



Univerzita  
Hradec Králové  
Filozofická  
fakulta

# **Základy kvantitativního výzkumu:** pro studenty humanitních oborů

**Kamil Janiš ml.**

Hradec Králové 2022

Původní skripta vznikla v roce 2014 pro Fakultu věrných politik, Slezské univerzity v Opavě, a to pod názvem: *Úvod do problematiky výzkumu I – základy kvantitativního výzkumu: pro studenty oboru Sociální patologie a prevence* a recenzovali je: doc. PhDr. Juraj Kalnický, Ph.D.; doc. RNDr. Tomáš Kopf, Ph.D.; Mgr. et Mgr. Marta Kolaříková, Ph.D.; PaedDr. Bronislava Štěpánková, Ph.D. Původní verzi nebylo přiděleno ISBN a nevzniklo v rámci žádného grantového projektu.

Nynější verze je mírnou upravou verze původní. Skripta v aktualizované verzi jsou zveřejněna v elektronické podobě a bez jejich komercializace.

© Kamil Janiš ml., 2022

*Skripta neprošla jazykovou úpravou.*

# Obsah

|                                                                       |                                        |
|-----------------------------------------------------------------------|----------------------------------------|
| Úvod.....                                                             | 4                                      |
| <b>1 Druhy výzkumů .....</b>                                          | <b>5</b>                               |
| 1.1 Porovnání kvantitativního a kvalitativního výzkumu .....          | 8                                      |
| <b>2 Etické aspekty výzkumu .....</b>                                 | <b>10</b>                              |
| <b>3 Kvantitativní výzkum .....</b>                                   | <b>14</b>                              |
| 3.1 Etapy kvantitativního výzkumu .....                               | 16                                     |
| 3.2 Výběr respondentů v kvantitativním výzkumu .....                  | 22                                     |
| 3.3 Formulace a testování hypotéz .....                               | 25                                     |
| 3.3.1 Hypotéza věcná, statistická, alternativní a nulová .....        | 28                                     |
| 3.3.2 Testování hypotéz .....                                         | 30                                     |
| 3.3.2.1 Test dobré shody chí-kvadrát .....                            | 30                                     |
| 3.3.2.2 Test nezávislosti chí-kvadrát pro kontingenční tabulku.....   | 34                                     |
| 3.3.2.3 Test nezávislosti chí-kvadrát pro čtyřpolní tabulku .....     | 37                                     |
| 3.3.2.4 Studentův t-test .....                                        | 39                                     |
| 3.3.2.5 Pearsonův koeficient korelace .....                           | 46                                     |
| 3.3.3 Chyba I. a II. druhu .....                                      | 48                                     |
| 3.4 Metody sběru dat v kvantitativním výzkumu .....                   | 48                                     |
| 3.4.1 Dotazník .....                                                  | 48                                     |
| 3.4.1.1 Otázky v dotazníku .....                                      | 50                                     |
| 3.4.1.2 Podoba a struktura dotazníku .....                            | 55                                     |
| 3.4.1.3 Rady a tipy při konstrukci a využití vlastního dotazníku..... | 56                                     |
| 3.5 Prezentace a popis získaných dat .....                            | 57                                     |
| 3.5.1 Třídění dat.....                                                | 58                                     |
| 3.5.2 Číselný popis nominálních (kategoriálních) dat .....            | 59                                     |
| 3.5.3 Číselný popis rozložení dat .....                               | 64                                     |
| <b>Závěr .....</b>                                                    | <b>70</b>                              |
| <b>Úkoly.....</b>                                                     | <b>Chyba! Záložka není definována.</b> |
| <b>Slovník doplňujících pojmů k výzkumu .....</b>                     | <b>71</b>                              |
| <b>Použitá literatura .....</b>                                       | <b>74</b>                              |
| <b>Seznam příloh .....</b>                                            | <b>76</b>                              |
| <b>Seznam použitých vzorců .....</b>                                  | <b>77</b>                              |

## Úvod

Skripta, která právě držíte v ruce, si kladou za cíl přiblížit Vám problematiku kvantitativního výzkumu, a to zejména v oblasti pedagogických a příbuzných sociálních disciplín.

Pokud se studenti a studentky rozhodnou ve své praktické části bakalářských prací realizovat kvantitativní výzkumná šetření, měl by jim tento text výrazně usnadnit její realizaci. Stejně tak by měl zvýšit kompetence v dané oblasti, které uplatní v praxi, a to při realizaci „drobných“ průzkumů, kde sice nejsou rozhodující částí statistické výpočty, ale dobrá formulace otázek, a tím získání i platných výsledků a relevantní možnosti interpretace.

Skripta s názvem *Základy kvantitativního výzkumu: pro studenty humanitních oborů* rozhodně nemohou nahradit publikace od Jana Hendla, Miroslava Chrásky, Petera Gavory, Michala Miovskeho a jiných, ale dávají elementární vhled do dané problematiky, který je nutné prohlubovat, a to nejen v rámci samostudia, ale i dalšího studia.

To ovlivňuje i relativně úzký výběr použité literatury, který Vás má dle uvedených odkazů nasměrovat na původní zdroje.

Výběr jednotlivých kapitol vychází z praktických potřeb bakalářských a diplomových prací. V textu nejsou prezentovány příklady standardizovaných testů, které mají vlastní specifický způsob vyhodnocení (např. motorické testy apod.).

Některé odstavce jsou psány relativně neformálním jazykem, který autor používá pro explicitnější vysvětlení podávané informace. Věnujte náležitou pozornost i poznámkám pod čarou! Nejsou tam zbytečně.

Autor

# 1 Druhy výzkumů

## Cíle kapitoly:

- *Student/ka si osvojí základní druhy výzkumů.*
- *Student/ka pochopí rozdíly mezi kvantitativním a kvalitativním výzkumem.*

## Klíčová slova:

*výzkum; základní výzkum; aplikovaný výzkum; teoretický výzkum; empirický výzkum; monodisciplinární výzkum; interdisciplinární výzkum; transdisciplinární výzkum; kvantitativní výzkum; kvalitativní výzkum*

V této kapitole jsou prezentovány vybrané druhy výzkumů. Jejich pochopení je základním východiskem k dalšímu studiu, a to nejen těchto skript, ale jakékoliv literatury k tématu. Kapitola 1 je koncipována jako krátký výkladový slovník.

Prvním termínem, který je nutné vymezit, je samotný termín **výzkum**. Gavora a kol. (2010)<sup>1</sup> definuje výzkum jako: „*souhrnný název pro vědeckou činnost. Je to aktivita specializovaných odborníků s příslušným vzděláním a kvalifikací s cílem budovat vědecké teorie. Při výzkumu se používají metody na registraci, zpracování a vyhodnocení zkoumaných jevů.*“

Vrána (1948, s. 154) definuje výzkum jako: „*původní hledání pravdy přímo z pramenů poznání, tj. bádání na základě pozorovaných, analyzovaných a tříděných faktů.*“

---

<sup>1</sup> <http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/veda-a-vyskum/vyskum.php?id=i1p1>

Průcha (2000, s. 181) prezentuje následující druhy pedagogických výzkumů:

Tab. č. 1 – Druhy pedagogických výzkumů

| <b>Druhy pedagogického výzkumu</b>  |                             |
|-------------------------------------|-----------------------------|
| <b><i>Rozlišovací kritérium</i></b> | <b><i>Druhy výzkumu</i></b> |
| stupeň obecnosti – konkrétnosti     | základní                    |
|                                     | aplikovaný                  |
| vztah k objektivní skutečnosti      | teoretický                  |
|                                     | empirický                   |
| základní paradigma                  | kvantitativní               |
|                                     | pozitivistický kvalitativní |
| zaměřenost na způsob využití        | akční výzkum                |
|                                     | strategicko-koncepční       |
| komplexnost objasňování             | monodisciplinární           |
|                                     | interdisciplinární          |
|                                     | transdisciplinární          |
| účelovost                           | deskriptivní                |
|                                     | diagnostický                |
|                                     | explorativní                |
|                                     | evaluační                   |
| prostředí realizace výzkumu         | laboratorní                 |
|                                     | terénní                     |
| výzkumné přístupy a metody          | experimentální              |
|                                     | observační                  |
|                                     | longitudinální/průřezový    |
|                                     | komparativní                |
|                                     | historický                  |
|                                     | prognostický                |
|                                     | interkulturní               |

Z množství uvedených druhů logicky vyplývá, že se s jednotlivými druhy setkáváme v rozdílné intenzitě.

Níže jsou specifikovány vybrané druhy pedagogických výzkumů (v obecné rovině se pochopitelně jedná i o výzkumy v jiných oborech). Průcha (2000, s. 182-187)

- **Základní výzkum** – jeho smyslem je v teoretické rovině objasňovat pedagogické problémy a oblasti. Výsledky základního výzkumu nemají bezprostřední využití pro praxi.
- **Aplikovaný výzkum** – výsledky aplikovaného výzkumu mají využití v oblasti pedagogické praxe.
- **Teoretický výzkum** – jedná se o výzkum, který využívá pouze teoretické metody. Zpravidla nepracuje s konkrétními daty. Pokud se zabývá pedagogickou realitou, tak na ní nahlíží z teoretického hlediska a výsledkem jsou teoretické závěry.
- **Empirický výzkum** – zkoumá se vždy nějaký konkrétní jev pedagogické reality. Jedná se o živé i neživé objekty. Výsledkem jsou konkrétní poznatky.
- **Monodisciplinární výzkum** – zaměřuje se na zkoumaný problém teoriemi a nástroji jedné vědecké disciplíny.
- **Interdisciplinární výzkum** – již z názvu vyplývá, že se na zkoumaný problém zaměřuje pohledem (teoriemi, přístupy, nástroji) dvou či více disciplín.
- **Transdisciplinární výzkum** – využívá se poznatků (teorií apod.) ze všech vědeckých disciplín, které se ke zkoumané problematice vztahují.

