačná štatistická analýza.....	44
5.4 Vzdialené hodnoty.....	45
5.5 Testovanie normality distribúcie premennej	48
6. Korelačná analýza	52
6.1 Korelačná analýza v programe jamovi	53
6.2 Interpretácia a zápis výsledkov korelačnej analýzy	58
7. Porovnávacie testy.....	61
7.1 Porovnávanie dvoch skupín (výberov).....	62
7.1.1 Testy pre dva nezávislé výbery a kontinuálne premenné v programe jamovi	65
7.1.2 Interpretácia a zápis výsledkov testov porovnania dvoch skupín	68
7.2 Porovnávanie dvoch meraní	70
7.2.1 Porovnávanie dvoch meraní v programe jamovi.....	71
7.2.2 Interpretácia a zápis výsledkov testov porovnania dvoch meraní.....	72
7.3 Jednofaktorová analýza rozptylu (porovnávanie 3 a viac skupín)	74

7.3.1 Jednofaktorová analýza rozptylu v programe jamovi	75
7.3.2 Interpretácia a zápis výsledkov jednofaktorovej analýzy rozptylov	80
7.3.3 Neparametrická verzia jednofaktorovej analýzy rozptylu.....	81
8. Vzťah dvoch nezávislých nominálnych premenných	85
8.1 Chí-kvadrát test pre dve nezávislé nominálne premenné v programe jamovi	86
8.2 Interpretácia a zápis výsledkov chí-kvadrát testu pre dve nezávislé nominálne premenné	89
9. Záverečné odporúčania.....	91
Záver.....	93
Referencie.....	94

Úvod

Tento učebný text je určený primárne pre študentov na katedre psychológie FF TU k predmetu Štatistika I. a II. a je nutným základom pre predmet Analýza psychologických dát. Uplatnenie možno nájsť aj u iných študentov, ktorých zaujímajú základy práce v programe jamovi. Cieľom tejto publikácie je ponúknuť čitateľovi jednoduchý návod pre dôležité funkcie v programe jamovi ako aj oboznámenie so základnou deskriptívnou a bivariačnou inferenčnou štatistikou. Ukážeme si základnú prácu s programom jamovi, od inštalácie, orientáciu v užívateľskom prostredí, cez vloženie dát a ich základný manažment. Druhá časť tejto publikácie je venovaná deskriptívnej štatistike. Uvedieme si, čo to je a na čo je dobrá. Následne si priblížime inferenčnú štatistiku, testovanie hypotéz a základné analýzy zamerané na skúmanie súvisu medzi dvomi premennými. Cieľom publikácie nie je poskytnúť vyčerpávajúci výklad všetkého k tejto téme, no práve naopak, efektívnym spôsobom priblížiť to dôležité a používané. V prípade, že by ste mali záujem ísť ďalej a dozvedieť sa viac, odporúčam využiť internet a vyhľadať si odpovede na vaše otázky v nespočetnom množstve článkov a prednášok (najmä v anglickom jazyku).

Prečo práve program jamovi? Analyzovať kvantitatívne výskumné dáta možno vo viacerých počítačových programoch. Niektoré z nich sú užívateľsky príjemné, no cenovo ťažko dostupné (napr. IBM SPSS). Niektoré sú zadarmo, no na druhú stranu pomerne zložité pre občasného (či jednorazového) používateľa, asi najznámejší je program R. V niektorých programoch sa spája príjemné s dostupným, ako napríklad PSPP, JASP a jamovi. Ak by ste mali záujem dozvedieť sa viac o práci v programe JASP, odporúčame pozrieť si aj učebnicu od Ballovej Mikuškovej (2021) alebo učebnicu od Hanáka (2016) pre program PSPP. Ako už názov tejto učebnice napovedá, budeme pracovať v počítačovom programe *jamovi* (The jamovi project, 2021). Tento program je voľne prístupný a zároveň užívateľsky príjemný. Má aj svoje nedostatky, teda občasné, drobné chybičky, no jeho autori na ňom stále pracujú a postupne zlepšujú a rozširujú funkcionality.

Momentálne čítate druhé, aktualizované vydanie tejto učebnice. Ako autor som rád za všetku pozitívnu spätnú väzbu od študentov a kolegov. Prvé vydanie však malo drobné nedostatky, ktoré sú v tejto verzii upravené. Motiváciou k aktualizácii boli tiež zmeny v programe jamovi. Ako som spomínal, vývojári na ňom stále pracujú. Ide o zmeny vo výzore a funkcionalite a aktualizované vydanie učebnice tieto zmeny zahŕňa (platí pre jamovi verziu 2.6.44).

1. Kvantitatívny výskum v psychológii

Kvantitatívny výskum, je metóda výskumu, ktorá je založená na meraní premenných prostredníctvom numerického systému. Merania sú ďalej analyzované prostredníctvom rôznych štatistických modelov alebo analýz s cieľom porozumieť, popísať a predikovať určitý fenomén (voľne podľa APA, n.d.). Už v názve tohto prístupu možno badať dva dôležité znaky – kvantum a kvantifikácia. Prvý pojem vyjadruje, že kvantitatívne výskumy sa spravidla uskutočňujú na desiatkach, stovkách či dokonca tisícoch účastníkov (Ilievová a kol., 2017). Druhý pojem je možno ešte dôležitejším pre pochopenie podstaty – skúmaný jav či fenomén meriame, t.j. určitým spôsobom ho kvantifikujeme. K takémuto, na číslach založenému prístupu, neoddeliteľne patrí štatistika.

Podľa Fielda (2018) výskumný proces prebieha v štyroch fázach. V prvej fáze sa objaví záujem o určitú problematiku. Ako náš záujem vznikol je individuálne, no vedie nás k túžbe dozvedieť sa viac o tejto problematike. V tejto fáze si kladieme výskumnú otázku a venujeme sa podstate javu, identifikujeme dôležité spojitosti – nabírame poznatky z doterajších výskumov. V ďalšej fáze formulujeme na základe získaných poznatkov určitú teóriu, ktorou vysvetľujeme podstatu javu. Vďaka nej dokážeme sformulovať výskumný problém a z neho vyplývajúce ciele nášho výskumu a na základe nich si stanovujeme určité predpoklady o psychologickom jave, ktoré odborne nazývame hypotézy. Tie chceme otestovať, a preto v tretej fáze zbierame dáta.

Zber dát je širokým pojmom, ktorého súčasťou je *výber meraných premenných* („Čo ideme merať?“) a ich *operacionalizácia* („Ako to ideme merať?“). Odpoveď na prvú otázku je daná povahou našich výskumných cieľov. Odpoveď na druhú otázku je zložitejšia. V rámci kvantitatívneho výskumu v psychológii meriame javy najčastejšie prostredníctvom psychodiagnostických metódik. Tie majú rôznu podobu – výkonové testy, postojové škály, osobnostné inventáre a podobne, no ich spoločným znakom je určitý číselný výsledok, s ktorým vieme následne pracovať. Hoci nejde o meranie ako také, meraním označujeme aj radenie do kategórií, napr. podľa pohlavia, vzdelania, rodinného stavu a pod. Dôležitým rozhodnutím pri zodpovedaní otázky, „Ako to ideme merať?“, je tiež výber výskumného dizajnu.

Korelačný výskum je výskumný dizajn, v ktorom skúmame súvis dvoch alebo viacerých premenných. Podstatou je, že pri meraní nijakým spôsobom priamo nezasahujeme, resp. neovplyvňujeme sledované premenné – jednoducho odmeriame a analyzujeme. Výhodou takéhoto prístupu je jednoduchosť uskutočnenia a možnosť merania pri veľkom množstve respondentov. Nevýhodou je, že nevieme spoľahlivo tvrdiť o smere efektu či konštatovať vplyv – vieme len potvrdiť súvis (koreláciu) premenných.

Experimentálny výskum je výskumný dizajn, v ktorom rovnako skúmame súvis premenných, no rozlišujeme závislú (cieľovú) premennú a skúmame, aký vplyv na ňu má nezávislá premenná, nazývaná aj prediktor. Pri tomto dizajne zasahujeme do výskumu a ovplyvňujeme mieru nezávislej premennej v snahe zistiť, či a aký efekt na závislú premennú má táto zmena (Walker, 2013). Pri experimentoch rozoznávame experimentálnu skupinu, pri ktorej bola uskutočnená zámerná zmena a kontrolnú skupinu, u ktorej nebol uskutočnený experimentálny zásah.

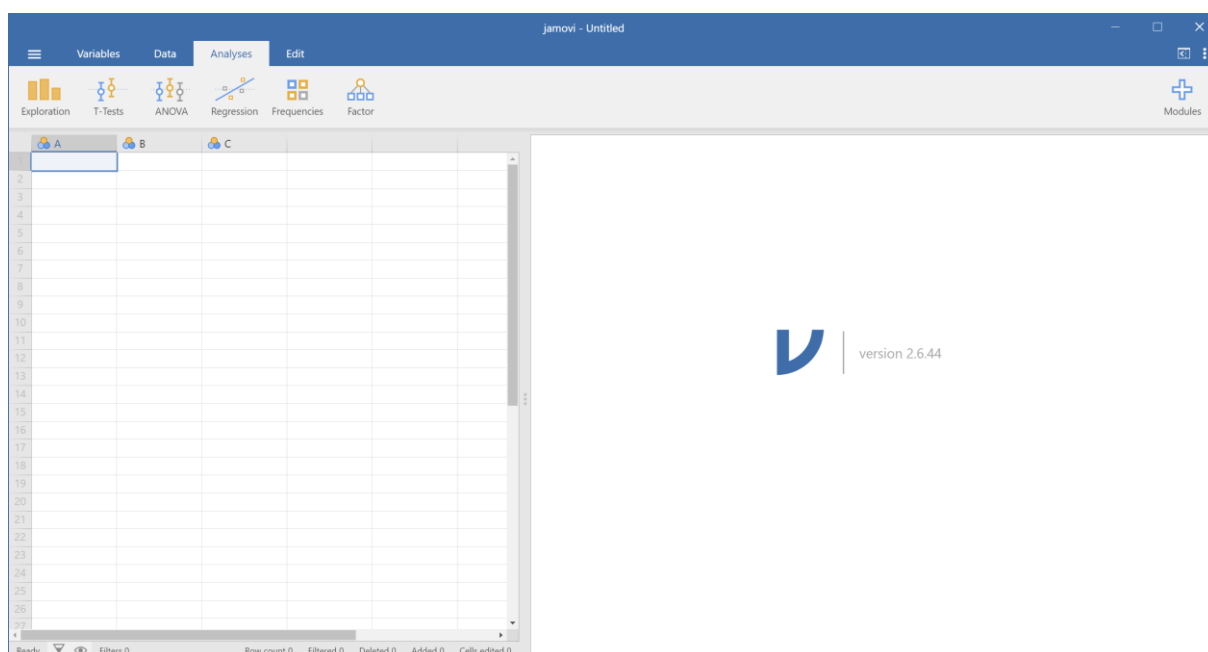
V rámci tretej fázy si musíme zodpovedať aj otázku „Na kom ideme merať?“. Pri tejto otázke je niekedy odpoveď skrytá priamo vo výskumnom probléme. Ak sa zaujíname o zvládanie záťaže u zamestnancov záchranej služby, je nám pomerne zrejmé, že dáta budeme zbierať u záchranárov. Vo všeobecnosti vyberáme takzvanú *cieľovú populáciu*, ktorú definujeme ako súbor všetkých možných skúmaných objektov. Populácia je tvorená *štatistickými jednotkami*. V psychológii sú nimi jednotliví ľudia, t.j. jednotlivci. Pre účely overenia hypotéz z tejto populácie vyberáme len niektoré štatistické jednotky, čím nám vzniká výskumný *výber* (alebo súbor), teda podmnožinu populácie (Chajdiak a kol. 1994).

Po tom, čo ukončíme zber dát, prichádza na rad posledná fáza výskumu – *analýza dát a interpretácia výsledkov* (Field, 2018). V tejto fáze využívame štatistické postupy pre analýzu dát. Výsledky analýzy ďalej interpretujeme, najmä s cieľom overiť platnosť nami stanovených hypotéz. Základmi analýzy psychologických dát sa venujeme v ďalších kapitolách.

2. Základy práce v programe jamovi

Pre začatie práce v programe jamovi si tento program musíme nainštalovať. Program môžeme stiahnuť zadarmo na stránke <https://www.jamovi.org/> pod záložkou *products*, kde zvolíme *jamovi Desktop*. Jamovi je možné stiahnuť v rôznych verziách. Na tomto programe jeho vývojári neustále pracujú (aspoň v čase písania tejto publikácie tomu tak bolo), a tak vám neviem povedať, ktorá verzia programu je aktuálne pre vás v ponuke. Program je však možné stiahnuť vo vydaní pre operačné systémy Windows 64-bitový, macOS, Linux a ChromeOS. V tejto publikácii používame operačný systém Windows, no zo skúsenosti je aj vydanie pre macOS užívateľsky totožné s tým pre Windows. Momentálne sú v ponuke dve verzie, jedna je stabilná (angl. *solid*), jedna je najnovšia (angl. *current*). Odporúčame vám zvoliť si stabilnú verziu, keďže tá si už prešla testovaním od užívateľov a prípadné chyby boli opravené. V tejto učebnici pracujeme s verziou 2.6.44.

Program sme si stiahli na pevný disk počítača a nainštalovali. Po spustení začíname v základnom prostredí, ktoré je zobrazené v Obrázku 2.1.



Obrázok 2.1: Základné rozhranie programu jamovi 2.6.44

V hornej lište tohto rozhrania máme šesť možností:

- tri vodorovné čiarky : cez túto možnosť sa dostaneme k ponuke vytvorenia nového dátového súboru (možnosť *New*); k možnosti otvorenia už rozpracovaného dátového súboru, v ktorom sme pracovali pred tým (možnosť *Open*); môžeme tiež importovať dátový súbor z iného programu (napr. tabuľku z tabuľkového procesora); nájdeme tu aj možnosť uložiť náš dátový súbor prostredníctvom možnosti *Save* alebo *Save as*, alebo možnosť exportovať dátový súbor pre iné programy (možnosť *Export*)
- lišta pre náhľad premenných v dátovom súbore *Variables* : v tejto lište budú zobrazené premenné, ktoré sú v našom dátovom súbore. Možno vďaka tomu vidieť celý

zoznam a rýchlo otvoriť nastavenie atribútov premenných. K týmto nastaveniam sa možno dostať aj inou podobne rýchlou cestou, ktorú využívame v tejto učebnici.

- lišta pre spracovanie dát *Data* : v tejto lište nájdeme možnosť vložiť, vystrihnúť, kopírovať hodnoty; možnosť vrátiť späť; možnosti nastavenia premenných (možnosť *Setup*); vypočítať nové premenné (možnosť *Compute*); transformovať hodnoty (možnosť *Transform*); pridať a odstrániť premenné alebo respondentov; nastaviť filtre – tieto úpravy dát priblížime v ďalších častiach práce
- lišta je určená pre štatistickú analýzu dát *Analyses* : nájdeme tu možnosti pre exploráciu dát, zisťovanie vzťahov medzi premennými či rozdielov medzi skupinami a i. Týmto štatistickým analýzom sa budeme venovať samostatne v jednotlivých kapitolách.
- lišta *Edit*  je určená pre úpravu výstupov analýz. Priamo vo výstupe v programe jamovi môžeme pridať popisy k výsledkom, upraviť formátovanie a pod. S touto kartou pracovať v tejto učebnici nebudeme, keďže výsledky neexportujeme a nezdierame, no prepisujeme do našej práce.
- zobrazenie alebo skrytie časti okna zobrazujúcej dáta 
- tri bodky pod sebou : kliknutím sa zobrazí kontextové okno v ktorom možno nastaviť priblíženie (*Zoom*), počet zobrazovaných desatinných miest, vzhľad grafov, jazykovú predvoľbu a nastavenie importu dát. V tejto časti je dôležité nastaviť počet desatinných miest. Štandardne sa mnohé číselné výsledky uvádzajú zaokrúhlené na 2 desatinné miesta, preto nastavíme *Number format* na „2 dp“ (dp – decimal points, teda desatinné miesta). Výnimkou je hodnota signifikancie, ktorá sa štandardne uvádza na 3 desatinné miesta, preto ponecháme *p-value format* nastavenie na „3 dp“.

Na ľavej strane je možné vidieť tabuľku, ktorá má pomenované stĺpce a očíslované riadky. Toto je základ pre zapísanie hodnôt z nášho výskumu. Na pravej strane je možné vidieť logo programu, no tento priestor slúži pre zobrazenie výsledkov uskutočnených analýz. Prvým krokom na začiatok je vytvorenie dátového súboru.

2.1 Dátový súbor

Dátový súbor možno jednoducho vysvetliť ako súbor, ktorý obsahuje dáta. Táto definícia je síce tautológiou, no väčšinou je s úsmevom prijímaná u študentov. V základe dátový súbor pozostáva (ako iné tabuľky) z dvoch hlavných elementov, a to:

- riadky: riadky v dátovom súbore vyjadrujú jednotlivé merania vo všeobecnosti. V rámci psychologického výskumu sú tieto merania najčastejšie jednotliví respondenti alebo participanti. Jednoducho povedané, čo riadok, to respondent/participant.
- stĺpce: stĺpce v dátovom súbore vyjadrujú jednotlivé *premenné*. Ide o charakteristiky prvkov štatistického súboru. Premenná môže nadobúdať viacero hodnôt, pričom je stanovené, akým spôsobom sú tieto hodnoty priradené (Hendl, 2006). V psychologickom výskume sa často zaujímate o rôzne socio-demografické charakteristiky, respondentom ponúkame položky psychodiagnostických metodík, ktoré potom ďalej spracovávame a prostredníctvom matematických operácií na základe

nich vytvárame nové charakteristiky (napr. priemerné skóre). Všetky tieto charakteristiky zastrešuje pojem *premenná*.

Premenné sa medzi sebou líšia na základe ich úrovne merania. Halama, Špajdel a Žitný (2013) uvádzajú 4 úrovne merania: nominálnu, ordinálnu, intervalovú a pomerovú. Zároveň však dodávajú, že posledné dve kategórie sú pomenovávané spoločným názvom *kardinálna premenná*. V programe jamovi sú tieto dve úrovne pomenované ako *kontinuálna* úroveň merania.

- *Nominálna premenná, nominálna úroveň merania:* ide o najnižšiu úroveň merania, kedy priradené číslo je len názvom pre určitý jav. Pri tejto úrovni merania s číslami nevieme robiť matematické operácie, ide len o pomenovanie. Príkladom môže byť spomenutá premenná *pohlavie*. V našom výskume sme sa respondentov pýtali, aké je ich pohlavie a mohli si vybrať jednu z dvoch možností – muž alebo žena. Pri prepise odpovedí respondentov môžeme odpoveď „muž“ číselne označiť ako 1 a odpoveď „žena“ ako 2, no aj naopak. Je dôležité, aby sme zvolený systém dodržali pri všetkých respondentoch.
- *Ordinálna premenná, ordinálna úroveň merania:* pri tejto úrovni merania znovu prideliť čísla k určitému javu, avšak tieto javy majú určitú hierarchiu a vieme ich na základe niečoho usporiadať. Príkladom môže byť *najvyššie dosiahnuté vzdelanie*. Respondenti si mohli zvoliť z možností základné, stredoškolské bez maturity, stredoškolské s maturitou, vysokoškolské – 1. stupňa (Bc.), vysokoškolské – 2. stupňa (Mgr., Ing.), vysokoškolské – 3. stupňa (PhD., CsC). Pri tejto úrovni merania sú síce jednotlivé možnosti usporiadané od najnižšieho po najvyššie, no medzi jednotlivými možnosťami nie sú rovnaké vzdialenosti. Vieme, že respondent, ktorý ukončil vysokoškolské vzdelanie 2. stupňa, má vyššie dosiahnuté vzdelanie ako respondent, ktorý označil stredoškolské bez maturity, no vzdialenosť medzi jednotlivými možnosťami nie je rovnaká. Na tejto úrovni merania nevieme robiť matematické operácie s hodnotami. Napr. aký je rozdiel medzi niekým, kto ukončil vysokoškolské vzdelanie 2. stupňa (hodnota 5) a niekým kto ukončil strednú školu bez maturity (hodnota 2). Vieme, že $5-2=3$, čiže by sme mohli povedať, že rozdiel je 3, čo v našom kódovaní znamená stredná škola s maturitou, a to nedáva zmysel.
- *Kontinuálna premenná, kontinuálna úroveň merania:* ide o spoločný názov pre intervalovú a pomerovú premennú, tiež nazývanú ako kardinálna premenná. Tento typ premennej používame pre premenné, kde dané číslo môžeme chápať ako číslo. Príkladom môže byť vek, výška, váha, ale aj premenné ako hrubé alebo priemerné skóre dotazníka, počet správnych odpovedí v teste a iné.

2.2 Vytvorenie dátového súboru

Pred akoukoľvek ďalšou prácou je potrebné vytvoriť si dátový súbor. Ak ste zbierali dáta metódou papier-pero, môžete definovať premenné a vkladať dáta priamo v programe jamovi. Odporúčam vám však začať v programe Microsoft Excel alebo podobnom tabuľkovom procesore. Napriek mnohým výhodám programu jamovi, pomenovanie premenných a vkladanie dát je rýchlejšie v Exceli. Takto vytvorený dátový súbor sa potom dá veľmi jednoducho presunúť do jamovi. Tí z vás, ktorí ste zbierali dáta prostredníctvom online

platformy, už pravdepodobne máte tabuľku, ktorú si môžete stiahnuť a po úpravách ju vložiť do jamovi.

2.2.1 Vytvorenie dátového súboru pri metóde papier-pero

Pokiaľ ste dáta zbierali spôsobom papier-pero, čaká vás množstvo príjemných chvíľ prepisovania hodnôt z papierových formulárov do počítača. Berte do úvahy počet hodnôt, ktoré musíte zadať. Ak ste mali vo svojom výskume 100 položiek a získali ste odpovede od 200 respondentov, budete musieť vložiť 20 tisíc hodnôt, a to zaberie nejaký čas.

V prvom kroku si otvorte tabuľkový procesor (napr. Excel) a vytvorte premenné. Stačí, ak do prvého riadku vložíte názvy, ako by ste chceli, aby sa premenné nazývali. Hoci sme písali, že v riadkoch sa nachádzajú respondenti, v tomto prípade nám prvý riadok posluží ako „hlavička“, v ktorej sú definované názvy premenných. Program jamovi z tohto prvého riadku vytvorí názvy premenných tak, aby sme ich nemuseli nastavovať ručne. Konkrétne názvy premenných sú už na vás a závisia na tom, aké položky ste zbierali. Ideálne je však nazvať premennú krátko, jednoslovne a výstižne, bez zbytočnej diakritiky alebo medzier, napríklad *pohlavie*, *vek*. V prípade položiek psychodiagnostických metodík je praktické použiť nejakú skratku vychádzajúcu z názvu metodiky a priradiť k nej číslo položky. V prípade prvej položky Inventára Veľkej Päťky (BFI-2), by sme mohli zvoliť názov BFI_1. Vo väčšine tabuľkových procesorov je možnosť automatického dopĺňania hodnôt. Táto funkcia vie značne urýchliť prácu. Do bunky zadáte skratku a číslo položky. Pravý spodný roh bunky má trochu hrubšie orámovanie, stačí kliknúť, potiahnuť v smere riadku a automaticky budú nastavené hodnoty. Za pár sekúnd tak viete pomenovať množstvo položiek (napr. BFI_1 až BFI_60).

Po vytvorení názvov pre jednotlivé premenné v prvom riadku prichádza na rad vkladanie hodnôt do jednotlivých buniek. Všetko, čo vkladáte, dávajte v číselnom formáte. Vopred si stanovte, aké číselné hodnoty pridelíte jednotlivým kategóriám premenných. Príkladom môže byť premenná pohlavie – zadefinujte si, akým číslom budú označení muži a akým ženy (napr. muž – 1, žena – 2). Pri nominálnych premenných je to na vás, no pri ordinálnych premenných (t.j. usporiadaných kategóriách) je dôležité, aby čísla boli zoradené tak, ako sú zoradené kategórie. Napríklad: základné vzdelanie – 1, stredoškolské bez maturity – 2, stredoškolské s maturitou – 3, vysokoškolské I. stupňa – 4, vysokoškolské II. stupňa – 5, vysokoškolské III. stupňa – 6.

2.2.2 Vytvorenie dátového súboru pri online zbere dát

Po ukončení online zberu vám väčšina platforiem umožní stiahnuť si súbor, ktorý otvoríte v tabuľkovom procesore. V prípade platformy Google Forms si môžete stiahnuť tabuľku Excel. Táto bude mať v každom stĺpci jednotlivé premenné. Prvý riadok tabuľky nesie názvy premenných – tie však majú znenie podľa toho, čo ste zadali pri tej ktorej položke online formulára. Odporúčam si tieto názvy premenovať spôsobom popísaným vyššie.

Odpovede jednotlivých respondentov budú v číselnom alebo slovnom formáte. Pokiaľ ste sa pýtali na vek, pričom respondenti odpovedali zadaním čísla, môžete s týmito odpoveďami priamo pracovať. V prípade položiek, pri ktorých respondenti vyberali jednu z možností (napr. Likertova škála), budú odpovede zapísané v podobe znenia tej ktorej možnosti (napr. „Veľmi nesúhlasím“). Takéto textové odpovede je potrebné zmeniť na čísla. V prípade nominálnych

premenných to nie je nutné – jamovi v tomto prípade dokáže pracovať s textovou hodnotou, no pri ordinálnych premenných je potrebné, aby boli hodnoty číselné. Formát odpovedí v stiahnutých dátach však závisí od konkrétnej platformy využitej pre zber dát a nemožno jednoducho generalizovať na všetky platformy.

Pokiaľ máme množstvo času a nič lepšie na práci, môžeme tieto hodnoty zameniť ručne, hodnotu po hodnote. Našťastie existujú možnosti automatického zamenenia hodnôt. V tabuľkovom procesore nájdite funkciu „Hľadať a nahradiť“ alebo „Nahradiť“ (aspoň takto je to pomenované v Microsoft Excel, kde je možné využiť klávesovú skratku Ctrl + H či Cmd + H na Mac). Vďaka tejto funkcii vieme hľadať určitý výraz či reťazec a zameniť ho za iný. V základe hľadá a nahrádza v celom hárku, no môžete vybrať len konkrétne stĺpce (označením stĺpca, resp. stĺpcov podržaním Shift a/alebo Ctrl), v ktorých chcete zmenu uskutočniť. V prípade spomenutej Likertovej škály by sme mohli dať hľadať „Veľmi nesúhlasím“ a zameniť to za 1, zvolením možnosti Nahradiť všetky (prípadne inak nazvanej možnosti v závislosti od konkrétneho tabuľkového procesora). Pri takomto zamieňaní hodnôt odporúčame pracovať opatrne a po častiach (označte si len stĺpce, pri ktorých máte rovnaký spôsob nahradenia hodnôt). Môže sa stať, že niektoré z použitých metodík majú rovnaké možnosti, no ich číselná hodnota je odlišná, napr. v jednom dotazníku máte možnosť Súhlasím, ktorú je potrebné zameniť za 4, no v inom dotazníku to má byť 6. Pochybenie v tomto kroku spôsobí, že vaše ďalšie výsledky budú neplatné. Dávajte si pozor a odporúčame vytvárať si „pevné“ kópie súborov, aby ste v prípade veľkého problému mohli využiť záložný súbor a nestratili ste hodiny práce.

2.3 Vloženie dátového súboru do jamovi a nastavenie premenných

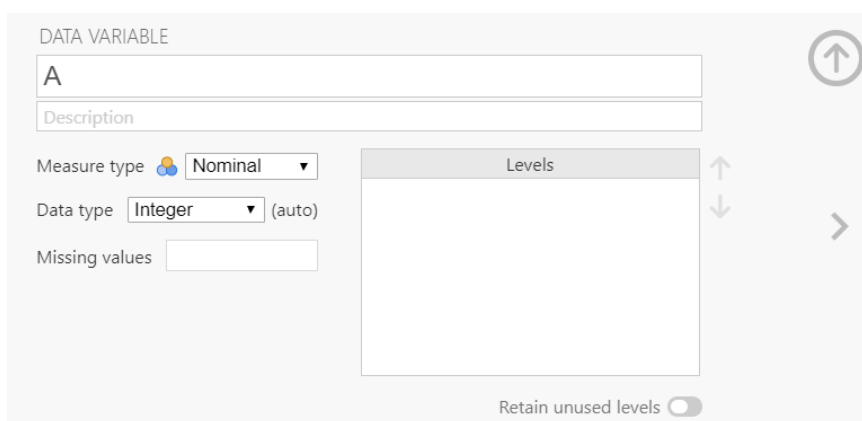
Keď už máte pripravený dátový súbor v tabuľkovom procesore, môžete dáta ľahko presunúť do programu jamovi. Uložte si dáta a otvorte program jamovi. V menu zvolte možnosť *Open*, nájdite svoj súbor a dajte otvoriť. Ak ste všetko nastavili správne, mali by sa vám dáta načítať aj s pomenovaním jednotlivých premenných. Obrázok 2.2 zobrazuje vzhľad časti dát v tabuľkovom procesore (konkrétne Microsoft Excel) a v programe jamovi po presune dát. Môžete si všimnúť, že pomenovania v prvom riadku sa stali pomenovaním premenných. Hneď, ako prenesiete dáta do programu jamovi, si tento súbor uložte (v menu zvolte možnosť *Save* a nastavte, kde a pod akým názvom ho uložiť).

	A	B	C	D	E
1	BFI_1	BFI_2	BFI_3	BFI_4	BFI_5
2	1	1	3	1	2
3	2	3	4	4	1
4	1	2	5	5	2
5	2	3	1	3	1
6	5	1	2	4	1

	 BFI_1	 BFI_2	 BFI_3	 BFI_4	 BFI_5
1	1	1	3	1	2
2	2	3	4	4	1
3	1	2	5	5	2
4	2	3	1	3	1
5	5	1	2	4	1

Obrázok 2.2: Dáta v tabuľkovom procesore a v programe jamovi

Ďalej je potrebné nastaviť atribúty jednotlivých premenných. Dvojklikom na názov premennej otvoríme nastavenie atribútov tejto premennej (obrázok 2.3).



Obrázok 2.3: Základný vzhľad rozhrania pre nastavenie atribútov premennej

V prvom okienku je názov premennej – ten máme nastavený už z tabuľky, ktorú sme vložili z tabuľkového procesoru, no v prípade potreby ho môžeme kedykoľvek zmeniť. V ďalšom okienku s názvom *Description* môžeme pridať vysvetľujúci popis k premennej, napríklad aby ste rozumeli aj neskôr, čo daná premenná vyjadruje. Najdôležitejším atribútom, ktorý je potrebné nastaviť je *Measure type* – úroveň merania premennej. Na výber máme možnosť *Nominal* pre nominálne premenné, *Ordinal* pre ordinálne premenné a *Continuous* pre kontinuálne premenné (popis je v časti 2.1). Nastavenie úrovne merania je dôležité pre niektoré štatistické analýzy, ktorým sa budeme venovať v ďalších častiach učebnice. Nie je však potrebné nastavovať každú položku osobitne. V prípade, že máte viacero položiek, pri ktorých chcete zmeniť nastavenia, môžete tak spraviť hromadne. Otvorte si nastavenia prvej z položiek, ktorú chcete upraviť, podržte na klávesnici tlačidlo Shift a kliknite na poslednú z položiek v rade, ktorú chcete upraviť. Vyberiete tým daný počet stĺpcov a zmeny robíte pre všetky tieto vybrané položky. V prípade, že by ste chceli naraz nastaviť viacero položiek, ktoré nie sú v rade za sebou, môžete tak urobiť tým, že namiesto Shift podržíte tlačidlo Ctrl a klikáte na názvy položiek pokiaľ nevyberiete všetky, o ktoré máte záujem.

Ďalším nastavením je *Data type*, prostredníctvom ktorého môžeme nastaviť, o aké dáta ide. Zadávať môžeme celé čísla – možnosť *Integer*, napríklad hodnotu, ktorú respondent zvolil na Likertovej škále alebo vek. Takéto premenné nazývame diskkrétne, pretože sú merané na úrovni celých čísiel. Opačným prípadom sú spojité premenné, pri ktorých sa môžu vyskytovať aj desatinné miesta (*Decimal*). Príkladom môže byť výška respondenta v metroch, napr. 1,76 metra. Iným príkladom môže byť priemerné skóre zo škály alebo dotazníka, napr. priemerné

skóre extravenzie a pod. Poslednou možnosťou je *Text*, kedy neplánujeme zadávať číselné hodnoty, ale text. Túto možnosť môžeme použiť napríklad pri premennej pohlavie – nemusíme zadávať číslo, ale muž alebo žena. Ideálnejšie je však zmeniť všetky informácie na čísla, keďže tieto čísla vieme v prípade potreby pomenovať. Toto nastavenie premennej nie je vo väčšine prípadov potrebné vopred meniť, keďže program jamovi ho nastaví automaticky (no pre istotu odporúčame skontrolovať). Pod nastavením typu dát, máme možnosť zadať hodnoty, ktoré v našom súbore vyjadrujú chýbajúce hodnoty – *Missing values*. Samozrejme, pokiaľ respondent nezodpovedal niektorú otázku alebo položku, môžeme nechať prázdnu bunku. Váš výskum ste ale mohli mať nastavený tak, že pri 4-stupňovej Likertovej škále (1 až 4) ste mali aj možnosť „Nechcem odpovedať“, ktorú ste pri prepise označili hodnotou napr. 0 alebo 99. Aby program jamovi vedel, že uvedené číslo vyjadruje chýbajúci údaj, môžete to nastaviť v tomto menu, kliknutím na okienko, pričom sa vám otvorí nové okno, v ktorom kliknete na [+ Add Missing Value](#) a nastavíte hodnotu, ktorá vyjadruje chýbajúci údaj.

Praktickým pomocníkom je aj možnosť pomenovať jednotlivé číselné hodnoty pri nominálnych alebo ordinálnych premenných. V časti *Levels* nájdeme číselné hodnoty obsiahnuté vo vybranej premennej či vybraných premenných. Kliknutím na číselnú hodnotu ju môžeme pomenovať. „Na pozadí“ nám stále zostane číselná hodnota, no v dátach a výsledkoch analýz sa nám zobrazia pomenovania. Ukážku nájdete v obrázku 2.4.

The image shows two screenshots of the Jamovi 'DATA VARIABLE' settings window. The top screenshot is for the variable 'pohlavie' (gender), set to 'Nominal' measure type and 'Integer' data type. The 'Levels' table lists 'muž' (1) and 'žena' (2). The bottom screenshot is for the variable 'vzdelanie' (education), set to 'Ordinal' measure type and 'Integer' data type. The 'Levels' table lists 'ZŠ' (1), 'SŠ bez maturity' (2), and 'SŠ s maturitou' (3).