Nejčastěji však výzkumy rozlišujeme podle zvoleného paradigmatu na – kvalitativní a kvantitativní.

To však může být limitujícím vnímáním pro zpracování empirických částí kvalifikačních prací studentů. Znalost dalších druhů může dopomoci k širšímu náhledu na tematické zaměření výzkumu a tím dospět i k novým zjištěním.

### 1.1 Porovnání kvantitativního a kvalitativního výzkumu

V níže uvedené tabulce jsou přehledně uvedeny rozdíly mezi kvantitativním a kvalitativním výzkumem.

Tab. č. 16 – Srovnání kvantitativního a kvalitativního výzkumu

| Kritérium                        | Výzkum                                                                                              |                                                                                                     |
|----------------------------------|-----------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------|
|                                  | Kvantitativní                                                                                       | Kvalitativní                                                                                        |
| <b><i>Filozofické kořeny</i></b> | pozitivismus                                                                                        | fenomenologie                                                                                       |
| <b><i>Cíle výzkumu</i></b>       | zevšeobecnění, vysvětlení jevu, získání objektivního důkazu, kontrolované prostředí, ověření teorie | získání vhledu, porozumění jevu, smyslu, chování lidí v přirozeném prostředí, vytvoření nové teorie |
| <b><i>Vztah k teorii</i></b>     | potvrzení, či vyvrácení teorie                                                                      | tvorba teorie                                                                                       |
| <b><i>Výzkumný charakter</i></b> | snaha o objektivnost                                                                                | subjektivita                                                                                        |



|                                          |                                                                                                                                                                                             |                                                                                                                                                       |
|------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Počet zkoumaných osob</b>             | preference velkého množství, reprezentativní vzorek                                                                                                                                         | klient, student, skupina, zařízení jsou reprezentativní ve smyslu specifčnosti                                                                        |
| <b>Postoj výzkumníka</b>                 | odstup od zkoumaných osob, jevů, reality; zobecnění                                                                                                                                         | vstup do reality, kontakt s participujícími osobami, jedinečnost, vcítění se                                                                          |
| <b>Plánování výzkumu</b>                 | na začátku důkladná, písemná příprava projektu podle dané struktury                                                                                                                         | pružné přizpůsobování, plán vzniká v průběhu práce, lze změnit zkoumané otázky a metody                                                               |
| <b>Průběh výzkumu</b>                    | plánovitě ověřování hypotéz a zjišťování kauzálních vztahů                                                                                                                                  | shromažďování velkého množství údajů o sledovaných jevech a jejich kontextech, ty se zaznamenávají a interpretují                                     |
| <b>Techniky a metody</b>                 | dotazník, testy, standardizované pozorování, strukturované interview, experimenty (manipulace s proměnnými) aj.                                                                             | dlouhodobý terénní výzkum, nestrukturované pozorování s různou mírou zúčastněnosti                                                                    |
| <b>Zpracování dat</b>                    | kvantitativní, počítačové, statistické – snaha o objektivní zobecnění a interpretaci dat                                                                                                    | kvalitativní kódování, analýza a interpretace na základě subjektivního porozumění                                                                     |
| <b>Výsledky, podoba závěrečné zprávy</b> | zobecnění výsledků na celý základní soubor (populaci), zjištění zákonitostí, výstižná a přehledná výzkumná zpráva (obsahuje výzkumný problém, metodologii, analýzu dat, diskusi k výsledkům | vysvětlování chování lidí v určitém kontextu s platností pro danou skupinu, jedince, zařízení; detailní, interpretativní nebo jen deskriptivní zpráva |

(Žumárová in Skutil a kol., 2011, s. 74-75)

## 2 Etické aspekty výzkumu

### Cíle kapitoly:

- *Student/ka dovede získané znalosti z oblasti etiky aplikovat při realizaci vlastního výzkumu.*

### Klíčová slova:

*etika; etické principy*

Etická rovina výzkumu neprobíhá pouze v přístupu k respondentům či jiným neživým subjektům. Jedná se i o rovinu přístupu k sobě samému, jiným odborníkům v daném oboru (ti, kteří budou náš výzkum číst) a v jisté souvislosti i širší laické veřejnosti.

Průcha (in Skutil a kol., 2011, s. 24-36) uvádí několik etických principů, které rozděluje na etické principy při přípravě výzkumu, při realizaci a při publikování výsledků výzkumu.

### Etické principy při přípravě výzkumu:

- **Výběr zkoumaných objektů nebo respondentů** – abychom docílili validity získaných dat, která je dána reprezentativností<sup>2</sup> objektů či respondentů, je zejména u kvantitativního výzkumu nutné zvolit správný druh výběru daných prvků. Reprezentativnosti nejlépe docílíme náhodným či stratifikovaným výběrem, a to zejména pokud chceme zkoumat nějakou realitu s celorepublikovou, krajskou aj. působností. Odpovídající (minimální) počet respondentů je možno stanovit (viz Chráska, 2007, s. 23-26). Někdy se však musíme z objektivních důvodů uchýlit k jinému druhu výběru. Výzkumník však musí na tento fakt upozornit, tím je dán impuls čtenáři výzkumu, že získaná data mají pravděpodobně nižší validitu<sup>3</sup>.

---

<sup>2</sup> Viz slovník doplňujících pojmů

<sup>3</sup> I když se u bakalářských prací nepředpokládá, že by počet respondentů dosahoval několika tisíc či stovek (od 500 výše), je dodržení výše uvedeného etického principu důležité. Pokud se např. chceme zabývat problematikou záškoláctví ve vybraném kraji a dotazníkové šetření provedeme pouze v jednom městě, a to z důvodu, že tam bydlíme, tak se nejedná o objektivní příčinu. Buď bychom měli

- **Zachování anonymity** – pokud se snažíme od respondentů získat nejpravdivější informace, zpravidla u dotazníku a výzkumu, je nutné jim zaručit anonymitu. Je nutné je o zachování anonymity informovat předem, ujistit je, že získané odpovědi slouží pouze pro daný účel. Ačkoliv se to může zdát jako samozřejmost, tak nesmíme zapomenout, že v některých případech je jistá možnost zpětné identifikace respondentů či organizace, kde výzkum probíhal. Průcha správně upozorňuje, že sice dodržíme anonymitu, ale v závěru poděkujeme řediteli daných škol apod. Někdy se výzkumník dostává do situace, kdy je etické zachovat anonymitu skupiny respondentů, ale zároveň je žádoucí daný soubor, co nejlépe charakterizovat. Zbytečně bychom také neměli zjišťovat osobní informace, které nepotřebujeme, resp. nikterak s nimi nechceme pracovat<sup>4</sup>.
- **Informovaný souhlas zkoumaných subjektů**<sup>56</sup> – znamená, že subjekty dobrovolně souhlasí se začleněním do výzkumu a zároveň znají záměry výzkumu, cíle výzkumu a využití získaných dat. Švaříček, Šedřová a kol. (2007, s. 46) poznamenávají, že někteří výzkumníci záměrně zatajují některé informace před respondenty, a to za účelem zvýšení validity získaných dat. Např. „výzkum na téma šikana je prezentován jako výzkum komunikace žáků.“ Prezentovaný příklad má relativně racionální základy, které mohou spočívat ve snaze prezentovat se lépe než je daná skutečnost. Ovšem stejní autoři se správně ptají: „*jestli takovému výzkumu můžeme věřit.*“ Informovaný souhlas může mít nejen podobu písemnou, ale i ústní či může být

---

problematiku záškoláctví zaměřit pouze na oblast vybraného města (a to již na počátku), nebo rozšířit počet respondentů o další města (např. prostřednictvím vícenásobného výběru).

<sup>4</sup> V současnosti se nabízí možnost využití tzv. online dotazníků. Tento způsob má však více záporů než kladů. Těžko docílíme reprezentativnosti. Nevíme, zda dotyčný splňuje kritéria apod. Podle autora má použití online dotazníku smysl pouze v případě, kdy zkoumaným souborem jsou jedinci, které nemůžeme ve společnosti lehce identifikovat. Může se jednat i o uzavřenou komunitu či lidi, kteří tráví svůj volný čas nějakým „zvláštním“ způsobem (např. geocaching). Mohou to být lidé, kteří se sdružují na serverech zaměřených na sebevraždy atd.

<sup>5</sup> Švaříček, Šedřová a kol. (2007) zmiňují ještě tzv. předpokládaný souhlas, který je v pedagogické realitě běžný. Jedná se o to, že ředitel/ka školy povolí vstup výzkumníkovi do školy a povolí mu realizaci výzkumu u žáků či učitelů.

<sup>6</sup> K výše uvedené poznámce je však nutné doplnit, že se opět částečně pohybuje na hraně etiky výzkumu, a to zejména u žáků a studentů mladších než 18 let. V podstatě byl souhlas vydán někým jiným než jejich zákonnými zástupci.

pořízen nějaký audiozáznam. Vždy záleží na povaze výzkumu (tématu výzkumu) atd.

Hendl (2008, s. 154) k zatajování cílů a okolností výzkumu doplňuje, že pokud se již rozhodujeme něco zatajit, měli bychom postupovat následovně:

- *„kdykoli je to možné, má výzkumník použít postup, který nevyžaduje zatajení skutečností;*
- *jestliže nelze použít alternativní metody, výzkumník se rozhoduje o zatajení podle vědecké a aplikované hodnoty výsledků studie;*
- *jestliže výzkumník musí něco zatajit, pak o tom informuje účastníky dostatečným způsobem, jakmile to bude možné.“*

#### **Etické principy při realizaci výzkumu:**

- **Zastupování výzkumného pracovníka** – výzkumník se rozhodne v nějaké fázi výzkumu (zpravidla sběr dat) pověřit jiné osoby, a to z důvodů objektivních či subjektivních. Pokud je tedy nutné, aby byl výzkumný pracovník zastupován, musí dotyčné zastupující osoby dostatečně proškolit, těžko však může kontrolovat jejich svědomitost, tedy zda postupují podle jeho instrukcí atd. Průcha (in Skutil a kol., 2011) opět vhodně poznamenává, že nelze říci, že by vždy bylo zastupování výzkumného pracovníka špatně. Záleží na povaze výzkumu a někdy se může jednat i o nutnost. Výzkumný pracovník by však měl nést hlavní zodpovědnost.

#### **Etické principy ve fázi publikování výzkumu:**

- Na tomto místě se autor domnívá, že není nutné nikterak zdůrazňovat problematiku plagiátorství a s ním souvisejícího správného citování, je to již samozřejmostí.
- **Nezkreslená prezentace výsledků výzkumu** – uvedený etický princip je v podstatě jasně srozumitelný, ale je nutné zdůraznit, že nesouvisí pouze s výzkumem. Vždy by měly být prezentovány výsledky v takové podobě, v jaké byly výzkumníkem zjištěny. To, že se nepotvrdí stanovené hypotézy, není chybou. Stejně tak by autor měl prezentovat

i názory jiných odborníků, které nejsou v souladu s jeho chápáním daného jevu.<sup>7</sup>

- **Opakované publikování téhož výzkumu** – eticky přípustné je, pokud je tentýž výzkum publikován vícekrát, ale vždy pro různou cílovou skupinu. Naše výsledky mohou mít širší společenský dopad a jejich publikování s vysokou mírou složitých statistických výpočtů by mohlo být nesrozumitelné pro širší veřejnost. Proto jej můžeme zjednodušit apod.

Hendl (2008) k etickým otázkám výzkumu poznamenává, že žádný výzkum nesmí ohrozit respondentovo fyzické a psychické zdraví.

V návaznosti na Hendlovu poznámku považuje autor za nutné doplnit následující. Pokud realizujete své empirické šetření ve školách (ale i různých volnočasových či sociálních zařízeních) není výjimkou, že vedení dotyčné instituce chce být s výsledky seznámeno. Jedná se o legitimní požadavek. Vy byste však měli vedení seznámit s celkovými výsledky tak, aby nebyla možná zpětná identifikace respondenta – např. rukopis v dotazníku, specifické vyjadřování v přepisu rozhovoru apod. Pokud zaručujeme anonymitu, tak ji nesmíme ničím ohrozit.

Ve světě, ale i v ČR existuje několik etických kodexů, které se váží k výzkumu či výzkumným pracovníkům, kteří výzkum realizují. Níže jsou uvedeny internetové odkazy, na kterých lze nalézt vybrané etické kodexy. Jednotlivé kodexy se mezi sebou výrazně neodlišují z hlediska obsahu, ale spíše z hlediska podrobnosti.

<http://www.une.edu.au/policies/pdf/codeofconductresearch.pdf>

[http://www.avcr.cz/o\\_avcr/zakladni\\_informace/dokumenty/eticky\\_kodex.html](http://www.avcr.cz/o_avcr/zakladni_informace/dokumenty/eticky_kodex.html)

<http://www.birmingham.ac.uk/Documents/university/legal/research.pdf>

[http://www.ucc.ie/research/rio/documents/CodeofGoodConductinResearch\\_000.pdf](http://www.ucc.ie/research/rio/documents/CodeofGoodConductinResearch_000.pdf)

---

<sup>7</sup> V této „činnosti“ vynikal český pedagog Otakar Kádner, který ve své publikaci Základy obecné pedagogiky (ve všech třech dílech) prezentuje jemu všechny dostupné informace. Dává na téma nahlédnout z různých úhlů a pojetí jednotlivých autorů, které cituje a parafrázuje.

### 3 Kvantitativní výzkum

#### Cíle kapitoly:

- *Student/ka dovede získané znalosti aplikovat při plánování kvantitativního výzkumu.*
- *Student/ka si osvojí základní druhy výběru respondentů.*
- *Student/ka po přečtení kapitoly dovede formulovat a testovat své hypotézy.*
- *Student/ka po přečtení kapitoly dovede třídit a popsat získaná data.*

#### Klíčová slova:

*etapy kvantitativního výzkumu; výběr respondentů; hypotéza; testování hypotéz; chí-kvadrát; t-test; Pearsonův koeficient korelace; dotazník; třídění dat; směrodatná odchylka; aritmetický průměr; modus; medián; koeficient kontingence; normovaná nominální variance*

*„Kvantitativní výzkum je označení pro přístup, jehož zdrojem má být pouze objektivní a co možná nejpřesnější zkoumání edukační reality, podobně jak v přírodních vědách.“ (Žumárová in Skutil a kol., 2011, s. 59)*

*„Kvantitativní přístupy k výzkumu v sociálních vědách v mnohém napodobují metodologii přírodních věd. Předpokládá se, že lidské chování můžeme do jisté míry měřit a předpovídat. Kvantitativní výzkum využívá náhodné výběry, experimenty a silně strukturovaný sběr dat pomocí testů, dotazníků nebo pozorování.“ (Hendl, 2008, s. 44)*

Kvantitativní výzkum stojí na základech pozitivistické či novopozitivistické filozofie, která vychází z daných skutečností. Nebere v úvahu žádná tázání. Zabývá se tedy fakty a jejich vztahy k faktům jiným. Nejedná se však pouze o pouhou deskripci, ale zabývá se i podmínkami, za kterých určitá fakta vznikají a dává je do souvislostí na základě principu podobnosti a následnosti (Chráška, 2007; Störig, 2000).

Smyslem, ale i cílem kvantitativního výzkumu je získat objektivní numerická data, která lze generalizovat. Snažíme se zjistit rozsah, frekvenci a intenzitu zkoumaného jevu (Gavora a kol., 2010)<sup>8</sup>.

Tab. č. 2 – Přednosti a nevýhody kvantitativního výzkumu

| <i><b>Přednosti</b></i>                                                                                             | <i><b>Nevýhody</b></i>                                                                                                    |
|---------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------|
| Testování a validizace teorií                                                                                       | Kategorie a teorie použité výzkumníkem nemusejí odpovídat lokálním zvláštnostem.                                          |
| Lze zobecnit na populaci.                                                                                           | Výzkumník může opominout fenomény, protože se soustřeďuje pouze na určitou teorii a její testování a ne na rozvoj teorie. |
| Výzkumník může konstruovat situace tak, že eliminuje působení rušivých proměnných, a prokázat vztah příčina-účinek. | Získaná znalost může být příliš abstraktní a obecná pro přímou aplikaci v místních podmínkách.                            |
| Relativně rychlý a přímočarý sběr dat.                                                                              | Výzkumník je omezen reduktivním způsobem získávání dat.                                                                   |
| Poskytuje přesná, numerická data.                                                                                   |                                                                                                                           |
| Relativně rychlá analýza dat (využití počítačů).                                                                    |                                                                                                                           |
| Výsledky jsou relativně nezávislé na výzkumníkovi.                                                                  |                                                                                                                           |
| Je užitečný při zkoumání velkých skupin.                                                                            |                                                                                                                           |

(Hendl, 2008, s. 47)

<sup>8</sup> <http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/veda-a-vyskum/kvantitativny-vyskum.php?id=i1p6>

### 3.1 Etapy kvantitativního výzkumu

V různých odborných publikacích lze nalézt rozdílná, ale ve své podstatě shodná členění etap kvantitativního výzkumu. Pro naše potřeby využijeme „rozetapování“ dle Gavory a kol. (2010)<sup>9</sup>:

Níže uvedené etapy nemusí na sebe navazovat chronologicky, tak jak jsou zde v pořadí uvedeny. Je sice předvídatelné, že nebudeme verifikovat hypotézy před jejich formulováním, avšak můžeme např. doplňovat či upřesňovat proměnné, a to na základě námi realizované pilotáže nebo předvýzkumu.