DATA VARIABLE							
pohlavie							
Description							
Measure type Nominal	<table border="1"> <thead> <tr> <th colspan="2">Levels</th> </tr> </thead> <tbody> <tr> <td>muž</td> <td>1</td> </tr> <tr> <td>žena</td> <td>2</td> </tr> </tbody> </table>	Levels		muž	1	žena	2
Levels							
muž		1					
žena	2						
Data type Integer							
Missing values							

DATA VARIABLE									
vzdelanie									
Description									
Measure type Ordinal	<table border="1"> <thead> <tr> <th colspan="2">Levels</th> </tr> </thead> <tbody> <tr> <td>ZŠ</td> <td>1</td> </tr> <tr> <td>SŠ bez maturity</td> <td>2</td> </tr> <tr> <td>SŠ s maturitou</td> <td>3</td> </tr> </tbody> </table>	Levels		ZŠ	1	SŠ bez maturity	2	SŠ s maturitou	3
Levels									
ZŠ		1							
SŠ bez maturity	2								
SŠ s maturitou	3								
Data type Integer									
Missing values									

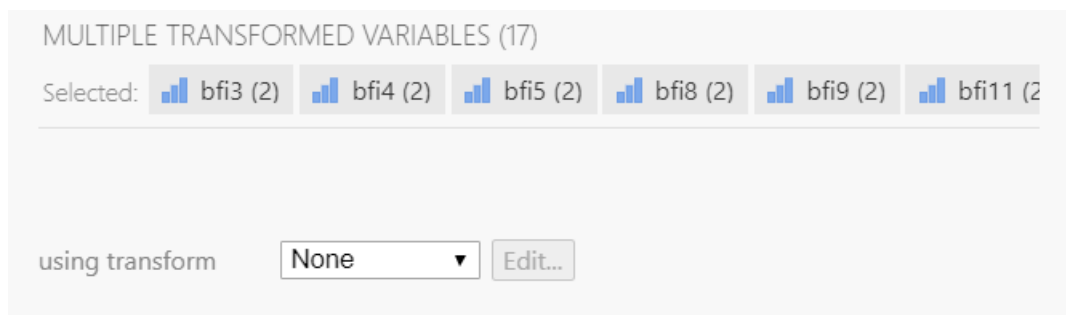
Obrázok 2.4: Ukážka nastaveného pomenovania hodnôt premennej

2.4 Transformácie hodnôt

Vložené dáta sú v „surovom“ stave. Toto čudné pomenovanie sa vyskytuje aj v anglickom jazyku tzv. „*raw data*“. Neznamená to, že ich je nutné tepelne upraviť (kreativite sa však medze nekladú), no dôležité je spraviť zopár úkonov, aby sme dáta dostali do stavu, kedy s nimi môžeme plne pracovať – hovoríme o transformácii dát či hodnôt. Čo to znamená? Ide o rôzne úpravy hodnôt, ktoré máme. Uvedieme dva príklady. Pravdepodobne najčastejšie využívaným je otočenie hodnôt premenných. Mnohé psychodiagnostické metodiky, najmä z oblasti sociálnej psychológie či psychológie osobnosti, využívajú takzvané *reverzné položky*. Ide o položky, ktoré sú významom opačné ako je meraný konštrukt. Napríklad, v prípade osobnostného dotazníka BFI-2 je polovica položiek formulovaná v smere meraného konštrukt a druhá polovica v opačnom smere. Konkrétne môžeme uviesť položku z domény Extraverzia: „Som niekto, kto je spoločenský, rád trávi čas s inými ľuďmi.“ Táto položka je určená pre meranie Sociability v rámci domény Extraverzia. Jej znenie je v smere meraného konštrukt, t.j. čím vyšší súhlas s touto položkou respondent vyjadrí, tým vyššiu mieru Sociability či celkovo Extraverzie by mal mať. No pre meranie tohto konštrukt v dotazníku BFI-2 nájdeme aj položku „Som niekto, kto je niekedy hanblivý, uzavretý.“ Zo znenia tejto položky vidíme, že je formulovaná v opačnom smere, to znamená, že čím vyšší súhlas s týmto tvrdením respondent vyjadrí, tým nižšiu mieru meraného konštrukt dosahuje. Čísla ale nepustia, t.j. pokiaľ sme pri prenose dát zamenili povedzme „Veľmi súhlasím“ za číselnú hodnotu 5, bude táto hodnota aj pri jednej aj pri druhej premennej. Z tohto dôvodu musíme otočiť hodnoty pri reverzných položkách vytvorením transformovaných premenných.

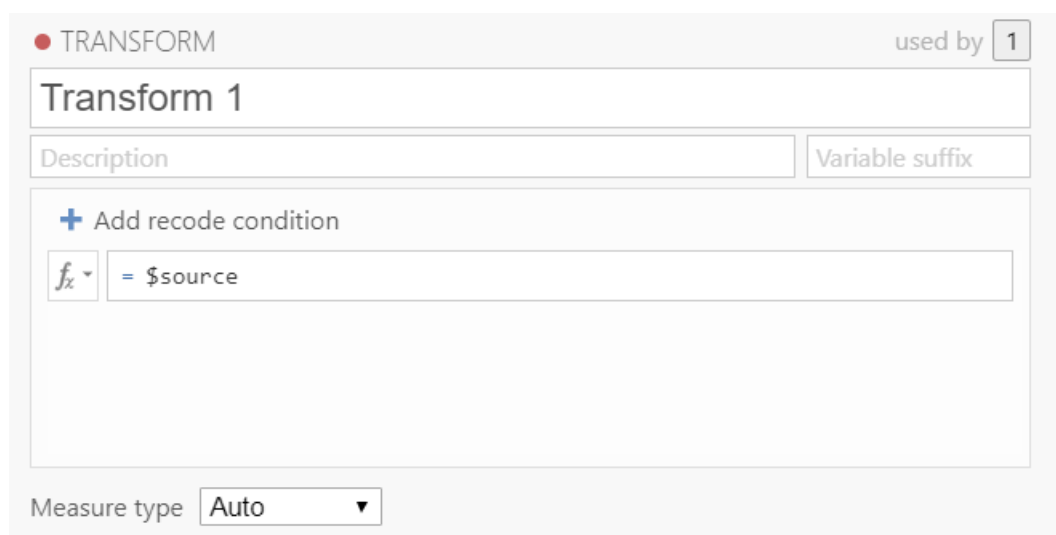
Druhým príkladom je vytvorenie premenných, v ktorých sme niektoré kategórie zlúčili. Využívame ho najmä, ak sme dáta zbierali detailnejšie, no neskôr nemáme záujem o takúto detailnú informáciu a chceme hodnoty zlúčiť. Príkladom môže byť premenná, ktorá vyjadruje najvyššie dosiahnuté vzdelanie – respondenti mali určiť svoje najvyššie dosiahnuté vzdelanie výberom jednej z možností: 1) základné, 2) stredoškolské bez maturity, 3) stredoškolské s maturitou, 4) vysokoškolské I. stupňa, 5) vysokoškolské II. stupňa, 6) vysokoškolské III. stupňa. Pre ďalšie analýzy však nepotrebujeme špecificky vymedzenú možnosť 6) a chceme ju zlúčiť s možnosťou 5. Takúto transformáciu nemusí uskutočniť nutne v programe jamovi. Pokiaľ sa pre ňu rozhodneme ešte pred vložením dát do programu jamovi, vieme tento úkon jednoducho spraviť aj v tabuľkovom procesore – v tomto prípade by sme, „5) vysokoškolské II. stupňa“ aj „6) vysokoškolské III. stupňa“, nahradili tou istou číselnou hodnotou (napr. 5), čím by sme ich zlúčili.

Transformáciu premenných môžeme v jamovi spraviť cez kartu *Data* zvolením možnosti *Transform*. Prvým krokom je označenie položiek, ktoré chceme transformovať tým istým spôsobom. Pri vašich dátach sa riadte pokynmi, ktoré sú v manuáloch metodík, ktoré ste použili. V prípade spomínaného dotazníka BFI-2 zvolíme všetky položky, ktoré podľa manuálu treba otočiť (výber položiek spravíte tak, že klikáte na názvy premenných v hlavičke tabuľky, pričom držíte Ctrl či Cmd). Po tom, ako ste označili všetky položky, zvolíte možnosť *Transform*. V dátovom súbore vám automaticky vzniknú nové premenné, ktoré budú niest transformované hodnoty a otvorí sa vám okno, v ktorom môžete nastaviť spôsob transformácie. Príklad zobrazujem v obrázku 2.5.



Obrázok 2.5: Okno pre nastavenie spôsobu transformácie

Aby sme vytvorili nový spôsob transformácie, musíme otvoriť výberové okno (v základe je v ňom popis *None*) a vybrať možnosť *Create New Transform...* Otvorí sa nám nové okno, ktoré zobrazujeme v obrázku 2.6. Toto okno ponúka viaceré nastavenia. Jedným z nich je pomenovanie nastavenia transformácie (*Transform 1*). Toto pomenovanie môžete ľubovoľne zmeniť na také, aké vám vyhovuje. Vytvorený spôsob transformácie sa dá opakovane použiť aj na iné položky. V našom prípade chceme nastaviť transformáciu pre Likertovu škálu s rozsahom od 1 až 5. Pomenujeme si túto transformáciu ako Likert 1-5. Ďalšou možnosťou je nastavenie popisu pre transformáciu (*Description*). Ak máte záujem, môžete pridať popis, v našom prípade ho nastavovať nebudeme. Užitočným je nastavenie prípony (*Variable suffix*). Ak túto hodnotu nenastavíte, novovytvorené transformované premenné budú mať názov pozostávajúci z názvu pôvodnej premennej spolu s názvom transformácie (napr. bfi1 Likert 1-5). Tento názov je zbytočne zložitý, a tak si môžete príponu nastaviť napr. na *R* ako rekódované a názov transformovanej premennej bude kratší. Najdôležitejšia je ale časť, kde nastavujeme spôsob transformácie. Tá môže byť jednoduchá ako pri otočení hodnôt alebo zložitejšia na základe logických výrazov tzv. podmieňovania. Posledným nastavením je *Measure type*, kde môžeme nastaviť, akú úroveň merania majú mať transformované položky. Pokiaľ nemáte záujem o konkrétne nastavenie, môžete ponechať automatické nastavenie, teda možnosť *Auto*.



Obrázok 2.6 Okno pre nastavenie transformácie premenných

Ak máme záujem o otočenie hodnôt, musíme vytvoriť vzorec, ktorý tieto hodnoty otočí. V prípade, že premenná, ktorú chceme otočiť, začína na hodnote 1 (napr. 5 bodová Likertova škála od 1 do 5), použijeme vzorec, kde k maximálnej hodnote škály pripočítame 1 a odčítame hodnotu zdrojovej (surovej) premennej. Hodnota surovej premennej, resp. originálnej hodnoty,

ktorú chceme transformovať je v jamovi označená ako *\$source*. V prípade, že máme škálu v rozmedzí od 1 do 5, vzorec bude vyzeráť takto $6 - \$source$ – to znamená, že ak niekto odpovedal 5, jeho transformovaná hodnota bude 1, ak niekto odpovedal 2, po transformácii to bude 4 atď. V prípade, že škála začína na 0, pracujete len s najvyššou hodnotou škály, t.j. nepričítavate 1, napr. pri škále od 0 do 4 bude vzorec vyzeráť $4 - \$source$. Príklad nastavenia pri 5 bodovej Likertovej škále s rozsahom od 1 po 5 uvádzame v obrázku 2.7.

The screenshot shows the 'TRANSFORM' panel in Jamovi. At the top, it says 'used by 1'. Below that is a text box labeled 'Likert 1-5'. Underneath is a 'Description' field with the letter 'R' next to it. There is a section titled '+ Add recode condition' with a dropdown menu showing 'fx' and a text box containing the formula '= 6 - \$source'. At the bottom, there is a 'Measure type' dropdown menu set to 'Auto'.

Obrázok 2.7: Príklad nastavenia jednoduchkej transformácie premennej

V prípade, ak by ste chceli spraviť zložitejšiu transformáciu, môžete k tomu využiť možnosť podmienkovania. Môže ísť o prípad, kedy máte premennú, pri ktorej chcete zlúčiť nejaké hodnoty alebo vytvoriť kategórie. Príkladom môže byť vytvorenie premennej, ktorá bude vyjadrovať vekové kategórie. Povedzme, že sme sa v našom dotazníku pýtali respondentov na ich vek v rokoch, no neskôr máme záujem pracovať s vekovými kategóriami. Vždy odporúčame merať na čo najvyššej úrovni. Niektorí z vás sa možno pýtajú, na čo zisťovať vek v rokoch, keď môžeme respondentom ponúknuť na výber jednu z kategórií, nech sa zaradia. Takýmto spôsobom však strácate množstvo informácie. Lepšie je mať detailnejšiu informáciu, ktorú potom transformujete na „hrubšiu“, pretože spätne to nejde. Ako teda na to? Pri nastavení transformácie musíme vytvoriť podmienku či podmienky. Kliknutím na *Add recode condition* otvoríme možnosť pre nastavenie podmienok (Obrázok 2.8).

The screenshot shows the 'Add recode condition' panel. It has a title '+ Add recode condition'. Below it are two rows of input fields. The first row has a dropdown menu with 'fx', a text box containing 'if \$source <= e.g. 2000', and another text box containing 'use e.g. 'Male'' with a close button 'X'. The second row has a dropdown menu with 'fx' and a text box containing 'else use e.g. \$source'. To the right of these rows are up and down arrow buttons.

Obrázok 2.8: Panel pre nastavenie podmienenej transformácie

Ako vidíte na obrázku 2.8, prvá kolónka obsahuje prednastavené *if \$source*, čo voľne znamená *ak je zdrojová hodnota* a za týmto môžeme nastaviť čo ďalej. Ak nám jedna takáto kolónka nestačí, môžeme pridať ďalšie, stlačením *Add recode condition*. Teraz musíme programu vysvetliť, ako má hodnoty transformovať. Vývojári jamovi nám rovno ponúkajú príklad pre lepšie intuitívne pochopenie nastavenia. V našom prípade by sme chceli transformovať

premennú vek (meraná na kontinuálnej úrovni), tak aby vznikla ordinálna premenná s tromi kategóriami: vek menší ako 30 rokov sa má stať hodnotou 1, medzi 30 (vrátane) až 50 (vrátane) sa má stať hodnotou 2, všetky ostatné hodnoty majú byť 3. V kolónke obsahujúcej *if* nastavíme podmienku a vo vedľajšej kolónke s *use* nastavíme, čo sa má stať, ak je podmienka splnená. Posledná kolónka je určená pre ostatné prípady (*else use*). Operátory („znamienka“), ktoré používame sú:

- > - väčšie ako hodnota
- < - menšie ako hodnota
- <= - menšie alebo rovné ako hodnota
- >= - väčšie alebo rovné ako hodnota
- == - rovné hodnote
- != - nie rovné hodnote

Tieto operátory sú vysvetlené aj priamo v programe po kliknutí na časť, kde sa píše. V rámci nášho príkladu by sme nastavili:

- if \$source > 50 use 3 – v tomto prípade sme nastavili strop ďalšiemu nastaveniu a hodnoty väčšie ako 50 sme transformovali za 3
- if \$source >= 30 use 2 – hodnoty, ktoré sú väčšie alebo rovné ako 30 zamení za 2
- if \$source < 30 use 1 – hodnoty menšie ako 30 zamení za 1
- else use – nemusíme nastavovať

V prípade tejto transformácie nastavíme aj *Measure type* na *Ordinal*. Nastavenie v programe jamovi zobrazujeme v obrázku 2.9. Po potvrdení uvidíme v dátach novú premennú „vek – R“, v stĺpci napravo od originálnej premennej „vek“. Alternatívnym nastavením je zložitejšia podmienka, kedy využívame operátory. Jedným z nich je *and* – v preklade „a“ tento operátor nám povolí v jednom kroku nastaviť 2 a viac podmienok, ktoré musia byť naplnené. Druhý operátor je *or* – v preklade „alebo“. Vďaka tomuto operátoru môžeme pridať viacero podmienok v jednom kroku s tým, že stačí, aby bol naplnený len jeden z nich, aby bola nastavená transformovaná hodnota. Alternatívne nastavenie transformácie by mohlo znieť takto (zobrazené v obrázku 2.10):

- if \$source < 30 use 1 – hodnoty menšie ako 30 zamení za 1
- if \$source >= 30 and \$source <= 50 use 2 – hodnoty, ktoré sú väčšie alebo rovné 30 a zároveň menšie alebo rovné 50 zamení za 2
- if \$source > 50 use 3 – hodnoty väčšie ako 50 zamení za 3

Pri každom nastavení odporúčame, aby ste si skontrolovali, či transformácia funguje tak ako ste plánovali – pozrite si zopár hodnôt originálnej premennej a hodnoty transformovanej premennej. Ak by transformácia nefungovala tak, ako ste plánovali, skúste sa znovu zamyslieť nad systémom transformácie a skontrolujte nastavené podmienky.

● TRANSFORM used by 1

vek kategorie

Description R

+ Add recode condition

f_x	if \$source > 50	use 3	×
f_x	if \$source >= 30	use 2	×
f_x	if \$source < 30	use 1	×

Measure type Ordinal

Obrázok 2.9: Príklad nastavenia transformácie premennej vek (jednoduchá podmienka)

● TRANSFORM used by 1

vek kategorie

Description R

+ Add recode condition

f_x	if \$source < 30	use 1	×
f_x	if \$source >= 30 and \$source <= 50	use 2	×
f_x	if \$source > 50	use 3	×

Measure type Ordinal

Obrázok 2.10: Príklad nastavenia transformácie premennej vek (zložená podmienka)

2.5 Matematické operácie

Zatiaľ sme si ukázali ako dostať dáta do programu jamovi, ako nastaviť premenné a základné možnosti transformácie. V ďalších častiach tohto učebného textu sa budeme venovať deskriptívnej štatistike, no skôr ako sa k tomu dostaneme, je potrebné, aby sme sa oboznámili s ďalším spôsobom vytvárania nových premenných. V psychologickom výskume, konkrétne v kvantitatívnom výskume, nemôžeme merať konštrukty priamo, no meriame ich nepriamo prostredníctvom položiek psychodiagnostických nástrojov. Tieto nástroje, nazývané aj škály, inventáre alebo dotazníky, obsahujú rôzny počet položiek, ktoré sú určené pre nepriame meranie, inak povedané „zachytenie“ cieľového konštruktu. Jednotlivé položky majú určitú výpovednú hodnotu a môžeme sa napríklad zaujímať o frekvenciu výskytu jednotlivých odpovedí v našom súbore. Pre ďalšie pokročilejšie analýzy využívame hodnoty, ktoré vznikajú sčítaním alebo spriemerovaním hodnôt viacerých položiek tak, aby sme dostali jednu hodnotu, ktorá prezentuje daný konštrukt. V prípade, že by sme sa zaujímali o mieru extravenzie, mohli by sme zvoliť položky z Inventára Veľkej Päťky 2. Pre meranie extravenzie je určených 12 položiek. Jednotlivé položky sú len nepriame indikátory, ktoré merajú konkrétnejšie prejavy

konštruktu, na základe jeho teoretického vymedzenia. Hodnotu extravenzie získame vypočítaním priemeru odpovedí. Podľa inštrukcie priloženej k inventáru je potrebné zohľadniť, že 6 položiek je formulovaných opačným smerom a je ich potrebné rekódovať (o reverzných položkách a ich transformácii sme informovali v predchádzajúcej podkapitole). Pri výpočtoch sa riadte inštrukciou, ktorá bola uvedená pri vami použitých metodikách. Najčastejšie sú však využívané sumárne skóre (spočítanie jednotlivých odpovedí na položky), nazývané tiež hrubé skóre a priemerné skóre (priemer odpovedí na položky). Priemerné skóre má výhodu a nachádza využitie pri Likertových škálach, pretože nie je závislé na počte položiek, zakaždým bude jeho hodnota v rozmedzí použitej Likertovej škály. Ďalšia výhoda je pri chýbajúcich hodnotách – ak respondent zodpovie len na 10 z 12 položiek, môžeme vypočítať jeho priemerné skóre (ak nám absencia určitého počtu odpovedí neprekáža). Sumárne skóre nachádza uplatnenie najmä pri položkách, kde je len jedna možnosť správna (napríklad výkonové testy, inteligenčné testy a iné) alebo pri položkách typu áno/nie. Keďže tieto položky majú po transformácii len 2 možné hodnoty 0 a 1, využíva sa skôr sumárne skóre. Oba spôsoby si ukážeme.

Pre vytvorenie novej vypočítanej premennej môžeme dvoj-kliknúť na hlavičku prázdneho stĺpca a zvoliť možnosť *NEW COMPUTED VARIABLE*, alebo môžeme ísť cez kartu *Data* a kliknúť na možnosť *Compute*. Otvorí sa nám okno podobné ako pri transformácii, teraz však nastavujeme vzorec, na základe ktorého budú vypočítané hodnoty novej premennej. V prvom kroku je užitočné nastaviť názov novej premennej. Názov by mal byť stručný a výstižný, napríklad *Extraverzia_HS*, *Extra_mean* a podobne. Dôležité je, aby ste sa vedeli ľahko zorientovať. Tak ako predtým, aj teraz môžete nastaviť popis pre novú premennú, no nie je to povinné. Najdôležitejšia časť je okienko na pravej strane, do ktorého píšete vzorec. Vzorec píšete s použitím názvov premenných. Na uľahčenie môžete kliknúť na tlačidlo *fx*, ktoré vám ponúkne premenné, ktoré sa nachádzajú vo vašich dátach ale aj iné automatické funkcie. Opätovne upozorňujem na zvýšenie pozornosti – ak zadáte zlý výpočet, budete pracovať s neplatnou premennou. Dobré si rozmyslite ako postupovať.

Ak by ste mali záujem o výpočet hrubého skóre, môžeme jednoducho vybrať položky alebo napísať ich názvy a oddeliť ich znamienkom plus, napr. $\text{premenná1} + \text{premenná2} + \dots + \text{premennáX}$. Taktiež môžete využiť funkciu *SUM()*, pričom do zátvorky napíšete alebo vyberiete premenné a oddelíte ich čiarkou, napr. *SUM(premenná1, premenná2, ... , premennáX)*. Výhodou tohto prístupu je možnosť počítat aj v prípade ak máte pri premenných chýbajúce hodnoty a to pridaním argumentu *ignore_missing = 1*, ktorý oddelíte čiarkou za poslednou premennou vo vzorci. Pokiaľ napíšete jednoduchý vzorec bez funkcie či automatickú funkciu *SUM()* bez tohto argumentu a v nejakom riadku máte chýbajúcu hodnotu, nebude v danom riadku vypočítaná žiadna hodnota. Zápis s argumentom pre počítanie s chýbajúcimi hodnotami môže vyzerat' napríklad takto *SUM(premenná1, premenná2, ... , premennáX, ignore_missing = 1)*.

Ak máte záujem o výpočet priemerného skóre, opäť môžete napísať vzorec ručne napr. $(\text{premenná1} + \text{premenná2} + \dots + \text{premennáX}) / X$; čiže zadáte štandardný vzorec výpočtu priemeru: súčet hodnôt delene ich počtom. Tak ako v predchádzajúcom prípade odporúčame však využiť automatickú funkciu *MEAN()*. Stačí do zátvoriek vložiť premenné a oddeliť ich čiarkou. Taktiež môžete nastaviť argument *ignore_missing = 1*, pokiaľ máte v riadkoch pre niektorých premenných chýbajúce hodnoty, no chcete s nimi pracovať. Výsledný vzorec bude v tvare *MEAN(premenná1, premenná2, ... , premennáX, ignore_missing = 1)*.

Ak počítate hrubé alebo priemerné skóre, dajte si pozor, aby ste použili už transformované reverzné položky (ak ste ich vypočítali). V obrázku 2.11 zobrazujeme nastavenia vypočítanej premennej Extra_priemer. Podľa manuálu k tomuto inventáru, máme otočiť položky číslo 11, 16, 26, 31, 36, 51 (v našom príklade sú tieto premenné pomenované *bfi* a číslo položky). Najskôr sme si tieto premenné transformovali (uvedené v predchádzajúcej podkapitole) a teraz vypočítame priemer pre každého respondenta. V tomto konkrétnom prípade sme si vybrali funkciu MEAN() a do zátvoriek sme nastavili premenné (vybrali sme ich zo zoznamu premenných v ponuke *fx*) a oddelili čiarkou. Všimnite si, že niektoré premenné sú bez úvodzoviek a niektoré v úvodzovkách. V úvodzovkách sú tie, ktorých názvy obsahujú medzeru. Výsledný vzorec vyzerá nasledovne: MEAN(bfi1, bfi6, `bfi11 - R`, `bfi16 - R`, bfi21, `bfi26 - R`, `bfi31 - R`, `bfi36 - R`, bfi41, bfi46, `bfi51 - R`, bfi56). Po napísaní vzorca máme automaticky vypočítané hodnoty v novej premennej Extra_priemer, vid' Obrázok 2.11.

COMPUTED VARIABLE

Extra_priemer

Description

Formula

f_x = MEAN(bfi1, bfi6, `bfi11 - R`, `bfi16 - R`, bfi21, `bfi26 - R`, `bfi31 - R`, `bfi36 - R`, bfi41, bfi46, `bfi51 - R`, bfi56)

Retain unused levels ☐

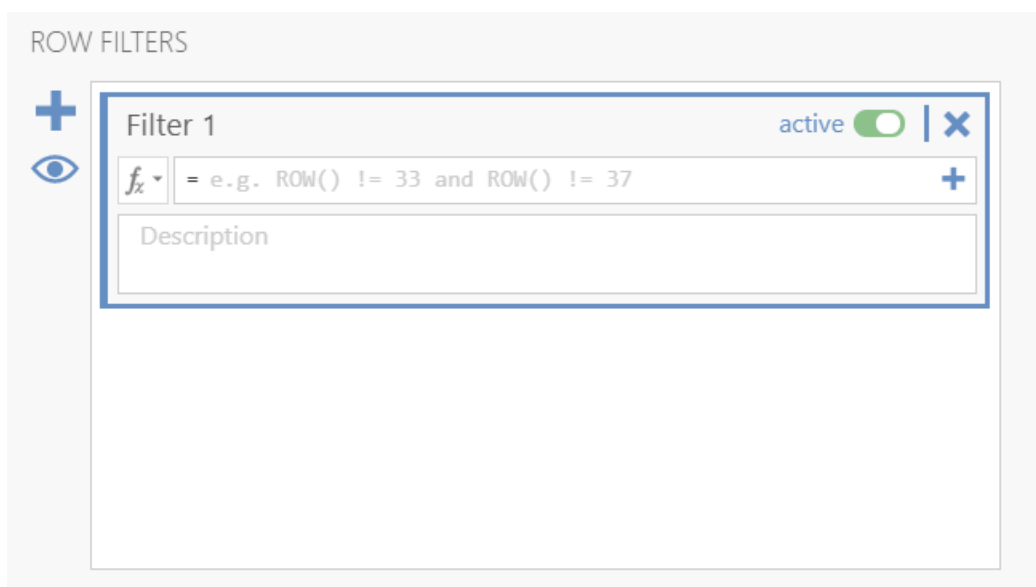
bfi36 - R	bfi41	bfi46	bfi51	bfi51 - R	bfi56	Extra_priemer
3	4	5	3	3	4	3.333
3	1	5	2	4	2	2.917
4	5	3	5	1	2	2.833
3	3	5	3	3	4	2.833
3	4	5	5	1	3	3.417
2	4	5	4	2	5	2.750

2.11 Ukážka nastavenia vypočítanej premennej

2.6 Filtrovanie dát

Poslednou zo základných funkcií v jamovi je filtrovanie dát. Nachádza uplatnenie v prípade, že z nejakého dôvodu nemáme záujem pracovať s celým dátovým súborom ale iba s respondentmi, participantmi či všeobecne prípadmi, ktoré spĺňajú nejakú podmienku. Nastavením filtra vylúčime z ďalších analýz tie prípady, ktoré nespĺňajú stanovenú podmienku. Podmienka môže byť jednoduchá – napríklad chceme pracovať len s tými, ktorí uviedli ako pohlavie muž, no môže byť aj zložená – napríklad muži, ktorí majú najvyššie dosiahnuté vzdelanie stredoškolské a majú viac ako 30 rokov.

Ak chceme nastaviť filter, môžeme tak spraviť cez kartu *Data* a vyberieme možnosť *Filters*. Zobrazí sa nám okno, v ktorom môžeme nastaviť filter (Obrázok 2.12). Pri prvom otvorení okna je automaticky vytvorený prvý filter *Filter 1*, a v prvom stĺpci dátového súboru sa zobrazí nová premenná *Filter 1*. V jednotlivých riadkoch dátového súboru je vidieť, či je daný riadok aktívny  alebo neaktívny . V okne nastavenia filtra máme rôzne možnosti. Veľké „+“ slúži pre prídanie ďalšieho filtra. Oko pod ním slúži pre zobrazenie alebo skrytie filtra (pozor, nie na vypnutie) v dátovom súbore, čo je užitočné, ak máme viacero filtrov a chceme si nechať zobrazené len tie, ktoré sú aktívne. Nastavený filter vieme podľa potreby vypnúť a znovu zapnúť, na čo slúži posuvné tlačidlo, ktoré je v základe v stave „active“ a svieti na zeleno. Stlačením tohto tlačidla vieme filter deaktivovať, čím prestane filtrovať dáta. Ak chceme filter natrvalo odstrániť, môžeme kliknúť na X. Asi najdôležitejším je samotný priestor pre nastavenie podmienky, na základe ktorej chceme dáta filtrovať. Ak chcete, môžete si nastaviť aj slovný popis, k čomu filter slúži (*Description*).



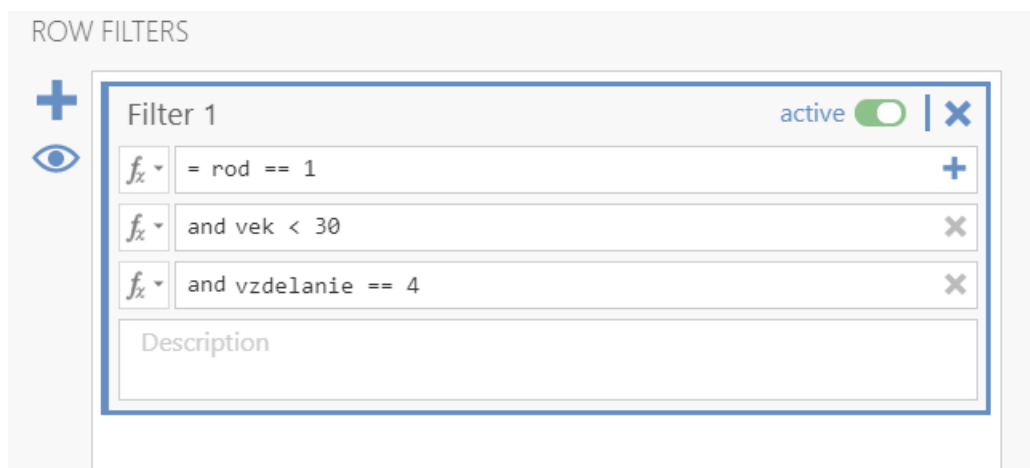
Obrázok 2.12: Základné okno pre nastavenie filtra prípadov (riadkov)

Výslednú podmienku môžeme nastaviť viacerými spôsobmi. Využívame na to princíp podobný tomu pri transformáciách. V základe, potrebujeme programu povedať, na akom princípe má rozhodnúť, či má byť riadok (prípád, respondent) aktívny alebo neaktívny pre ďalšie analýzy. Napríklad, chceme nechať aktívne len ženy v našom dátovom súbore. Takúto informáciu nesie premenná „rod“, pri ktorej máme v rámci nášho príkladu hodnotu 1 pre ženy a 2 pre mužov. Podmienka bude $rod = 1$, t.j. chceme nechať aktívne riadky, v ktorých platí, že premenná hodnota v premennej *rod* je rovná 1. Dajte si pozor, v zápise sú dve = za sebou. Po dopísaní podmienky si môžete všimnúť, že aktívne zostanú len riadky, kde je podmienka naplnená (Obrázok 2.13).

	Filter 1	rod	vzdelanie	vek
1		žena	3	22
2		muž	2	18
3		muž	4	27
4		žena	5	32
5		žena	4	24

Obrázok 2.13 Príklad vzhľadu dát po nastavení filtra $\text{rod} == 1$

Ak by sme chceli filtrovať na základe viacerých podmienok, môžeme tak spraviť viacerými spôsobmi. Jedným z nich je vytvorenie ďalšieho filtra, pridanie ďalšieho riadku do jedného filtra (malé „+“ v okne, kde nastavujeme podmienku), prípadne priamo v riadku. To, aký spôsob zvolíte, je už na vás a závisí aj od toho, čo konkrétne potrebujete. Pre ukážku chceme pracovať len s prípadmi, kedy ide o ženy, ktoré majú menej ako 30 rokov a najvyššie dosiahnuté vzdelanie „Vysokoškolské I. stupňa“ (hodnota 4). Priamo v jednom riadku by sme mohli napísať „ $\text{rod} == 1$ and $\text{vek} < 30$ and $\text{vzdelanie} == 4$ “. Alebo môžeme túto podmienku nastaviť vo viacerých riadkoch (obrázok 2.14 a 2.15).



2.14 Nastavenie filtra s viacerými podmienkami

	Filter 1	F1 (2)	F1 (3)	rod	vzdelanie	vek
1	✓	✓	✗	žena	3	22
2	✗	✓	✗	muž	2	18
3	✗	✓	✓	muž	4	27
4	✓	✗	✗	žena	5	32
5	✓	✓	✓	žena	4	24

Obrázok 2.15 Príklad vzhľadu dát po nastavení filtra $\text{rod} == 1$ and $\text{vek} < 30$ and $\text{vzdelanie} == 4$

Pri práci s filtrami si dajte pozor, aby ste ich *vypli (deaktivovali) prípadne odstránili, keď ich už nepotrebujete*.

3. Základná deskriptívna štatistika

Dáta, s ktorými pracujeme, poskytujú detailnú informáciu, ktorú je však nemožné uchopiť a pracovať s ňou. Pre možnosť zachytiť pomery v dátach využívame deskriptívnu štatistiku, ktorá nám umožňuje nazrieť na tieto dáta globálnejšie. Za pojmom deskriptívna štatistika sa skrýva množstvo ukazovateľov, s ktorými sa stýkame takmer dennodenne a v mnohých prípadoch si to ani neuvedomujeme. O niektorých z nich, ktoré používame pri psychologickom výskume a spracovávaní výsledkov, si povieme v nasledujúcich častiach. Rozdeľujeme ich na dve základné časti a to miery centrálnej tendencie, resp. miery polohy a miery variability.

3.1 Miery polohy

Miery polohy nám poskytujú predstavu o centrálnej tendencii ostatných hodnôt, zastupujú a reprezentujú ostatné hodnoty. Pri práci s psychologickými dátami používame najčastejšie aritmetický priemer, no využitie nachádza aj medián a modus.

Aritmetický priemer je miera centrálnej tendencie, ktorá je v bežnom živote asi najčastejšie využívanou. Aritmetický priemer poznáme už zo základnej školy. Väčšina z nás ho využívala napríklad pre výpočet známky, ktorú dostaneme na vysvedčení z predmetov. Ak sme chceli vedieť, akú známku budeme mať z matematiky, sčítali sme si jednotlivé známky a vydělili počtom týchto známok. Výsledkom bola priemerná známka z predmetu, vďaka čomu sme mali informáciu, ako na tom zhruba sme. Priemer teda vypočítame ako súčet všetkých hodnôt premennej, ktorý vydělíme počtom týchto hodnôt.

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$$

Ak by sme z matematiky mali za polrok 6 známok, napríklad: 2, 3, 2, 1, 3, 1, priemer známok by sme vypočítali ako:

$$\begin{aligned}\bar{x} &= \frac{2 + 3 + 2 + 1 + 3 + 1}{6} \\ \bar{x} &= \frac{12}{6} \\ \bar{x} &= 2\end{aligned}$$

Zistili by sme, že na vysvedčení pravdepodobne uvidíme známku 2. Už pri takto nízkom počte prvkov (jednotlivých známok) je priemer užitočným ukazovateľom a vieme vďaka nemu zhodnotiť aj porovnať pomery v dátach. Náš spolužiak by mohol mať známky 2, 1, 2, 3, 2, 1 – na prvý pohľad vidíme, že známky sa mierne líšia, a vďaka priemeru zistíme, že spolužiak má o niečo lepšie známky ako my, jeho priemer by bol cirka 1,83 (no dvojku na vysvedčení bude mať pravdepodobne aj on). Pri práci s priemerami je však potrebné uvedomiť si, že priemer priemerov nemusí byť nutne priemer celku. Ak by sme chceli vypočítať priemernú známku v skupine kamarátov, mohli by sme spočítať priemerné známky našich kamarátov a vydeliť ich

počtom kamarátov. V prípade, že by všetci kamaráti, s ktorých známami by sme pracovali, mali rovnaký počet známok, získame platnú informáciu. Ak sa však počet známok líši, nezískame platnú informáciu. V takomto prípade musíme vypočítať *vážený aritmetický priemer*. Slovo „vážený“ v tomto prípade neznamena oslovenie, vyjadrujúce našu mieru úcty k priemeru, no znamená, že pri výpočte prihliadame na počet prvkov – vyvažujeme vzhľadom na počet prvkov. Vážený priemer vypočítame tak, že jednotlivé priemery vynásobíme počtom hodnôt, z ktorých boli získané, tieto hodnoty spočítame a vydáme celkovým počtom hodnôt. Rozdiel medzi priemerom priemerov a váženým aritmetickým priemerom uvádzame nižšie. V tomto príklade sme sa rozhodli zistiť priemernú známku v skupine 10 kamarátov. Poprosili sme ich, aby nám povedali, akú majú priemernú známku z matematiky (zaokrúhlenú na 2 desatinné miesta) ako aj počet známok, ktoré majú. Podľa informácií v tabuľke vidíme, že počet známok sa u jednotlivých kamarátov líši. Ak by sme v tomto prípade len vypočítali priemer priemerov, získali by sme inú hodnotu ako v prípade platného postupu – počítania váženého aritmetického priemeru.

Kamarát	Počet známok	Priemerná známka	Priemerná známka krát počet známok
Veronika	4	1,25	5
Peter	6	1,83	10,98
Patrik	6	2	12
Ivana	6	2,33	13,98
Vlada	6	1,33	7,98
Dominika	7	1,71	11,97
Juraj	4	2,25	9
Peto	5	2,6	13
Eva	7	1,14	7,98
Michal	5	3,2	16
Priemer priemerov		1,964	
Vážený aritmetický priemer			1,927

Uvádzame tiež zopár dôležitých informácií ohľadom priemeru.

Súčet odchýlok všetkých hodnôt od aritmetického priemeru sa rovná nule: ak by sme z hodnôt vypočítali priemer, potom tento priemer odčítali od jednotlivých hodnôt a tieto rozdiely sčítali, získame hodnotu 0. Máme hodnoty: 2, 3, 2, 1, 3, 1. Priemer je 2. Rozdiely od priemeru sú: 0, 1, 0, -1, 1, -1. Ich súčet je 0.

Hodnota priemeru je najbližšou hodnotou k všetkým ostatným hodnotám, na základe ktorých bol priemer vypočítaný. Odbornejšie (no zložitejšie) vyjadrené, súčet druhých mocnín odchýlok všetkých hodnôt od aritmetického priemeru je menší ako súčet druhých mocnín odchýlok všetkých hodnôt od akejkoľvek inej hodnoty. To znamená, že pokiaľ dáme všetky odchýlky od priemeru na druhú (vynásobíme ich sebou samým) a tieto hodnoty sčítame, získame menšie číslo, ako keby sme pracovali s akoukoľvek inou hodnotou mimo priemeru. Odchýlky od priemeru sme vypočítali v predchádzajúcom bode, teraz ich umocníme: 0, 1, 0, 1, 1, 1. Súčet je 4. Vyskúšame si spraviť odchýlky od inej hodnoty, napr. 2,1. Odchýlky budú: -

0,1; 0,9; -0,1; -1,1; 0,9; -1,1. Umocnením získame hodnoty: 0,01; 0,81; 0,01; 1,21; 0,81 a 1,21. Súčet je 4,07 – táto hodnota je vyššia, ako keď sme použili odchýlky od priemeru.

Ak jednotlivé hodnoty, na ktorých bol vypočítaný priemer, nahradíme ich priemerom, ich priemer zostane rovnaký. To znamená, že ak nahradíme naše pôvodné hodnoty ich priemerom a vypočítame priemer, hodnota priemeru sa nezmení.

Aritmetický priemer je ovplyvniteľný vzdialenými alebo extrémnymi hodnotami. To znamená, že ak by sme mali v súbore hodnôt netradične nízku alebo vysokú hodnotu, priemer sa zmení a je menej výpovedný pre väčšinu dát. Príkladom môže byť rad hodnôt 1, 2, 1, 1, 5. Hodnota 5 je v tomto prípade vzdialená od zvyšku. Ak vypočítame priemer z týchto hodnôt, získame hodnotu 2, bez tejto vzdialenej hodnoty by bol priemer 1,25.

Aritmetický priemer aritmetických priemerov častí súboru nie je aritmetickým priemerom celého súboru – túto informáciu sme popísali v príklade vyššie.

Priemer, ako deskriptívny údaj, nemôžeme používať v prípade, ak je rozloženie dát viacvrcholové. To znamená, že je viacero hodnôt, ktoré majú rovnakú a zároveň najčastejšiu početnosť. Z hľadiska rozloženia dát priemer tiež nepoužívame, ak sú hodnoty veľmi zošikmené alebo je počet hodnôt extrémne malý. Priemer nemôžeme použiť, pokiaľ nemeríme na kontinuálnej úrovni, nemôžeme ho použiť pri ordinálnych alebo nominálnych premenných (napr. priemerné najvyššie dosiahnuté vzdelanie alebo priemerné pohlavie).

Medián je prostrednou hodnotou v usporiadanom rade hodnôt. Medián poskytuje informáciu o povahe dát takým spôsobom, že polovica hodnôt v tomto rade je menšia ako medián alebo rovná mediánu a polovica hodnôt je rovná alebo vyššia ako medián. Ak máme rad čísel, ktorý usporiadame od najmenšej po najväčšiu, tak číslo uprostred je medián. Pri nepárnom počte hodnôt nájdeme medián tak, že v usporiadanom rade hodnôt sledujeme číslo na pozícii $(n-1)/2+1$, pričom „n“ vyjadruje počet hodnôt. Ak by sme mali 7 hodnôt: 1, 1, 2, 3, 3, 4, 4, tak medián nájdeme na pozícii $(7-1)/2+1$, t.j. štvrtá hodnota v rade, v našom prípade hodnota 3. V prípade párneho počtu hodnôt je mediánom aritmetický priemer hodnôt, ktoré sa nachádzajú na pozíciách $n/2$ a $n/2+1$. Napríklad pri 8 hodnotách 1, 1, 2, 3, 3, 4, 4, 4 vypočítame medián ako priemer hodnôt na štvrtej $(8/2)$ a piatej $(8/2+1)$ pozícii, v našom prípade hodnota 3. Medián nie je ovplyvnený extrémnymi hodnotami, keďže nie je vypočítaný na základe všetkých konkrétnych hodnôt ale na základe ich počtu. Pre príklad môžeme zameniť jednu z hodnôt 4 v predchádzajúcom rade za, povedzme, 15. Pri aritmetickom priemere by takáto extrémna hodnota spôsobila výraznú zmenu (2,75 verzus 4,125), no táto extrémna hodnota nič nemení na prostrednej hodnote, t.j. medián je stále 3. Medián nachádza uplatnenie aj pri deskripcii ordinálnych premenných. Môžeme sa zaujímať o najvyššie dosiahnuté vzdelanie u respondentov, pričom by mali na výber možnosti 1) Základná škola, 2) Stredná škola bez maturity, 3) Stredná škola s maturitou, 4) Vysoká škola I. stupňa a 5) Vysoká škola II. stupňa a viac. Pokiaľ mediánom tejto premennej bude hodnota 3, zistíme, že polovica respondentov má dosiahnuté nižšie ako stredoškolské vzdelanie s maturitou alebo práve toto vzdelanie a polovica má dosiahnuté stredoškolské vzdelanie s maturitou alebo vyššie. Medián nemôžeme použiť pri deskripcii nominálnych premenných, keďže jednotlivé kategórie nominálnym premenných nie je možné usporiadať.

Modus pravdepodobne všetci poznáme, najmä vďaka ich hitu Ty, ja a môj brat („odrazové sklíčka, sklíčka dotykov, svetia vo tme, ty si taká fajn...“). V rámci deskriptívnej štatistiky je

modus možno menej obľúbený, no vyjadruje hodnotu, ktorá je zastúpená v najvyššej frekvencii (najčastejšie). Použiteľný je pri všetkých úrovniach merania (nominálnej, ordinálnej a kontinuálnej), no málokedy je priamo pomenovávaný. Príkladom môže byť deskripcia výskumného súboru, kde popíšeme, že väčšina respondentov boli ženy. V takomto prípade je modusom premennej rod hodnota žena (prípadne číslo, ktorým sme túto nominálnu kategóriu označili). V rade čísel 1, 1, 1, 2, 2, 2, 2, 3, 3, sa 1 vyskytuje trikrát, 2 štyrikrát a 3 dvakrát. Modusom je teda hodnota 2, lebo má najvyššiu frekvenciu

3.2 Miery variability

Miery variability či miery disperzie ponúkajú doplňujúcu informáciu k mieram polohy. Informujú nás o tom, ako hodnoty premennej variujú, aké majú rozloženie. Pre dôkladnejšie informovanie o povahe hodnôt používame miery polohy a variability spoločne. Pravdepodobne najznámejším a najpoužívanším párom je priemer a štandardná odchýlka. Vo výskumných štúdiách ste sa už určite stretli s M (z anglického *mean*, priemer) a SD (z anglického *standard deviation*, štandardná odchýlka). Postupne uvedieme krátku informáciu o používaných mierach variability.

Variačné rozpätie

Variačné rozpätie, po anglicky *range*, je najjednoduchšou mierou variability. Vyjadruje rozdiel medzi maximálnou hodnotou premennej (x_{max}) a minimálnou hodnotou premennej (x_{min}). Jej matematické vyjadrenie vyzerá takto:

$$R = x_{max} - x_{min}$$

Dajte si však pozor: veľké R je používané najmä pre označenie viacnásobného korelačného koeficientu, preto ak sa s ním stretnete vo výsledkoch nejakej štúdie, pravdepodobne nejde o variačné rozpätie. Vo všeobecnosti sa len málokedy variačné rozpätie používa pre účely popisu variability dát. Je málo informatívne – ak je variačné rozpätie veku respondentov 30, tak vieme, že medzi najmladším a najstarším respondentom je rozdiel 30 rokov, avšak nevieme, aké sú pomery. Môže ísť o situáciu, kedy má najmladší respondent 18 rokov a najstarší 48 rokov, no rovnakú hodnotu dostaneme aj pri zásadne inom súbore, napríklad vo veku 43 až 73 rokov. Taktiež neponúka informáciu o rozložení hodnôt okolo centrálnej hodnoty a je ovplyvnené extrémnymi hodnotami. Príkladom môže byť súbor, kde je väčšina respondentov vo veku 18 až 30 rokov, v takom prípade je hodnotou variačného rozpätia 12, no stačí, že pribudne jeden 50 ročný respondent a hodnota sa rázne zmení na 32, len kvôli jednej, odľahlej hodnote. V praxi je lepšie uviesť informáciu o minimálnej a maximálnej hodnote, čo je pre nás ako aj pre čitateľa informatívnejšie.

Rozptyl

Rozptyl, po anglicky *variance*, je na pozadí viacerých pokročilejších štatistických analýz. Vyjadruje rozloženie hodnôt okolo priemeru, avšak v jednotkách na druhú, čiže je ťažko interpretovateľný. Vzorec pre jeho výpočet vyzerá nasledovne:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Slovne by sme tento vzorec vyjadrili ako *priemer štvorcov odchýlok všetkých hodnôt od priemeru*. Keby sme chceli rozptyl vypočítať ručne, v prvom kroku by sme vypočítali priemer všetkých hodnôt $\bar{x}$. Potom by sme zobrali prvú hodnotu a odčítali od nej priemer a výsledok by sme dali na druhú (to je ten štvorec v definícii). Prečo na druhú? Ak si pamätáte z hodín matematiky, tak akékoľvek číslo na druhú je nutne pozitívne. Ak by sme to nespravili, získali by sme negatívne čísla pri tých hodnotách, ktoré sú menšie ako priemerná hodnota. Ak ste pozorne čítali vlastnosti priemeru, tak viete, že súčet odchýlok hodnôt od priemeru týchto hodnôt sa rovná nule – výsledkom by bola nula. Následne, po vypočítaní štvorcov odchýlok pre všetky hodnoty, vypočítame ich priemer, a *voilà*, úspešne sme vypočítali rozptyl. Výsledná hodnota je na druhú, takže je ťažko interpretovateľná – napríklad rozptyl veku 256 rokov² je prinajmenšom podivné vyjadrenie. Stačí však spraviť jednu matematickú operáciu, konkrétne odmocnenie, a vznikne nám...

Štandardná odchýlka

Štandardná odchýlka, prípadne smerodajná odchýlka, anglicky *standard deviation*, je vypočítaná prostredníctvom odmocnenia rozptylu. Vzorec pre výpočet vyzerá takto:

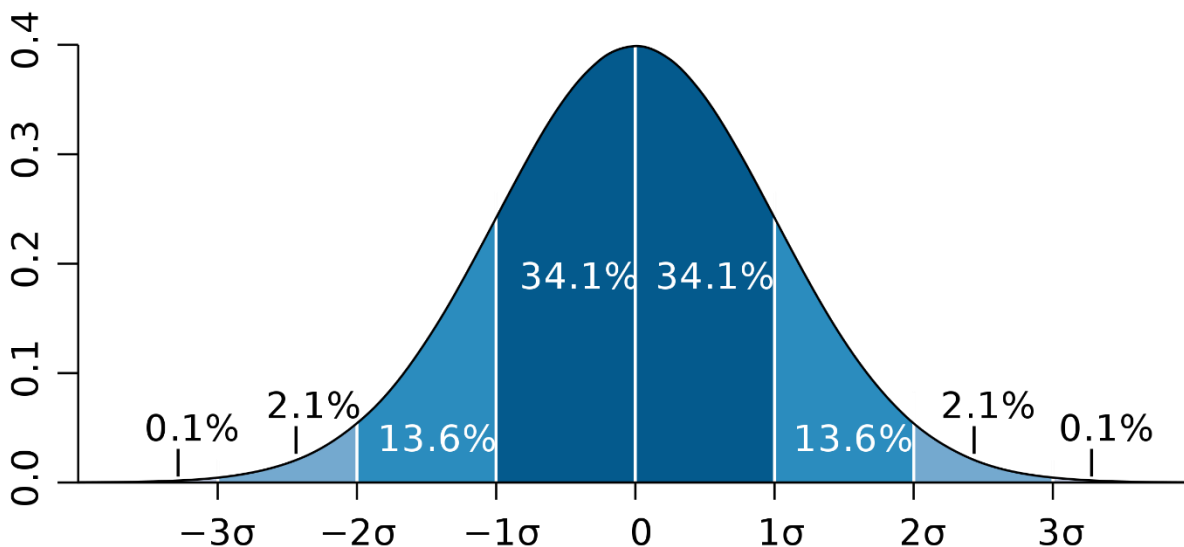
$$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$$

Ak by sme ho chceli vyjadriť slovne, povedali by sme, že ide o *odmocninu priemeru štvorcov odchýlok všetkých hodnôt od priemeru*, ale to už vieme na základe definície rozptylu.

V psychologickej literatúre sa štandardná odchýlka najčastejšie označuje ako SD. Vysoko pravdepodobne ste sa pri čítaní psychologickej literatúry s touto skratkou stretli, keďže je uvádzaná snád' v každej výskumnej štúdii. Prečo je taká populárna? Na čo nám štandardná odchýlka slúži? *Štandardná odchýlka* je dobrou priateľkou miery centrálnej tendencie *priemeru*. On je priemerný, ona štandardná, a napriek tomu pre nás ako výskumníkov má tento pár veľký význam. Štandardná odchýlka totižto vyjadruje rozloženie hodnôt okolo priemeru. V prípade distribúcie, ktorá je podobná *normálnej* distribúcii dát, platí pravidlo 68-95-99,7, ktoré je pre nás informatívne, a vďaka nemu máme lepšiu predstavu o povahe dát, resp. distribúcii hodnôt analyzovanej premennej:

- približne 68 % hodnôt premennej leží v oblasti medzi 1 SD pod priemerom a 1 SD nad priemerom
- ak sa presunieme na oblasť -2 SD až +2 SD, nájdeme tam približne 95 % hodnôt
- takmer všetky hodnoty (99,7 %) sa nachádzajú v rozmedzí -3 SD až +3 SD

Grafické zobrazenie tohoto pravidla nájdete v obrázku 3.1. Môžete si všimnúť, že priemer tejto distribúcie je 0, a potom sa tam nachádzajú nejaké iné znaky σ (malá sigma) – netreba sa báť, ide o takzvanú *populačnú štandardnú odchýlku*, no toto nás ďalej nemusí zaujímať, berme to ako SD. Vidíme, že medzi priemerom a 1 SD či už plus alebo mínus sa nachádza 34,1 % hodnôt, medzi 1 SD a 2 SD sa nachádza 13,6 % hodnôt, medzi 2 SD a 3 SD sa nachádza 2,1 % hodnôt.



Obrázok 3.1 Normálna distribúcia a oblasti štandardných odchýlok (zdroj: M. W. Toews - Own work, based (in concept) on figure by Jeremy Kemp, on 2005-02-09, CC BY 2.5, <https://commons.wikimedia.org/w/index.php?curid=1903871>)

Vďaka tomuto pravidlu si vieme prakticky predstaviť, ako sú dáta rozložené. Priemer ako deskriptívny ukazovateľ je dôležitý, no sám o sebe nedostačujúci. Predstavte si, že sme robili dva výskumy na samostatných súboroch (v každom výskumnom súbore boli jedineční respondenti). V oboch súboroch sa nám podarilo vyzbierať odpovede od 300 respondentov. Zamerali sme sa na všeobecnú dospelú populáciu – oslovovali sme dospelých ľudí s prosbou o zapojenie sa do výskumu. Chceme sa pozrieť na to, aký je priemerný vek prvom a druhom súbore (ako na to v programe jamovi si ukážeme v ďalšej časti). Úplnou náhodou zistíme, že v oboch súboroch je priemerný vek $M_{vek} = 41,35$ roka (veľmi nepravdepodobné, ale pre účely príkladu ideálne). Čo je však rozdielne, je práve štandardná odchýlka. V prvom súbore je to napríklad $SD = 2,70$ a v druhom $SD = 7,80$. Na základe štandardnej odchýlky vidíme, ako sú dáta vzdialené od priemeru. Zatiaľ čo v prvom súbore máme väčšinu respondentov v rozmedzí od 38,65 rokov ($M - 1 SD$) až 44,05 rokov ($M + 1 SD$), a takmer všetkých respondentov v rozsahu 35,95 rokov ($M - 2 SD$) až 46,75 rokov ($M + 2 SD$), v druhom súbore sú tieto rozpätia širšie, 33,55 až 49,15 roka v prípade jednej štandardnej odchýlky, a 25,75 až 56,95 roka v prípade dvoch štandardných odchýlok. Štandardná odchýlka je pre naše zhodnotenie pomerov premenných dôležitou deskriptívnou jednotkou. Na základe rovnakého alebo veľmi podobného priemeru by sme mohli usúdiť, že súbory sú v tomto prípade veľmi podobné, avšak na základe veľmi rozdielnej štandardnej odchýlky vidíme, že distribúcia veku je pomerne rozdielna v týchto dvoch súboroch. Túto informáciu majte na pamäti nielen v prípade práce na vlastnom výskume, ale aj pri rešerši výskumnej literatúry – zaujímajte sa o deskriptívu či už pri popise výskumného súboru ale aj pri meraných konštruktoch, získate tak lepšiu a kritickejšiu pohľad na výskum a možné závery z neho plynúce.

Medzikvartilové rozpätie

Medzikvartilové rozpätie, po anglicky *interquartile range*, v skratke *IQR*, je posledná miera variability, ktorej sa budeme venovať. Pojem rozpätie sme použili hneď na začiatku tejto podkapitoly, šlo o hodnotu, kde sme z maximálnej hodnoty odpočítali minimálnu hodnotu.

V prípade medzikvartilového rozpätia pôjde o podobný princíp, avšak ako už názov napovedá, pôjde o rozpätie medzi kvartilmi. Čo sú to tie kvartily a ako vypočítame medzikvartilové rozpätie? Kvartil je jeden z *koeficientov relatívneho umiestnenia*. Tieto koeficienty nám hovoria o relatívnej pozícii určitej hodnoty v usporiadanom rade hodnôt. Najčastejšie sa používajú percentily, decily a kvartily. S percentilmi a decilmi sa budete stretávať napríklad pri psychodiagnostike. Prostredníctvom psychodiagnostickej metodiky získate určitú hodnotu pre daného človeka. Ako je však na tom v porovnaní s inými ľuďmi? Vďaka normám k metodike, ktoré sú získané prostredníctvom testovania mnohých ľudí, získate odpoveď. Testovaný človek môže mať hodnotu, ktorá sa rovná povedzme 65 percentilu – to znamená, že 65 % ľudí na základe ktorých boli počítané normy, má rovnakú alebo nižšiu hodnotu ako je nami zistená hodnota a naopak 35 % ľudí má rovnakú alebo vyššiu hodnotu. Poďme naspäť ku kvartilom. Ľudskejšou rečou, kvartil je štvrtina. Predstavme si, že máme 100 hodnôt. Týchto 100 hodnôt usporiadame od najmenšej po najväčšiu a chceme ich rozdeliť na 4 rovnaké časti. Slovo rovnaké znamená, že tieto časti budú mať rovnaký počet hodnôt. V prípade 100 hodnôt je to jednoduché: $100 / 4 = 25$. Každá časť má obsahovať 25 hodnôt. Hodnotu 100 používame aj z dôvodu, že si to vieme jednoducho preniesť na percentá – každá časť má obsahovať 25 % hodnôt. Na to, aby sme rozdelili usporiadaný rad hodnôt na štyri rovnaké časti, potrebujeme 3 hodnoty. Číslo, ktoré vyjadruje prvý kvartil, vymedzuje prvých 25 % hodnôt, číslo ktoré vyjadruje druhý kvartil, vymedzuje prvých 50 % hodnôt. Ak ste pozorne čítali o mediáne, iste viete, že je to vlastne druhý kvartil. Posledné, tretie číslo, vymedzí prvých 75 % hodnôt, pričom nám zostane najvyšších 25 % hodnôt. Pre grafické zobrazenie si pozrite Obrázok 3.2.



Obrázok 3.2 Grafické zobrazenie medzikvartilového rozpätia

Ako môžete vidieť, medzikvartilové rozpätie vyjadruje rozpätie medzi prvým a tretím kvartilom. Ako ho teda vypočítať? Podobne ako rozpätie, avšak v tomto prípade nepoužijeme hodnoty maxima a minima, ale od hodnoty tretieho kvartilu (Q_3) odčítame hodnotu prvého (Q_1).

$$IQR = x_{Q3} - x_{Q1}$$

Tak ako štandardná odchýlka patrí k priemeru, medzikvartilové rozpätie používame ako mieru variability k mediánu. Vyjadruje nám rozpätie stredných 50 % hodnôt. No podobne ako pri variačnom rozpätí, ak chcete uviesť detailnejšiu deskriptívu, uveďte radšej hodnotu prvého a tretieho kvartilu.

3.3 Frekvenčná analýza

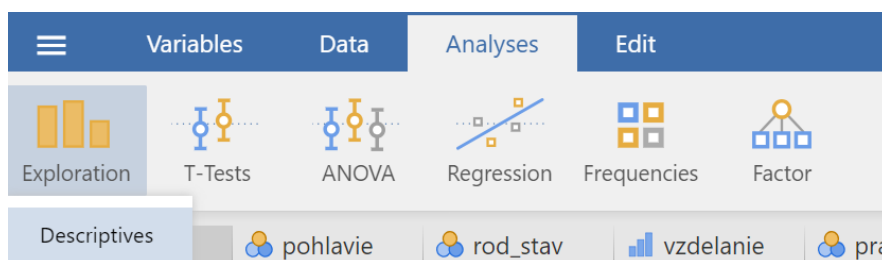
Frekvenčná analýza je pomerne jednoduchá záležitosť. Vďaka tejto analýze zisťujeme frekvenciu, či inak povedané početnosť výskytu určitého javu. Napríklad chceme vedieť, aký je počet mužov a žien v našom výskumnom súbore alebo koľko ľudí ukončilo určitý stupeň vzdelania. Frekvenčnú analýzu uplatňujeme najmä na nominálne alebo ordinálne premenné. Pri popise premenných meraných na kontinuálnej úrovni preferujeme skôr miery centrálnej

tendencie a variability, ktoré sú uvedené v predchádzajúcich častiach. Pri popise početnosti určitého javu používame buď to absolútnu hodnotu, napríklad v našom výskumnom súbore bolo 32 respondentov, ktorí uviedli ako najvyššie dosiahnuté vzdelanie „stredoškolské bez maturity“. Uvádzanie konkrétnych hodnôt patrí k štandardu, no rovnako je vhodné doplniť túto hodnotu aj percentuálnym podielom. S percentami sme sa už mnohokrát stretli a ich pochopenie je veľmi intuitívne. Čitateľ vďaka percentuálnemu vyjadreniu ľahko pochopí pomery v našich dátach. Nezabudnite preto uvádzať nielen konkrétne hodnoty ale aj ich percentuálny podiel.

3.4 Deskriptívna analýza v programe jamovi

V tejto časti si konkrétne ukážeme, ako spraviť deskriptívnu štatistiku v programe jamovi. Našťastie pre nás je to naozaj jednoduché. Domnievame sa, že najlepšie sa to naučíme na konkrétnom prípade. Deskriptívna štatistická analýza nachádza uplatnenie hneď na začiatku analýzy dát, keď sa chceme dozvedieť niečo o našom výskumnom súbore. Tieto informácie sú zaujímavé nielen pre nás, ale tvoria dôležitú časť vedeckých publikácií – nájdete ju v každej dobrej kvantitatívnej výskumnej práci a rovnako je tomu aj v prípade vašich záverečných prác. Poďme teda na to!

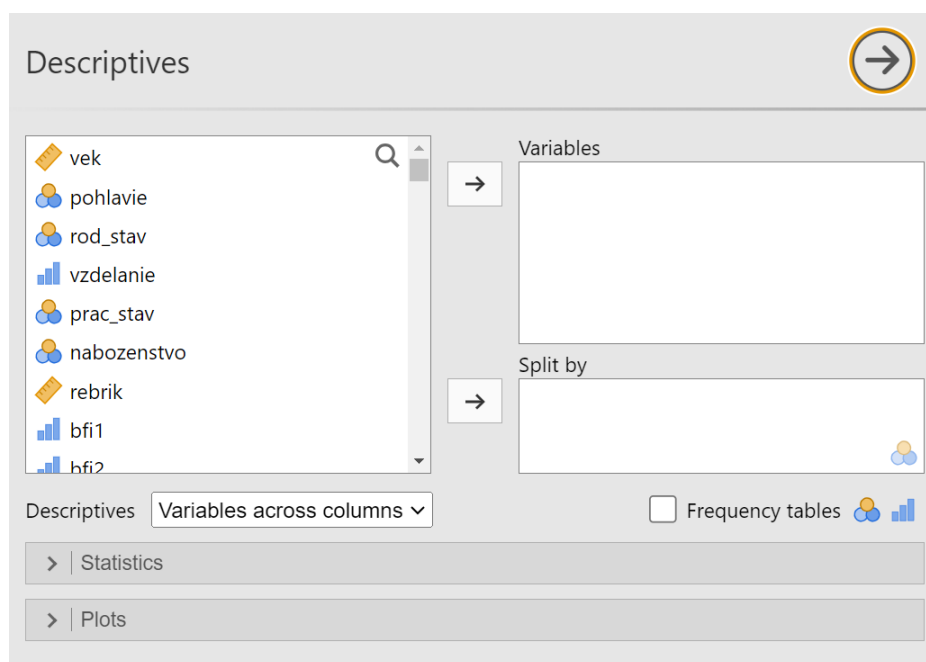
V prvom rade si v jamovi otvoríme kartu *Analyses*, kde nájdeme ikonu *Exploration*, klikneme a zvolíme *Descriptives* (Obrázok 3.3).



Obrázok 3.3 Umiestnenie ponuky deskriptívnej štatistickej analýzy

Po kliknutí sa zmení používateľské rozhranie, pričom pribudne nové okno, v ktorom sa bude odohrávať celé čaro deskriptívnej analýzy, všetko „pod jednou strechou.“ Rovnako tak nám v časti pre výsledky hneď pribudla tabuľka, ktorá neobsahuje žiadne výsledky. To dôležité z deskripcie si ukážeme v tejto časti, no vrátíme sa sem aj pri ďalších častiach tejto učebnice.

Ponuku pre deskriptívnu analýzu zobrazuje Obrázok 3.4. V prvom rade si môžeme všimnúť okno, v ktorom sa nachádzajú premenné v našom súbore. Hneď vedľa je okno s názvom *Variables* a pod ním je okienko s názvom *Split by*. Do okna *Variables* presúvame premenné, ktoré chceme analyzovať. Okienko *Split by* plní veľmi šikovnú funkciu – ak do neho dáme premennú, všetky výpočty, ktoré si nastavíme, sa budú robiť samostatne pre skupiny definované touto premennou. Táto funkcia je povolená len pre nominálne a ordinálne premenné. Nižšie si môžete všimnúť začiarokovacie políčko *Frequency tables* – vďaka nej nám bude vypočítaná frekvenčná analýza. Jamovi ponúka aj možnosť dvojakeho zobrazenia výsledkov pomocou výberového okna. V základe (*Variables across columns*) budú premenné v stĺpcoch a deskriptívne ukazatele v riadkoch, alebo naopak pri zvolení možnosti *Variables across rows*. Nižšie sa nachádzajú dve časti, *Statistics* a *Plots*.



Obrázok 3.4 Vzhľad okna deskriptívnej analýzy

V tejto podkapitole si ukážeme len časť *Statistics*, k *Plots* sa dostaneme neskôr. Keď si rozklikneme túto časť, otvorí sa nám množstvo možností, ktoré môžeme označiť, pričom viaceré sú už v základe zvolené (Obrázok 3.5). Priblížime, čo ktorá možnosť znamená:

- **Sample Size** – informácie o našom súbore, respektíve o respondentoch, ktorí spĺňajú kritéria nastaveného filtrovania. *N* vyjadruje celkový počet, *Missing* znamená, že sa nám zobrazí počet respondentov, ktorým chýba hodnota v analyzovanej premennej.
- **Percentile Values** – výpočet hodnôt určitého percentilu alebo hodnôt, ktoré rozdeľujú usporiadaný rad hodnôt nejakej premennej. Ak by sme zvolili, *Cut points for 4 equal groups*, získali by sme hodnoty kvartilov. Ak by sme zvolili, *Percentiles*, získame hodnoty konkrétnych percentilov.
- **Central Tendency** – tu si nastavujeme miery centrálnej tendencie, v základe sú zvolené *Mean* (priemer) a *Median* (medián). Na výber máme aj *Mode* (modus) alebo *Sum* (súčet všetkých hodnôt premennej), ten však pre naše účely nepoužívame.
- **Dispersion** – tu si nastavujeme miery variability, inak povedané disperzie. V základe je označené prakticky všetko, čo potrebujeme – *Std. deviation* (štandardná odchýlka), *Minimum* (najnižšia hodnota), *Maximum* (najvyššia hodnota). Zaujímať nás môže aj *IQR* (medzikvartilové rozpätie). Ďalšie možnosti pre naše účely nepotrebujeme.
- **Distribution a Normality** – výpočty vhodné pre posúdenie distribúcie hodnôt premennej, viac si k tomu povieme pri časti o zhodnotení normality dát, v časti o inferenčnej štatistike.
- **Mean Dispersion** – výpočet doplnujúcich informácií k priemeru, ktoré pre naše účely nepotrebujeme.
- **Outliers** – jamovi nám uvedie čísla riadkov najnižších a najvyšších hodnôt premennej. Táto možnosť môže byť užitočná pre rýchle zhodnotenie takzvaných vzdialených hodnôt – tejto problematike sa venujeme v kapitole 5.4.

Obrázok 3.5 Nastavenia v časti Statistics

Po presunutí niektorej z premenných do okienka *Variables* sa nám vo vedľajšom okne jamovi objavia výsledky deskriptívnej alebo frekvenčnej analýzy. Ak by sme chceli bližšie popísať náš súbor, môžeme sa zamerať základné socio-demografické premenné. Najskôr môžeme charakterizovať náš súbor z hľadiska celkového počtu respondentov a ich priemerného veku. Do okienka *Variables* si vložíme premennú vek, výsledok je zobrazený v Obrázku 3.6. Vidíme tabuľku, v ktorej sú deskriptívne ukazatele písané v riadkoch (*Variables across columns*).

Descriptives	
	vek
N	281
Missing	0
Mean	42.43
Median	41
Standard deviation	15.49
Minimum	18
Maximum	72

Obrázok 3.6 Výsledok deskriptívnej analýzy v programe jamovi

Ak by sme chceli informáciu o veku uviesť špecificky pre mužov a ženy, pridáme do okienka *Split by* aj premennú pohlavie. Výsledok je zobrazený v Obrázku 3.7. Ako vidíte, výsledky sú samostatne v riadkoch pre mužov a ženy. Takýmto spôsobom to funguje či máte dve alebo viac skupín (kategórií premennej).

Descriptives		
	pohlavie	vek
N	muž	138
	žena	143
Missing	muž	0
	žena	0
Mean	muž	42.59
	žena	42.28
Median	muž	41.00
	žena	41
Standard deviation	muž	15.70
	žena	15.34
Minimum	muž	18
	žena	18
Maximum	muž	72
	žena	70

Obrázok 3.7 Výsledok deskriptívnej analýzy s použitím funkcie Split by v programe jamovi

Pre zistenie počtosti jednotlivých kategórií premennej, využijeme funkciu *Frequency tables*. Túto funkciu však používame len pre nominálne alebo ordinálne premenné. V Obrázku 3.8 zobrazujeme výsledok frekvenčnej analýzy pre premenné pohlavie a najvyššie dosiahnuté vzdelanie. Pre účely získania detailnejších hodnôt pre percentá, sme v nastavení zmenili počet desatinných miest na tri. V tejto tabuľke vidíme stĺpec s názvom *Levels*, v ktorý nesie informáciu o jednotlivých hodnotách (kategóriách) premennej. Keďže sme si tieto hodnoty v dátovom súbore pomenovali, vidíme ich názvy. V druhom stĺpci s názvom *Counts* sa nachádza informácia o absolútnom počte. Tretí stĺpec *% of Total* vyjadruje percentuálny pomer jednotlivých kategórií. Posledný stĺpec *Cumulative %* vyjadruje takzvané kumulatívne percento. Uplatnenie nachádza pri ordinálnych premenných, kde percentuálna hodnota vyjadruje, koľko respondentov označilo danú odpoveď alebo nižšiu. V rámci nášho príkladu môžeme vidieť, že stredoškolské vzdelanie s maturitou alebo nižšie označilo 82,2 % respondentov ($7,1 + 34,5 + 40,6 = 82,2$).

Frequencies of pohlavie			
Levels	Counts	% of Total	Cumulative %
muž	138	49.1 %	49.1 %
žena	143	50.9 %	100.0 %

Frequencies of vzdelanie			
Levels	Counts	% of Total	Cumulative %
ZŠ	20	7.1 %	7.1 %
SŠ bez maturity	97	34.5 %	41.6 %
SŠ s maturitou	114	40.6 %	82.2 %
VŠ I. stupeň	8	2.8 %	85.1 %
VŠ II. stupeň	42	14.9 %	100.0 %

Obrázok 3.8 Výsledok frekvenčnej analýzy premenných pohlavie a vzdelanie

Takýmto spôsobom sa dokážeme dopracovať k rôznym informáciám o premenných. Nižšie uvádzame príklad zápisu informácií o výskumnom súbore:

Náš výskumný súbor tvorilo 281 respondentov vo veku od 18 do 72 rokov ($M = 42,34$; $SD = 15,49$). Priemerný vek mužov ($N = 138$) bol $42,59$ ($SD = 15,70$) a žien ($N = 143$) bol $42,28$ ($SD = 15,34$) roka. Väčšina respondentov uviedla ako najvyššie dosiahnuté vzdelanie stredoškolské s maturitou ($N = 114$; 40,6 %), ďalej stredoškolské bez maturity ($N = 97$; 34,5 %), vysokoškolské II. stupňa ($N = 42$; 14,9 %), základné ($N = 20$; 7,1 %) a vysokoškolské I. stupňa ($N = 8$; 4,2 %). Z hľadiska rodinného stavu, väčšina respondentov bola v manželskom zväzku (48,8 %), za ktorými nasledovali slobodní respondenti (23,1 %), rozvedení (12,5 %), respondenti vo vzťahu (10,7 %) a ovdovení respondenti (5 %).

Samozrejme, zápis a množstvo informácií, ktoré uvediete závisí od toho, čo je pre vás dôležité, no zaužívané je uvedenie počtu respondentov, mužov a žien; priemeru a štandardnej odchýlky veku respondentov; frekvenčnej analýzy najvyššieho dosiahnutého vzdelania. V prípade, že chcete uviesť frekvenčnú analýzu premennej, ktorá má viacero kategórií, napr. 4 a viac, môžete využiť aj tabuľku. Niekedy sa pre grafické zobrazenie frekvencie využíva napr. koláčový/kruhový graf, jeho využitie však zvážte – naozaj je prínosné, ak pridám graf, nestačí informácia v texte alebo tabuľke? Osobne tieto grafy nepovažujeme za potrebné a preferujeme zápis v texte alebo tabuľke. Dajte si tiež pozor, ak uvediete číselnú informáciu v tabuľke, neuvádzajte ju duplicitne v texte (ak by sme napríklad prezentovali výsledok frekvenčnej analýzy premennej rodinný stav v tabuľke, neuviedli by sme ho v texte). Ako ste si mohli v príklade všimnúť, štatistické skratky a symboly uvádzame *kurzívou*.