- 1. Volba výzkumného tématu** – volba výzkumného tématu je poměrně důležitým krokem, náměty lze získat na základě studia odborné literatury, vlastních zkušeností nebo diskuzí s odborníky. Vždy je nutné si k danému tématu nastudovat dostatek relevantní literatury (zda vůbec existuje). Zároveň by téma mělo odpovídat i našim zkušenostem, kompetencím a mělo by být i časově zvládnutelné. Vliv na výběr tématu mají i finanční limity, dostupnost a zajímavost tématu.

*„Každým tématem se někdo zabývá.“* (z přednášky prof. PhDr. Jiřího Mareše, CSc. na PdF UP v Olomouci)

*Příklad příliš širokého výzkumného tématu: Šikana<sup>10</sup>.*

*Příklad relativně správně formulovaného výzkumného tématu: Životní styl žáků 2. stupně základní škol.*

Výběr tématu je opravdu odpovědným krokem, který by neměl být určován tím, že kvalifikační práce jsou vnímaným nutným zlem, a proto si vybíráme téma, které je subjektivně „lehké“, mnohokrát zpracované apod. Téměř v každém tématu se dá nalézt nová dimenze, podívat se na známé téma jiným úhlem pohledu a vybrat si i jiný druh výzkumu (viz výše).

---

<sup>9</sup> <http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/veda-a-vyskum/etapy-prace.php?id=i1p8>

<sup>10</sup> V podstatě jednoslovné výzkumné téma nebývá výjimkou, resp. s takovými návrhy v některých případech přicházejí studenti za pedagogy. Z uvedeného příkladu je patrné, že nelze určit, čím se bude student zabývat v oblasti šikany. Budete to šikana na pracoviště? Kyberšikana? ...



- 2. Stanovení výzkumného problému** – výzkumný problém má konkrétnější podobu než výzkumné téma. V rámci výzkumného tématu může být formulováno více výzkumných problémů.

Příklad výzkumného problému (k výše uvedenému tématu): *Faktory ovlivňující životní styl žáků základních škol.*

Pokud je výzkumný problém formulován v tázací podobě, jedná se o výzkumnou otázku.

Výzkumná otázka by měla odpovídat úrovni práce, pro kterou je formulována. Nikdy by se na ni nemělo dát odpovědět pouze „ano“ či „ne“. Měla by začínat slovy „jak“ nebo „co“. Být přesně formulována. Neměla by být do jisté míry kontroverzní či příliš široká<sup>11</sup> (Skutil in Skutil a kol, 2011).

*Příklad kontroverzní výzkumné otázky: Co způsobuje, že jsou blondýnky hloupé?*

*Příklad příliš široké výzkumné otázky: Co je možné zrealizovat pro inkluzi romské populace?*

*Příklad relativně správně formulované výzkumné otázky: Jak středoškolští studenti tráví čas v online prostoru?*

*„Dobrý projekt se vyznačuje vysokou kompatibilitou účelu výzkumu, teorie, výzkumných otázek, metod, výběrových strategií a postupů pro zajištění validity výsledků. Všímáme si těchto situací:*

- *Jestliže výzkumné otázky, k nimž získáte odpovědi, nemají přímo vztah k účelu studie, pak se pravděpodobně musí změnit výzkumné otázky.*
- *Jestliže výzkumné otázky nejsou propojené s teorií, pak není jisté, že odpovědi budou mít nějakou hodnotu. Potom je zapotřebí navrhnout jinou teorii nebo upravit výzkumné otázky.*

---

<sup>11</sup> Na druhou stranu musí být dostatečně široká, abychom nevyřadili určité aspekty (Švaříček, Šedřová a kol., 2007)

- *Jestliže metody a výběrové strategie nepovedou k zodpovězení výzkumných otázek, budeme sbírat nová data, rozšíříme výběr nebo upravíme výzkumné otázky.*“ (Hendl, 2008, s. 145)

### 3. Stanovení proměnných ve výzkumu – rozeznáváme několik typů proměnných, které by se měly objevovat již při formulaci našeho výzkumného problému (výzkumné otázky).

a) Zpravidla hovoříme o proměnné závislé a nezávislé. Závisle proměnná závisí na nezávisle proměnné, tedy pokud dojde ke změně nezávisle proměnné, nastane změna u závisle proměnné. Předpokládáme, že tam existuje nějaký vztah, který se rozhodneme zkoumat (Hendl, 2009).

*Příklad: Rychlost žáka 9. třídy v běhu na 100 m je závislá na jeho tělesné výšce.*

*Závisle proměnnou je v tomto případě rychlost, nezávisle proměnnou tělesná výška.*

Pochopitelně nemůžeme měnit výšku jednoho jedince během dne, týdne. Předpokládáme však, že výška má příčinnou souvislost s rychlostí v běhu na 100 m. Proto je rozlišován tzv. *přirozený a manipulativní prediktor* u nezávisle proměnných. Přirozený ovlivnit nemůžeme. Jedná se právě o výšku, pohlaví apod. Manipulativní prediktor ano. Jednalo by se např. o délku spánku apod. (Hendl, 2009).

Hendl (2009) doplňuje ještě tzv. rušivou (matoucí) proměnnou, kterých může být více a mají vztah k námi zkoumaným nezávisle a závisle proměnným. K výše uvedenému příkladu by takových matoucích proměnných mohlo být více – např. trénovanost jedince, jeho tělesná hmotnost aj.

Rovněž se dá použít termín intervenující proměnné. Tedy takové, které jsou těžko zachytitelné, ale přesto nějakým způsobem ovlivňují námi

získaná data. Ve vztahu k výsledkům vzdělávání žáků různých typů škol může tímto být – prostředí v širším významu, roční doba, osobnost učitele apod.

**b)** Proměnné můžeme rozdělit dle typu měření, které zvolíme. Tím pak rozlišujeme následující proměnné:

**nominální (kategorální)** – uvedený typ proměnné se objevuje nejčastěji v dotaznících vlastní konstrukce v BP. Odpovědi lze roztrždit do kategorií. Příkladem nominální proměnné je: muž/žena; stupeň vzdělání apod. (Žumárová in Skutil a kol., 2011).

Se získanými údaji nemůžeme zacházet úplně jako s čísly. Krásný příklad uvádí Hendl (2009, s. 50). Pokud bychom například kvalitativnímu znaku přiřadili číslo (muž – 1; žena – 2), spočítali bychom aritmetický průměr, tak získáme naprostý nesmysl.

**Co můžeme „počítat“:** *sčítat a odčítat četnosti; chí-kvadrát; Fisherův test; procenta; koeficient kontingence; nominální varianci aj.* (Chráska, 2007).

**ordinální** – jedná se o typ proměnné, kdy jednotlivé výsledky lze seřadit, např. výsledky v závodech (na základě dosaženého času) apod.

**Co můžeme „počítat“:** *medián; kvartilovou odchylku; Spearmanův koeficient pořadové korelace; U-test; Wilcoxonův test aj.* (Chráska, 2007).

**intervalové** – data získaná intervalovým měřením vyjadřují, jaké jsou mezi získanými výsledky rozdíly. Intervalové měření nemá přirozený nulový bod. Jedná se např. o výsledky získané v didaktickém testu (Chráska, 2007, s. 36-37).

**Co můžeme „počítat“:** *aritmetický průměr, směrodatnou odchylku, Studentův t-test, párový t-test apod.* (Chráska, 2007).

**poměrové** – hodnoty vyjadřují množství vlastnosti, kterou zjišťují. Existuje zde přirozená nula. Příkladem poměrového měření je např. věk dítěte (Chráska, 2007, s. 37).

**Co můžeme „počítat“:** *v podstatě vše, co bylo výše uvedeno. Určité riziko však představuje využití postupů pro data nominální a ordinální.*

Poměrové a intervalové měření (proměnné) označujeme jako měření metrické, získáváme tedy metrická data.

**4. Formulování hypotéz** – hypotézám je věnována samostatná subkapitola, autor skript se domnívá, že to přispěje ke zpřehlednění a lepšímu pochopení.

**5. Stanovení výzkumných metod, výběr hotových či vlastních (samostatně) konstruovaných výzkumných nástrojů** – nejčastěji se objevují dotazníky vlastní konstrukce, rozhovory s vlastními otázkami apod. Práce s hotovými nástroji<sup>12</sup> vyžaduje určité kompetence v dané oblasti (u psychologických testů i příslušné vzdělání). U hotových nástrojů je striktně stanovena administrace (postup). Pokud se rozhodneme využít např. dotazník vlastní konstrukce (a to se zpravidla rozhodneme), je nutné zajistit elementární míru validity a praktičnosti<sup>13</sup>.

Validita znamená platnost. O dobré validitě hovoříme v případě, že se skutečně měří to, co se měřit má (Chráska, 2007).

Výše uvedená věta zní možná jednoduše a samozřejmě, ale je nutné si ji připomenout již ve fázi stanovení výzkumného problému. Vše poté souvisí s formulovanými otázkami (v případě dotazníku, rozhovoru).

Významnou vlastností je praktičnost. Významnou z toho důvodu, že vystihuje přesně to, co se od výzkumu v BP i DP očekává. Takový výzkum má být jednoduchý,

---

<sup>12</sup> Jedná se o nástroje, kde je již zajištěna validita a reliabilita.

<sup>13</sup> Elementární mírou se rozumí vztah k úrovni práce, kterou představuje práce bakalářská.

finančně přijatelný<sup>14</sup>, časově nenáročný<sup>15</sup>, není vyžadována nějaká vysoká erudice výzkumníka, snadno proveditelný (Chráska, 2007).