4. Vnútna konzistencia

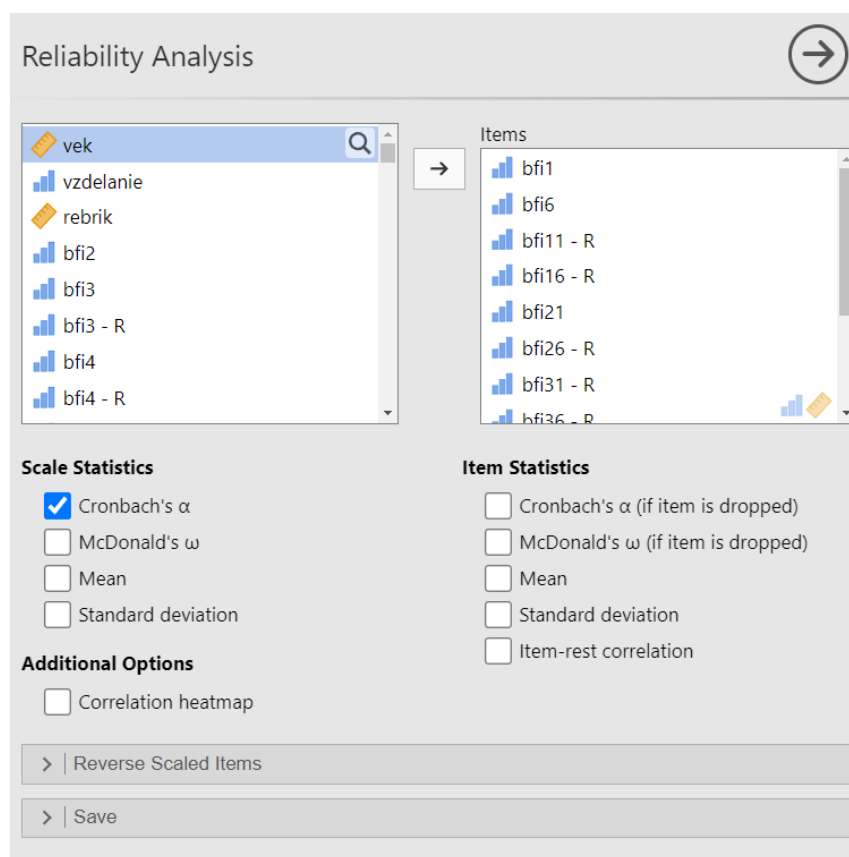
Meranie v psychológii prebieha prevažne prostredníctvom rôznych testov, dotazníkov alebo škál. Takýto spôsob merania je nutne spojený s určitou chybou merania, ktorá je spojená s presnosťou psychodiagnostických nástrojov. Posudzovaním toho, ako dobre tieto nástroje fungujú, sa zaoberá samostatná vedná disciplína *psychometria*. Psychometrické analýzy a postupy sú nad rámec tejto učebnice. Pre záujemcov odporúčame publikáciu ohľadom princípov psychodiagnostiky od Halamu (2011). Jedna z vecí, ktorá sa však bežne testuje a nájdete ju aj vo výskumných prácach, je *reliabilita* prípadne konkrétnejšie *vnútorná konzistencia*. Ak si vo výskumných štúdiách pozriete časť, v ktorej autori popisujú použité metodiky, stretnete sa s hodnotením reliability, vnútornej konzistencie (angl. *internal consistency*) alebo konkrétnymi pojmami ako Cronbachova alfa, α . V tejto časti si priblížime, čo to tá vnútorná konzistencia je a ako ju vypočítať v programe jamovi.

Psychodiagnostické metodiky sú určené na meranie určitej konkrétnej (širšej či užšej) oblasti. Jednotlivé položky by teda rovnako mali merať túto oblasť. Či ju aj naozaj merajú, je otázkou *validity* psychodiagnostických nástrojov. Reliabilita nám na toto neposkytne odpoveď, no dozvieme sa, ako *súdržne* teda *konzistentne* tieto položky merajú. Je viacero spôsobov, ktorými sa reliabilita posudzuje – napríklad zisťujeme vzťah medzi výsledkami dvoch meraní v rôznom čase (test-retest reliabilita) alebo výsledkami paralelných metód. Keďže bežne nemáme tieto informácie dostupné, používame koeficienty vnútornej konzistencie, najčastejšie už spomenutú Cronbachovu alfu (Cronbach, 1951). Veľmi jednoducho povedané, koeficienty vnútornej konzistencie vyjadrujú mieru toho, ako silno spolu položky súvisia, aký silný je medzi položkami vzťah. Súvis medzi položkami je dôležitá vlastnosť pri metodikách, kde predpokladáme, že odpovede respondentov sú spôsobené meraným konštruktom na pozadí. Ak je vzťah medzi položkami nízky, je otázne, či je odpoveď na položky daná tým istým konštruktom, alebo ide o viacero konštruktov.

Výpočet koeficientov vnútornej konzistencie je v jamovi jednoduchý. V karte *Analyses* zvolíme možnosť *Factor* a vyberieme *Reliability Analysis*. Tak ako v prípade deskriptívnej analýzy sa nám zobrazí stredové okno, v ktorom nastavíme analýzu. Podobne ako pri iných analýzach máme okienko, v ktorom máme na výber položky v našom súbore. Vyberieme položky, pri ktorých chceme zisťovať vnútornú konzistenciu a presunieme ich do okienka *Items*. Automaticky sa nám vo výsledkovej časti zobrazia výsledky. Pri výbere položiek si dajte pozor na to, čo robíte. Položky vyberajte tak, ako sú radené v metodike, napr. podľa vyhodnocovacieho kľúča – vkladajte tie položky, s ktorými ste počítali hrubé alebo priemerné skóre. Ak máte metodiku, ktorá má viacero samostatných faktorov, pri ktorých sa nedá vypočítať spoločný faktor, nedávajte všetky položky naraz! Príkladom môže byť v úvodných častiach spomenutý Inventár Veľkej Päťky 2. Tento osobnostný inventár má 5 samostatných faktorov: extravézia, prívetivosť, svedomitosť, negatívna emocionalita a otvorenosť. Ak chceme zistiť vnútornú konzistenciu, musíme spraviť 5 samostatných analýz, pre každý faktor zvlášť. Nič také ako vnútorná konzistencia Inventára Veľkej Päťky 2 nie je, t.j. nemôžete dať analyzovať všetky položky naraz. V prípade, že má metodika aj reverzne skórované položky, vkladajte ich transformovanú podobu. Ak by ste sa pomýlili a vložili reverznú položku, jamovi vás na to upozorní, keď zistí negatívny vzťah medzi určitou položkou a zvyšnými položkami. Avšak, tak ako pri všetkom ostatnom, dajte si čas a venujte nastaveniu pozornosť. Pod okienkami s položkami si môžete všimnúť začiarokovacie políčka. V základom nastavení je zvolené len *Cronbach's α* . V podstate je to všetko, čo potrebujeme. Obrázok 4.1 zobrazuje

nastavenie, v ktorom zisťujeme vnútornú konzistenciu domény Extraverzia. Do analýzy sme vložili už vopred transformované reverzné položky.

Na výber máme McDonaldovu omegu ako druhý koeficient vnútornej konzistencie. Pre naše účely ale bežne postačuje Cronbachova alfa. Priamo v nastavení tejto analýzy si môžeme dať vypočítať priemer a štandardnú odchýlku, štatistiku pre položky, či vykresliť graf korelácií medzi položkami – tieto dopĺňajúce výsledky bežne nepotrebujeme. Ako je na Obrázku 4.1 vidieť, nižšie sa nachádza záložka s názvom *Reverse Scaled Items*. Slúži na to, aby sme určili, ktoré z vybraných položiek sú reverzne skórované a majú byť transformované. Znovu pozor! Ak sme vybrali už raz transformované položky, *nenastavujeme* ich opätovnú transformáciu – vrátili by sme ich tým do pôvodného (surového) stavu. Posledná záložka *Save* nám umožňuje uložiť sumárne (hrubé) alebo priemerné skóre. Nevýhodou je, že dané skóre je viazané na analýzu, a ak ju odstránime, odstráni sa aj vypočítaná premenná z dát. Dá sa to obísť tak, že vypočítanú premennú skopírujeme do prázdneho stĺpca, no tento postup hodnotíme ako zbytočne náročný, preto odporúčame držať sa postupu uvedeného v kapitole 2.5.



Obrázok 4.1 Nastavenie výpočtu Cronbachovej alfy

Výsledok pri základnom nastavení vyzerá úplne jednoducho (Obrázok 4.2).

Scale Reliability Statistics	
Cronbach's α	
scale	0.76

Obrázok 4.2 Výsledok analýzy vnútornej konzistencie pri predvolenom nastavení

Získali sme číslo 0,76 – ako však túto hodnotu interpretovať? Odpoveď nie je úplne jednoduchá. Hodnota Cronbachovej alfy je závislá od vzájomných vzťahov medzi položkami (čím silnejšie, tým lepšie) ale aj od počtu položiek – to znamená, že v prípade metodík s nízkym počtom položiek (napr. 4) môžeme získať nízku hodnotu napriek uspokojivým vzťahom medzi položkami. Naopak, v prípade veľkého počtu položiek môžeme získať vysokú hodnotu, napriek tomu, že medzi niektorými položkami je veľmi slabý vzťah. Interpretácia tiež závisí od toho, čo daná metodika meria. Pokiaľ ide o veľmi úzko vymedzenú oblasť, očakávame „vyššiu“ hodnotu alfy, v prípade širších oblastí, napr. osobnostné charakteristiky, nám postačuje aj „nižšia“ hodnota. Čo je to tá vyššia hodnota a čo nižšia? Vo všeobecnosti považujeme hodnoty Cronbachovej alfy väčšie ako 0,7 za dobré, 0,8 za veľmi dobré a 0,9 za výborné. Hodnoty pod 0,7 sa pri metodikách s nízkym počtom položiek niekedy považujú za dostačujúce, pokiaľ nie sú nižšie ako 0,6. Ide však len veľmi hrubé rozdelenie, ktoré na základe vyššie spomenutého ťažko uplatniť na každú oblasť psychológie. Niekedy autori používajú vyjadrenie *akceptovateľná*, prípadne *vyhovujúca* vnútorná konzistencia, na základe ich zhodnotenia faktorov vstupujúcich do výpočtu, t.j. počet položiek a šírka meranej oblasti. V našom prípade by sme rovnako zhodnotili vnútornú konzistenciu domény Extraverzia za vyhovujúcu.

Zhodnotenie vnútornej konzistencie uvádzame najčastejšie pri popise použitých metodík. Napríklad:

Pre meranie piatich osobnostných faktorov sme použili Inventár Veľkej Päťky 2 od autorov Halama et al. (2020). Tento inventár obsahuje celkovo 60 položiek, 12 pre každú doménu: Extraverzia, Prívetivosť, Svedomitosť, Negatívna emocionalita a Otvorenosť. Inventár má spoločný základ „*Som niekto, kto...*“ a pokračuje jednotlivými položkami, ktoré sú vo forme krátkych výrokov, napr. „*je spoločenský, rád trávi čas s inými ľuďmi.*“ alebo „*sa niekedy správa nezodpovedne.*“ Respondenti vyjadrujú mieru súhlasu s týmito položkami prostredníctvom 5-bodovej Likertovej škály, od „*Veľmi nesúhlasím*“ po „*Veľmi súhlasím*“. Vnútorná konzistencia jednotlivých domén je vyhovujúca: Extraverzia $\alpha = 0,76$; Prívetivosť $\alpha = 0,78$; Svedomitosť $\alpha = 0,84$; Negatívna emocionalita $\alpha = 0,83$ a Otvorenosť $\alpha = 0,80$.

Čo robiť v prípade, že nami použitá metodika vykazuje slabú vnútornú konzistenciu ($\alpha < 0,60$)? Odpoveď závisí od toho, čo sme použili. Pokiaľ sme použili metodiku, ktorá je adaptovaná pre použitie u nás, a autori slovenskej adaptácie uvádzajú akceptovateľnú vnútornú konzistenciu, môžeme sa zamyslieť, či rozdiel môže spočívať v špecifikách nášho súboru, resp. populácie, na ktorú sme zamerali náš výskum. V takom prípade je vhodná určitá polemika v rámci diskusie zistení a najmä limitácií výskumu. Ak sme použili metodiku, ktorá nie je na Slovensku adaptovaná alebo sme si vytvorili vlastnú, môžeme sa pozrieť, či sú vzťahy medzi položkami slabé vo všeobecnosti, alebo je nejaká položka prípadne položky, ktorých prítomnosť výrazne znižuje vnútornú konzistenciu. Príčinou môže byť zlý preklad či nevhodná formulácia. V jamovi si v časti *Item Statistics* zvolíme možnosť *Cronbach's alpha (if item is dropped)*. Táto funkcia nám dá informáciu o tom, aká by bola Cronbachova alfa, keby vynecháme určitú položku. Vďaka tomu zistíme, či by odstránenie niektorej položky spôsobilo zvýšenie vnútornej konzistencie na prijateľnú úroveň. Pokiaľ by sme takú identifikovali, môžeme zvážiť jej nepoužitie pre výpočet hrubého alebo priemerného skóre (jej odstránenie). Pokiaľ však takúto položku neidentifikujeme, musíme sa zmieriť s nízkou vnútornou konzistenciou a *hrať s tým, čo máme* – pokračovať, no pri diskusii výsledkov zohľadniť tento fakt. V Obrázku 4.3 je zobrazený výsledok analýzy. Pri tejto 4 položkovej metodike sme zistili vzhľadom na nízky

počet položiek veľmi dobrú vnútornú konzistenciu $\alpha = 0,77$. Výsledky doplnujúcej analýzy ukázali, že odstránenie ktorejkoľvek z prvých troch položiek by viedlo k zníženiu α – hovoríme, že tieto položky práve prispievajú vnútornej konzistencii. Odstránenie poslednej položky by viedlo k pomerne výraznému zvýšeniu α . V prípade, že by sme mali slabú vnútornú konzistenciu, stálo by za zváženie túto položku odstrániť, no v našom prípade to nie je potrebné, keďže v základe je naša α vyhovujúca.

Item Reliability Statistics	
	If item dropped
	Cronbach's α
shs_1	0.63
shs_2	0.62
shs_3	0.71
shs_4 - R	0.85

Obrázok 4.3 Zmena Cronbachovej alfy v prípade odstránenia niektorej položky.

5. Inferenčná štatistika

Zatiaľ sme si ukázali základy práce s dátovým súborom, základy práce v programe jamovi, základy deskriptívnej analýzy a povedali sme si niečo o vnútornej konzistencii. V nasledujúcich častiach však príde to „pravé čaro“ analýzy psychologických dát. V psychologických kvantitatívnych výskumoch mnohokrát očakávame súvis určitých premenných. Pojem súvis je veľmi široký, no práve preto vystihuje širokú paletu prípadov. Napríklad predpokladáme, že existuje vzťah medzi vekom a extraverciou alebo rozdiel v negatívnej emocionalite u mužov a žien, alebo efekt najvyššieho dosiahnutého vzdelania na schopnosť kritického myslenia. Aj keď ide o rôzne analýzy a rôzne pomenovania, na pozadí je stále hľadanie a potvrdzovanie súvisu medzi premennými. Naše očakávanie nás vedie k túžbe skúmať danú oblasť a overiť naše predpoklady. Tieto sa týkajú určitej populácie, teda celej množiny ľudí, ktorí spĺňajú nejaké podmienky, napr. populácia dospelých ľudí, prípadne užšie definované množiny, napr. populácia pracovníkov záchrannej služby, deti v predškolskom veku a iné. Keďže populácia obsahuje veľké množstvo ľudí a my nemáme možnosť zapojiť do výskumu každého jedného človeka, robíme výber, na základe ktorého vyvodzujeme závery pre celú cieľovú populáciu. Tento proces sa nazýva štatistická *inferencia*, z čoho pramení názov *inferenčná štatistika*. Našou snahou je teda nie len skúmať pomery v dátach, ale zistiť, či môžeme zistené pomery zovšeobecniť na celú populáciu. Týmto sa dostávame ku konceptu *testovania signifikancie nulovej hypotézy*, v skratke NHST z anglického *null hypothesis significance testing*.

5.1 Testovanie signifikancie nulovej hypotézy

Vedieť, či môžeme naše zistenie zovšeobecniť na cieľovú populáciu znie fajn a ak ste niekedy aspoň mierne prebehli výsledky v kvantitatívnych výskumných štúdiách, určite ste sa stretli s pojmami ako signifikancia, štatistická významnosť alebo anglicky *significance*. Práve tieto pojmy sú u začínajúcich študentov psychológie veľmi honosne vyslovované pri prezentáciách výskumných štúdií, ktoré spracovali, hoci pravdepodobne (skoro) nikto nevie, čo to vlastne znamená. Slovo významné alebo signifikantné v nás automaticky evokuje pocit, že je to niečo „významné“ teda dôležité. V mnohých prípadoch tomu tak je, v mnohých nie. Na to, aby sme vedeli, čo to znamená významnosť, musíme si najskôr priblížiť, čo je to nulová hypotéza a na čom je postavené testovanie jej signifikancie.

Nulová hypotéza je ústredným pojmom v rámci NHST. Je to základný predpoklad, ktorý je vopred platný, pokiaľ dôkazy neukážu inak. Ako sme už spomínali, v našom výskume máme nejaké predpoklady, ktoré chceme overiť – tzv. výskumné hypotézy, nazývané aj meritórne alebo alternatívne hypotézy. Väčšinou sa tieto hypotézy týkajú *prítomnosti súvisu* medzi premennými. Nulová hypotéza má znenie v súvislosti s tým, aký test robíme ako aj s tým, ako znie naša výskumná hypotéza, takže ťažko hovoriť za všetky prípady, no zjednodušene, nulová hypotéza hovorí o *neprítomnosti javu* v populácii. Ak predpokladáme, že v populácii *existuje* vzťah medzi premennými, nulová hypotéza hovorí, že *neexistuje*. Ak hovoríme, že sa dve skupiny (napr. muži a ženy) *líšia* v miere nejakej vlastnosti, nulová hypotéza tvrdí, že sa *nelíšia*. Keďže je nulová hypotéza a priori platná, my sme tí, ktorí musia dokázať náš predpoklad. Z tohto dôvodu sa naše výskumné hypotézy nazývajú aj alternatívnymi – sú alternatívou voči nulovej hypotéze. Výnimkou sú výskumné hypotézy, ktoré sa rovnajú nulovej hypotéze – napríklad predpoklad neprítomnosti rozdielu medzi skupinami. Ak by sme chceli byť veľmi kritický, takéto hypotézy sú problematické v tom, že nulová hypotéza je vopred platná, na čo ju teda potvrdzovať? Ešte prísnejšie povedané, NHST nám konceptovo neumožňuje overiť

platnosť nulovej hypotézy, predpoklad o jej platnosti je zadefinovaný v rámci tohto konceptu (Field, 2018). Napriek tomu sa s takýmito výskumnými hypotézami môžeme stretnúť.

Každá štatistická analýza, ktorú robíme má určitý číselný výsledok, ktorý nazývame *koeficient*. Nulová hypotéza pri mnohých testoch hovorí o tom, že tento koeficient sa v populácii rovná 0. Vďaka štatistickej analýze sa v našom výskume dopracujeme k nejakej hodnote koeficientu, ktorý sa nerovná nule. Toto zistenie je však platné len pre náš výskumný súbor a nehovorí o pomeroch v populácii – tam predsa a priori platí nulová hypotéza. V tomto momente vstupuje do hry štatistická významnosť, signifikancia. Veľmi zjednodušene ale výstižne povedané, vďaka nej vieme, *aká je pravdepodobnosť získať v našom výskumnom súbore/výbere takúto hodnotu koeficientu ak v populácii platí nulová hypotéza*. Táto definícia v sebe nesie dva kľúčové aspekty: výskumný súbor/výber a hodnotu koeficientu. Signifikancia výsledku je teda závislá od veľkosti koeficientu a počtu ľudí, na ktorých bola táto analýza vykonaná. Počet ľudí, na ktorých bola analýza vykonaná môžeme chápať ako nejakú váhu dôkazov. Ak by sme mali v dvoch prípadoch rovnakú hodnotu koeficientu ale výrazne iný počet ľudí, hodnota signifikancie by bola iná. A zas, ak by sme mali v dvoch prípadoch rovnaký počet ľudí, ale rozdielny koeficient, hodnota signifikancie by sa líšila.

V rámci definovania štatistickej významnosti považujeme za dôležité upozorniť na dva *nesprávne* spôsoby chápania signifikancie, ktorými sa niekedy implicitne poníma. Štatistická významnosť výsledku nám v základe *nehovorí* o nejakej praktickej významnosti výsledku – predstavme si výskum, do ktorého sa zapojil veľký počet mužov a žien, napr. $N > 300$ z každej skupiny. V tomto výskume sme participantom ponúkli parené buchty (bol dostatok pre každého) a zistili sme, že muži zjedli v priemere 4,3 parenej buchty, zatiaľ čo ženy 4,2 parenej buchty. Prostredníctvom štatistickej analýzy, ktorú priblížime v inej časti učebnice, by sme zistili, že muži zjedli o jednu desatinu buchty viac (povedzme, že jedno naozaj malé sústo). Je tento rozdiel veľký, dôležitý, prakticky významný? Asi ani nie, no napriek tomu by nám mohol vyjsť ako štatisticky významný. Štatistická významnosť nám rovnako nehovorí o tom, aká je pravdepodobnosť získať daný výsledok v populácii – t.j. *nevieme* povedať, ako pravdepodobné je, že v populácii naozaj muži zjedli v priemere o 1/10 parenej buchty viac. Vďaka štatistickej významnosti vieme v podstate len to, ako veľmi pravdepodobné je pri počte našich respondentov získať hodnotu koeficientu, ktorú sme získali, ak v populácii platí nulová hypotéza. Pokiaľ je tá pravdepodobnosť *veľmi nízka*, naše uvažovanie pokračuje ďalej – určite platí nulová hypotéza, ak je takto nízka pravdepodobnosť získať takúto hodnotu koeficientu? V rámci takéhoto uvažovania sme pomerne asertívny a nepochybujeme o sebe, ale o platnosti nulovej hypotézy. Hovoríme, že nulovú hypotézu *zamietame*, no možno lepšie povedané, *sponchybnujeme* jej platnosť a prijímame našu alternatívnu hypotézu, v ktorej sme predpokladali súvis medzi premennými. Naopak, ak je vysoká pravdepodobnosť získať nami zistenú nenulovú hodnotu koeficientu pri platnosti nulovej hypotézy, nepochybujeme o jej platnosti a zamietame našu alternatívnu hypotézu, v ktorej sme predpokladali súvis medzi premennými.

Čo znamená malá pravdepodobnosť? Koľko je to? Takáto hodnota sa volá *hladina významnosti alfa*. Konvenčne zaužívanou hodnotou je 5 %, ale používajú sa aj prísnejšie hodnoty ako 1 % či 0,1 %. Určite ste sa už niekedy stretli s tabuľkou, v ktorej bola informácia $p < 0,05$, alebo $p < 0,01$, alebo $p < 0,001$, čo znamená, že výsledky boli signifikantné na úrovni 5 %, 1 % alebo 0,1 %, a teda nasvedčovali, že autori výskumu pochybovali o platnosti nulovej hypotézy a prijímali stanovené alternatívne hypotézy (vlastné predpoklady).

5.2 Alternatívne hypotézy

V predchádzajúcej časti sme spomenuli alternatívne hypotézy. Ide o hypotézy, ktoré hovoria o prítomnosti niečoho a sú v protismere nulovej hypotézy. Pri každej jednej analýze si ukážeme, akým spôsobom môžeme formulovať hypotézy, preto v tejto časti len krátko spomenieme dva druhy alternatívnych hypotéz: dvojstranná/dvojsmerná a jednostranná/jednosmerná hypotéza. Dvojstranná/dvojsmerná hypotéza (*angl. two-tailed hypothesis*) je hypotéza, v ktorej len vyjadrujeme predpoklad existencie súvisu medzi premennými – napríklad predpokladáme, že existuje korelácia medzi premennými alebo že sa porovnávané skupiny líšia. Pri jednostrannej/jednosmernej hypotéze (*angl. one-tailed hypothesis*) nepredpokladáme len súvis medzi premennými, no bližšie ho špecifikujeme – napríklad predpokladáme pozitívnu koreláciu medzi premennými alebo očakávame, ktorá z porovnávaných skupín má vyššiu hodnotu. Rozdiel medzi týmito hypotézami tkvie najmä vo výpočte štatistickej významnosti. Vráťme sa k príkladu ohľadom porovnania skonzumovaných parených buchiet u mužov a žien. Ak predpokladáme, že je rozdiel v množstve zjedených buchiet medzi mužmi a ženami, stanovili sme si dvojstrannú hypotézu. Dá sa povedať, že to hráme na obe strany – je nám „jedno“ (nie naozaj), či sú to muži alebo ženy, kto zje viac buchiet. Nulová hypotéza by tvrdila, že nie je rozdiel. Ak bude rozdiel pri našom počte participantov dostatočne veľký, získame hodnotu štatistickej významnosti, ktorá je nižšia ako alfa (konvenčne 5 %), čím pochybujeme o platnosti nulovej hypotézy, a teda prijímame našu alternatívnu hypotézu. V takomto prípade však rozdeľujeme nami zvolených 5 % na obe strany, a preto výsledok bude signifikantný, iba ak ide o hodnoty koeficientu menej pravdepodobné ako 2,5 % z každej strany. Ak by sme stanovili jednosmernú hypotézu, napr. že muži zjedia viac ako ženy, nulová hypotéza by stále tvrdila opak, v tomto prípade by to ale bolo, že „muži nezjedia viac“ – čo zahŕňa aj prípad, keď ženy zjedia viac. Ak by sme zistili, že ženy zjedia o dosť viac ako muži (je tam rozdiel), štatistická významnosť bude počítaná s touto otázkou na pozadí: aká je pravdepodobnosť získať takýto výsledok v našom súbore, ak v populácii nemá byť rozdiel alebo ženy zjedia viac? Ak ste pozorne čítali, pravdepodobnosť bude veľmi vysoká (blížiaca sa 100 %), keďže nulová hypotéza zahŕňa aj tento jav. Z hľadiska NHST by sme tvrdili, že nepochybujeme o nulovej hypotéze. Ale veď ten rozdiel tam je, nie? Áno je, no v opačnom smere ako sme očakávali, preto sa tým ďalej nezaobráme. Ak by sme mali stanovenú dvojsmernú hypotézu, mohli sme potenciálne získať signifikantný výsledok. Z nášho pohľadu je to veľmi tvrdé zatváranie si očí pred niečím. Áno, musíme zamietnuť našu alternatívnu hypotézu, ale aspoň v diskusii sa môžeme vyjadriť, že tento rozdiel medzi skupinami (či všeobecne súvis medzi premennými) tam je, len v opačnom smere. Z tohto dôvodu niektorí autori (napr. Field, 2018) pri pojednaní o tejto problematike navrhujú používanie dvojsmerných hypotéz. Výhodou jednosmernej hypotézy je, že celú našu hladinu významnosti venujeme len na jednu stranu, a teda pokiaľ výsledok bude v smere nášho očakávania, môže byť signifikantný aj pri nižšej hodnote koeficientu, keďže nás zaujíma 5 % najextrémnejších hodnôt v porovnaní s 2,5 % pri dvojsmernej hypotéze. Ktorý typ alternatívnej hypotézy si vyberiete, je na vás. Pri koreláciách si ukážeme, že overenie jednosmernej hypotézy je v podstate rovnaké ako dvojsmernej, no líši sa výpočet štatistickej významnosti výsledku. Detailnejší popis nájdete v kapitole 6.1.

5.3 Bivariačná a multivariačná štatistická analýza

Existuje veľké množstvo spôsobov analýzy dát, ktoré pokrývajú širokú paletu výskumných cieľov zameraných na zistenie súvisu medzi premennými. V základe ich rozdeľujeme na dve kategórie – *bivariačné* a *multivariačné* analýzy. Bivariačná analýza pracuje s dvomi premennými – zisťujeme vzájomný súvis dvoch premenných. Ako už názov napovedá,

multivariačná analýza rieši vzájomný súvis troch a viacerých premenných. Výhodou bivariačnej analýzy je jej jednoduchosť – pre výpočet ako aj pre porozumenie. Jednoduchosť však nie je na škodu a v mnohých výskumoch bivariačná analýza sprevádza výsledky multivariačnej analýzy, ktorá je zložitejšia pre výpočet ako aj porozumenie, no dokáže komplexnejšie uchopiť skúmanú problematiku. Príkladom multivariačnej analýzy môže byť analýza kovariancie, viacnásobná regresná analýza, exploratórna faktorová analýza a množstvo iných... V tejto učebnici sa budeme venovať len bivariačnej analýze a to konkrétne:

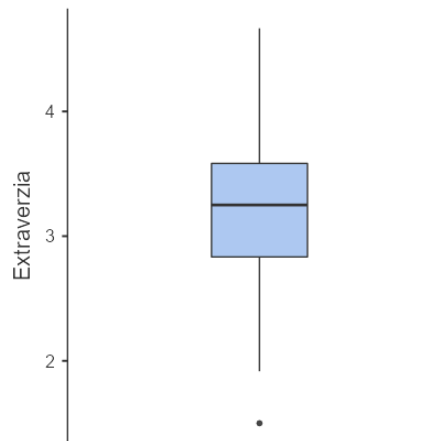
- vzťah medzi dvomi premennými - korelácia
- porovnávanie dvoch skupín
- porovnávanie dvoch meraní
- porovnávanie troch a viacerých skupín
- chí-kvadrát teste pre súvis dvoch nominálnych premenných

Podme na to!

5.4 Vzdialené hodnoty

Pre zabezpečenie spoľahlivejších výsledkov je potrebné, aby sme sa pozreli na premenné aj z hľadiska vzdialených hodnôt. Pri korelačnej analýze ale aj porovnávacích testoch, ktoré pracujú s kontinuálnymi premennými, sa stretneme s pojmami *parametrické* a *neparametrické* testy. V kvantitatívnych výskumoch, ktoré pracujú s výskumnými súbormi rádovo v stovkách participantov, sa najčastejšie využívajú parametrické testy. Tieto testy pri výpočte koeficientu využívajú priemer a štandardnú odchýlku. Ako sme už v časti o mierach polohy a variability načrtli, tieto ukazovatele sú *citlivé* v prípade prítomnosti vzdialených hodnôt alebo narušenia normálnej distribúcie hodnôt premennej.

Vzdialené hodnoty sú hodnoty premennej, ktoré sú výrazne vzdialené od ostatných hodnôt, t.j. hodnoty, ktoré sú výrazne väčšie alebo menšie ako väčšina hodnôt. Čo však znamená výrazne vzdialené? Jeden zo spôsobov definovania je vzdialenosť väčšia ako 1,5 násobok medzikvartilového rozpätia (*IQR*) od prvého kvartilu (*Q1*) v smere nižšie alebo vyššie od tretieho kvartilu (*Q3*). Podľa tohto pravidla môžeme v programe jamovi najrýchlejšie zistiť prítomnosť vzdialených hodnôt prostredníctvom *krabičkového grafu* (angl. *boxplot*). Ten si môžeme jednoducho vykresliť pri deskriptívnej analýze, kde v časti *Plots* zvolíme *Boxplot* teda „krabičkový graf“. Nájde v ňom vyobrazené medzikvartilové rozpätie prostredníctvom krabičky, v ktorej je vodorovná čiara znázorňujúca medián. Pod a nad krabičkou nájdeme kolmé úsečky. Spodná vyjadruje rozsah 1,5 násobku medzikvartilového rozpätia od prvého kvartilu nižšie ($Q1 - 1,5 \cdot IQR$), vrchná zas naopak, od tretieho kvartilu vyššie ($Q3 + 1,5 \cdot IQR$). V prípade, že sa v premennej nevyskytujú vzdialené hodnoty, sú tieto úsečky vyjadrením rozsahu od minimálnej hodnoty k maximálnej. Ak sa v premennej nachádzajú vzdialené hodnoty, budú v grafe zobrazené prostredníctvom bodov. Príklad krabičkového grafu pre premennú Extraverzia zobrazuje Obrázok 5.1. Ako je možné vidieť, v spodnej časti grafu sa nachádza jeden bod. Ide o hodnotu, ktorá je od hodnoty prvého kvartilu vzdialená viac ako 1,5 násobok medzikvartilového rozpätia. Čo robiť s vzdialenými hodnotami?



Obrázok 5.1 Krabičkový graf s jednou vzdialenou hodnotou

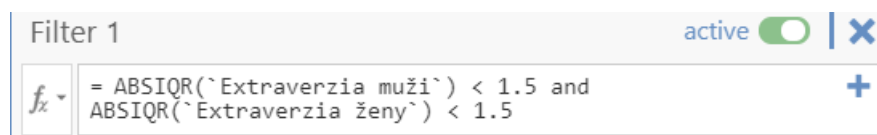
Odpoveď na túto otázku nie je jednoznačná. V prvom rade musíme posúdiť, či je táto hodnota *reálne* dosiahnuteľná. Takéto posúdenie robíme už na začiatku práce s premennými. Pozrieme si minimum a maximum a zhodnotíme, či premenné neobsahujú hodnoty, ktoré by *nemali* obsahovať. Príkladom môže byť vek. Ak respondenti vyplňali online dotazník, mohlo sa stať, že náhodou preťukli a uviedli vek 256 rokov. Pri takejto hodnote je jasné, že je nereálna. Najlepšie je hodnotu odstrániť, a tak jeden z respondentov nebude mať uvedený vek. Môže sa stať, že sa vám do výskumu zapojí niekto, kto nie je z vašej plánovanej cieľovej populácie. Ak sa zameriavate na neskorú dospelosť, no zapojí sa niekto, kto má 18 rokov, treba ho odstrániť úplne, pretože nespadá do vami stanovenej cieľovej populácie. Iným príkladom nereálnej hodnoty môže byť chyba na vašej strane. Rovnako ako respondentovi, tak aj vám sa mohlo stať, že ste nechtiac ťukli do klávesnice, a prepísali ste hodnotu (na toto veľký pozor!). V prípade, že viete dohľadať originálnu odpoveď, môžete tak spraviť, ak nie, hodnotu vymažte. Ak sú hodnoty reálne dosiahnuteľné, postup je zložitejší. Musíme sa rozhodnúť, čo s takýmito hodnotami spraviť. Jeden pohľad môže byť, že tieto hodnoty sa môžu v populácii vyskytnúť, a preto je otázne, či máme takéto hodnoty odstrániť. Na druhú stranu, parametrické testy sú citlivé na vzdialené hodnoty, a preto výsledky môžu byť týmito hodnotami posunuté. Pokiaľ máme vysoký počet respondentov a odstránenie pár hodnôt „nepocítíme“, môžeme tieto hodnoty odstrániť alebo odfiltrovať. Pokiaľ máme nízky počet respondentov (napr. 20), odstránenie zopár hodnôt je citelné. V takomto prípade môžeme voľiť neparametrické testy, pre ktoré odľahlé hodnoty nie sú problémom a pracovať s celým súborom.

Pred každou analýzou si môžeme nastaviť filter, tak aby sme v analyzovaných premenných alebo skupinách vyfiltrovali vzdialené hodnoty. V jamovi na to môžeme využiť šikovnú zabudovanú funkciu `ABSIQR()`. Táto funkcia vypočíta vzdialenosť jednotlivých hodnôt od „krabičky“ v jednotkách medzikvartilového rozpätia. Ak sa hodnota v danom riadku nachádza v rozmedzí medzikvartilového rozpätia, funkcia vráti hodnotu 0, ak sa nachádza mimo, vráti hodnotu vzdialenosti v IQR a v absolútnej hodnote. Vďaka tomu vieme nastaviť odfiltrovanie hodnôt, ktorých absolútna vzdialenosť od krabičky je väčšia ako 1,5. Ak pracujeme s viacerými premennými, napríklad riešime koreláciu medzi dvomi premennými, zadáme príkaz pre obe premenné v jednom filtri. Nastavenie filtra pre premenné Extraverzia a Prívetivosť zobrazujeme v obrázku 5.2. Keďže filter ponechá hodnoty, ktoré spĺňajú podmienku, musíme ho nastaviť tak, aby hodnota vrátená funkciou `ABSIQR` bola menšia ako 1,5 (ponecháme len hodnoty, ktoré nie sú vzdialené).



Obrázok 5.2 Nastavenie filtra pre ponechanie hodnôt, ktoré nie sú vzdialené v premenných Extraverzia a Prívetivosť

V prípade, keď budeme porovnávať dve alebo viac skupín, je potrebné, aby sme odfiltrovali vzdialené hodnoty pre jednotlivé skupiny zvlášť. Určitá hodnota nemusí byť vzdialenou, ak berieme súbor ako celok, no môže byť vzdialenou v konkrétnej skupine. V takomto prípade si najskôr musíme vytvoriť nové vypočítané premenné (*Computed variable*, kap. 2.3). Pri výpočte využijeme logickú funkciu *IF()*. Vďaka nej vieme nastaviť, že novovytvorená premenná má obsahovať len hodnoty pôvodnej premennej z vybranej skupiny. Pri tejto funkcii musíme najskôr zadať podmienku, a potom to, čo sa má stať, ak je podmienka naplnená a ak nie je naplnená: *IF([podmienka], [čo ak je naplnená], [čo ak nie je naplnená])*. Ak by sme chceli preniesť do novej premennej len hodnoty premennej Extraverzia u mužov, nastavenie bude vyzeráť takto: *IF(pohlavie == 1, Extraverzia)*. V tomto prípade sme neuviedli, čo sa má stať, ak podmienka nie je naplnená – nespraví sa nič, bunka v danom riadku zostane prázdna, a to presne chceme. Potom spravíme tento úkon pre druhú skupinu – vytvoríme novú vypočítanú premennú, v ktorej zmeníme príkaz: *IF(pohlavie == 2, Extraverzia)*. Pri nastavovaní príkazu si môžete pomôcť a názvy premenných vložiť cez tlačidlo *f_z*. V našom prípade sme si vytvorili dve premenné, „Extraverzia muži“, ktorá nesie hodnoty u mužov a „Extraverzia ženy“, ktorá nesie hodnoty u žien. Následne môžeme nastaviť filter s funkciou *ABSIQR()* pre obe premenné. Príklad uvádzame v Obrázku 5.3. Nezabudnite, že takto vytvorené premenné nám slúžia len pre nastavenie filtra a pri analýze už budeme pracovať s originálnou premennou.



Obrázok 5.3 Nastavenie filtra pre ponechanie hodnôt, ktoré nie sú vzdialené v premennej Extraverzia, samostatne pre mužov a ženy

Po nastavení sa vyfiltrujú respondenti, ktorí mali vzhľadom na ich skupinu odľahlú hodnotu v danej premennej. Ak si dáte znovu vykresliť krabičkový graf, môže sa stať, že pribudnú nové hodnoty, ktoré sú označené bodom, teda sú vzdialené. Odstránením niektorých hodnôt sa zmenili pomery v dátach a môže sa stať, že hodnoty, ktoré pôvodne neboli vzdialené sa po odfiltrovaní stanú odľahlými. V takomto prípade už neopakujeme čistenie a pracujeme so súborom tak, ako je, t.j. čistenie premenných od odľahlých hodnôt robíme len raz.

Pri riešení odľahlých hodnôt môžete vopred identifikovať takéto hodnoty pre všetky analýzy, ktoré budete robiť a odstrániť respondentov ako takých. Získate tak finálny súbor respondentov, s ktorým budete robiť ďalšie analýzy. Na druhú stranu, najmä v prípade väčšieho počtu hypotéz a premenných vstupujúcich do analýz, tak môžete prísť o väčší počet respondentov. Domnievame sa, že ak má respondent odľahlú hodnotu pri jednej z viacerých premenných, je zbytočné odstrániť ho zo súboru. Ideálnejšie je pred každou analýzou samostatne analyzovať vzdialené hodnoty v konkrétnych premenných, vyfiltrovať ich a pokračovať analýzou. Po

dokončení hlavnej analýzy (napr. po overení hypotézy), *nezabudnite* filter vymazať, prípadne ho prerobte pre účely ďalšej analýzy.

Znovu však opakujeme, že problematika odstraňovania vzdialených hodnôt, najmä v sociálno-psychologickom výskume, je komplikovaná. Príkladom je už práve spomenutá osoba s veľmi nízkou mierou extravenzie (Obrázok 5.1). Na jednu stranu by sme mohli túto osobu pri práci s premennou Extraverzia vyfiltrovať, pretože je vzdialená, no rovnako môžeme viesť polemiku, že táto osoba nie je nereálna a naozaj v populácii existujú ľudia s veľmi nízkou extravenziou. Čo ak by sme mali omnoho vyšší počet respondentov v našom výskume? Ak by sme pracovali nie v stovkách, ale tisícoch či desaťtisícoch, je možné, že by táto hodnota bola frekventovanejšia a nebola by vzdialenou. Vyfiltrovaním tejto osoby tak odstránime niekoho, kto reprezentuje populáciu. Pre účely ďalších výpočtov pracujeme v tejto učebnici so všetkými respondentmi – neodstraňujeme vzdialené hodnoty.

5.5 Testovanie normality distribúcie premennej

Po odstránení odľahlých hodnôt pokračujeme. Pre posúdenie, či je premenná normálne distribuovaná, môžeme zvoliť viaceré štatistické testy. Jeden z nich, Shapiro-Wilkov, je priamo implementovaný v programe jamovi. Tento test nám, podobne ako iné inferenčné štatistické testy, vráti koeficient, ktorý však neinterpretujeme, no sledujeme jeho štatistickú významnosť. Nulová hypotéza tohto testu je, že premenná je v populácii normálne distribuovaná. Vypočítaná hodnota štatistickej významnosti nám teda hovorí, aká je pravdepodobnosť získať distribúciu, ktorú máme my, pri našom počte respondentov, v prípade že platí nulová hypotéza. Ak je hodnota signifikancie nízka ($p < 0,05$), berieme to tak, že nulová hypotéza neplatí a premenná nie je normálne distribuovaná. Takéto zhodnotenie je problematické. Dôvod sme si už spomenuli: výpočet signifikancie je závislý aj od počtu respondentov, na ktorých je analýza uskutočnená. Čím je počet vyšší, tým nižšia hodnota štatistickej významnosti nám vyjde. Pokiaľ máme vysoký počet respondentov, môže sa stať, že aj pri dostatočne normálne distribuovanej premennej získame štatisticky významný výsledok a mali by sme voliť neparametrický test. Naopak, ak máme veľmi nízky počet respondentov, Shapiro-Wilkov test môže vyjsť nesignifikantný napriek tomu, že distribúcia testovanej premennej nie je ideálna pre použitie parametrického testu. Z tohto dôvodu je vhodné na normalitu distribúcie nahliadať komplexnejšie prostredníctvom viacerých dôkazov. Osobne sa najviac prikláňame k hodnoteniu tzv. quantile-quantile grafov, ktoré sa zaužívané nazývajú Q-Q grafy či anglicky *Q-Q plots*, no môžeme sa zamerať aj na hodnoty šikmosti (*angl. skewness*) a špicatosti (*angl. kurtosis*), ktoré nás informujú o povahe distribúcie, alebo si distribúciu môžeme nechať graficky zobrazit' prostredníctvom histogramu.

Spôsob, akým sú hodnoty premennej distribuované, vieme číselne vyjadriť prostredníctvom spomenutých koeficientov šikmosti a špicatosti. Šikmosť nám hovorí o tom, či sú hodnoty rovnomerne distribuované z oboch strán, od minimálnej hodnoty k priemeru a od priemeru k maximálnej hodnote. Špicatosť nám jednoducho vypovedá o tom, aký je rozdiel v hustote najčastejšie zastúpených hodnôt v porovnaní s ostatnými menej zastúpenými hodnotami, inak povedané o hrúbke krajov (chvostov) distribúcie. Pri posudzovaní normality distribúcie sa primárne zameriavame na šikmosť, no ideálne je, aby oba koeficienty boli čo najbližšie k nule, ideálne medzi -0,5 a 0,5, no akceptovateľné je aj rozmedzie -1 až 1.

Prakticky si ukážeme zhodnotenie normálnej distribúcie pri premenných, s ktorými budeme neskôr pracovať pri korelačnej analýze. V jamovi si kartu *Analyses* zvolíme možnosť

Exploration a ďalej *Descriptives*. Nastavenie tejto analýzy sme preberali v časti o deskriptívnej analýze. Teraz sa zameriame na šikmosť, špicatosť distribúcie, overenie normality prostredníctvom Shapiro-Wilkovho testu a v časti *Plots* si zapneme *Histogram*, *Density* a *Q-Q*. Zapneme si tiež možnosť *N*, aby sme zistili počet respondentov, ktorí vstupujú do analýzy. Pre prehľadnosť výsledkov sme zatiaľ ostatné možnosti vyplí. Do okienka *Variables* si prenesieme premenné, ktoré chceme analyzovať a získame požadované výsledky. Nastavenie pre tento príklad zobrazujeme v Obrázku 5.4.

Obrázok 5.4 Nastavenie analýzy pre posúdenie normality distribúcie premennej

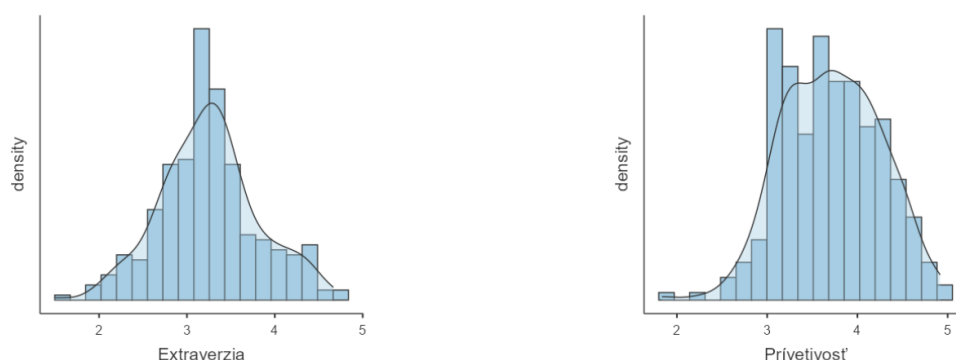
Získame tabuľku, v ktorej budú výsledky pre každú premennú (Obrázok 5.5).

Descriptives		
	Extraverzia	Prívetivosť
N	281	281
Missing	0	0
Skewness	0.11	-0.10
Std. error skewness	0.15	0.15
Kurtosis	0.10	-0.25
Std. error kurtosis	0.29	0.29
Shapiro-Wilk W	0.99	0.99
Shapiro-Wilk p	0.043	0.037

Obrázok 5.5 Výsledok analýzy šikmosti, strmosti a Shapiro-Wilkovho testu

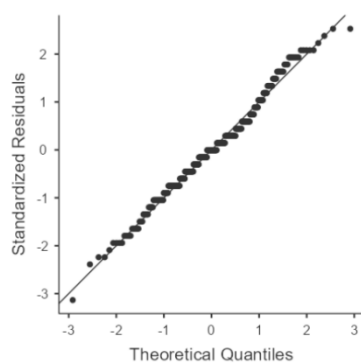
Postupne zhodnotíme šikmosť, strmosť a signifikanciu Shapiro-Wilkovho testu. Pri premennej Extraverzia sme zistili hodnotu šikmosti 0,11 a hodnotu špicatosti -0,10. Obe tieto hodnoty sú pomerne blízke nule a nachádzajú sa v spomenutom pásme -0,5 až 0,5. Podobne je tomu aj v prípade premennej Prívetivosť. Pri Shapiro-Wilkovom teste sa zameriavame na hodnotu

signifikancie, ktorá sa zaužívané označuje ako p . Ako sme v úvode tejto podkapitoly spomenuli, ak získame hodnotu nižšiu ako 0,05, pochybujeme o tom, že je dodržaná normálna distribúcia premennej. Pri oboch premenných je táto hodnota nižšia ako 0,05. Ak by sme vyvodili záver len na základe tohto ukazovateľa, prišli by sme k záveru, že normálna distribúcie nie je dodržaná a volili by sme neparametrický test. Avšak hodnota significancie môže byť zapríčinená nie tak problematickou distribúciou hodnôt tejto premennej ale vysokým počtom respondentov ($N = 281$), ktorí vstupujú do tejto analýzy. Test je „citlivý“ aj na prakticky zanedbateľný rozdiel od normálnej distribúcie, preto vyjde štatisticky významný pri $p < 0,05$. Distribúciu hodnôt premenných máme graficky vyobrazenú prostredníctvom histogramov s vykreslenou krivkou hustoty dát. Tieto grafy zobrazujeme v Obrázku 5.6.

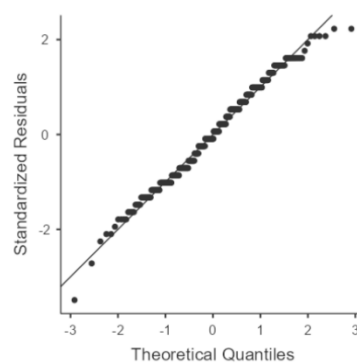


Obrázok 5.6 Histogram pre vybrané premenné

Histogram nám zobrazuje frekvenciu skupín susediacich hodnôt. Ako na prvý pohľad vidíme, krivky pri jednotlivých premenných sa veľmi nepodobajú na Gausovu krivku normálnej distribúcie (Obrázok 3.1). Napriek tomu vidíme podobný trend. Hodnoty okolo priemeru sú zastúpené najčastejšie a čím ďalej od priemeru ideme, tým klesá frekvencia výskytu týchto hodnôt. Tento trend je pomerne dobre viditeľný pri premennej *Extraverzia*. Pri premennej *Prívetivosť* je širšie pásmo hodnôt, ktoré majú veľmi podobnú frekvenciu výskytu. Histogram je šikovník pomocník pre znázornenie distribúcie hodnôt, avšak pre posúdenie normality distribúcie je nápomocnejší Q-Q graf. Tieto grafy obsahujú priamku, ktorá vyjadruje normálnu distribúciu a jednotlivé body (krúžky) znázorňujú hodnoty jednotlivých respondentov. Posúdenie týchto grafov je naozaj jednoduché. Pravidlom je, aby *skoro všetky body boli pretnuté priamkou, prípadne ležali v jej tesnej blízkosti*. Grafy pre obe premenné zobrazujeme v Obrázku 5.7. Ako je možné vidieť, naozaj väčšina týchto bodov je pretnutá priamkou. V praxi sa môžete stretnúť s tým, že na začiatku a konci môžu byť niektoré body vzdialenejšie od priamky, avšak pokiaľ ich je len zopár, nie je to problém. Problém je v prípade, že sa viaceré body vzdiaľujú od priamky a vzniká tak oblý rad, môžeme ho nazvať „bruškom“, prípadne sa okolo priamky tvorí z týchto bodov krivka, ktorá sa striedavo z jednej a druhej strany vzdiaľuje od priamky. Pokiaľ je tento odklon graficky výrazný, môžeme zvážiť použitie neparametrického testu.



Extraverzia



Prívetivosť

Obrázok 5.7 Q-Q grafy pre vybrané premenné

Aký je teda záver? Štatisticky významný Shapiro-Wilkov test naznačuje, že dáta nie sú normálne distribuované, avšak tento výsledok môže byť spôsobený vysokým počtom respondentov. Dôkazy založené na hodnote šikmosti v spojitosti so zhodnotením Q-Q grafov poukazujú na to, že môžeme použiť parametrický test.

Takéto posúdenie robíme aj v prípade porovnávania dvoch alebo viacerých skupín, ktorým sa budeme venovať v kapitole 7. Pri týchto testoch však posudzujeme normalitu testovanej premennej samostatne v jednotlivých skupinách. Stačí, ak pri analýze vložíme do okienka *Split by* premennú, ktorá definuje porovnávané skupiny.

6. Korelačná analýza

Korelačná analýza slúži pre zisťovanie vzťahov medzi dvomi premennými meranými na kontinuálnej, prípadne ordinálnej úrovni. Je asi najčastejšie využívanou analýzou, či už v záverečných prácach alebo kvantitatívnych výskumných štúdiách. Korelačnú analýzu používame v prípade, ak očakávame, že s narastajúcou hodnotou jednej premennej bude narastať alebo klesať hodnota druhej premennej. Výsledkom korelačnej analýzy je korelačný koeficient a jeho štatistická významnosť. Korelačný koeficient vyjadruje tesnosť vzťahu dvoch premenných, no častejšie je používané vyjadrenie *sila vzťahu*. Najčastejšie používaným korelačným koeficientom je Pearsonov, no používa sa aj Spearmanov a Kendalov korelačný koeficient. Výber koeficientu súvisí s povahou premenných. Pearsonov korelačný koeficient je parametrický. Predpokladom pre použitie tohto koeficientu je normálna distribúcia oboch premenných, ktoré dávame do vzťahu ako aj to, že obe premenné majú byť merané na *kontinuálnej* úrovni. V prípade, že jedna alebo obe premenné nespĺňajú podmienku normálnej distribúcie alebo kontinuálnej úrovne merania, používame neparametrické koeficienty – Spearmanov alebo Kendalov. Pokiaľ predpokladáme normálnu distribúciu meraných *kontinuálnych* premenných v populácii a máme dostatočne veľký počet ľudí, na ktorých analýzu uskutočňujeme (napr. viac ako 100), môžeme si výber zjednodušiť a použiť Pearsonov korelačný koeficient. V takomto prípade je vysoká pravdepodobnosť, že naše dáta sú dostatočne podobné normálnej distribúcii. Pokiaľ máme jednu alebo obe premenné merané na *ordinálnej* úrovni, alebo sa nám *nepotvrdila* normálna distribúcia použitých premenných, prípadne máme nízky počet ľudí vstupujúcich do analýzy (napr. menej ako 30), vyberieme neparametrický korelačný koeficient. Teraz si ukážeme, ako by mali vyzeráť hypotézy týkajúce sa korelácií, ako ich štatisticky vyhodnotiť a ako výsledky zapísať a interpretovať.

Vo všeobecnosti stanovujeme hypotézu ako predpoklad o prítomnosti štatisticky významného vzťahu:

H: Predpokladáme štatisticky významný vzťah medzi X a Y.

Pri stanovení hypotézy môžeme no nemusíme používať explicitné vyjadrenie predpokladania, čiže hypotéza môže znieť aj:

H: Existuje štatisticky významný vzťah medzi X a Y.

Dôležité je, aby sme explicitne zapísali, čo očakávame a medzi ktorými premennými.

Nulová hypotéza v tomto prípade hovorí o neprítomnosti vzťahu a jej znenie by bolo:

H_0 : Neexistuje vzťah medzi X a Y.

V oboch prípadoch sme stanovili dvojsmernú hypotézu. Ako sme v predchádzajúcej časti uviedli, pri jednosmernej hypotéze je špecifikovaný aj smer vzťahu. Pokiaľ očakávame priamu úmeru, teda že so zvyšujúcou sa hodnotou jednej premennej sa zvyšuje hodnota druhej premennej, môže našu hypotézu zapísať takto:

H: Predpokladáme štatisticky významný pozitívny vzťah medzi X a Y.

Nulová hypotéza má v tomto prípade znenie:

H_0 : X a Y nemajú medzi sebou pozitívny vzťah.

V prípade nepriamej úmery, teda že so zvyšujúcou sa hodnotou jednej premennej klesá hodnota druhej premennej, uvádzame negatívny:

H: Predpokladáme štatisticky významný negatívny vzťah medzi X a Y.

Nulová hypotéza má v tomto prípade znenie:

H0: X a Y nemajú medzi sebou negatívny vzťah.

Pre praktickú ukážku zameníme X a Y za niečo špecifické a ukážeme si ako takéto hypotézy analyzovať a interpretovať. Vo vašich záverečných prácach nezabúdajte na to, že hypotézy majú byť pri ich formulácii odôvodnené. Vaše hypotézy uvádzate vo výskumnej časti práce a mali by nasledovať za výskumným problémom a cieľmi. Pokúste sa pri ich formulácii explicitne pripomenúť čitateľovi, na základe čoho máte také očakávanie. Pre ukážku v tejto učebnici však uvádzame len hypotézy, ktoré si následne vyhodnotíme.

H1: Predpokladáme štatisticky významný vzťah medzi extraverziou a prežívaným šťastím.

H2: Prívetivosť a prežívané šťastie budú spolu štatisticky významne, pozitívne korelovať.

H3: Predpokladáme, že negatívna emocionalita bude štatisticky významne, negatívne korelovať s prežívaným šťastím.

H4: Najvyššie dosiahnuté vzdelanie bude štatisticky významne a pozitívne súvisieť so svedomitosťou.

Ako ste si mohli všimnúť, prvá hypotéza je dvojsmerná a ostatné sú jednosmerné. Ak chceme analyzovať dáta a potvrdiť či vyvrátiť nami stanovené, meritórne/alternatívne hypotézy, musíme najskôr zvoliť vhodný korelačný koeficient. V prípade prvých troch hypotéz pracujeme s premennými, ktoré sú *kontinuálne*. Vzhľadom na počet respondentov v našom dátovom súbore a na fakt, že tieto premenné by mali byť v našej cieľovej populácii približne normálne distribuované, mohli by sme zvoliť Pearsonov korelačný koeficient. Pri poslednej hypotéze pracujeme aj s ordinálnou premennou, preto by bolo vhodné použiť Spearmanov korelačný koeficient. Ak chceme pristúpiť k výberu koeficientu zodpovednejšie, mali by sme zhodnotiť normálnu distribúciu týchto premenných (podkapitola 5.5).

6.1 Korelačná analýza v programe jamovi

Korelačnú analýzu v jamovi nájdeme v záložke *Analyses*, kde zvolíme možnosť *Regression* a *Correlation Matrix*. Základné okno obsahuje dve okienka. Ako zvyčajne, v jednom nájdeme premenné v našom súbore a do druhého okienka presunieme premenné, medzi ktorými chceme zisťovať vzťah. Ikonky v rohu tohto okienka poukazujú na to, že môžeme vkladať kontinuálne a ordinálne premenné. Pod týmito okienkami nájdeme ďalšie nastavenia, ktoré zobrazujeme v Obrázku 6.1.

Correlation Coefficients	Additional Options
☒ Pearson	☒ Report significance
☐ Spearman	☐ Flag significant correlations
☐ Kendall's tau-b	☐ N
	☐ Confidence intervals
	Interval: 95 %
Hypothesis	Plot
☒ Correlated	☐ Correlation matrix
☐ Correlated positively	☐ Densities for variables
☐ Correlated negatively	☐ Statistics

Obrázok 6.1 Základné nastavenie pri korelačnej analýze

V časti *Correlation Coefficients* volíme korelačný koeficient, pričom môžeme zvoliť aj viacero naraz. V tejto učebnici budeme pracovať s Pearsonovým a Spearmanovým korelačným koeficientom. V ďalšej časti *Hypothesis* nastavujeme naše očakávanie ohľadom vzťahu. Možnosť *Correlated* vyjadruje dvojstrannú hypotézu – očakávame len vzťah bez určenia smeru. Ďalšie dve možnosti *Correlated positively* a *Correlated negatively* vyjadrujú jednostrannú hypotézu, očakávanie pozitívnej alebo negatívnej korelácie. Toto nastavenie má za cieľ uľahčenie výpočtu štatistickej významnosti výsledku. Vyskytuje sa aj pri niektorých ďalších bivariačných testoch. V praxi však niekedy študentom skôr komplikuje situáciu ako pomáha.

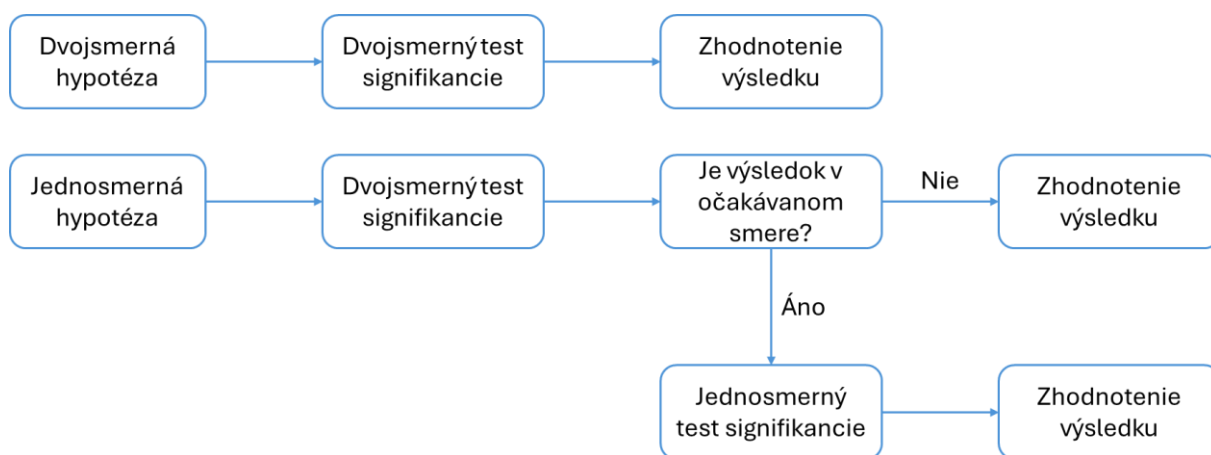
Prvý problém je v nepozornosti – pokiaľ overujeme viacero hypotéz, ktoré majú rôzne smery, môže sa stať, že ponecháme nesprávne nastavenie. Vysvetlíme, prečo je to problém. V prípade, že by sme pri overení hypotézy o prítomnosti negatívnej korelácie medzi premennými ponechali nastavenie *Correlated positively*, významnosť výsledku by bola počítaná pre alternatívnu hypotézu v opačnom smere, ako sme mienili. Vypočítaná sila vzťahu bude rovnaká pri všetkých nastaveniach, líšiť sa bude hodnota štatistickej významnosti.

Povedzme, že by bol vzťah medzi premennými $-0,30$ – neskôr uvádzame, že ide o stredne silný negatívny vzťah. V tomto hypotetickom príklade nám pri nastavení *Correlated* (dvojsmerná hypotéza) vyjde hodnota signifikancie nižšia ako $0,001$. Ak by sme chceli zistiť hodnotu signifikancie pre jednosmerný test, stačí túto hodnotu vydeliť dvomi, čím sa rovnako dopracujeme k zaokrúhlenej hodnote $p < 0,001$. Túto hodnotu získame aj pri nastavení *Correlated negatively*. Získali by sme teda štatisticky významný výsledok a našu hypotézu potvrdili.

Ak by sme nesprávne nastavili *Correlated positively*, získaná hodnota by bola po zaokrúhlení $p = 1$, čo je rozhodne viac ako $0,05$, preto by sme pri nekritickom hodnotení výsledku konštatovali štatisticky nevýznamný výsledok a našu hypotézu nesprávne nepotvrdili. Možno sa pýtate „prečo?“ Pripomeňme si informácie, ktoré sme uviedli v kapitolách 5.1 a 5.2. Výsledok dvojsmerného testu signifikancie nám povie, aká je pravdepodobnosť získať takúto koreláciu pri našich dátach, ak by v populácii korelácia nemala byť – pravdepodobnosť je teda menej ako $0,1\%$ ($p < 0,001$). Pri správnom nastavení pre tento príklad (*Correlated negatively*) nám signifikancia povie, aká je pravdepodobnosť získať takúto hodnotu pri platnosti nulovej hypotézy „vzťah nie je negatívny“ – a opäť, táto pravdepodobnosť je menej ako $0,1\%$ ($p < 0,001$). Nesprávne nastavenie (*Correlated positively*) nám dá výsledok, aká je pravdepodobnosť

získať takúto hodnotu, ak v populácii „vzťah nie je pozitívny“ – hodnota -0,30 rozhodne nie je pozitívna, a preto je pravdepodobnosť blízka 100 % ($p = 1$).

Druhý problém nastáva v ojedinelej situácii, keď sme si stanovili jednosmernú hypotézu, napr. predpokladali sme negatívny vzťah, no „z nejakého“ dôvodu je realita opačná a vzťah medzi premennými je pozitívny (a má aj svoju vecnú významnosť). Upozorňujeme, že sme nespravili metodologickú chybu alebo nesprávne výpočty – prosté vzťah medzi premennými je v populácii opačný, ako sme očakávali. V takomto prípade by nám opäť pri jednosmernom teste signifikancie (*Correlated negatively*) vyšla signifikancia $p > 0,05$, pričom hypotéza by sa nám nepotvrdila. Problém však nastáva pri ďalšej interpretácii vzťahu. Ak by sme zostali pri jednosmernom testovaní, mohlo by nám ujsť dôležité zistenie, že síce sa hypotéza nepotvrdila a vzťah nie je negatívny, no práve naopak – je pozitívny a k tomu štatisticky významný ($p < 0,05$) z hľadiska dvojsmerného testu signifikancie. Z tohto dôvodu odporúčame pri overovaní jednosmernej hypotézy začať dvojsmerným testom a ak je výsledok v smere hypotézy (napr. hypotéza o negatívnom vzťahu a identifikovaný negatívny vzťah), prepnúť na jednosmerný test. Ak by výsledok bol v opačnom smere (napr. hypotéza o negatívnom vzťahu a identifikovaný pozitívny vzťah), ponechať a reportovať signifikanciu pre dvojsmernú hypotézu. Pre lepší prehľad uvádzame nasledujúci diagram.



V časti *Additional Options* môžeme zvoliť ďalšie informácie, ktoré chceme zobraziť. V základe zvolené *Report significance* vyjadruje, že chceme mať uvedenú hodnotu štatistickej významnosti – áno chceme, je to pre nás dôležitá informácia. *Flag significant correlations* nám prostredníctvom hviezdíčiek označí hodnoty korelačných koeficientov, ktoré sú signifikantné na úrovni $p < 0,05$ (*), $p < 0,01$ (**) a $p < 0,001$ (***). Je to praktické v prípade, že sme naraz zvolili viacero premenných, kedy nám jamovi automaticky vypočíta vzájomné vzťahy medzi všetkými zvolenými premennými. V prípade dvoch premenných ide o jednu koreláciu. V prípade 5 premenných už získame 10 korelačných koeficientov a hodí sa nám aj takéto zvýraznenie. Ak zvolíme *N*, získame informáciu o počte ľudí, na ktorých bola analýza vykonaná. Táto informácia sa zide najmä v prípade, že niektorí respondenti majú chýbajúce údaje alebo máte nastavené nejaké filtrovanie dát (napr. odstránené vzdialené hodnoty). Možnosť *Confidence intervals* nám vypočíta konfidenčné intervaly korelačných koeficientov, toto je však téma nad rámec tejto učebnice. V poslednej časti *Plots* si môžeme nechať graficky znázorniť korelačnú maticu. Takéto grafické znázornenie však v záverečných prácach neuvádzame a nebudeme sa mu ani venovať.

Prakticky si ukážeme overenie hypotézy H1. V prípade oboch premenných sme potvrdili normálnu distribúciu premenných, preto volíme Pearsonov korelačný koeficient r . Otvoríme si korelačnú analýzu a do okienka presunieme premennú Extraverzia a SHS_priemer. Nastavenie *Hypothesis* necháme na *Correlated*, keďže ide o dvojsmernú hypotézu. Výsledok, ktorý sa nám zobrazí, nájdete v Obrázku 6.2.

Correlation Matrix		Extraverzia	SHS_priemer
Extraverzia	Pearson's r	—	
	df	—	
	p-value	—	
	N	—	
SHS_priemer	Pearson's r	0.48	—
	df	279	—
	p-value	< .001	—
	N	281	—

Obrázok 6.2 Výsledok korelačnej analýzy pre H1

Výsledok pre zhodnotenie prvej hypotézy máme už zobrazený v Obrázku 6.2. V prvom rade sledujeme hodnotu štatistickej významnosti p -value. V tomto prípade sme zistili, že $p < 0,001$ (ak by sme zmenili nastavenia, môžeme sa dozvedieť aj konkrétnejšiu hodnotu, v tomto prípade je hodnota p dokonca menšia ako jedna milióntina). Úplne nám však postačuje informácia, ktorú máme v základe. Čo vďaka tomu vieme? Naša alternatívna hypotéza hovorí o prítomnosti vzťahu medzi extravertiou a prežívaným šťastím. Nulová hypotéza v tomto prípade hovorí, že v populácii tento vzťah nie je – korelačný koeficient by mal byť rovný 0. Hodnota štatistickej významnosti nám hovorí o tom, aká je pravdepodobnosť získať hodnotu korelačného koeficientu 0,48, pri 281 respondentoch, ak v populácii nie je vzťah medzi premennými. Hodnota štatistickej významnosti je nižšia ako konvenčne prijímaná hodnota 0,05, dokonca výrazne nižšia. Takto nízka pravdepodobnosť ($p < 0,001$, t.j. menej ako 0,1 %) nás vedie k pochybovaniu o platnosti nulovej hypotézy a k presvedčeniu, že nami stanovená hypotéza platí, t.j. v populácii existuje vzťah medzi týmito dvomi premennými. Tento systém je prakticky rovnaký pre všetky alternatívne hypotézy týkajúce sa korelácií. V prípade jednosmerných hypotéz však nulová hypotéza zahŕňa aj opačný ako nami stanovený vzťah. Napríklad, ak očakávame pozitívny vzťah, nulová hypotéza zahŕňa neprítomnosť vzťahu a akýkoľvek negatívny vzťah. A naopak, ak očakávame negatívny vzťah, nulová hypotéza zahŕňa neprítomnosť vzťahu a akýkoľvek pozitívny vzťah. Dajte si preto pozor na nastavenie hypotézy (*Hypothesis*) v jamovi. Pri analýze si taktiež doprajte čas na rozmýšľanie a nemýľte si korelačný koeficient s jeho štatistickou významnosťou! Nanešťastie sa neraz stalo, že študenti interpretovali korelačný koeficient ako štatistickú významnosť a naopak – ich zistenia a všetko na to nadväzujúce je automaticky neplatné!

Tento postup použijeme aj pri overení hypotéz H2 a H3. V oboch prípadoch sme začali dvojsmerným testom signifikancie (*Correlated*). Oba zistené vzťahy boli v očakávanom smere, takže sme pokračovali samostatne jednosmerným testom. V prípade H2 nastavíme *Correlated positively* a pri H3 zas *Correlated negatively*. V prípade H4 volíme neparametrický

Spearmanov korelačný koeficient (jedna premenná je ordinálna). Výsledky z programu jamovi uvádzame v Obrázku 6.3.

Correlation Matrix			
		SHS_priemer	Prívetivosť
SHS_priemer	Pearson's r	—	
	df	—	
	p-value	—	
	N	—	
Prívetivosť	Pearson's r	0.31	—
	df	279	—
	p-value	< .001	—
	N	281	—

Note. H_a is positive correlation

Correlation Matrix			
		SHS_priemer	Negatívna emocionalita
SHS_priemer	Pearson's r	—	
	df	—	
	p-value	—	
	N	—	
Negatívna emocionalita	Pearson's r	-0.49	—
	df	279	—
	p-value	< .001	—
	N	281	—

Note. H_a is negative correlation

Correlation Matrix			
		Svedomitosť	vzdelanie
Svedomitosť	Spearman's rho	—	
	df	—	
	p-value	—	
	N	—	
vzdelanie	Spearman's rho	0.11	—
	df	279	—
	p-value	0.076	—
	N	281	—

Correlation Matrix			
		Svedomitosť	vzdelanie
Svedomitosť	Spearman's rho	—	
	df	—	
	p-value	—	
	N	—	
vzdelanie	Spearman's rho	0.11	—
	df	279	—
	p-value	0.038	—
	N	281	—

Note. H_a is positive correlation

Obrázok 6.3 Ukážka výsledkov korelačných analýz pre H2 až H4 v programe jamovi

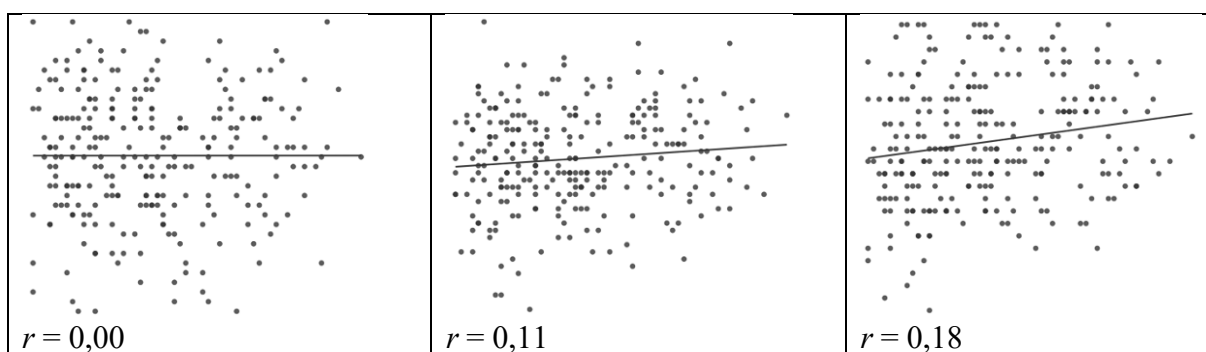
Ako je na obrázku vidieť, vzťah prežívaného šťastia s prítelivosťou a negatívnu emocionalitou je štatisticky významný, pretože hodnota p je menšia ako 0,05 – tieto hypotézy sa nám potvrdili. Pre zaujímavosť sme v Obrázku 6.3 uviedli aj výsledok dvojsmernej signifikancie pre hypotézu H4. Ak by táto hypotéza bola dvojsmernou, výsledok by nebol štatisticky významný, pretože $p = 0.076$ je viac ako 0,05. Keďže je hypotéza jednosmerná, výsledná jednosmerná signifikancia je polovičná voči dvojsmernej $p = 0.038$, čo je menej ako 0,05 a hypotéza sa nám potvrdila.