- 6. Sběr dat** – v této etapě můžeme narážet na různé překážky, se kterými jsme při plánování výzkumu nepočítali. Rozhodně bychom měli mít dostatek času na sběr. Připraveny záložní alternativy.

V kvantitativním výzkumu je sběr dat samostatnou částí výzkumu. V rámci kvalitativního výzkumu se sběr dat prolíná s jejich zpracováváním a vyhodnocováním (Hendl, 2008).

- 7. Zpracování údajů (dat)** – u této etapy máme v podstatě tři možnosti. První možností je sednout si na zem do pokoje, rozložit si kolem sebe dotazníky (nejčastější metoda při kvantitativním výzkumu), hodinu brečet a poté ručně třídit jednotlivé údaje. Tato možnost má výhodu pouze v tom, že k tomu nepotřebujeme nic více než tužku a papír. Hlavním rizikem je, že se v datech postupně ztrácíme a roste pravděpodobnost chyby.

Druhá možnost je v podstatě stejná, jen s tím rozdílem, že si nesesedneme na zem, ale k počítači, brečet však stále můžeme. Zpravidla využíváme tabulkový editor Microsoft Office, data tak mnohem přehledněji uspořádána. Podmínkou však je, abychom si přehledně nastavili jednotlivé excelovské listy (sloupce, řádky). Můžeme využívat i statistické funkce, které poskytuje zmiňovaný program (tabulky četností apod.). Je však nutné mít pokročilejší znalosti daného programu. Zejména u sestavování tabulek četností jsou „excelovské“ kroky relativně nelogické a i nápověda je spíše matoucí než napovídající (vysvětlíme v příloze č. 2). Další výhodou je kompatibilita s některými statistickými softwary.

Třetí a poslední možnost je mít náš dotazník či jiný záznamový arch v takové formální podobě, že můžeme využít dokumentační scannery a

---

<sup>14</sup> Pamatujte na to, že ne vše se dá řešit pomocí internetu a z pohodlí domova. Můžete očekávat i to, že specifické skupiny obyvatel po Vás mohou požadovat určitý finanční obnos.

<sup>15</sup> Nejedná se jen o fázi sběru dat, ale i o část, kdy data vyhodnocujete.

ve spolupráci s odpovídajícími programy se nám výsledky přímo zaznamenávají.

Prezentaci získaných dat se budeme věnovat v samostatné subkapitole.

## 8. Verifikace hypotéz – verifikaci hypotéz je věnována samostatná subkapitola.

Nedílnými etapami výzkumu je pilotáž a předvýzkum. Někdy jsou tyto pojmy zaměňovány nebo chápány jako synonyma. Jedná se však o dvě rozdílné činnosti, které jsou důležité při každém výzkumu (BP, DP aj.). Jsou realizovány před samotným sběrem dat.

**Pilotáž** slouží k tomu, abychom se o našem problému dozvěděli více, pronikli do určitých zákonitostí a zjistili, zda má vůbec význam danou oblast zkoumat<sup>16</sup>. Může se jednat o rozhovory či pozorování. Na základě pilotáže je možné zpřesnit naše hypotézy, zpřesnit výzkumný problém (Chráska, 2007).

**Předvýzkum.** Hlavním cílem předvýzkumu je ověřit, zda používáme správnou metodu pro sběr dat. Neméně podstatné je rovněž zjistit (zpravidla u dotazníku) srozumitelnost otázek, náročnost otázek, jaký čas zabere jeho vyplňování, zda respondenti chápou dané otázky tak, jak jsme je skutečně mysleli. Předvýzkum se realizuje na malém vzorku respondentů z cílové skupiny. Předvýzkum je realizován i u jiných metod než je dotazník (Bartošová, Skutil in Skutil a kol., 2011).

Pokud se pilotáži a předvýzkumu „vyhneme“, můžeme očekávat problémy při sběru a následném vyhodnocování dat. Dosud realizovaná práce může v podstatě začít znovu, ale zpravidla již s jinými respondenty.

### 3.2 Výběr respondentů v kvantitativním výzkumu

Výběr respondentů je dalším odpovědným krokem k úspěšné realizaci výzkumu. Rozlišujeme několik druhů výběru prvků do našeho výzkumu. Jaký druh

---

<sup>16</sup> Zdali je vůbec reálné se danou problematikou zabývat.

výběru zvolíme, opět záleží na tom, co chceme zjistit, jaké máme definované proměnné apod. Každopádně bychom i v BP a DP měli usilovat o co největší reprezentativnost<sup>17</sup> daného vzorku.

Punch (2008b) zmiňuje, že právě studenti často pracují s omezeným časem, zdroji a možnostmi, které jim mohou znesnadnit přístup k některým lidem (skupinám).

Základní soubor je tvořen všemi prvky (jedinci), kteří zapadají do jeho vymezení. Základní soubor je vymezen místem (např. ČR, MSK apod.), časem (např. rok 2011), věkem, pohlavím, vzděláním a dalšími proměnnými. To vše poté nazýváme znaky základního souboru (Žumárová in Skutil a kol., 2011).

Zpravidla nemůžeme výzkum realizovat na všech prvcích základního souboru<sup>18</sup>. Výzkum tedy realizujeme na zmenšeném počtu respondentů. Takový soubor označujeme jako soubor výběrový. Výběrový soubor musí být zmenšeným obrazem základního souboru, tedy musí být proporcionálně stejný. Tím plní předpoklad reprezentativnosti.

Zvolený výběr by měl mít určitou logiku, a to ve vztahu k výzkumným otázkám, tématu apod. Z tohoto důvodu nelze říci, který druh výběru je přímo lepší či horší, vždy to závisí na našem výzkumném projektu (Punch, 2008b). Poslední odstavec trochu rozvíňme. Často se stává, že studenti svou práci nazvou např. *Problematika kyberšikany u žáků 2. stupně ZŠ na Broumovsku*. Základní soubor tedy tvoří všichni žáci 2. stupně ZŠ na Bruntálsku. Student/ka však rozdává své dotazníky pouze na ZŠ ve městě Bruntál. Předpokládejme, že výzkumná část bude naplňovat hlavní cíl práce. Zcela jistě při tomto postupu splněn nebude.

Níže je uveden základní výčet druhů výběru, který je dostačující pro cíl těchto skript. Chráška (2007); Punch (2008b); Žumárová in Skutil a kol. (2011):

- 1. Dostupný výběr** – tento druh výběru by se dal nazvat jako „chytnu toho, kdo šel okolo“. S tímto druhem se lze často setkat v BP i DP, kdy jsou např. v případě dotazníku osloveni známí, kamarádi, dotazník je rozdán na pracovišti<sup>19</sup>. Výsledky výzkumu, kde byl zvolen uvedený výběr, je nutné

---

<sup>17</sup> Viz slovník doplňujících pojmů.

<sup>18</sup> U velmi specifických a malých skupin ano – viz níže.

<sup>19</sup> Rozdání dotazníku na pracovišti nemusí být ničím negativním, ale výsledky je nutné vztahovat pouze k danému pracovišti. Je nutné výzkum dostatečně propojit s teorií, což může být obtížné.

dávat do souvislosti pouze s danou skupinou. Důležité je téma výzkumu, protože uvedený typ výběru nemusí být opodstatněný. Výsledky jsou pak významně devalvovány.

2. **Náhodný výběr** – dalo by se říci, že se jedná o ideální druh výběru respondentů ve vztahu k zobecnitelnosti a reprezentativnosti. O jejich zařazení rozhoduje náhoda. V praxi jej však těžko uplatňujeme, a to z důvodu uvedených dále. Představit si jej může tak, že všechny členy základního souboru „hodíme“ do klobouku a náhodně z takového osudí losujeme respondenty.
3. **Záměrný výběr** – výzkumník volí na základě svých zkušeností, teorie, vlastního úsudku prvky, které splňují znaky základního souboru. Rozeznáváme tři typy záměrného výběr: Anketní výběr – respondenti se sami hlásí. Anketní výběr je u některých výzkumů v podstatě jediný možný, protože nemáme jinou možnost, jak určit jedince dle znaků základního souboru. Jako příklad uveďme BP jedné studentky, která se zabývala problematikou mužů na rodičovské dovolené. Zvolila tedy inzerát, aby se jí přihlásili. Za anketní výběr můžeme považovat i online dotazník. Výběr průměrných jednotek – zde je vyžadována vysoká odbornost výzkumníka, který vybere respondenty jako typické reprezentanty základního souboru, což může být ve svém důsledku velice těžko obhajitelné. Kvótní výběr – respondenty vybíráme na základě určitých znaků. V základním souboru je poměrné zastoupení prvků dle pohlaví, vzdělání apod. Na základě determinovaných znaků určujeme kvóty. Např. pokud by znakem mělo být vzdělání, tak víme, že v základním souboru je např. 30 % lidí se základním vzděláním, 55 % lidí se středoškolským vzděláním a 15 % má vzdělání vysokoškolské. Výběrový soubor musí kopírovat toto procentuální rozdělení. Příkladem jsou např. i nejrozumnější volební průzkumy.
4. **Stratifikovaný výběr** – jedná se o výběr, kdy je základní soubor tvořen několika dalšími podskupinami. Jednotlivé podskupiny nemusí být proporcionálně zastoupeny, tak jak jsou v základním souboru. Vybírá se



pomocí náhodného výběru z jednotlivých podskupin. Chráska (2007) uvádí příklad rozdělení učitelů do podskupin dle délky praxe.