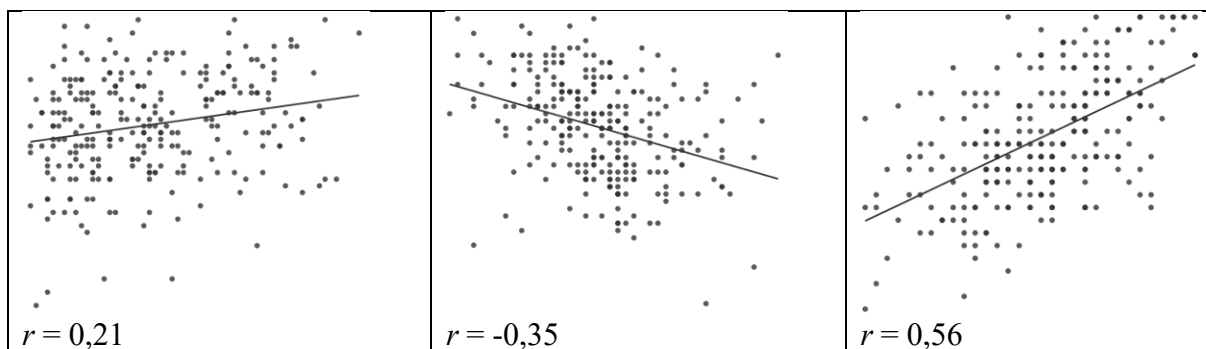
6.2 Interpretácia a zápis výsledkov korelačnej analýzy

Z hľadiska NHST je pre nás dôležitá najmä informácia o tom, či pochybujeme alebo nepochybujeme o platnosti nulovej hypotézy, teda či potvrdzujeme alebo zamietame nami stanovenú alternatívnu hypotézu. Z tohto dôvodu sme sa zatiaľ venovali len hodnote štatistickej významnosti p . Avšak korelačná analýza má aj iný výsledok, ktorý je pre nás zaujímavý, a to práve korelačný koeficient. Korelačný koeficient je štandardizovaná hodnota s rozsahom od -1 (absolútny negatívny vzťah) až 1 (absolútny pozitívny vzťah). Ako sme už v úvode kapitoly spomenuli, vyjadruje silu vzťahu medzi týmito premennými. Negatívne hodnoty hovoria o zápornom vzťahu, pozitívne hodnoty vypovedajú o kladnom vzťahu. Nižšie uvádzame ako interpretovať rôzne hodnoty korelačného koeficientu (Cohen, 1988). Ignorujúc znamienko korelačného koeficientu (pracujeme s absolútnou hodnotou), približné delenie je:

- 0 – 0,1 – žiaden vzťah, vzťah blízky nule, prakticky zanedbateľný vzťah
- 0,1 – 0,2 – veľmi slabý vzťah
- 0,2 – 0,3 – slabý vzťah
- 0,3 – 0,5 – stredne silný vzťah
- 0,5 – 0,7 – silný vzťah
- 0,7 a viac – veľmi silný vzťah

Takéto hodnotenie je len rámcové, no napomáha nám pri interpretácii hodnôt. Vždy však záleží aj od oblasti, ktorú riešite. Pri niektorých oblastiach psychologického výskumu môže byť aj slabý vzťah *zaujímavý*, zatiaľ čo pri iných môže byť stredne silný vzťah „slabým“ (napr. vzťah výsledku psychodiagnostickej metodiky s kritériom, kde by sme očakávali silnejší vzťah). Koreláciu dvoch premenných môžeme graficky znázorniť prostredníctvom *bodového* grafu. V tomto grafe jednotlivé body vyjadrujú priesečníky hodnôt dvoch premenných pre každého jedného respondenta a môžeme vidieť, či a aký trend je v dátach. K týmto bodom je pridaná aj priamka, ktorá je umiestnená tak, aby bola k všetkým bodom najbližšie. Pri jednotlivých grafoch uvádzame aj hodnotu Pearsonovho korelačného koeficientu.





Obrázok 6.4 Grafické znázornenie vzťahov dvoch premenných (bodový graf s priamkou)

Zápis výsledkov korelačnej analýzy je jednoduchý – potrebujeme uviesť hodnotu korelačného koeficientu (Pearsonovo r alebo Spearmanovo ρ [rhó]) a jeho štatistickú významnosť (p). V rámci slovného hodnotenia skomentujeme štatistickú významnosť, silu a smer vzťahu. Vhodné je, aby sme na začiatok uviedli aj to, ktorý korelačný koeficient sme použili a prečo. Nižšie uvádzame odpovede na nami stanovené hypotézy.

Pre overenie hypotéz H1 až H3 sme na základe dodržania predpokladov pri všetkých testovaných premenných použili Pearsonov korelačný koeficient.

H1: Predpokladáme štatisticky významný vzťah medzi extraverziou a prežívaným šťastím.

Hypotéza H1 sa potvrdila. Zistili sme štatisticky významný, stredne silný, pozitívny vzťah ($r = 0,48$; $p < 0,001$) medzi extraverziou a prežívaným šťastím.

H2: Prívetivosť a prežívané šťastie budú spolu štatisticky významne, pozitívne korelovať.

Hypotéza H2 sa potvrdila. Zistili sme štatisticky významný, stredne silný, pozitívny vzťah ($r = 0,31$; $p < 0,001$) medzi prívetivosťou a prežívaným šťastím.

H3: Predpokladáme, že negatívna emocionalita bude štatisticky významne, negatívne korelovať s prežívaným šťastím.

Hypotéza H3 sa potvrdila. Negatívna emocionalita a prežívané šťastie spolu štatisticky významne, negatívne korelujú. Tento vzťah je stredne silný ($r = -0,49$; $p < 0,001$).

H4: Najvyššie dosiahnuté vzdelanie bude štatisticky významne a pozitívne súvisieť so svedomitosťou.

Pre overenie hypotézy H4 sme zvolili Spearmanov neparametrický korelačný koeficient, keďže premenná najvyššie dosiahnuté vzdelanie je ordinálna. Naša hypotéza sa potvrdila. Medzi premennými sme zistili štatisticky významný, slabý, pozitívny vzťah ($\rho = 0,11$; $p = 0,038$).

Pri uvádzaní výsledkov teda nezapadnite napísať, aký koeficient ste použili, či sa hypotéza potvrdila alebo nie, hodnotu korelačného koeficientu a štatistickej významnosti. Nezapadnite tieto informácie dôkladne slovne zhodnotiť. A pre istotu, znovu zopakujeme: *nezmýľte si štatistickú významnosť so silou vzťahu!*

V prípade, že máte viacero hypotéz, ktoré sa zaoberajú vzťahom nejakej premennej s inými premennými, môžete zvoliť zápis výsledkov do tabuľky. Treba zvážiť, či je praktickejšie

zapísať výsledky v texte alebo do tabuľky. Ak uvádzate číselné výsledky do tabuľky, neuvádzate ich duplicitne v texte, no stále je nutné popísať výsledky – interpretovať hodnoty. Príkladom môžu byť naše hypotézy H1, H2 a H3. Spoločnou premennou vo všetkých troch hypotézach je prežívané šťastie. Predtým, než sa vyjadríme k hypotézam samostatne, by sme mohli čitateľa informovať, že konkrétne výsledky uvádzame v tabuľke. Potom zhodnotiť hypotézy bez uvedenia číselných hodnôt a uviesť tabuľku. Príklad takejto tabuľky uvádzame nižšie. Osobne by sme v tomto prípade uvádzanie výsledkov v tabuľke nepoužili, keďže ide len o výsledky troch analýz, no uváženie je na vás. V prípade, že by sme uvádzali výsledky pre viacero analýz so spoločnou premennou/premennými, bolo by to vhodné a pre čitateľa prehľadnejšie.

Tabuľka X

Výsledky korelačnej analýzy pre hypotézy H1 až H3

	Prežívané šťastie		
	<i>N</i>	<i>r</i>	<i>p</i>
Extraverzia	281	0.48	< 0.001
Prívetivosť	281	0.31	< 0.001
Negatívna emocionalita	281	-0.49	< 0.001

Poznámka.

V niektorých výskumných štúdiách sa môžete stretnúť s tým, že autori používajú Pearsonovo *r* aj pre vyjadrenie súvisu medzi ordinálnymi premennými, či dokonca aj medzi nominálnymi. Môže to vytvárať chaos v tom, čo kedy použiť. Matematika v pozadí týchto analýz je vec kúzelná (čo iné povedať, ak sa v tom nevyznáme) – niekedy sa k tomu istému výsledku dá dopracovať prostredníctvom iných analýz. Naopak, v niektorých prípadoch ide len o zjednodušené informovanie o pomeroch v dátach, najmä ak autori primárne používajú komplexnejšie analýzy a súvis medzi párami premenných len hrubo dopĺňa tieto komplexné výsledky.

7. Porovnávacie testy

V našich výskumoch sa často stretávame s potrebou porovnávať. Možno najčastejšie využívaným je test pre porovnanie dvoch skupín, no porovnáваме aj dve merania alebo našu hodnotu voči konkrétnej hodnote. V tejto časti si postupne ukážeme testy pre dva nezávislé výbery a testy pre dve párové merania.

Testy pre dva nezávislé výbery (angl. *independent samples*) využívame vtedy, ak potrebujeme zistiť, či sa líšia pomery v dvoch výberoch. Slovo výber je v tomto zmysle možno lepšie zameniť za skupina – ak porovnáваме dve skupiny v miere určitého konštrukt, ktorý je meraný na kontinuálnej úrovni. *Nezávislé* v tomto prípade vyjadruje, že respondenti sa v jednom a druhom výbere neopakujú. Ak sa na vec pozrieme skrz skupiny, dáva to vcelku zmysel – ak by sme porovnávali mužov a ženy, nie je možné, aby niekto označil, že je muž aj žena. V rámci názvu zas dáva zmysel, že možno chceme porovnať dva výbery z populácie, takže podmienkou je, aby boli od seba nezávislé, t.j. respondenti sa neopakujú v jednom a druhom výbere.

Testy pre párové merania (angl. *paired samples*) využívame v prípade, že chceme porovnať dve kontinuálne premenné na tom istom výbere. V tomto prípade sa nelíšia respondenti – tí sú rovnakí, no líšia sa premenné. Tieto analýzy používame najčastejšie v prípade, ak meriame nejaký konštrukt dvakrát v čase (longitudinálny dizajn), pred nejakou udalosťou či intervenciou a po nej (pretest a posttest) alebo dve premenné, ktoré sú merané rovnakým spôsobom no v inom kontexte: miera obľúbenosti sladkého alebo slaného pečiva, aktuálna životná spokojnosť verzsus očakávaná za rok, a iné.

Pri potulkách výskumnými štúdiami alebo programom jamovi sa môžete stretnúť aj s jednovýberovým testom (angl. *one sample*). Tomuto testu sa v tejto učebnici nebudeme detailnejšie venovať, no v prípade záujmu ho po prečítaní tejto kapitoly hravo zvládnete. Tento test používame, ak chceme premennú v našom výskume porovnať voči konkrétnej hodnote. Táto hodnota nepochádza z nášho výskumu – môže ísť o hodnotu nejakého kritéria alebo odhadovanú hodnotu v populácii. V tomto prípade pracujeme len s jedným výberom – tým našim, a porovnáваме ho voči nejakej konkrétnej hodnote, ktorá nepochádza z nášho dátového súboru. V jamovi ho nájdete v časti *T-Tests* pod názvom *One Sample T-Test*.

V závere tejto kapitoly sa oboznámime s porovnávacími testami pre 3 a viac skupín, ktorý sa nazýva jednofaktorová analýza rozptylu, v anglickej literatúre nazývaná *one-way ANOVA*. Tento test využijete v prípade, že by ste chceli nejakú kontinuálnu premennú porovnať medzi tromi alebo viacerými skupinami, napr. odbor štúdia či profesie, rôzne experimentálne skupiny a pod.

Predpokladom pre použitie spomenutých testov je kontinuálna úroveň merania porovnávaných premenných a v prípade parametrických testov aj ich normálna distribúcia. Pri porovnávaní dvoch výberov potrebujeme zistiť, či je táto premenná normálne distribuovaná v oboch skupinách, pri jednovýberovom len či je normálne distribuovaná v našom výbere. V prípade porovnávania dvoch meraní zistíme, či je *rozdiel* medzi prvou a druhou premennou normálne distribuovaný. Ak sa nám nepodari potvrdiť normálnu distribúciu, môžeme použiť neparametrickú alternatívu.

V nasledujúcich častiach si pre jednotlivé testy ukážeme, ako môžu vyzerat' hypotézy, ako ich overiť v jamovi a ako výsledky interpretovať a zapísať.

7.1 Porovnávanie dvoch skupín (výberov)

Pre skúmanie rozdielov medzi dvomi skupinami používame parametrický Studentov alebo Welchov t-test, prípadne neparametrický Mann-Whitneyho U-test. Ako sme už na začiatku spomenuli, pre parametrické testy je dôležité naplnenie podmienky normálnej distribúcie porovnáwanej premennej, ktorú nazývame *závislá* premenná, v oboch porovnávaných skupinách. Pokiaľ táto podmienka nie je naplnená, môžeme použiť neparametrickú verziu porovnávacieho testu. Zjednodušene povedané, základný rozdiel medzi t-testom a U-testom tkvie v tom, že parametrické testy pracujú s priemerom, zatiaľ čo neparametrické testy s priemerným poradím. Neparametrické testy teda nie sú závislé na normálnej distribúcii a zvládnu aj prípady odľahlých hodnôt. Nevýhodou je, že nepracujú s hodnotami ako takými, a teda ich priemerom a štandardnou odchýlkou, no len poradím hodnôt.

V prvom kroku si môžeme predstaviť hypotézy, ktoré by sme analyzovali prostredníctvom porovnávacích testov medzi skupinami. V prípade alternatívnych hypotéz vyjadrujeme, že existuje rozdiel v nejakej premennej medzi určitými skupinami. Podobne ako pri korelačnej analýze rozoznávame dvojstranné a jednostranné hypotézy. Pri dvojstranných očakávame, že pomery sa v jednej a druhej skupine líšia. Pri jednostranných hypotézach naše očakávanie špecifikujeme aj o to, pri ktorej skupine očakávame vyššie alebo nižšie hodnoty. Príkladom môže byť dvojstranná hypotéza:

H: Predpokladáme, že existuje rozdiel v [závislá premenná] medzi skupinou X a Y.

Nulová hypotéza v tomto prípade vyjadruje neprítomnosť rozdielu:

H_0 : Neexistuje rozdiel v [závislá premenná] medzi skupinou X a Y.

Alebo jednosmerná:

H: Predpokladáme, že skupina X dosahuje vyššiu úroveň [závislej premennej] v porovnaní s Y.

Pri ktorej nulová hypotéza vyjadruje:

H_0 : Skupina X nedosahuje vyššiu úroveň [závislej premennej] v porovnaní s Y.

Rovnako ako pri koreláciách bude výsledkom analýzy koeficient a jeho štatistická významnosť. V prípade parametrických testov je výsledkom buď to Studentov alebo Welchov koeficient t . V prípade dvojsmerného testu, nulová hypotéza hovorí, že t sa má v populácii rovnať 0. My však tieto skupiny porovnáme a zistíme že t sa nerovná 0. Štatistická významnosť nám povie, aká je pravdepodobnosť získať takú hodnotu t pri danom počte *stupňov voľnosti*, ak v cieľovej populácii platí nulová hypotéza. Ak je táto pravdepodobnosť veľmi nízka, resp. nižšia ako nami stanovaná hladina významnosti (konvenčne 0,05), pochybujeme o platnosti nulovej hypotézy a prijímame našu alternatívnu hypotézu. V predchádzajúcej vete sme použili pojem *stupne voľnosti* (angl. *degrees of freedom*, skratka DF alebo df). Ide o počet elementov, ktoré sa podieľajú na výpočte koeficientu a sú dôležité pre výpočet štatistickej významnosti koeficientu t . Našťastie pre nás, jamovi to všetko vypočíta za nás, a preto sa tým nemusíme až tak veľmi trápiť.

Ako príklad si overíme nasledujúcu hypotézu

H1: Predpokladáme, že ženy v porovnaní s mužmi dosahujú vyššiu mieru prívetivosti.

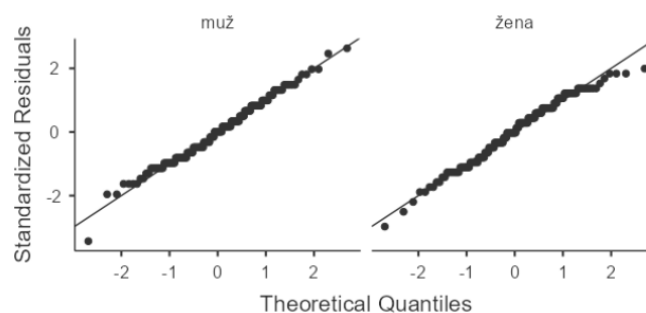
Na začiatok by sme mohli očistiť premennú Prívetivosť od odláhlych hodnôt. Postupujeme samostatne pre skupinu mužov a žien. Pre pripomenutie postupu si pozrite podkapitolu 5.4. Ako sme už spomenuli, v tejto učebnici pracujeme s dátami tak ako sú a odláhlé hodnoty neodstraňujeme. Pokračujeme overením normálnej distribúcie tejto premennej v jednotlivých skupinách. Postup je rovnaký ako v prípade korelačnej analýzy, s tým rozdielom, že použijeme aj funkciu *Split by*. V karte *Analyses* si zvolíme *Exploration* a *Descriptives*. Do okienka *Variables* si vložíme závislú premennú a do *Split by* dáme premennú, ktorá vyjadruje porovnávané skupiny. Pri nastavení analýzy si necháme prednastavené možnosti, no pridáme šikmosť a špicatosť (*Skewness* a *Kurtosis*), zapneme test normality (*Shapiro-Wilk*) a z grafov si dáme vykresliť Q-Q grafy (*Q-Q*). Výsledkom bude tabuľka obsahujúca deskriptívne údaje a výsledok Shapiro-Wilkovho testu, Q-Q graf. Všetky tieto informácie budú samostatne pre mužov a ženy, keďže tieto dve skupiny obsahuje premenná pohlavie. Pri hodnotení postupujeme rovnako ako v kapitole 6.1. Sledujeme najmä šikmosť distribúcie a Q-Q grafy. Keďže máme vyšší počet respondentov, výsledky Shapiro-Wilkovho testu môžu byť signifikantné aj napriek pomerne dobrej distribúcii. V Obrázku 7.1 uvádzame výsledky zo spustenej analýzy. Už v tomto momente môžeme vidieť, aké sú pomery medzi mužmi a ženami. Ak sme predpokladali, že ženy sú prívetivejšie v porovnaní s mužmi, očakávame vlastne, že priemer prívetivosti v skupine žien je vyšší ako u mužov. Stačí sledovať výsledok deskriptívnej analýzy a zistíme, či sme na dobrej ceste k overeniu hypotézy. V riadku *Mean* vidíme, že priemer prívetivosti v skupine mužov je 3,58, zatiaľ čo u žien je to 3,85 – čo je vyššia hodnota ako u mužov. Takže sme na dobrej ceste, lebo dáta ukazujú trend v súlade s našou hypotézou. Ale pozor! Samotné porovnanie priemerov NIE je dostačujúce k overeniu hypotézy! Pre overenie potrebujeme uskutočniť štatistický test, ktorý popisujeme v ďalšej časti.

Rozdielna situácia by nastala v prípade, ak by nám už výsledky deskriptívnej štatistiky ukázali opačný trend ako sme očakávali. Skúsme si predstaviť, že by muži mali vyššiu hodnotu priemeru prívetivosti v porovnaní so ženami. V takomto prípade automaticky vieme, že naša hypotéza H1 sa nepotvrdila – nezistili sme očakávaný trend. V takomto prípade, podobne ako sme popísali pri koreláciách, odporúčame zvoliť dvojsmerný test signifikancie.

Descriptives		
	pohlavie	Prívetivosť
N	muž	138
	žena	143
Missing	muž	0
	žena	0
Mean	muž	3.58
	žena	3.85
Median	muž	3.58
	žena	3.83
Standard deviation	muž	0.51
	žena	0.54
Minimum	muž	1.83
	žena	2.25
Maximum	muž	4.92
	žena	4.92
Skewness	muž	0.00
	žena	-0.29
Std. error skewness	muž	0.21
	žena	0.20
Kurtosis	muž	0.26
	žena	-0.38
Std. error kurtosis	muž	0.41
	žena	0.40
Shapiro-Wilk W	muž	0.99
	žena	0.98
Shapiro-Wilk p	muž	0.342
	žena	0.070

Obrázok 7.1 Výsledok šikmosti, špicatosti a Shapiro-Wilkovho testu pre premenné Extraverzia a Prívetivosť, samostatne pre mužov a ženy

Ako možno vidieť, jednotlivé výsledky sú samostatne pre mužov a ženy. Šikmosť a špicatosť je v oboch skupinách v rámci normy. Shapiro-Wilkov (Shapiro-Wilk p) test nám vyšiel štatisticky nevýznamný ($p > 0,05$) v oboch skupinách. Pre finálne rozhodnutie skontrolujeme aj Q-Q grafy, ktoré sú znázornené v Obrázku 7.2.



Obrázok 7.2 Q-Q grafy pre premennú Prívetivosť, samostatne pre mužov a ženy

Ako možno vidieť, v oboch prípadoch ležia takmer všetky body na priamke. Z hľadiska všetkých dôkazov sme sa rozhodli pre použitie parametrického testu – t-testu. Či vybrať Studentov alebo Welchov t-test posúdime na základe výsledku Levenovho testu homogenity rozptylov.

7.1.1 Testy pre dva nezávislé výbery a kontinuálne premenné v programe jamovi

Všetko, čo potrebujeme, nájdeme v karte *Analyses*, kde zvolíme *T-Tests* a vyberieme *Independent Samples T-Test*. Otvorí sa nám okno, ktoré zobrazujeme v Obrázku 7.3.

Obrázok 7.3 Okno pre nastavenie analýzy testu pre dva nezávislé výbery

Ako pri iných analýzach máme okienko, v ktorom sa nachádzajú všetky premenné v našom súbore. Vedľa nájdeme okienko *Dependent Variables*, kam vkladáme závislé premenné, teda tie, v ktorých chceme hľadať rozdiel. Pod ním je okienko *Grouping Variable*, do ktorého vložíme premennú, ktorá určuje dve skupiny. Dajte si pozor. Pokiaľ máte premennú, ktorá má viac ako dve hodnoty (viac ako dve skupiny) a vložíte ju do tohto okienka, jamovi vám vráti chybu – nevie totiž, ktoré skupiny má porovnať. Je preto potrebné, aby ste premennú transformovali alebo zapli filter tak, aby zostali len dve aktívne skupiny (vyfiltrujete respondentov, ktorí majú inú skupinu ako tie, ktoré chcete porovnávať). V časti *Tests* si môžeme vybrať, ktorý test chceme vypočítať.

V časti *Hypothesis* vyberáme, akú máme hypotézu. Dvojstranná hypotéza je nastavená v základe (*Group 1 ≠ Group 2*). Ak máme stanovenú jednostrannú hypotézu a výsledky deskriptívnej analýzy sú v súlade s našim očakávaním, volíme podľa našej hypotézy. Group 1 je skupina, ktorá má v premennej nižšiu hodnotu, Group 2 je tá, ktorá má vyššiu. V našom prípade Group 1 znamená muži (číselne ich máme označených ako 1) a Group 2 sú ženy (označené ako 2). Pri overovaní našej hypotézy H_1 nastavíme *Group 1 < Group 2*, keďže očakávame, že ženy majú v priemere vyššiu prívetivosť.

V časti *Missing values* nastavujeme, ako sa má jamovi vysporiadať s chýbajúcimi hodnotami závislej premennej. Pokiaľ naraz analyzujete viaceré závislé premenné, pri ktorých máte aj chýbajúce hodnoty, môžete nastaviť *Exclude cases analysis by analysis*, kedy analýza prebehne s počtom respondentov, ktorí majú prítomné hodnoty pri jednotlivých závislých premenných – počty respondentov sa tak môžu meniť od analýzy k analýze. Ak zvolíte *Exclude cases listwise*, budú respondenti, ktorí majú chýbajúcu hodnotu v čo i len jednej zo závislých premenných, vylúčení zo všetkých analýz, ktoré máte nastavené – počet respondentov bude rovnaký pri každej analýze.

Dôležitá časť je *Additional Statistics*, kde si môžeme zvoliť vypočítanie ďalších informácií. *Mean difference* nám vráti hodnotu rozdielu medzi priermi závislej premennej v jednej a druhej skupine. *Effect size* nám vypočíta hodnotu miery efektu, inak povedané praktickú veľkosť rozdielu – Cohenovo d pre parametrické testy a poradovo-biseriálny korelačný koeficient pre Mann-Whitneyho U-test. Zvolením *Descriptives* získame deskriptívnu štatistiku pre jednotlivé skupiny a závislé premenné (tú už ale poznáme ešte z exploračie dát, no neublíži, ak to máme „pokope“).

V poslednej časti *Assumption Checks* si môžeme zvoliť Shapiro-Wilkov test a zobrazenie Q-Q grafu pre posúdenie normality. My sme však tento krok už spravili jednotlivo pre každú skupinu zvlášť, takže tieto možnosti nepotrebujeme, no v prípade ak sme sa rozhodli pre t-test, musíme vybrať, ktorý použijeme. Rozhodnutie spravíme na základe výsledku Levenovho testu pre rovnosť rozptylov. Studentov t-test má predpoklad, že rozptyly hodnôt okolo priemeru sú v oboch skupinách rovnaké. Vzorec, prostredníctvom ktorého sa počíta Studentov t-test, obsahuje spoločnú štandardnú odchýlku. Ak by sme nepotvrdili predpoklad rovnosti rozptylov, volíme Welchov t-test, ktorého vzorec je prispôsobený pre túto situáciu. V jamovi zvolíme *Homogeneity test*, čím získame výsledok Levenovho testu. Ten má na pozadí nulovú hypotézu, ktorá hovorí, že testované skupiny majú rovnaké rozptyly. V prípade, že nám štatistická významnosť tohto testu vyjde vyššia ako 0,05, konštatujeme, že rovnosť rozptylov je dodržaná (nepochybujeme o nulovej hypotéze) a pre zisťovanie rozdielov medzi skupinami volíme Studentov t-test. V prípade, že by bol výsledok Levenovho testu štatisticky významný ($p < 0,05$), konštatujeme, že rovnosť rozptylov nie je dodržaná a volíme Welchov t-test.

Ako teda na to? V prvom rade pomaly. Stretávame sa s tým, že si študenti mýlia výsledky Shapiro-Wilkovho testu s Levenovym testom a ešte horšie s výsledkom t-testu. Shapiro-Wilkov a Levenov test *nehovoria nič o rozdieloch priemerov medzi skupinami*. Prostredníctvom nich sa len dopracovávame k tomu, ktorý test vybrať. V našom prípade sme usúdili dodržanie podmienky normality distribúcie pre použitie t-testu. V ďalšom kroku sledujeme výsledok Levenovho testu, ktorý nám povie, či je medzi skupinami dodržaná rovnosť rozptylov. Výsledok zobrazujeme v Obrázku 7.4.

Homogeneity of Variances Test (Levene's)				
	F	df	df2	p
Prívetivosť	0.87	1	279	0.351

Note. A low p-value suggests a violation of the assumption of equal variances

Obrázok 7.4 Výsledok Levenovho testu

Pri výsledku sledujeme štatistickú významnosť, ktorá je v tomto prípade vyššia ako 0,05 ($p = 0,351$). Ako sme už uviedli, pokiaľ je hodnota $p > 0,05$, hovoríme o dodržaní rovnosti rozptylov v porovnávaných skupinách a pre porovnanie skupín volíme Studentov t-test. V tomto momente si vypneme Levenov test, aby sme nemali zbytočne veľa výsledkov a môžeme ísť analyzovať. V časti *Tests* zvolíme *Student's*. V časti *Hypothesis* zvolíme jednosmerný test signifikancie. V H1 predpokladáme, že ženy sú prívetivejšie, takže zvolíme $Group\ 1 < Group\ 2$. Výsledky zobrazujeme v Obrázku 7.8. V tomto prípade je $p < 0,001$, teda je určite menšie ako 0,05, kedy považujeme výsledok za štatisticky významný. Pýtame sa: aká je pravdepodobnosť získať takýto rozdiel medzi skupinami pri tomto počte stupňov voľnosti, ak v populácii nie je rozdiel alebo majú muži vyššiu prívetivosť? H1 je jednostranná, preto nulová hypotéza zahŕňa aj opačný prípad voči tomu, ktorý sme vyjadrili v hypotéze. Pravdepodobnosť je nižšia ako 0,1 %, a preto pochybujeme o platnosti nulovej hypotézy a prijímame našu hypotézu. Hypotéza H1 sa potvrdila.

Independent Samples T-Test

Independent Samples T-Test						
		Statistic	df	p	Effect Size	
Prívetivosť	Student's t	-4.29	279.00	< .001	Cohen's d	-0.51

Note. H_a: μ muž < μ žena

Group Descriptives						
	Group	N	Mean	Median	SD	SE
Prívetivosť	muž	138	3.58	3.58	0.51	0.04
	žena	143	3.85	3.83	0.54	0.04

Obrázok 7.5 Výsledok Studentovho t-testu s mierou efektu a deskriptívnou štatistikou (prívetivosť a pohlavie)

Ak by sme nezistili rovnosť rozptylov, teda výsledok Levenovho testu by bol štatisticky významný ($p < 0,05$), zvolili by sme Welchov t-test. Výsledok je zobrazený v Obrázku 7.6. Ako možno vidieť, výsledok sa v tomto prípade v podstate nezmenil – Hypotéza H1 sa potvrdila.

Independent Samples T-Test						
		Statistic	df	p	Effect Size	
Prívetivosť	Welch's t	-4.29	278.90	< .001	Cohen's d	-0.51

Note. H_a: μ muž < μ žena

Obrázok 7.6 Výsledok Studentovho t-testu s mierou efektu (prívetivosť a pohlavie)

Ak by sme v prípade tejto hypotézy usúdili nedodržanie normálnej distribúcie závislej premennej v porovnávaných skupinách, mohli by sme zvoliť neparametrický Mann-Whitneyho test. Výsledok je zobrazený v Obrázku 7.7. Opäť sledujeme štatistickú významnosť testu (p) a hodnotíme, že H1 sa potvrdila.

Independent Samples T-Test					
		Statistic	p	Effect Size	
Prívetivosť	Mann-Whitney U	6975.50	< .001	Rank biserial correlation	0.29

Note. H_a: μ muž < μ žena

Obrázok 7.7 Výsledok Mann-Whitneyho U-testu s mierou efektu (prívetivosť a pohlavie)

7.1.2 Interpretácia a zápis výsledkov testov porovnania dvoch skupín

Podobne ako pri korelačnej analýze sa pri interpretácii výsledkov zaujímame o štatistickú významnosť ako aj mieru efektu, teda nejaké štandardizované vyjadrenie veľkosti rozdielu medzi skupinami. Toto sme získali vďaka tomu, že sme si dali vypočítať *Effect size*. Ako ste si v obrázkoch 7.5 a 7.6 mohli všimnúť, máme tam aj hodnotu Cohenovho *d*. Ide o mieru efektu, ktorá berie v úvahu veľkosť rozdielu priemerov v porovnávaných skupinách vzhľadom na štandardnú odchýlku premennej. Vďaka tomu je štandardizovaná a môžeme porovnávať veľkosť rozdielov s inými výskumami a stabilným (zaužívaným) spôsobom ju interpretovať. Vyjadruje, o koľko štandardných odchýlok sa skupiny medzi sebou líšia. V prípade Mann-Whitneyho U-testu používame poradovo-biseriálny korelačný koeficient *r* (*Rank biserial correlation*). Ten na rozdiel od Cohenovho *d* nevyužíva priemer, preto je vhodný pre situáciu, kedy sa nepotvrdí normálna distribúcia. Ako tieto hodnoty interpretovať? Podobne ako v prípade korelačného koeficientu, interpretácia je len približná, no zaužívané delenie uvádzame v tabuľke nižšie (Cohen, 1988):

Veľkosť rozdielu	Cohenovo d	Poradovo-biseriálny koeficient r
Zanedbateľný	< 0,2	< 0,1
Malý	0,2 – 0,5	0,1 – 0,3
Stredný	0,5 – 0,8	0,3 – 0,5
Veľký	> 0,8	> 0,5

Pri zápise výsledkov overovania hypotéz je dôležité informovať o tom, aký test sme použili a prečo. Ďalej zapíšeme, či sme našu hypotézu potvrdili alebo nie. Za tým zapíšeme výsledky testu a deskriptívu pre jednotlivé skupiny. Ak sme použili t-test, mali by sme uviesť hodnotu *t*, *df*, *p* a Cohenovho *d*. Formát zápisu t-testu je: $t_{(df)} = \text{„hodnota t“}$. Ku každej skupine uvedieme počet respondentov *N*, priemer *M* a štandardnú odchýlku *SD*. Ak máte záujem, môžete uviesť aj hodnotu mediánu. Ak sme použili Mann-Whitneyho U-test, zapisujeme hodnotu *U* a *p*. Ku každej skupine zapíšeme počet respondentov *N*, medián *Mdn* hodnoty prvého a tretieho kvartilu (*Q1* a *Q3*), ktoré používame pre deskripciu premenných, ktoré nemajú normálnu distribúciu. Tie si však musíme vypočítať zvlášť (*Explore a Descriptives*).. Pokiaľ chceme, môžeme zapísať aj priemer a štandardnú odchýlku.

Nižšie uvádzame možnú podobu zápisu výsledkov jednotlivých testov. Pri zápise sa riad'te tým, čo sme spomenuli vyššie. Konkrétna formulácia viet je už na vás, no domnievame sa, že kreatívne vyžitie si treba dopriať v iných častiach práce, tu je dôležité, aby bolo jasné, čo a ako ste zistili.

Studentov t-test:

Pre overenie hypotézy H1 sme použili Studentov t-test, keďže sme potvrdili normálnu distribúciu a rovnosť rozptylov prívetivosti v oboch skupinách. Hypotéza H1 sa potvrdila. Zistili sme, že ženy ($N = 143$; $M = 3,85$; $SD = 0,54$) dosahujú štatisticky

významne vyššiu mieru prívetivosti ($t_{(279)} = -4,29$; $p < 0,001$) v porovnaní s mužmi ($N = 138$; $M = 3,58$; $SD = 0,51$). Tento rozdiel je stredne veľký ($d = -0,51$).

Welchov t-test:

Pre overenie hypotézy H1 sme použili Welchov t-test, keďže sme potvrdili normálnu distribúciu, no nepotvrdili sme predpoklad rovnosti rozptylov prívetivosti v oboch skupinách. Hypotéza H1 sa potvrdila. Zistili sme, že ženy ($N = 143$; $M = 3,85$; $SD = 0,54$) dosahujú štatisticky významne vyššiu mieru prívetivosti ($t_{(278,90)} = -4,29$; $p < 0,001$) v porovnaní s mužmi ($N = 138$; $M = 3,58$; $SD = 0,51$). Tento rozdiel je stredne veľký ($d = -0,51$).

Mann-Whitneyho U-test:

Podmienka normálnej distribúcie testovanej premennej nebola v porovnávaných skupinách naplnená, preto sme pre overenie hypotézy použili Mann-Whitneyho U-test. Hypotéza H1 sa potvrdila. Zistili sme, že ženy ($N = 143$; $Mdn = 3,83$; $Q1 = 3,42$; $Q3 = 4,25$) dosahujú štatisticky významne vyššiu mieru prívetivosti ($U = 6975,50$; $p < 0,001$) v porovnaní s mužmi ($N = 138$; $Mdn = 3,58$; $Q1 = 3,17$; $Q3 = 3,92$). Tento rozdiel je malý ($r = 0,29$).

Ak máte nižší počet respondentov, na základe ktorých overujete hypotézu, môže sa stať, že vám aj *prakticky zaujímavý* výsledok vyjde štatisticky nevýznamný. Podobne ako pri koreláciách, to nič nemení na tom, že sa vám hypotéza nepotvrdila, no v rámci diskusie výsledkov môžete spomenúť napríklad to, že ste zistili malý až stredne veľký rozdiel, no nebol štatisticky významný. Pokiaľ zistíte zanedbateľný alebo veľmi malý rozdiel a nebol signifikantný, interpretujete to ako neprítomnosť rozdielu. Ak by ste zistili zanedbateľný rozdiel, ktorý by bol vďaka vysokému počtu stupňov voľnosti štatisticky významný, prihliadnite na jeho veľkosť pri interpretácii – síce sme zistili štatisticky významný rozdiel, no jeho veľkosť je prakticky nevýznamná/zanedbateľná.

Ak by ste mali viacero závislých premenných, ktoré porovnávame medzi rovnakými skupinami tým istým testom, môžeme číselné výsledky zapísať do tabuľky. V takomto prípade v texte uvedieme, že výsledky sú v danej tabuľke a do textu už hodnoty *neuvádzame*, no stále *musíme* uviesť slovné zhodnotenie, podobne ako v príklade vyššie. Nižšie uvádzame možný vzhľad tabuľky. Ak by ste mali v jednej tabuľke výsledky Studentovho a Welchovho t-testu, nie je to problém. Z popisu pri hypotézach to čitateľ pochopí, no môžete mu to pripomenúť aj poznámkou v poslednom riadku tabuľky.