**5. Vícenásobný výběr** – tento druh výběru je realistický v BP. Výše jsme zmiňovali příklad tématu BP s Bruntálskem. Nemusíme objíždět všechny školy, ale začneme právě výběrem škol (náhodným) – dále již vybíráme jednotlivé žáky na školách (které jsme vylosovali), kteří splňují znaky základního souboru. Postupujeme tedy od nejvyššího řádu až k jednotlivým prvkům.

**6. Vyčerpávající výběr** – jedná se o výběr, kdy jsou respondenty všichni členové (prvky) základního souboru. Základní soubor je tak malý, že nám nečiní žádné obtíže oslovit všechny. Např. jedná se o učitele na jedné škole.

U jednotlivých výběrů, které jsou zvoleny, záleží, jak prezentujeme získaná data. Ani jeden z výběrů není primárně úplně špatně, jak již bylo uvedeno výše. Vše se odvíjí od našeho tématu, výzkumného problému, výzkumných otázek apod.

Výše chybí možná informace, kterou byste rádi věděli. Jak velký má být soubor? Obecnou radou je, že čím více, tím lépe, alespoň v kvantitativním výzkumu.

Disman (2008) uvádí, abychom mohli tvrdit, že všechny vrány jsou černé, musíme pozorovat všechny vrány.

Existují i postupy pro výpočet velikosti výběrového souboru, které však v BP či DP využijeme zřídka, prakticky nikdy.

### **3.3 Formulace a testování hypotéz**

Formulace hypotéz je stejně odpovědnou částí výzkumu jako všechny ostatní. Hypotézy by neměly být vystřeleny tzv. od boku, bez rozmyslu. Tedy jen proto, že je nutné, aby tam byly (většinou byly).

*„Hypotézou se rozumí předpoklad, tvrzení, podmíněný výrok o vztazích mezi dvěma či více proměnnými. Chápe se jako odpověď na výzkumnou otázku, musí být*

ověřitelná a měla by obsahovat dvě alternativy (platí – neplatí). Hypotézy tvoří základ kvantitativně orientovaných výzkumů.“ (Žumárová in Skutil a kol., 2011, s. 60)

Hendl (2008, s. 39) hypotézu definuje: „*Predikce nebo odhad vztahu, který existuje v reálném světě za určitých podmínek.*“

Výše bylo uvedeno, že hypotéza by neměla být vystřelena od boku, měla by být formulována na základě teorie, vlastní praxe apod. Neměla by být ani předem „jasná“, tedy taková, kdy před započatým výzkumem již se 100% či 99% jistotou víme, jak to dopadne. Příkladem takové hypotézy je: „*Policisté nosí ve své profesi častěji zbraň než knihovníci.*“<sup>20</sup> I když je uvedený příklad sám o sobě nesmyslný, protože srovnávat nošení zbraně mezi knihovníky a policisty není dobrým výzkumným tématem, tak má ad absurdum demonstrovat výše uvedenou poznámku.

Pro lepší pochopení bude níže uveden **vhodný příklad** Gavory a kol. (2010)<sup>21</sup>: „*Žáci s výborným prospěchem z matematiky mají lepší postoj k vyučování tohoto předmětu než žáci se slabým prospěchem.*“

V tomto případě tak „jasno“ dopředu nemáme. Stejný autor doporučuje, aby výzkumník svou hypotézu rozšířil (resp. doplnil další hypotézy) a zaměřil se např. na více ročníků, na osobnost učitele, třídní klima apod. (vše ve vztahu k původní hypotéze (tématu). Rozpracováním hypotézy odkrývá další dimenze zkoumané problematiky.

Pravidla pro formulaci hypotéz (Chráska, 2007, s. 17):

1. „*Hypotéza je tvrzení, které je vyjádřeno oznamovací větou.*“<sup>22</sup>
2. *Hypotéza musí vyjadřovat vztah mezi dvěma proměnnými (pokud se nejedná o vyjádření vztahů, není možno hovořit o vědecké hypotéze). Proto musí být hypotéza vždy formulována jako tvrzení o rozdílech, vztazích nebo následcích*<sup>23</sup>.

---

<sup>20</sup> I když dostat románem Vojna a mír po hlavě není nic příjemného.

<sup>21</sup> <http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/hypotezy/ako-vznikaju.php?id=i7p1>

<sup>22</sup> Viz výše

<sup>23</sup> Je možné stanovit i hypotézu o jednom parametru. Jedná se především o parametry dobře teoreticky ukotvené. Příkladem takto stanovené parametrické hypotézy by mohla být taková, která se týká určitých dovedností, klíčových kompetencí apod. na určitém vývojovém stupni jedince.

3. *Hypotézu musí být možno empiricky ověřit. Proměnné, které vystupují, musí být měřitelné (byť např. jen na základě kategorizace).“*

Poslední bod je obzvláště důležitý, nezřídka se stává, že je hypotéza formulována technicky dobře, ale proměnné jsou neměřitelné, resp. téměř neměřitelné.

Hypotéza musí být jasná, stručná a výstižná (Gavora a kol., 2010).

**Nejčastější chyby** při formulaci hypotéz v podstatě vycházejí z nedodržení výše uvedeného. Pokud si tedy zapamatujete, jak má být hypotéza formulována, měli byste se vyhnout všem níže uvedeným nedostatkům.

Chráška (2007, s. 18) zmiňuje následující nedostatky:

1. Nejsou vyjádřeny vztahy mezi proměnnými, i když jistý vztah může v hypotézách nalézt, nemusí být jednoznačně pochopen. Jako **příklad** uvádí: „*Dívky si nejraději barví vlasy.*“
2. Hypotéza obsahuje neurčitou formulaci. Obecný **příklad** takové hypotézy: „*jev A někdy vyvolává jev B*“. A někdy taky nevyvolává.

**Konkrétní příklad:** *Pití alkoholu někdy vyvolává zvracení.*

Mohli bychom ještě doplnit, že se často objevují formulace hypotéz kombinující výše uvedené nedostatky. Popř. začínají slovy: Více než 50 % ...; Většina dotázaných ... apod. Není výjimkou ani věta obsahující: Větší polovina ... . Je tedy nutné nad hypotézami dostatečně popřemýšlet. Špatná formulace nevychází ani tak z neznalosti, ale spíše ze zbrklosti. Plánování výzkumu a jeho příprava není závod.

**Příklady špatně formulovaných hypotéz:**

*Většina lidí pije alkohol.*

*Více než 50 % dotázaných má problémy s kouřením.*

*Respondenti mají doma počítač.*

### 3.3.1 Hypotéza věcná, statistická, alternativní a nulová

#### Hypotéza věcná

Hypotézy, které formulujeme pro náš výzkum, nazýváme jako hypotézy věcné. Jedná se o hypotézy, jejichž příklady jsou uvedeny výše. V těchto hypotézách jsou tedy vyjádřeny *rozdíly, vztahy nebo následky*.<sup>24</sup>

Hypotézy vyjadřující rozdíly jsou nejčastěji používanými hypotézami v BP. Využívají se v nich slova jako: častěji, více apod. Příkladem takové hypotézy je: *Žáci 9. tříd v Ostravě mají více neomluvených hodin ve škole než žáci 9. tříd v Opavě.*

Hypotézy vyjadřující vztahy se formulují v případě, že chceme zjišťovat míru a směr vztahu. Směrem je myšlen vztah pozitivní či negativní. Příkladem takové hypotézy je: *Mezi dosaženým stupněm vzdělání a výší příjmu je pozitivní vztah.*

Hypotézy vyjadřující následky využívají spojení slov: jak-tak, čím-tím. Příkladem takové hypotézy je: *Pokud se zvýší množství společných pohybových aktivit žáků, tak se zlepší třídní klima.*

Jednotlivé proměnné, které se vyskytují v hypotéze, je nutné *operacionalizovat*, tedy vyjádřit tak, aby je bylo možné přesně změřit. Chráska (2007) uvádí příklad, že vědomost žáka je možné operacionalizovat jako výsledek didaktického testu.

#### Hypotéza statistická

Zjednodušeně je statistickou hypotézou myšlena taková hypotéza, která již obsahuje statisticky měřitelné termíny. Tedy to, podle čeho vlastně bude ověřovat hypotézu věcnou. Zda podle bodů, četností zaznamenaných odpovědí, zaznamenaných projevů chování apod. Věcnou hypotézu tedy převádíme na hypotézu statistickou, aby ji bylo možno statisticky ověřit.