Tabuľka X

Výsledky porovnania mužov a žien v prívetivosti a otvorenosti

Závislá premenná	Muži			Ženy			Výsledok t-testu			
	<i>M</i>	<i>SD</i>	<i>N</i>	<i>M</i>	<i>SD</i>	<i>N</i>	<i>t</i>	<i>df</i>	<i>p</i>	<i>d</i>
Prívetivosť	3,58	0,51	137	3,85	0,54	143	-4,29	279	< 0,001	-0,51
Otvorenosť	3,52	0,64	138	3,49	0,58	143	0,38	279	0,705	0,05

Poznámka.

Na záver znovu upozorňujeme, *nemýľte si štatistickú významnosť s veľkosťou rozdielu!*

7.2 Porovnávanie dvoch meraní

Táto časť je určená pre všetkých z Vás, ktorí sa zameriavate na zmenu v určitej kontinuálnej premennej u tých istých respondentov či participantov. Testy pre porovnávanie dvoch meraní najčastejšie nachádzajú uplatnenie pri výskumoch, ktoré sa zaoberajú zmenou:

- v čase – výskumník odmeria určitú premennú v čase 1 a v čase 2 (napr. o hodinu, týždeň či rok) a skúma, či nastala zmena
- po intervencii či nejakej udalosti – rovnako ide o zmenu v čase, no v tomto prípade sa v dizajne počíta s tým, že by zmena mala nastať na základe nejakej udalosti – napr. experimentálneho zásahu, intervencii, liečbe a pod.
- podľa špecifikácie – meriame to isté v rovnakom čase no s iným obsahom, napr. zaujímame sa o subjektívne hodnotenie vzťahu s matkou a otcom a porovnávame či existuje rozdiel

Pri porovnávaní dvoch meraní nás vlastne zaujíma, či je priemerný rozdiel medzi dvomi spojitými premennými dostatočne veľký na to, aby bol pri veľkosti nášho výberu štatisticky významný. Podobne ako pri testovaní rozdielov medzi dvomi skupinami si môžeme stanoviť dvojstrannú alebo jednostrannú hypotézu. V tomto prípade ale ťažšie uviesť nejakú všeobecnú formu (pravdou je, že nám žiadna všeobecná formulácia nenapadá), preto si ukážeme fiktívny príklad a veríme, že pochopíte princíp. V našom fiktívnom výskume sme sa zaujímali, či nastane zmena v pociťovanom strese u študentov po absolvovaní skúšky. V rámci konzistencie a zredukovania efektu ostatných premenných sme výskum uskutočnili pred a po skúške z toho istého predmetu v rámci jedného termínu. Aby sme respondentov príliš nezaťažili „vypytovaním sa“ (najmä pri čakaní na skúšku), pre meranie pociťovaného stresu sme použili jednoduchú otázku: „*Ako veľmi sa cítite byť stresovaný/á?*“. Respondenti odpovedali prostredníctvom 11 bodovej škály od 0 – *Vôbec necítim stres* po 10 – *Cítim neznesiteľný stres*. V rámci nášho výskumu predpokladáme, že nastane zmena, preto by sme mohli stanoviť dvojstrannú hypotézu:

H1: Predpokladáme, že nastane štatisticky významná zmena v pociťovanom strese u respondentov v čase pred a po skúške.

alebo

H1: Predpokladáme, že pociťovaný stres pred skúškou a po skúške sa bude štatisticky významne líšiť.

Môžeme však predpokladať, že po skúške stres zo študentov „opadne“, a preto nastane zníženie pociťovaného stresu:

H1: Predpokladáme, že pociťovaný stres po skúške bude štatisticky významne nižší ako pred skúškou.

Do nášho výskumu sa zapojilo 84 študentov, ktorí sa zúčastnili na skúške. Dáta máme, môžeme ísť overiť hypotézu. Na začiatok potrebujeme vedieť, či použiť parametrický (Studentov t) alebo neparametrický (Wilcoxonov) test. Rozhodnutie, ktorý test vybrať, robíme ako pri iných, už spomenutých analýzach. V tomto prípade však nekontrolujeme odľahlé hodnoty a neposudzujeme distribúcie samostatných premenných ale ich rozdielu. V našom príklade sa nebudeme zaujímať o distribúciu premenných *stres pred skúškou* a *stres po skúške*, no vytvoríme si novú premennú, ktorá vyjadruje rozdiel týchto dvoch premenných. V jamovi si

vytvoríme novú vypočítanú premennú (*New computed variable*), ktorú nazveme *rozdiel v strese* a zadáme tam prvú premennú, od ktorej odčítame druhú premennú (alebo naopak): *stres pred skúškou* – *stres po skúške*. Ak chceme odstrániť vzdialené hodnoty, pracujeme s touto premennou. Ak chceme vybrať, či použiť parametrický alebo neparametrický test, posudzujeme túto premennú. Pokiaľ nemáte v pláne riešiť vzdialené hodnoty a pre posúdenie normality vám stačí výsledok Shapiro-Wilkovho testu alebo Q-Q graf, nemusíte robiť nič a prejsť rovno k analýze v jamovi.

7.2.1 Porovnávanie dvoch meraní v programe jamovi

Ak chceme, môžeme si na začiatok vypočítať premennú vyjadrujúcu rozdiel medzi dvomi meraniami, tak ako sme uviedli v predchádzajúcom odseku. Túto premennú môžeme následne analyzovať – identifikovať a odfiltrovať vzdialené hodnoty, posúdiť distribúciu (šikmost, špicatosť, Shapiro-Wilkov test, Q-Q graf). Premenná v našom príklade neobsahuje žiadne vzdialené hodnoty a na základe komplexného posúdenia sme sa dostali k záveru, že môžeme použiť parametrický test, no v rámci interpretácie si ukážeme oba testy.

V jamovi nájdeme test pre porovnanie dvoch meraní v karte *T-Tests*, kde zvolíme možnosť *Paired Samples T-Test*. Otvorí sa nám okno, ktoré zobrazujeme v obrázku 7.8.

The screenshot shows the 'Paired Samples T-Test' dialog box in the jamovi software. The interface is divided into several sections:

- Variables:** On the left, a list of variables includes 'stres pred skúškou', 'stres po skúške', and 'rozdiel v strese'. An arrow points from this list to the 'Paired Variables' box on the right, which is currently empty.
- Tests:** This section contains checkboxes for 'Student's' (checked), 'Bayes factor', and 'Wilcoxon rank'. A 'Prior' value of '0.707' is entered next to the Bayes factor option.
- Additional Statistics:** This section includes checkboxes for 'Mean difference', 'Confidence interval' (set to 95%), 'Effect size', and another 'Confidence interval' (also set to 95%). There are also checkboxes for 'Descriptives' and 'Descriptives plots'.
- Hypothesis:** This section has three radio button options: 'Measure 1 ≠ Measure 2' (selected), 'Measure 1 > Measure 2', and 'Measure 1 < Measure 2'.
- Missing values:** This section has two radio button options: 'Exclude cases analysis by analysis' (selected) and 'Exclude cases listwise'.
- Assumption Checks:** This section at the bottom includes checkboxes for 'Normality test' and 'Q-Q Plot'.

Obrázok 7.8 Okno pre nastavenie analýzy rozdielu dvoch meraní

Podobne ako pri porovnávaní dvoch skupín tu nájdeme možnosť zvoliť parametrický test alebo neparametrický test v časti *Tests*. V časti *Hypothesis* volíme, či ide o dvojstrannú hypotézu (je rozdiel) alebo jednostrannú hypotézu: $Measure\ 1 > Measure\ 2$ vyjadruje, že prvá premenná v páre je väčšia ako druhá, a zas $Measure\ 1 < Measure\ 2$ vyjadruje opak (prekvapivo...) V časti *Additional Statistics* si môžeme zapnúť *Mean difference*, vďaka čomu budeme vedieť priemerný rozdiel (tú istú hodnotu získame vypočítaním priemeru vypočítanej premennej rozdielu). Čo nás určite zaujíma je *Effect size*, vďaka čomu budeme vedieť posúdiť praktickú veľkosť rozdielu. Zapnúť môžeme aj *Descriptives*. V prípade, že by sme vopred neriešili normalitu distribúcie, môžeme si zapnúť aj *Normality test* a *Q-Q Plot*. Je to však to isté, ako keď tieto analýzy spravíme na vypočítanej premennej rozdielu (v našom príklade *rozdiel v strese*). Aby sme spustili analýzu, musíme do okienka *Paired Variables* presunúť premenné, ktoré chceme porovnávať. Pri tejto analýze nebudeme mať premenné pod sebou, ale vedľa seba – v páre. Dajte si pozor, aby ste premenné vložili správne, inak sa vám analýza nespustí. Výsledok analýzy nájdete v Obrázku 7.9.

Paired Samples T-Test

Paired Samples T-Test									
			statistic	df	p	Mean difference	SE difference	Effect Size	
stres pred skúškou	stres po skúške	Student's t	6.093	83.000	< .001	1.071	0.176	Cohen's d	0.665

Descriptives					
	N	Mean	Median	SD	SE
stres pred skúškou	84	6.917	7.000	1.909	0.208
stres po skúške	84	5.845	6.000	2.045	0.223

Obrázok 7.9 Výsledok Studentovho t-testu pre dve merania

V tomto prípade bola pri analýze nastavená dvojstranná hypotéza. Rovnako ako pri predchádzajúcich analýzách, sledujeme hodnotu štatistickej významnosti p a porovnáваме túto hodnotu s nominálnou hodnotou, ktorá je zaužívaná 0,05. Vidíme, že $p < 0,001$. To znamená, že rozdiel medzi dvomi meraniami je štatisticky významný, tým pádom sa naša hypotéza potvrdila. Ak by sme mali stanovenú špecifickejšiu hypotézu, v ktorej by sme predpokladali, že pociťovaný stres sa po skúške zníži, potvrdila by sa nám, keďže priemerný stres pred skúškou je vyšší ako po skúške – túto informáciu by sme už zistili pri deskriptíve premenných, ale ak sme v nastavení zvolili *Descriptives*, nájdeme ju aj vo výsledku analýzy. Ak by ste pracovali s premennými, pri ktorých nemožno použiť parametrický test – napríklad by premenné neboli kontinuálne ale ordinálne, mali by ste veľmi nízky počet respondentov, alebo by nebola dodržaná normalita distribúcie rozdielu medzi meraniami – môžete použiť neparametrický Wilcoxonov test. Stačí, ak ho zvolíte v ponuke *Tests*.

7.2.2 Interpretácia a zápis výsledkov testov porovnania dvoch meraní

Interpretácia a zápis výsledkov je veľmi podobný ako v prípade porovnávaní dvoch skupín, ktoré nájdete v časti 7.1.2. Ako sme v predchádzajúcom odseku uviedli, pre vyhodnotenie výsledku najskôr sledujeme štatistickú významnosť. V našom prípade je rozdiel medzi meraniami pred a po skúške štatisticky významný. Druhá vec, ktorá nás nielen môže, ale aj má zaujímať, je praktická veľkosť rozdielu. Opäť ide o Cohenovo d v prípade parametrického testu

alebo poradovo-biseriálny korelačný koeficient v prípade neparametrického testu. Interpretácia je rovnaká ako pri porovnávaní dvoch skupín.

Zápis výsledku by mal obsahovať informáciu o tom, aký test ste použili, hodnoty koeficientov, deskriptívne údaje a vyjadrenie sa k potvrdeniu alebo nepotvrdeniu hypotézy:

Na základe dodržania podmienky normálnej distribúcie rozdielu medzi meraniami sme pre overenie hypotézy H1 použili Studentov t-test. Hypotéza H1 sa potvrdila. Zistili sme, že priemerný stres u študentov pred skúškou ($M = 6,92$; $SD = 1,91$) je vyšší ako po skúške ($M = 5,85$; $SD = 2,05$). Tento rozdiel je stredne veľký a štatisticky významný ($M = 1,07$; $t_{(83)} = 6,09$; $p < 0,001$; $d = 0,67$).

Ak by ste pre overenie hypotézy použili neparametrický Wilcoxonov test, zápis bude veľmi podobný, no uvediete aj medián, hodnoty prvého a tretieho kvartilu, a na miesto Cohenovho d uvediete poradovo-biseriálny koeficient r .

Na základe nedodržania podmienok pre použitie parametrického testu sme pre overenie H1 použili neparametrický Wilcoxonov test. Hypotéza H1 sa potvrdila ($W = 958,00$; $p < 0,001$). Zistili sme, že priemerný stres u študentov pred skúškou ($Mdn = 7,00$; $Q1 = 5,00$; $Q3 = 8,00$; $M = 6,92$; $SD = 1,91$) je štatisticky významne vyšší ako po skúške ($Mdn = 6,00$; $Q1 = 5,00$; $Q3 = 7,00$; $M = 5,85$; $SD = 2,05$). Tento rozdiel je veľký, poradovo-biseriálny koeficient $r = 0,85$.

Môžete si všimnúť, že pri použití oboch testov sa hypotéza potvrdila. Rozdiel je pri hodnotení veľkosti rozdielu – pri parametrickom teste sme konštatovali stredne veľký rozdiel, zatiaľ čo pri neparametrickom až veľký rozdiel. Podobne ako pri zápise výsledkov predchádzajúcich analýz je možné formulovať zápis výsledkov rôzne, dôležité však je, aby ste uviedli všetky potrebné informácie.

7.3 Jednofaktorová analýza rozptylu (porovnávanie 3 a viac skupín)

Analýza rozptylu, anglicky *analysis of variance*, v skratke *ANOVA*, patrí do skupiny často používaných analýz v psychologických výskumoch. Pri čítaní kvantitatívnych výskumných štúdií ste sa už pravdepodobne niekedy s týmto názvom alebo skratkou stretli (ak teda sledujete aj časť popisujúcu použité štatistické analýzy). V tejto učebnici si konkrétne ukážeme takzvanú jednoduchú alebo jednofaktorovú analýzu rozptylu, ktorá nachádza uplatnenie aj pri výskumoch v rámci záverečných prác. Pre tých z vás, ktorí máte záujem dozvedieť sa viac aj o pokročilejších analýzach rozptylu, odporúčame učebnicu od Špajdela (2020), ktorá je sprievodcom týchto analýz v programe SPSS, no ak pochopíte princíp, tieto analýzy môžete uskutočniť aj v programe jamovi.

Chceme porovnávať mieru kontinuálnej premennej medzi 3 alebo viacerými skupinami, prečo teda jednofaktorová analýza rozptylu? V úvode do bivariačnej štatistiky sme spomenuli, že skúmame vzťah dvoch premenných. V tomto prípade ide znovu o vzťah či súvis dvoch premenných – jedna z nich je kontinuálna a druhá je nominálna s tromi alebo viacerými kategóriami (skupinami), ktorá sa nazýva *faktor*, z čoho je odvodený názov jednofaktorová ANOVA. Jednoducho povedané, vďaka tejto analýze vieme overiť, či medzi dvomi takýmito premennými existuje štatisticky významný súvis alebo nie. Matematika za touto analýzou je zložitejšia, no tak ako v prípade porovnávania dvoch skupín, získame výsledok v podobe koeficientu, stupňov voľnosti a štatistickej významnosti koeficientu. Jednofaktorová analýza rozptylu nám však v základe neodpovie na to, či sú štatisticky významné rozdiely medzi všetkými skupinami, alebo medzi ktorými skupinami konkrétne sú rozdiely, odpovie nám na to, či existuje súvis medzi kontinuálnou a nominálnou premennou – zistíme, či sa nejakým spôsobom skupiny medzi sebou líšia. Pokiaľ získame štatisticky významný výsledok, vieme, že existuje súvis medzi premennými a môžeme pokračovať ďalej, porovnávaním jednotlivých skupín navzájom, pomocou takzvaných *post-hoc* testov. Tieto testy fungujú na princípe už spomínaných testov pre dve skupiny, no pri výpočte berú v úvahu fakt, že porovnáваме viacero skupín. Toto je dôležité z hľadiska zníženia rizika falošne signifikantných výsledkov – výpočet štatistickej významnosti týchto testov je prísnejší.

Hypotézy, pri ktorých využívame jednofaktorovú analýzu rozptylu, nemajú stanovený smer, predpokladáme v nich len súvis premenných. Uvádzame príklad možných hypotéz:

H: Predpokladáme, že existuje štatisticky významný súvis medzi [závislá premenná] a [faktor].

H: Predpokladáme, že existuje štatisticky významný rozdiel medzi skupinami [faktor] v miere [závislá premenná].

Pre ilustráciu uvádzame príklad (v podstate fiktívneho) výskumu, v ktorom sme sa zaujímali o to, či existuje rozdiel v sebaúcte u ľudí na základe ich rodinného stavu. Určite sa tejto problematike venovalo mnoho štúdií, no nám stačí, že máme dáta a môžeme si ukázať analýzu. Naším cieľom bolo potvrdiť hypotézu:

H1: Predpokladáme, že existuje rozdiel v miere sebaúcty jednotlivcov na základe ich rodinného stavu.

Konkrétne sme sa zaujímali o slobodných ľudí, ľudí vo vzťahu, ľudí v manželskom zväzku a rozvedených ľudí. Do nášho výskumu sa zapojilo celkovo 503 respondentov, z toho 135 (26,8 %) bolo slobodných, vo vzťahu bolo 165 (32,8 %) respondentov, 120 (23,9 %) bolo zosobašených a 83 (16,5 %) respondentov uviedlo, že sú rozvedení.

7.3.1 Jednofaktorová analýza rozptylu v programe jamovi

Na začiatku sme overili normalitu distribúcie premennej sebaúcta, samostatne pre jednotlivé skupiny (viď. kapitola 5.5). Na základe komplexného posúdenia – s prihliadnutím na kontinuálnu povahu premennej, šikmosť a špicatosť, Q-Q grafy, sme usúdili, že podmienka normálnej distribúcie je dodržaná a rozhodli sme sa použiť parametrický test.

V jamovi môžeme uskutočniť jednofaktorovú analýzu rozptylu prostredníctvom dvoch funkcií. Obe nájdeme v karte *Analyses* pod možnosťou *ANOVA*. Prvá funkcia, ktorú môžeme využiť sa priamo nazýva *One-Way ANOVA*, druhá má všeobecný názov *ANOVA*. Prostredníctvom druhej funkcie môžete spraviť aj pokročilejšie analýzy rozptylu. Výhodou prvej je jej čisté zamerania sa na jednofaktorovú analýzu rozptylu, no chýba v nej jednoduchá možnosť výpočtu mier efektu. V druhej táto možnosť je, no nemožno zvoliť koeficient pre prípad nedodržania podmienky rovnosti rozptylov. Najskôr začneme funkciou *ANOVA*, kde si overíme homogenitu rozptylov závislej premennej. Pokiaľ sa nám rovnosť rozptylov potvrdí, zostaneme v tejto funkcii. Pokiaľ nám Levenov test vyjde signifikantný, využijeme obe funkcie.

V jamovi si v karte *Analyses* otvoríme možnosť *ANOVA* a zvolíme *ANOVA*. Otvorí sa nám okno, ktoré zobrazujeme v Obrázku 7.10. V hlavnej časti opäť vidíme okienko s premennými v dátovom súbore a vedľa neho dve ďalšie. Do prvého s názvom *Dependent Variable* vkladáme závislú kontinuálnu premennú (v našom prípade je to Sebaúcta). Do okienka nižšie *Fixed Factors* vkladáme premennú, ktorá vyjadruje skupiny, takzvaný faktor. Nižšie si môžeme zvoliť rôzne miery efektu. Pri analýze rozptylu sa využívajú najmä dve miery efektu. Stretnúť sa môžeme s η^2 (eta na druhú, anglicky *eta squared*) alebo s ω^2 (omega na druhú, anglicky *omega-squared*). Pri štúdiu literatúry možno narazíte aj na pojmy ako epsilon na druhú (epsilon squared, ε^2), parciálna eta na druhú (partial eta-squared, η_p^2) či parciálna omega na druhú (partial omega-squared, ω_p^2). Tieto sa používajú pri pokročilejších analýzach rozptylu. Pre naše účely nám postačí ω^2 , ktorej využívanie odporúča aj Field (2018). V ďalších častiach sú rôzne nastavenia – v tomto momente sú pre nás dôležité časti *Assumption Checks* a *Post Hoc Tests*.

ANOVA

rodstav
sebaúcta

Dependent Variable

Fixed Factors

Model Fit

Overall model test

Effect Size

η^2 partial η^2 ω^2

> | Model

> | Assumption Checks

> | Contrasts

> | Post Hoc Tests

> | Estimated Marginal Means

Obrázok 7.10 Základný vzhľad funkcie ANOVA v jamovi

Pre začatie analýzy si do okienka *Dependent Variable* vložíme závislú premennú, v našom prípade „sebaúcta“. Do okienka *Fixed Factors* vložíme premennú „rodstav“. Automaticky sa vo výsledkovej časti objaví výsledky. Tie nás v tomto momente zatiaľ nezaujímajú. Otvoríme si časť *Assumption Checks* a zvolíme *Homogeneity test*. Týmto sa nám vypočíta Levenov test rovnosti rozptylov. Sledujeme jeho výsledok. Zvýšte však pozornosť. Výsledok analýzy rozptylu a Levenovho testu sa vcelku podobajú, uistite sa, že sledujete správnu tabuľku. Výsledok z príkladu nájdete v Obrázku 7.11.

Assumption Checks

Homogeneity of Variances Test (Levene's)			
F	df1	df2	p
2.476	3	499	0.061

Obrázok 7.11 Výsledok testu rovnosti rozptylov – Levenov test

V našom prípade sme získali signifikanciu Levenovho testu vyššiu ako 0,05 ($p = 0,061$), čo svedčí o tom, že rozptyly v skupinách nie sú štatisticky významne odlišné, preto môžeme pre výpočet použiť jednofaktorovej analýzy rozptylu využiť Fisherov koeficient. Ak by bol výsledok Levenovho testu štatisticky významný, bolo by vhodnejšie použiť Welchov koeficient, ktorý však cez túto funkciu nezískame (postup ukážeme neskôr). Po zhodnotení Levenovho testu ho vypneme. V tomto momente si v prvej časti funkcie *ANOVA* zapneme mieru efektu ω^2 a sledujeme výsledok hlavnej analýzy, ktorý nájdete v Obrázku 7.12.

ANOVA

ANOVA - sebaúcta

	Sum of Squares	df	Mean Square	F	p	ω^2
rodstav	502.127	3	167.376	7.355	< .001	0.037
Residuals	11356.346	499	22.758			

Obrázok 7.12 Výsledok jednofaktorovej analýzy rozptylu (Fisherov koeficient F)

V prvom rade sledujeme štatistickú významnosť. Ak je hodnota nižšia ako 0,05, konštatujeme, že efekt skupín je štatisticky významný. To znamená, že sme zistili štatisticky významný súvis týchto premenných, a teda miera kontinuálnej premennej sa líši v súvislosti so skupinami. V našom prípade je $p < 0,001$, to znamená, že sa naša hypotéza potvrdila. Štatisticky významný výsledok jednofaktorovej analýzy rozptylu však automaticky neznamená, že sa jednotlivé skupiny navzájom signifikantne líšia a nehovorí nám ani o tom, ktoré skupiny konkrétne sa medzi sebou líšia. Pokiaľ by bola štatistická významnosť vyššia ako 0,05, nebol by preukázaný štatisticky významný efekt, naša hypotéza by sa nepotvrdila a analýzu by sme v tomto kroku ukončili. My však máme signifikantný výsledok, a preto pokračujeme ďalej.

Je čas zistiť, ktoré skupiny sa medzi sebou líšia. V časti *Post Hoc Tests* nájdeme dve okienka. V jednom sa nachádzajú premenné, ktoré máme nastavené ako faktory. Keďže máme len jeden faktor, budeme tam mať len jednu premennú, ktorú presunieme do pravého okienka. Pri nastavení *Correction* ponecháme zvolenú možnosť *Tukey* a zapneme tiež možnosť *Cohen's d*, aby sme mali vypočítanú mieru efektu pri testoch rozdielov medzi jednotlivými skupinami. Týmto nastavením nám vo výsledkoch pribudne tabuľka, v ktorej nájdeme výsledky viacerých testov. Ich počet závisí od toho, koľko máme skupín. Ak máme 3 skupiny, získame 3 výsledky (prvá vs. druhá, prvá vs. tretia a druhá vs. tretia skupina). V našom prípade získame 6 výsledkov. Výsledok zobrazujeme v Obrázku 7.13.

Post Hoc Comparisons - rodstav

Comparison		Mean Difference	SE	df	t	P _{Tukey}	Cohen's d
rodstav	rodstav						
Slobodní	- Vo vzťahu	-2.448	0.554	499.000	-4.423	< .001	-0.513
	- Manželstvo	-1.567	0.599	499.000	-2.618	0.045	-0.328
	- Rozvedení	-2.283	0.665	499.000	-3.431	0.004	-0.479
Vo vzťahu	- Manželstvo	0.882	0.572	499.000	1.541	0.414	0.185
	- Rozvedení	0.166	0.642	499.000	0.258	0.994	0.035
Manželstvo	- Rozvedení	-0.716	0.681	499.000	-1.051	0.719	-0.150

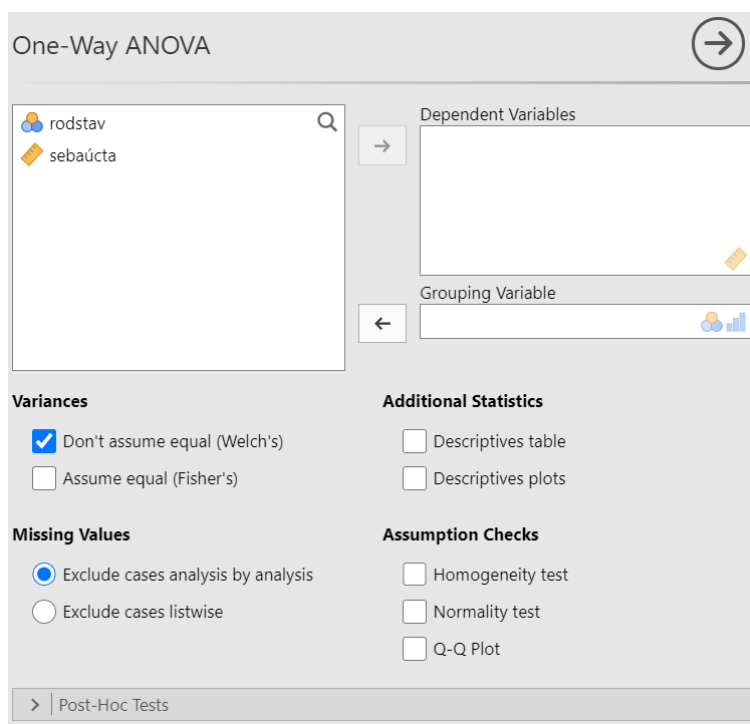
Note. Comparisons are based on estimated marginal means

Obrázok 7.13 Výsledky post-hoc testov s Tukeyho korektúrou štatistickej významnosti

Tieto výsledky už vieme poľahky zhodnotiť, keďže ide o porovnávanie dvoch skupín. Sledujeme hodnoty štatistickej významnosti (p_{Tukey}) v jednotlivých riadkoch, aby sme zistili, ktoré skupiny sa navzájom štatisticky významne líšia. Tam, kde je hodnota nižšia ako 0,05, konštatujeme štatisticky významný rozdiel. Môžete si všimnúť, že v našom príklade sme zistili štatisticky významný rozdiel v sebaúcte medzi slobodnými a tými vo vzťahu, manželstve alebo

rozvedenými, no nezistili sme štatisticky významné rozdiely pri ostatných porovnávaniach. Pri porovnávacích testoch odporúčame začať deskriptívou, aby sme vedeli, aké sú priemery v jednotlivých skupinách (kto má viac, kto má menej?), no dozvieme sa to aj z hodnoty Cohenovho d . Ak má prvá skupina nižšiu hodnotu závislej premennej, Cohenovo d bude negatívne, ak vyššiu, d bude pozitívne.

Späť o dva kroky – čo ak je výsledok Levenovho testu štatisticky významný (nie je dodržaná podmienka homogenity rozptylov)? Usmejeme sa (nie že by to bolo potrebné pre analýzu, ale že vraj to pomáha tak nejak všeobecne). Ako sme v úvode spomínali, funkcia *ANOVA* nevie vypočítať Welchov koeficient, ktorý je vhodnejší v prípade, že skupiny nemajú homogénne rozptyly. Už ju ale máme spustenú, takže si aspoň zapíšeme hodnotu miery efektu ω^2 a spustíme si funkciu *One-Way ANOVA* (Obrázok 7.14). Postup je veľmi podobný. Závislú premennú vložíme do okienka *Dependent Variables*, premennú definujúcu skupiny vložíme do okienka *Grouping Variable*. Môžete si všimnúť, že je v základe zvolená možnosť *Don't assume equal (Welch's)* v časti *Variances*. Toto presne chceme, keďže sme zistili, že rozptyly jednotlivých skupín nie sú podobné.



Obrázok 7.14 Základný vzhľad funkcie One-Way ANOVA v jamovi

Po vložení premenných sledujeme výsledok (Obrázok 7.15). Opäť nás zaujíma hodnota štatistickej významnosti a nič sa nemení na tom, že výsledok je štatisticky významný ak je $p < 0,05$. V tomto prípade nás už neprekvapí, že výsledok je štatisticky významný – zistili sme to pár riadkov dozadu, a teraz znovu len s použitím iného koeficientu. V prípade veľkého počtu respondentov v skupinách a veľmi nízkej hodnoty štatistickej významnosti sa oboma testami, resp. koeficientmi dostaneme k rovnakému výsledku. V rámci exaktnosti je však vhodné brať v úvahu potrebu rovnosti rozptylov pri Fisherovom koeficiente, a v prípade výraznej nerovnosti použiť Welchov koeficient – a to najmä v prípade nižšieho počtu respondentov v skupinách.

One-Way ANOVA (Welch's)				
	F	df1	df2	p
sebaúcta	6.264	3	252.018	< .001

Obrázok 7.15 Výsledok jednofaktorovej analýzy rozptylu (Welchov koeficient F)

Keďže sme zistili štatisticky významný efekt, tak ako predtým pokračujeme testami rozdielov medzi jednotlivými skupinami. V časti *Post-Hoc Tests* nájdeme Games-Howellov test, ktorý je podobný ako predtým použitý Tukeyho, no je prispôsobený pre nerovnosť rozptylov medzi skupinami. Zvolíme teda *Games-Howell (unequal variances)* a označíme aj ostatné možnosti v časti *Statistics*. Automaticky nám pribudnú výsledky testov, ktoré zobrazujeme v Obrázku 7.16. Vďaka hviezdikám ľahko vidíme, pri ktorých porovnávaniach boli zistené štatisticky významné rozdiely. Z hľadiska záveru nastala jedna zmena v porovnaní s Tukeyho post-hoc testom – rozdiel v sebaúcte medzi slobodnými a zosobašenými nevyšiel ako štatisticky významný.

Games-Howell Post-Hoc Test – sebaúcta					
		Slobodní	Vo vzťahu	Manželstvo	Rozvedení
Slobodní	Mean difference	—	-2.448 ***	-1.567	-2.283 **
	t-value	—	-4.121	-2.519	-3.427
	df	—	261.534	250.472	204.565
	p-value	—	< .001	0.059	0.004
Vo vzťahu	Mean difference		—	0.882	0.166
	t-value		—	1.634	0.281
	df		—	262.584	176.150
	p-value		—	0.361	0.992
Manželstvo	Mean difference			—	-0.716
	t-value			—	-1.159
	df			—	180.342
	p-value			—	0.654
Rozvedení	Mean difference				—
	t-value				—
	df				—
	p-value				—

Note. * p < .05, ** p < .01, *** p < .001

Obrázok 7.16 Výsledky post-hoc testov s Games-Howellovou korektúrou štatistickej významnosti

Ako možno vo výsledkoch vidieť, nemáme tu informáciu o praktickej veľkosti rozdielu, teda Cohenovom *d*. Musíme si ho vypočítať inak. Ak vás baví nastavovanie filtrov v jamovi, môžete si nastaviť filter tak, aby ste mali vždy aktívne len dve skupiny – tie, ktoré chcete porovnávať a postupne počítat' Welchove t-testy (kapitola 7.1.1), z ktorých si vezmete informáciu o Cohenovom *d*. Avšak pozor – vo výsledkoch zapisujeme hodnoty Games-Howellovho t-testu, z analýzy Welchovým t-testom si vezmeme len hodnotu Cohenovho *d*. A samozrejme, na konci nezabudnite vypnúť filter.

Druhým spôsobom je ručný výpočet prostredníctvom nasledujúceho vzorca:

$$d = \frac{M \text{ prvej skupiny} - M \text{ druhej skupiny}}{\sqrt{\frac{(SD \text{ prvej skupiny})^2 + (SD \text{ druhej skupiny})^2}{2}}}$$

Vieme si predstaviť tú radosť na vašej tvári! Jeden spôsob lepší ako druhý! Haha... Osobne by sme postupovali výpočtom cez jamovi s nastavovaním filtrov, no vieme si pomôcť aj Excelom či iným tabuľkovým procesorom. Otvorte si nový súbor tabuľkového procesora. Aby sme sa nemýlili, do prvého riadku si dajte samostatne do buniek M1, SD1, M2, SD2 a d. Do druhého riadku budete vkladať hodnoty, ktoré v jamovi jednoducho zistíte tak, že si pri nastavení funkcie *One-Way ANOVA* zapnete v časti *Additional Statistics* možnosť *Descriptives table*. Zaujímajú nás priemery (M) a štandardné odchýlky (SD) v jednotlivých skupinách. Do bunky E2 v tabuľkovom procesore vložte nasledujúci text:

=ROUND((A2-C2)/SQRT((B2^2+D2^2)/2);3)

Príklad, v ktorom sme zadali hodnoty M1 a SD1 zo skupiny „Slobodní“ a hodnoty M2 a SD2 zo skupiny „Vo vzťahu“ nájdete v Obrázku 7.17.

	A	B	C	D	E
1	M1	SD1	M2	SD2	d
2	29.067	5.503	31.515	4.609	-0.482

Obrázok 7.17 Nastavenie výpočtu Cohenovho d v programe Excel

Dajte si pozor na znamienko oddelujúce desatinné miesta. Jamovi používa bodku, no u nás používame čiarku. Ak máte slovenskú lokalizáciu tabuľkového procesora, použite čiarku, inak vám to nebude fungovať (my máme anglickú lokalizáciu, preto používame bodku, podobne ako v jamovi). Kto však chce, môže počítať ručne pre ten pocit nostalgie základ školských čias. Naši učitelia matematiky by boli na nás hrdí! Tak či tak, máme Cohenovo d pre prvé porovnanie, dopočítame ďalšie a môžeme ísť interpretovať a zapisovať výsledky.

7.3.2 Interpretácia a zápis výsledkov jednofaktorovej analýzy rozptylov

Tak ako pri iných testoch, aj pri jednofaktorovej analýze rozptylu uvádzame hodnotu koeficientu, stupne voľnosti, hodnotu štatistickej významnosti a mieru efektu. Hodnoty ω^2 môžeme zhruba interpretovať týmto spôsobom (Cohen, 1988):

- 0,010 až 0,059 – malý efekt
- 0,060 až 0,139 – stredne veľký efekt
- 0,140 a viac – veľký efekt

Ak sme zistili štatisticky významný výsledok, uvádzame aj výsledky jednotlivých post-hoc testov, kde znovu uvádzame tieto informácie. Cohenovo d interpretujeme rovnako, ako v prípade porovnávania dvoch skupín (viď. 7.1.2). Zápis výsledkov je rôzny. Pokojne môžete všetko zapísať v riadku, no ideálnejší je zápis časti výsledkov do tabuľky. V rámci nášho príkladu by mohol zápis vyzeráť napríklad nasledovne:

H1: Predpokladáme, že existuje rozdiel v miere sebaúcty jednotlivcov na základe ich rodinného stavu.

Pre overenie hypotézy H1 sme použili jednofaktorovú analýzu rozptylu, keďže sa naplnili predpoklady pre použitie parametrického testu. Na základe dodržania rovnosti rozptylov v porovnávaných skupinách sme zvolili Fisherov koeficient. Hypotéza H1 sa potvrdila. Zistili sme, že jednotlivci sa štatisticky významne líšia v miere sebaúcty na základe ich rodinného stavu ($F_{(3; 499)} = 7,36$; $p < 0,001$). Tento efekt je malý ($\omega^2 = 0,04$). Jednotlivé skupiny sme porovnávali prostredníctvom Tukeyho post-hoc testu. Zistili sme, že slobodní jednotlivci dosahujú štatisticky významne nižšiu mieru sebaúcty ako jednotlivci s partnerom, zosobášení a rozvedení jednotlivci. V prvom prípade ide o stredne veľký rozdiel, v ďalších dvoch o malý rozdiel. Deskriptívu premennej sebaúcta pri jednotlivých skupinách uvádzame v Tabuľke X. Výsledky post-hoc testov uvádzame v Tabuľke Y.