Věcná hypotéza: *Žáci 9. tříd v Ostravě mají celkově více neomluvených hodin ve škole než žáci 9. tříd v Opavě.*

---

<sup>24</sup> Níže jsou informace parafrázovány a kombinovány z publikací: Skutil a kol., 2011; Chráska, 2007; Gavora a kol., 2010

Statistická hypotéza: *Celkový počet neomluvených hodin žáků 9. tříd v Ostravě je vyšší než u žáků 9. tříd v Opavě.*

### **Hypotéza nulová a alternativní**

Statistická hypotéza není verifikována přímo, ale pomocí nulové a alternativní hypotézy. Nejjednodušším vysvětlením je, že *nulová hypotéza* říká – nejsou rozdíly, vztahy, následky. Alternativní říká, že rozdíly, vztahy či následky jsou. Dle formulace alternativní hypotézy hovoříme buď o oboustranném, či jednostranném testu. Alternativní hypotéza pro oboustranný test nám říká, že  $a \neq b$  (tedy, že rozdíly jsou, ale neříká jakým směrem). Pro jednostranný test je alternativní hypotéza formulována buď jako  $a > b$ , nebo  $a < b$ . V případě, že testujeme naši věcnou hypotézu uvedenou výše, je naše alternativní hypotéza jednostranná ve vztahu k hypotéze věcné.

Nulová hypotéza je tedy to, co považujeme za přijatelné bez dalšího zkoumání.

Problematika využití jednostranných či oboustranných testů je komplexně složitější a není rozhodně jedno, který test využíváme. V případě využití např. t-testu je patrné, že docházíme k rozdílným výsledkům (jiná hodnota kritického oboru, p-hodnota) (blíže viz Hendl, 2009)<sup>25</sup>.

O využití jednostranného či oboustranného testu má význam uvažovat např. v případě z-testu či t-testu. Uveďme pouze, že chí-kvadrát sleduje vždy jednostrannou závislost.

Všechny níže uvedené konkrétní příklady hypotéz jsou testovány jednostranně.

Za určitých okolností hypotézy formulovány být nemusí. Jedná se o situace, kdy máme pouze jednu proměnnou a výzkum je deskripcí určité situace, jevu apod. Je však nutné formulovat výzkumné otázky (Gavora a kol., 2010).<sup>26</sup> To však výzkumníkovi (studentovi) situaci rozhodně neulehčuje a nejsou nikterak sníženy nároky na kvalitu dotazníku. Statistické vyhodnocení je tou jednodušší částí výzkumu.

---

<sup>25</sup> Viz příloha č. 3

<sup>26</sup><http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/hypotezy/hypotezy-neformuluju.php?id=i7p7>

### 3.3.2 Testování hypotéz

Po správném naformulování hypotéz, sesbírání potřebných dat a vůbec všem potřebném, nezbyvá nic jiného než naše hypotézy otestovat. Jak bylo uvedeno v kapitole 2, nejčastěji se v BP vyskytují dotazníky s nominálními proměnnými. V této souvislosti se níže zaměříme především na statistické metody pro analýzu nominálních dat. Zmíníme však i metody pro data metrická (t-test a Pearsonův koeficient korelace).

Dnes již pochopitelně existují sofistikované statistické programy (SPSS, STATISTICA aj.), které významně práci výzkumníka ulehčují a eliminují riziko početní chyby. Jmenované programy, resp. jejich pořízení, se pohybuje řádově v desítkách tisíc a jejich ovládání není rozhodně jednoduchou záležitostí. Alternativu představuje program Microsoft Excel, ovšem nastavení funkce (konkrétně chí-kvadrát) je časově náročnější téměř tak stejně jako u ručního počítání.

Než se pustíte do čtení dalších odstavců a subkapitol, tak si nachystejte: *tužku, papíry a kalkulačku*.

Asi každého napadla otázka – *Proč se to učit počítat ročně, když na to máme software?* „Ruční“ výpočet u některých metod vede k pochopení, jak k výsledku dospějeme, čím je ovlivněn a rovněž nám to může pomoci k lepší interpretaci výsledku.

#### 3.3.2.1 Test dobré shody chí-kvadrát

Test dobré shody chí-kvadrát patří k nejužívanějším a nejelementárnějším testům ke zjištění statistické významnosti. K nejužívanějším patří z toho důvodu, že v dotaznících vlastní konstrukce zjišťujeme zpravidla četnosti jednotlivých nabízených odpovědí (nominální data). V sociálních vědách neexistuje příliš mnoho možností, kde bychom využívali stopky, didaktické testy, pedometry apod. Avšak jistá metrická data získat můžeme.

Smyslem testu dobré shody chí-kvadrát je ověřit, zda existují statisticky významné rozdíly mezi námi zaznamenanými četnostmi (tzv.  $P$  – pozorovaná četnost)

a četnostmi očekávanými (tj.  $O$  – očekávaná četnost), někdy též nazvanými jako četnosti teoretické.

$$\chi^2 = \sum \frac{(P - O)^2}{O}$$

(Chráška, 2007, s. 72) [1]

V modelovém příkladu 1 jsme navštívili 93 žáků sedmých tříd a zjišťovali jsme oblíbenost organizačních forem výuky. Zjišťujeme, zda je mezi organizačními formami výuky statisticky významný rozdíl.

Tab. č. 3 – modelový příklad 1

| Organizační forma výuky | P – pozorovaná četnost |
|-------------------------|------------------------|
| Frontální výuka         | 17                     |
| Skupinová výuka         | 23                     |
| Párová výuka            | 25                     |
| Individuální výuka      | 16                     |
| Týmová výuka            | 12                     |
| $\Sigma$                | 93                     |

Celkem máme 93 respondentů. Nyní si musíme určit naše očekávané četnosti. V tomto případě zpravidla očekáváme stejné zastoupení odpovědí u jednotlivých organizačních forem výuky. Lze si však očekávané četnosti určit téměř libovolně, avšak musíme pamatovat, že součet očekávaných četností musí být shodný s počtem pozorovaných četností. V našem příkladu tedy 93.

Tab. č. 4 – pokračování modelového příkladu 1

| Organizační forma výuky | P – pozorovaná četnost | O – očekávaná četnost | P-O                                                                                   | (P-O) <sup>2</sup>                                                                    | $\frac{(P-O)^2}{O}$ |
|-------------------------|------------------------|-----------------------|---------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|---------------------|
| Frontální výuka         | 17                     | 18,6                  | -1,6                                                                                  | 2,56                                                                                  | 0,14                |
| Skupinová výuka         | 23                     | 18,6                  | 4,4                                                                                   | 19,36                                                                                 | 1,04                |
| Párová výuka            | 25                     | 18,6                  | 6,4                                                                                   | 40,96                                                                                 | 2,20                |
| Individuální výuka      | 16                     | 18,6                  | -2,6                                                                                  | 6,76                                                                                  | 0,36                |
| Týmová výuka            | 12                     | 18,6                  | -6,6                                                                                  | 43,56                                                                                 | 2,34                |
| $\Sigma$                | 93                     | 93                    |  |  | 6,08                |

V našem modelovém příkladu jsme předpokládali, že očekávaná četnost bude u všech organizačních forem stejná. Vydělili jsme tedy celkový počet respondentů 93 číslem 5, což je celkový počet organizačních forem výuky<sup>27</sup>. Následně jsme postupovali dle tabulky č. 4. Dle vztahu 1 jsme zjistili, že hodnota testového kritéria chí-kvadrát je rovna 6,08<sup>28</sup>. Tato hodnota nám sama o sobě nic neřekne. Námi vypočtenou hodnotu testového kritéria musíme porovnat s hodnotou kritickou.

Potřebujeme však znát ještě dva údaje. Prvním z nich je tzv. počet stupňů volnosti. Počet stupňů volnosti<sup>29</sup> určíme jednoduše ze vztahu  $f = k - 1$ , kde  $f$  je počet stupňů volnosti a  $k$  je počet nabízených možností. V modelovém příkladu je tedy  $f = 5 - 1$  a počet stupňů volnosti je roven 4. Ještě je potřeba zvolit tzv. hladinu významnosti, nejčastěji 0,05 (5 %), 0,01 (1 %) či 0,001 (0,1 %) (Hendl, 2009).

*„Hladina významnosti je pravděpodobnost, že neoprávněně (nesprávně) odmítneme nulovou hypotézu.“* (Chráska, 2007, s. 72)

Jinými slovy řečeno. Pokud zjistíme statisticky významné rozdíly na hladině významnosti 0,05, můžeme říci, že naše riziko omylu rozhodnutí o zamítnutí nulové hypotézy je nižší než 5 %, při hladině významnosti 0,01 je nižší než 1 % a při hladině významnosti 0,001 je nižší než 0,1 %. Hladinu významnosti volíme dle závažnosti výzkumu (Chráska, 2007).

Nyní již známe všechny potřebné informace k určení, zda v našem modelovém příkladu jsou či nejsou statisticky významné rozdíly. Hladinu významnosti jsme si zvolili 0,05, známe hodnotu testového kritéria 6,08 a víme, že počet stupňů volnosti je roven 4. V příloze č. 1 je uvedena tabulka kritických hodnot. Pro 4 stupně volnosti a zvolenou hladinu významnosti 0,05 je kritická hodnota rovna 9,488.

Pokud je námi vypočtená hodnota nižší než hodnota kritická (na zvolené hladině významnosti), konstatujeme, že mezi proměnnými neexistují žádné statisticky významné rozdíly a nemůžeme zamítnout nulovou hypotézu. Pokud by námi

---

<sup>27</sup> Logicky by to byly jednotlivé možnosti v dotazníku u jednotlivých otázek.

<sup>28</sup> Předpokládám, že všichni víme, jak jsme k danému číslu dospěli.