Tabuľka X

Deskriptíva premennej sebaúcta na základe rodinného stavu

Skupina	<i>N</i>	<i>M</i>	<i>SD</i>
Slobodní	135	29,07	5,50
Vo vzťahu	165	31,52	4,61
Manželstvo	120	30,63	4,42
Rozvedení	83	31,35	4,27

Poznámka.

Tabuľka Y

Výsledok Tukeyho post-hoc testov pre premennú sebaúcta medzi skupinami podľa rodinného stavu

		Tukeyho post-hoc testy			
		<i>t</i>	<i>df</i>	<i>p</i>	Cohenovo <i>d</i>
Slobodní	- Vo vzťahu	-4,423	499	<,001	-0,51
	- Manželstvo	-2,618	499	0,045	-0,33
	- Rozvedení	-3,431	499	0,004	-0,48
Vo vzťahu	- Manželstvo	1,541	499	0,414	0,19
	- Rozvedení	0,258	499	0,994	0,04
Manželstvo	- Rozvedení	-1,051	499	0,719	-0,15

Poznámka.

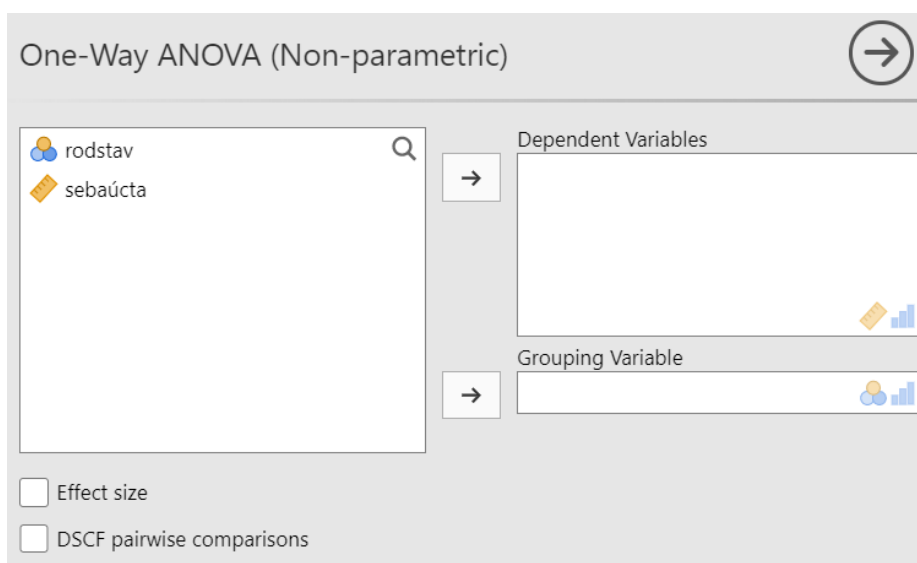
7.3.3 Neparametrická verzia jednofaktorovej analýzy rozptylu

V predchádzajúcej časti sme si ukázali jednofaktorovú analýzu, ktorá má ako predpoklad použitia normálnu distribúciu závislej premennej ako aj to, že táto závislá premenná je

kontinuálna. Čo však v prípade, že tieto podmienky nie sú naplnené? Tak ako aj pri porovnávaní dvoch skupín máme neparametrický Mann-Whitneyho U-test, tak aj pri jednofaktorovej analýze môžeme použiť Kruskal-Wallisov test, niekedy nazývaný aj neparametrická jednofaktorová ANOVA.

Ak ste ešte nezačali, prípadne neukončili zber dát, odporúčame vám, aby ste rátali s tým, že jednotlivé skupiny musia byť dostatočne zastúpené, jednak z hľadiska analýzy dát ale aj pre možnosť zovšeobecnenia vašich zistení na cieľovú populáciu. Ak sa dá, snažte sa mať dostatočný počet respondentov/participantov – pokiaľ je predpoklad, že vaša závislá premenná je konceptuálne normálne distribuovaná v cieľovej populácii, je v prípade vyššieho počtu respondentov vyššia šanca, že dáta budú normálne distribuované. Posudzovanie normality robte tiež komplexne, tak ako sme uviedli v časti o testovaní normality distribúcie. V praxi sa pri kvantitatívnych výskumoch, ktoré majú vyšší počet respondentov (rádovo desiatky, stovky) len málokedy stretávame s využívaním neparametrických testov, pokiaľ ich závislá premenná nie je priamo ordinálna (ide však len o subjektívny pohľad). Pokiaľ sa však po uvážení rozhodnete pre neparametrický test, ukážeme si ako na to. Ak ste pochopili princíp toho, čo sme riešili doteraz, nebude to pre vás žiaden problém.

V jamovi si otvoríme možnosť *ANOVA* a v časti *Non-Parametric* zvolíme *One-Way ANOVA* s podnadpisom *Kruskal-Wallis*. Otvorí sa nám pomerne jednoduché okno, v ktorom ľahko nastavíme analýzu (Obrázok 7.18).



Obrázok 7.18 Okno pre analýzu neparametrickej analýzy rozptylu (Kruskal-Wallisov test)

Podobne ako pred tým, do okienka *Dependent Variables* vložíme závislú premennú (v našom prípade „sebaúcta“) a do *Grouping Variable* vložíme premennú, ktorá definuje skupiny. Zvolíme si aj *Effect size* a v prípade, že zistíme štatisticky významný súvis premenných, zvolíme *DSCF pairwise comparisons*. Vďaka prvej možnosti získame informáciu o epsilon na druhú (ϵ^2), čo je ďalšia miera efektu (podobne ako omega na druhú). Vďaka druhej možnosti nám bude vypočítaný Dwass-Steel-Critchlow-Flignerov post-hoc test porovnávajúci jednotlivé skupiny navzájom. Vzhľad výsledkov nájdete v Obrázku 7.19.

V prvom rade sledujeme výsledok Kruskal-Wallisovho testu, v ktorom sa zameriame na štatistickú významnosť. Ak je $p < 0,05$, konštatujeme, že sa potvrdil štatisticky významný súvis

týchto dvoch premenných. V takomto prípade zhodnotíme aj praktickú veľkosť efektu. Delenie malého, stredného a veľkého efektu je rovnaké ako pri ω^2 , ktoré sme uviedli v predchádzajúcej časti. Následne môžeme spustiť post-hoc testy a sledujeme signifikanciu pri jednotlivých porovnávaniach. V našom prípade sme zistili štatisticky významný rozdiel medzi skupinou slobodných a ľudí vo vzťahu, slobodných a rozvedených. Ostatné porovnávanie neboli štatisticky významné. Záver je teda podobný tomu, čo sme zistili pri výsledku porovnávania prostredníctvom Games-Howellovho testu. Na záver potrebujeme zistiť mieru efektu týchto rozdielov – použijeme poradovo-biseriálny koeficient r . Asi najľahšou cestou je využitie automatického výpočtu pri porovnávaní dvoch skupín Mann-Whitneyho U-testom (nezabudnite, že pre výpočet porovnania dvoch skupín musíte nastaviť filter tak, aby ste pracovali len s dvomi skupinami). Pozor však, „požičiame“ si len mieru efektu, rozdiely medzi skupinami hodnotíme prostredníctvom post-hoc testu.

One-Way ANOVA (Non-parametric)

Kruskal-Wallis				
	χ^2	df	p	ε^2
sebaúcta	15.376	3	0.002	0.031

Dwass-Steel-Critchlow-Fligner pairwise comparisons

Pairwise comparisons - sebaúcta			
		W	p
Slobodní	Vo vzťahu	5.105	0.002
Slobodní	Manželstvo	2.917	0.166
Slobodní	Rozvedení	4.065	0.021
Vo vzťahu	Manželstvo	-2.189	0.409
Vo vzťahu	Rozvedení	-0.415	0.991
Manželstvo	Rozvedení	1.500	0.714

Obrázok 7.19 Výsledok Kruskal-Wallisovho testu s post-hoc testami

Zápis výsledkov je podobný ako pri parametrickej verzii. Pri výsledku Kruska-Wallisovho testu uvádzame χ^2 , df. Pri deskriptívnych hodnotách sa však prevažne zameriame na medián a *IQR* prípadne hodnoty *Q1* a *Q3*.

H1: Predpokladáme, že existuje rozdiel v miere sebaúcty jednotlivcov na základe ich rodinného stavu.

Pre overenie hypotézy H1 sme použili neparametrický Kruskal-Wallisov test, keďže sa nám nenaplnili podmienky pre použitie parametrického testu. Hypotéza H1 sa potvrdila. Zistili sme štatisticky významný súvis sebaúcty s rodinným stavom ($\chi^2 = 15,38$; $df = 3$, $p = 0,002$). Tento efekt je malý ($\varepsilon^2 = 0,03$). Zistili sme, že slobodní jednotlivci dosahujú štatisticky významne nižšiu mieru sebaúcty ako jednotlivci s partnerom a rozvedení jednotlivci. Tieto rozdiely sú malé. Deskriptívu premennej sebaúcta pri jednotlivých skupinách uvádzame v Tabuľke X. Výsledky post-hoc testov uvádzame v Tabuľke Y.

Tabuľka X*Deskriptíva premennej sebaúcta na základe rodinného stavu*

Skupina	<i>N</i>	<i>Mdn</i>	<i>Q1</i>	<i>Q3</i>	<i>M</i>	<i>SD</i>
Slobodní	135	29,00	26	34	29,07	5,50
Vo vzťahu	165	31,00	28	35	31,52	4,61
Manželstvo	120	30,50	28	34	30,63	4,42
Rozvedení	83	31,00	29	34	31,35	4,27

*Poznámka.***Tabuľka Y***Výsledok post-hoc testov pre premennú sebaúcta medzi skupinami podľa rodinného stavu*

			<i>W</i>	<i>p</i>	Poradovo-biseriálny koeficient <i>r</i>
Slobodní	-	Vo vzťahu	5,11	0,002	0,24
	-	Manželstvo	2,92	0,166	0,15
	-	Rozvedení	4,07	0,021	0,23
Vo vzťahu	-	Manželstvo	-2,19	0,409	0,11
	-	Rozvedení	-0,42	0,991	0,02
Manželstvo	-	Rozvedení	1,50	0,714	0,09

Poznámka.

8. Vzťah dvoch nezávislých nominálnych premenných

Doposiaľ sme spoločne zvládli základy inferenčnej štatistiky, korelácie dvoch kontinuálnych či ordinálnych premenných a porovnávanie skupín a meraní. V tejto poslednej časti učebnice sa dostávame k analýze, ktorá má za účel overiť vzťah dvoch nominálnych premenných. V literatúre je táto analýza nazývaná ako Pearsonov chí-kvadrát test (angl. *Pearson's chi-square test*). Názov vychádza z názvu koeficientu chí na druhú χ^2 , ktorý je však používaný pri viacerých komplexnejších analýzach, takže sa s ním môžete stretnúť aj mimo skúmania vzťahu dvoch nominálnych premenných, čomu sa venujeme v tejto časti.

Pre Pearsonov chí-kvadrát test je potrebné, aby boli premenné na sebe nezávislé, t.j. nejde o test-retest a každá kombinácia bola zastúpená aspoň 5 prípadmi. Program jamovi nám umožňuje aj analýzu opakovaného merania na úrovni nominálnych premenných, no tieto analýzy sú nad rámec tejto učebnice. Analýza funguje na princípe porovnania *očakávanej* frekvencie jednotlivých kombinácií s tým, ako sú jednotlivé kombinácie zastúpené. Čo sú to kombinácie a ich očakávaná frekvencia? Dve nominálne premenné, ktorých vzťah analyzujeme, môžu mať dve alebo viac kategórií. Uvedieme príklad, pri ktorom sú obe premenné dichotómne, teda majú len dve kategórie. Rozhodli sme sa skúmať, či sa muži a ženy líšia v tom, či obľubujú alebo neobľubujú parené buchty. Respondenti mali na výber pohlavie (muž/žena) a otázku „Máte rád/rada parené buchty?“ s možnosťou odpovede „Nie“ alebo „Áno“. Celkovo sa do výskumu zapojilo 100 respondentov. Úplnou náhodou sa do výskumu zapojilo presne 50 mužov a 50 žien, a ešte väčšia sranda je, že 50 ľudí uviedlo, že majú radi parené buchty a 50, že nie. Tieto počty sú vhodné pre najľahší možný príklad. Ak sa zamyslíte nad kombinatorikou, tak prídete na to, že môžeme získať 4 možné kombinácie: muž/áno, žena/áno, muž/nie, žena/nie. Ak by tieto dve premenné nemali žiaden súvis, očakávali by sme, že zastúpenie jednotlivých kombinácií bude rovnaké, teda frekvencia ich výskytu je totožná. Ak si spomeniete na nulovú hypotézu (ako je na ňu možné zabudnúť), výsledok frekvenčnej analýzy by podľa nej vyzeral takto:

Máte rád/rada parené buchty?			
Pohlavie	Nie	Áno	Celkovo
Muž	25	25	50
Žena	25	25	50
Celkovo	50	50	100

Ako vidíte, v každej kategórii sa nachádza rovnaký počet ľudí, to znamená, že rovnaký počet mužov má a nemá v obľube parené buchty, rovnaký počet žien má a nemá rado parené buchty, rovnaký počet ľudí, ktorí obľubujú buchty sú muži a ženy a rovnaký počet ľudí, ktorí nemajú radi parené buchty je mužov a žien. Tento príklad je samozrejme veľmi ideálny, v oboch premenných sú obe kategórie zastúpené presne pol na pol. Očakávaná frekvencia výskytu jednotlivých kombinácií by sa líšila v prípade, že by kategórie premenných neboli rovnako zastúpené v dátach – nemôžeme očakávať, že bude mať 25 mužov rado buchty a 25 mužov nemá rado buchty, ak by sa celkovo do výskumu zapojilo len 35 mužov. V prípade, že by sa do výskumu zapojilo 40 mužov a 60 žien a 43 ľudí by odpovedalo Nie a 57 Áno, očakávané frekvencie by vyzerali takto:

Pohlavie	Máte rád/rada parené buchty?		
	Nie	Áno	Celkovo
Muž	17,20	22,80	40,00
Žena	25,80	34,20	60,00
Celkovo	43,00	57,00	100,00

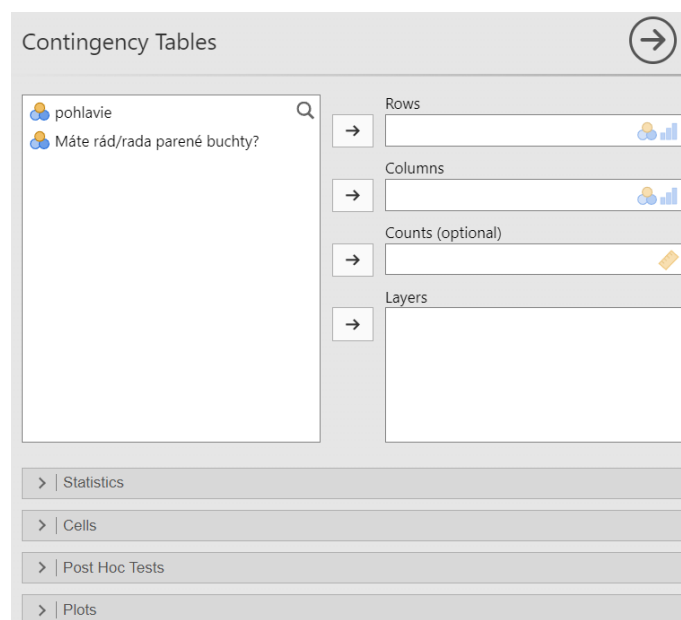
Výsledkom analýzy bude opäť koeficient (χ^2), stupne voľnosti a štatistická významnosť koeficientu. Pri analýze sa porovnáva frekvencia jednotlivých kombinácií v našich dátach s očakávanou frekvenciou. Koeficient vyjadruje, ako veľmi sa odlišuje očakávanie voči realite v našich dátach, a na základe počtu stupňov voľnosti je vypočítaná hodnota p , ktorá nám povie, či je táto odlišnosť štatisticky významná alebo nie. Nadviažeme na uvedený príklad a overíme nasledovnú hypotézu:

H1: Predpokladáme štatisticky významný vzťah pohlavia a obľuby parených buchiet.

Ako sme už uviedli, do výskumu sa zapojilo 40 mužov a 60 žien.

8.1 Chí-kvadrát test pre dve nezávislé nominálne premenné v programe jamovi

V jamovi tento test nájdete v ponuke *Frequencies*, časť *Contingency Tables*, kde zvolíme *Independent Samples χ^2 test of association*. Otvorí sa nám okno, ktoré zobrazujeme v Obrázku 8.1. Opäť v ňom nájdeme okienko s našimi premennými a štyri okienka. Najviac nás zaujímajú prvé dve. Do okienka *Rows* vkladáme premennú, ktorá bude umiestnená ako riadky a do *Columns* zas premennú, ktorá bude tvoriť stĺpce. Na základe tohto nastavenia nám bude vypočítaná frekvencia jednotlivých kombinácií. To ako premenné vložíte je v podstate na vašej preferencii, na výsledku to v základe nič nezmení. Ďalšie dve okienka nepoužívame: *Counts* slúži pre vloženie frekvencie kombinácií v prípade, že nemáme „surové“ dáta a *Layers* slúži na pridanie premennej, na základe ktorej bude analýza rozdelená na samostatné časti (vrstvy).



Obrázok 8.1 Základný vzhľad okna pre nastavenie chí-kvadrát testu dvoch nezávislých nominálnych premenných

V časti *Statistics* nájdete viacero možných nastavení (Obrázok 8.2).

The screenshot shows a 'Statistics' dialog box with the following sections:

- Tests:**
 - ☒ χ^2
 - ☐ χ^2 continuity correction
 - ☐ Likelihood ratio
 - ☐ Fisher's exact test
 - ☐ z test for difference in 2 proportions
- Hypothesis:**
 - ☒ Group 1 $\neq$ Group 2
 - ☐ Group 1 > Group 2
 - ☐ Group 1 < Group 2
- Nominal:**
 - ☐ Contingency coefficient
 - ☐ Phi and Cramer's V
- Comparative Measures (2x2 only):**
 - ☐ Odds ratio
 - ☐ Log odds ratio
 - ☐ Relative risk
 - ☐ Difference in proportions
 - ☒ Confidence intervals
 - Interval: %
 - Compare:
- Ordinal:**
 - ☐ Gamma
 - ☐ Kendall's tau-b
 - ☐ Mantel-Haenszel

Obrázok 8.2 Možnosti nastavenia rôznych koeficientov

Pri skúmaní vzťahu dvoch nominálnych premenných využívame možnosť χ^2 , ktorá je v základe nastavená. Pre niektorých z vás môže byť potrebný aj Fisherov exaktný test (*Fisher's exact test*), ktorý sa využíva v prípade, že niektorá/niektoré kombinácie sú zastúpené menej ako 5 prípadmi. Podmienkou použitia štandardného chí-kvadrát testu je to, aby sa každá kombinácia v dátach vyskytla aspoň 5-krát. V prípade nedodržania tejto podmienky je výsledok chí-kvadrát testu nespoľahlivý a je potrebné použiť Fisherov exaktný test, ktorý si vie poradiť aj s nízkou frekvenciou výskytu niektorých kategórií. Pre výpočet miery efektu, v tomto prípade sily vzťahu dvoch nominálnych premenných, si zvolíme možnosť *Phi and Cramer's V*. Ostatné nastavenia a analýzy sú nad rámec tejto učebnice, no záujemcom odporúčame voľne dostupnú knihu o analýzach v jamovi od Navarro a Foxcroft (2019), ktorá túto problematiku rozoberá detailnejšie.

V ďalšej karte *Cells* máme možnosť zvoliť si zobrazenie frekvencie kombinácii v našich dátach (*Observed Counts*) ale aj očakávanej frekvencie (*Expected Counts*). Nastavíme si tiež zobrazenie percent pre riadky, stĺpce a celok (označíme *Row*, *Column* a *Total*). V poslednej karte *Plots* môžeme nastaviť grafické znázornenie, to však nepotrebujeme.

Pri takomto nastavení získame 3 tabuľky s hodnotami. Prvá tabuľka obsahuje informáciu o frekvencii zastúpenia jednotlivých kategórií. V druhej tabuľke nájdeme výsledok chí-kvadrát testu a v poslednej nájdeme miery efektu (Obrázok 8.3).

Contingency Tables

Contingency Tables

pohlavie		Máte rád/rada parené buchty?		Total
		nie	áno	
muž	Observed	11	29	40
	Expected	17.20	22.80	40.00
	% within row	28 %	73 %	100 %
	% within column	26 %	51 %	40 %
	% of total	11 %	29 %	40 %
žena	Observed	32	28	60
	Expected	25.80	34.20	60.00
	% within row	53 %	47 %	100 %
	% within column	74 %	49 %	60 %
	% of total	32 %	28 %	60 %
Total	Observed	43	57	100
	Expected	43	57	100
	% within row	43 %	57 %	100 %
	% within column	100 %	100 %	100 %
	% of total	43 %	57 %	100 %

χ^2 Tests

	Value	df	p
χ^2	6.53	1	0.011
N	100		

Nominal

	Value
Phi-coefficient	0.26
Cramer's V	0.26

Obrázok 8.3 Výsledok chí-kvadrát testu dvoch nezávislých nominálnych premenných

V prvej tabuľke vidíme, ako sú zastúpené jednotlivé kategórie v riadkoch *Observed* ako aj očakávané zastúpenie v prípade platnosti nulovej hypotézy v riadku *Expected*. V ďalších riadkoch môžeme sledovať frekvenciu v percentách voči danému riadku, stĺpcu a celku. V našom príklade je očakávané, aby 17,2 mužov odpovedalo nie, no v našich dátach je ich nižší počet, len 11. Naopak mužov, ktorí odpovedali áno, je viac ako je očakávané. U žien je situácia opačná, viac žien, ako je očakávané, odpovedalo nie a menej odpovedalo áno. Voľnejšie povedané, muži majú radšej parené buchty ako ženy. Aby sme však mohli potvrdiť našu hypotézu, musíme sledovať štatistickú významnosť chí-kvadrátu – $p = 0,011$, čo je menej ako 0,05, teda súvis premenných je štatisticky významný a naša hypotéza sa potvrdila. Keďže ide o vzťah dvoch premenných, môžeme zhodnotiť aj jeho silu. Pri dvoch nominálnych premenných, ktoré majú len dve kategórie (tak ako v našom príklade), môžeme použiť phi alebo Cramerovo V, keďže oba koeficienty sú totožné. V prípade ak by ste pracovali s nominálnymi premennými, ktoré majú viac ako 2 kategórie, použite Cramerovo V. Interpretácia Cramerovho V sa líši v závislosti od počtu stupňov voľnosti Cramerovho V. Pozor, nemýľte si ich so stupňami voľnosti chí-kvadrát testu. Výpočet je ľahký, vezmeme si počet kategórii tej

nominálnej premennej, ktorá ich má najmenej a odpočítame 1. V našom prípade majú obe premenné 2 kategórie, čiže $2-1=1$. Ak by ste mali dve premenné, jedna by mala 4 kategórie, druhá 3, výpočet by vyzeral takto $3-1=2$. Interpretácia veľkosti efektu pre 1, 2, 3, 4 a 5 stupňov voľnosti je takáto:

Stupne voľnosti Cramerovho V	Malý efekt	Stredný	Veľký
1	0,10	0,30	0,50
2	0,07	0,21	0,35
3	0,06	0,17	0,29
4	0,05	0,15	0,25
5	0,04	0,13	0,22

8.2 Interpretácia a zápis výsledkov chí-kvadrát testu pre dve nezávislé nominálne premenné

Ako som už uviedol, pre informáciu o tom, či existuje súvis medzi premennými, musíme sledovať štatistickú významnosť. V rámci interpretácie a zápisu teda uvádzame hodnotu χ^2 , stupne voľnosti a hodnotu p . Pre zhodnotenie veľkosti efektu uvádzame Cramerovo V, ktoré interpretujeme na základe počtu stupňov voľnosti Cramerovho V. Pre úplnosť výsledkov je vhodné uviesť tabuľku, v ktorej uvedieme početnosť kombinácií v našom výskume ako aj očakávanú početnosť. V rámci interpretácie nezabudnite aj na popísanie informácií z tejto tabuľky. Príklad:

H1: Predpokladáme štatisticky významný vzťah pohlavia a obľuby parených buchiet.

Pre overenie hypotézy H1 sme použili Pearsonov chí-kvadrát test. Hypotéza H1 sa potvrdila, zistili sme štatisticky významný vzťah pohlavia a obľuby parených buchiet ($\chi^2_{(1)} = 6,53$; $p = 0,011$). Na základe Cramerovho V hodnotíme tento vzťah ako slabý ($V = 0,26$). Na základe posúdenia očakávaných a zistených frekvencií sme zistili, že muži vo všeobecnosti viac preferujú parené buchty v porovnaní so ženami. Očakávanú a zistenú frekvenciu jednotlivých kombinácií uvádzame v Tabuľke 1.

Tabuľka 1

Očakávaná a zistená frekvencia kombinácií pohlavia a obľuby parených buchiet

Pohlavie	Frekvencia	Máte rád/rada parené buchty?		Celkovo
		Nie	Áno	
Muž	Zistená	11	29	40
	Očakávaná	17,20	22,80	40
Žena	Zistená	32	28	60
	Očakávaná	25,80	34,20	60
Celkovo	Zistená	43	57	100
	Očakávaná	43	57	100

Poznámka.

Hypotéza sa potvrdila, výsledok sme zapísali. Pri tomto teste však stále ide o celkový vzťah premenných. Na základe zistenej a očakávanej frekvencie vidíme rozdiel, ktorý vieme interpretovať, no ak chceme hovoriť o štatistickej významnosti rozdielu medzi očakávaním a zistením, musíme si vypočítať Pearsonov reziduál, čiže štandardizovaný rozdiel medzi týmito frekvenciami. Vypočítame ho pomerne jednoducho:

$$\text{Pearsonov reziduál} = \frac{\text{zistená frekvencia} - \text{očakávaná frekvencia}}{\sqrt{\text{očakávaná frekvencia}}}$$

Ak by sme chceli zistiť, či sa štatisticky významne líši zistená a očakávaná frekvencia mužov, ktorí nemajú radi buchty, vzorec by vyzeral takto:

$$\text{Pearsonov reziduál pre muž/nie} = \frac{11 - 17,2}{\sqrt{17,2}} \cong -1,49$$

Ak je v absolútnej hodnote Pearsonov reziduál väčší ako 1,96, ide o štatisticky významný rozdiel pri $p < 0,05$, ak je väčší ako 2,58 tak $p < 0,01$, a ak je väčší ako 3,29 tak $p < 0,001$. V našom prípade rozdiel medzi očakávanou a zistenou frekvenciou nie je tak veľký, aby bol štatisticky významný. Takýto záver vyvodíme aj pri ostatných kombináciách: muž/áno (1,30), žena/nie (1,22), žena/áno (-1,06). Hoci je teda celkový vzťah štatisticky významný, nezistili sme štatisticky významné rozdiely v očakávanej a zistenej frekvencii pri jednotlivých kombináciách. Nemusíme smútiť, celkový záver – významný vzťah dvoch premenných platí, no tieto detailnejšie zistenia spolu so slabou silou vzťahu musíme mať na mysli pri diskusii zistení.

V prípade, že je súvis medzi premennými štatisticky významný, nemusíme Pearsonove reziduály počítať ručne, ale v časti *Post Hoc Tests* zvolíme *Pearson residuals* a Jamovi nám ich vypočíta. Pre lepšiu prehľadnosť dokonca zvýrazní reziduály, ktoré sú štatisticky významné pri približne $p < 0,05$ (ak chcete, môžete si hladinu upraviť na základe hodnôt uvedených vyššie). Výsledok zobrazujeme na Obrázku 8.4.

Post Hoc Test (Pearson Residuals)		
pohlavie	Máte rád/rada parené buchty?	
	nie	áno
muž	-1.49	1.30
žena	1.22	-1.06

Obrázok 8.4 Výsledok Pearsonových post-hoc testov v jednotlivých kategóriách

9. Záverečné odporúčania

Na záver, máme niekoľko (podľa nás) dôležitých postrehov, ktoré sme nazbierali za tých pár rokov vlastného výskumu a konzultácií so študentmi.

Čaro výskumu je pred zberom dát – dajte si záležať na príprave toho, čo idete zbierať. Váš výskum je len tak dobrý, ako dobré sú jeho časti – teoretický koncept, opodstatnenie predpokladov, metodológia a použité psychodiagnostické metodiky. Skúmať sa dá všeličo, no dôležité je, aby to dávalo zmysel (malo koncept) a výskum bol vhodne spravený. S odstupom času a najmä po zbere dát či pri diskusii výsledkov vám napadne zopár vecí, ktoré ste mohli spraviť inak, čo je prirodzené. Ak však zabudnete na niečo, čo je esenciálne, už to ľahko (alebo skôr vôbec) nenapravíte. Vopred majte pripravené a pevne stanovené ciele, výskumné hypotézy, metodiky ako aj to, ako budete dáta analyzovať. Vďaka tomu možno odhalíte nedostatky a napadnú vám veci, ktoré môžete vo vašom výskume zlepšiť.

Výskumný výber a cieľová populácia – váš výskum má priniesť zistenia pre určitú širšiu alebo užšiu vymedzenú cieľovú populáciu (napr. všeobecná populácia dospelých verzus členovia záchranárskych jednotiek). Pri užšiem vymedzenej populácii, povedzme záchranároch, nebudete zbierať dáta u pekárov – pomerne jasná chyba. No dajte si záležať na tom, aby ste pri zbere dát zachytili cieľovú populáciu, čo najviac sa dá. Pri všeobecnej populácii dospelých je potrebné osloviť a prizvať do výskumu rôznych ľudí z danej populácie. Najčastejším príkladom je, že hovoríme o širokej populácii na základe vysokoškolských študentov – pre výskum možno najprístupnejšej skupine ľudí. Takáto chyba je vytykána aj výskumom publikovaným v rôznych kvalitných vedeckých časopisoch. Dôležité je aj to, aby ste pri interpretácii a diskusii výsledkov zohľadnili, aký efekt na zovšeobecnenie zistení môžu mať charakteristiky vášho výskumného výberu.

Nepotvrdené predpoklady – možno sa vám stane, že niektoré predpoklady, teda vami stanovené hypotézy sa nepotvrdia. Dôkladne skontrolujte, či ste z technického hľadiska spravili všetko správne a ak áno, tak sa s tým musíte vysporiadať. Nebojte sa toho. V diskusii máte možnosť uviesť možné dôvody, prečo sa vám hypotéza nepotvrdila. Aj nepotvrdená hypotéza je zistenie! Najmä sa v návale paniky nedopustíte hrubého prešľapu – vytvárania nových „predpokladov“. Ak ste šli do zberu dát so štyrmi hypotézami, skončíte so štyrmi hypotézami. Nepridávajte hypotézy len kvôli tomu, aby sa vám niečo potvrdilo. Vaša záverečná práca nie je o potvrdených hypotézach, je to dôkaz o vašej schopnosti nazbierať a syntetizovať poznatky, formulovať výskumný problém, ciele, zvoliť správny metodologický postup, analyzovať dáta a diskutovať zistenia. Nikde nie je spomenuté niečo ako „minimálny počet potvrdených hypotéz“. Samozrejme, ak ste pri práci s dátami narazili na niečo zaujímavé, prípadne vám niečo napadlo, čo by stálo za preverenie, môžete to uviesť, no jasne vymedzte, že sú to dopĺňajúce zistenia nad rámec hypotéz.

Na číslach záleží – pri korelačnej analýze sme uviedli, aby ste si nemýlili signifikanciu s korelačným koeficientom, čo platí aj pri ostatných analýzach. Dajte si čas a uistite sa, že pozeráte na správne hodnoty. Stáva sa to na seminároch, skúške a nanešťastie aj v záverečných prácach. Študent si pomýli napr. hodnotu p s hodnotou r , a tak štatisticky nevýznamný, silou prakticky nulový vzťah ($r = 0,04$; $p = 0,532$) sa stane štatisticky významným, silným vzťahom...

Formálny zápis nie je „formalitka“ – v učebnici sme uviedli príklady zápisu výsledkov. Samozrejme je to len jedna z možností zápisu, konkrétna formulácia sa môže líšiť, no dôležité

je, aby váš zápis obsahoval všetky dôležité číselné hodnoty ako aj popis toho, čo tie hodnoty znamenajú. „Hypotéza H1 sa potvrdila. $r = 0,35$ a $p < 0,001$ “ nestačí.

Nebojte sa pokročilejších analýz – odporúčanie sa týka najmä študentov magisterského stupňa štúdia. Táto učebnica je základom, vďaka ktorému snád' naberie istotu v základoch a zároveň odvahu pokračovať ďalej. Pri diplomových prácach sa častejšie riešia komplexnejšie výskumné problémy, ktoré možno analyzovať komplexnejšími analýzami. Možnože vám názov „hierarchická lineárna regresia“ alebo „analýza kovariancie“ znie desivo, no ak ste pri štúdiu literatúry narazili na analýzu, ktorú by ste mohli využiť, „googlite“ – na internete určite nájdete množstvo zaujímavých návodov, no nebojte sa spýtať školiteľa alebo vyučujúceho.

Záver

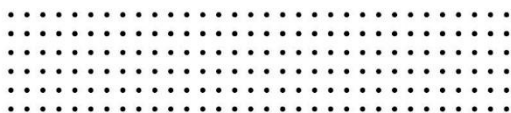
Kvantitatívny výskum je úzko spätý s analýzou dát, ktorá je vďaka pokroku v informačných technológiách čoraz dostupnejšia a užívateľsky priateľskejšia. Jedným z príkladov je program jamovi, ktorý je zadarmo dostupný pre kohokoľvek a zároveň ponúka širokú paletu štatistických analýz. Táto učebnica je určená pre úplných začiatočníkov v oblasti analýzy psychologických dát a práce v programe jamovi. Čitateľovi ponúka základný prehľad toho dôležitého pri analýze psychologických dát. Využitie nájde pri vyhodnocovaní dát pri kvantitatívnom výskume v rámci záverečných prác, no dúfame, že motivuje čitateľov k ďalšiemu rozširovaniu obzorov

Prajeme vám veľa príjemných chvíľ v programe jamovi, pri analýze dát a interpretácii výsledkov, aj keď vieme, že o tom s radosťou v budúcnosti nebudete hovoriť vašim vnúčatám...

Referencie

- American Psychological Association. (n.d.). Quantitative research. In APA dictionary of psychology. <https://dictionary.apa.org/quantitative-research>
- Ballová Mikušková, E. (2021). Štatistika pre začiatočníkov: základy štatistických analýz pre študentov učiteľstva. Nitra: Univerzita Konštantína Filozofa.
- Chajdiak, J., Rublíková, E., Gudába, M. (1994). Štatistické metódy v praxi. Bratislava: Statis.
- Cohen, J. (1988). Statistical Power Analysis for the Behavioral Sciences. 2nd. Lawrence Erlbaum.
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16(3), 297–334. <https://doi.org/10.1007/bf02310555>
- Field, A. (2018). Discovering statistics using IBM SPSS Statistics. London: Sage.
- Halama, P. (2011). Princípy psychologické diagnostiky. Trnava: Filozofická fakulta TU v Trnave.
- Halama, P., Špajdel, M., Žitný, P. (2013). Štatistika pre psychológov s využitím softwaru SPSS I. Deskriptívna štatistika. Trnava.
- Hanák, R. (2016). Dátová analýza pre sociálne vedy. Bratislava: Ekonóm.
- Hendl, J. (2006). Přehled statistických metod spracování dat. Analýza a metaanalýza dat. Praha: Portál.
- Ilievová, Ľ., Boroňová, J., Lajdová, A., Uričková, A., Žitný, P. (2017). Výskum v ošetrovatel'stve: krok za krokom. I. časť. Trnava: Typi Universitatis Tyrnaviensis a VEDA, vydavateľstvo SAV.
- Navarro, D. J., Foxcroft, D. R. (2019). learning statistics with jamovi: a tutorial for psychology students and other beginners. (Version 0.70). DOI: 10.24384/hgc3-7p15 [Available from url: <http://learnstatswithjamovi.com>]
- Špajdel, M. (2020). Sprievodca analýzou rozptylu v psychologickom výskume. Trnava: Filozofická fakulta TU v Trnave.
- The jamovi project (2021). *jamovi* (Version 1.6) [počítačový softvér]. Stiahnuté z: <https://www.jamovi.org>
- Walker, I. A. N. (2013). Výzkumné metody a statistika. Praha: Grada.

Filozofická fakulta TU v Trnave



Sprievodca základmi analýzy
psychologických dát
v programe Jamovi

2026©Michal Kohút