<sup>29</sup> Stupeň volnosti se dá jednoduše popsat tak, že příslušný počet značí, do kolika polí očekávaných četností si můžeme zvolit libovolnou hodnotu. V naší tabulce si můžeme libovolnou hodnotu zvolit do 4 polí, ale do pátého pole již musíme dosadit takovou hodnotu, aby celkový součet byl roven 93 (tedy v našem případě).



vypočtená hodnota byla vyšší než hodnota kritická (na zvolené hladině významnosti), konstatujeme, že existují statisticky významné rozdíly a zamítáme nulovou hypotézu.

Níže si uveďme příklad, jak potvrdit či nepotvrdit platnost naší věcné hypotézy. Formulovali jsme si následující věcnou hypotézu: *Opilých žáků 9. tříd je významně více ke dni předání vysvědčení než na Silvestra.*

Tab. č. 5 – modelový příklad 2

| Den                | Počet opilých žáků |
|--------------------|--------------------|
| Předání vysvědčení | 85                 |
| Silvestr – 31.12.  | 31                 |
| $\Sigma$           | 116                |

Následně postupujeme stejně jako v tab. č. 6.

Tab. č. 6 – pokračování modelového příkladu 2

| Den                | P – pozorovaná četnost | O – očekávaná četnost | P-O | (P-O) <sup>2</sup> | $\frac{(P-O)^2}{O}$ |
|--------------------|------------------------|-----------------------|-----|--------------------|---------------------|
| Předání vysvědčení | 85                     | 58                    | 27  | 729                | 12,57               |
| Silvestr – 31.12.  | 31                     | 58                    | -27 | 729                | 12,57               |
| $\Sigma$           | 116                    | 116                   |     |                    | 25,14               |

Zjistili jsme, že hodnota testového kritéria je rovna 25,14. Dopočítáme počet stupňů volnosti dle výše uvedeného vztahu, tedy pro tento příklad je počet stupňů volnosti 1, zvolili jsme hladinu významnosti 0,05. Nyní porovnáme naši vypočtenou hodnotu s hodnotou kritickou. Kritická hodnota pro 1 stupeň volnosti a zvolenou hladinu významnosti je rovna 3,841. Zjišťujeme, že námi vypočtená hodnota je vyšší než hodnota kritická. Tedy, že mezi proměnnými existují statisticky významné rozdíly. Víme, že zaznamenané četnosti (pozorované) odpovídají naší formulované věcné hypotéze (blíže viz formulace o vztahu u jednostranného a oboustranného testu). Můžeme tedy konstatovat, že se námi formulovaná hypotéza potvrdila.

Zjištění statisticky významných rozdílů neznamena „automatické“ potvrzení věcné hypotézy. Vždy je nutné se „podívat“, zda zaznamenané četnosti odpovídají našemu jednostrannému vztahu! Uvedené bývá častou chybou. Výsledky neinterpretujte „slepě“.

Pokud nebyly zjištěny statisticky významné rozdíly, můžeme konstatovat, že nelze odmítnout nulovou hypotézu a námi formulovaná věcná hypotéza se nepotvrdila. Uvedené platí pro danou hladinu významnosti (viz níže – „chyby“).

Ve výše uvedených příkladech byly vždy očekávané četnosti rovnoměrně rozděleny. Taková situace však není využitelná vždy. Uvedený postup ke stanovení očekávaných četností je využitelný pouze za předpokladu, že se ptáme zpravidla jedné homogenní skupiny (z hlediska věku, pohlaví, vzdělání apod.) na otázku, u které nelze určit jinak teoretické rozdělení odpovědí. Příkladem, kdy lze určit očekávané četnosti jinak, by byla např. otázka zaměřující se na druhy vyzkoušených drog. Nelze očekávat, že by kokain měl stejnou četnost jako marihuana. K určení očekávaných četností nám mohou pomoci nějaké celostátní výzkumy, každopádně reprezentativní výzkumy, dle kterých můžeme poměrově rozdělit očekávané četnosti.

Posledně uvedenému odstavci věnujte pozornost. Jelikož se pohybujeme v oblasti statistiky a kvantitativního výzkumu. Nemůžeme tam ty čísla „střílet“, jak nás napadne. Nezapomínejte, že cílem kvantitativního výzkumu je mj. ověřit teorii, která mohla být vytvořena na základě výzkumů jiných. Máme-li tedy data k celostátním údajům a snažíme se je komparovat s naším lokálním či regionálním výzkumem, tak teoretické rozprostření četností musíme respektovat.

### **3.3.2.2 Test nezávislosti a homogenity pro kontingenční tabulku**

Kontingenční tabulka je někdy označena jako tabulka se dvěma vstupy. Tedy dvěma proměnnými (např. pohlaví ve vztahu k počtu utracených peněz za „něco“). Je využito testu dobré shody chí-kvadrát (porovnání mezi pozorovanými a očekávanými četnostmi). Naše výpočty se budou lišit pouze u výpočtu očekávané četnosti  $O$  a počtu stupňů volnosti.

Pojmy nezávislost a homogenita mají v tomto případě vliv zejména na formulaci nulové hypotézy. V případě nezávislosti<sup>30</sup> tedy nulovou hypotézu formulujeme tak, že se proměnné vzájemně neovlivňují, tedy, že mezi proměnnou  $A$  a  $B$  není žádná souvislost.

---

<sup>30</sup> Prokázání závislosti je poměrně obtížné.

Homogenita předpokládá, že: „pravděpodobnostní rozdělení kategoriální proměnné  $B$  je stejné v různých populacích, které jsou identifikovány faktorem  $A$ .“ (Hendl, 2009, s. 318). Takovým příkladem je i modelový příklad č. 3 (viz tab. č. 7).

Formulujme si opět věcnou hypotézu, kterou chceme testovat: *Dívky utratí za mobilní telefon více peněz než chlapci.*

Tab. č. 7 – modelový příklad 3

| Pohlaví                       | Utracené peníze |              |               | Marginální četnosti ( $n_r$ ) |
|-------------------------------|-----------------|--------------|---------------|-------------------------------|
|                               | 0 – 200 Kč      | 201 – 300 Kč | 301 a více Kč |                               |
| Chlapci                       | 25              | 45           | 16            | 86                            |
| Dívky                         | 48              | 45           | 23            | 116                           |
| Marginální četnosti ( $n_s$ ) | 73              | 90           | 39            | Celková četnost $n = 202$     |

V tab. č. 7 máme zaznamenány pozorované četnosti  $P$ . Nyní si musíme určit očekávané četnosti  $O$ . V kontingenční tabulce se očekávané četnosti vypočítávají ze vztahu dle vzorce [2].

$$O = \frac{n_r \cdot n_s}{n}$$

(Chráška, 2007, s. 77) [2]

Očekávanou četnost musíme dopočítat pro jednotlivá pole tabulky, kterých je 6. Ve vztahu 2 je  $n_r$  tzv. marginální četnost v řádku,  $n_s$  marginální četnost ve sloupci a  $n$  celková četnost.

Pro pole, kde jsme zaznamenali četnost 25 (25 chlapců utratí za měsíc 0 až 200 Kč), vypočítáme očekávanou četnost následně. Příslušná marginální četnost v řádku je rovna 86, marginální četnost ve sloupci 73 a celková četnost 202. Dosadíme do vzorce a výsledná hodnota očekávané četnosti pro dané pole je 31,08.

$$O = \frac{86 \cdot 73}{202} = 31,08$$

Pro pole s četností 48 (dívky, které utratí 0 – 202 Kč za měsíc) je hodnota očekávané četnosti rovna 41,92.

$$O = \frac{116 \cdot 73}{202} = 41,92$$

Pro lepší přehlednost a jasné pochopení je níže naznačeno, které marginální četnosti se vztahují k vybraným polím tabulky. Pro dvě pole již máme očekávanou četnost dopočítanou, nyní je tedy nutné dopočítat i zbývající čtyři pole.

Obr. č. 1 – chí-kvadrát; výpočet očekávaných četností

| Pohlaví                         | Utracené peníze |              |               | Marginální četnosti ( $n_{.}$ ) |
|---------------------------------|-----------------|--------------|---------------|---------------------------------|
|                                 | 0 – 200 Kč      | 201 – 300 Kč | 301 a více Kč |                                 |
| Chlapci                         | 25              | 45           | 16            | 86                              |
| Dívky                           | 48              | 45           | 23            | 116                             |
| Marginální četnosti ( $n_{.}$ ) | 73              | 90           | 39            | Celková četnost<br>$n = 202$    |

Pro pole s četností 45 u chlapců je očekávaná četnost  $O = 38,32$

Pro pole s četností 45 u dívek je očekávaná četnost  $O = 51,68$

Pro pole s četností 16 u chlapců je očekávaná četnost  $O = 16,60$

Pro pole s četností 23 u dívek je očekávaná četnost  $O = 22,40$

Následně postupujeme stejným způsobem jako v předchozím příkladu. V jednotlivých polích tabulky provedeme výpočet:  $P - O$ ,  $(P - O)^2$  a  $\frac{(P - O)^2}{O}$

Výpočet testového kritéria získáme součtem hodnot, dle vzorce č. 1. Jednotlivé hodnoty jsou uvedeny v tab. č. 8. V závorce je jako první uvedena hodnota očekávané četnosti a jako druhá hodnota vypočtená ze vztahu  $\frac{(P - O)^2}{O}$

Tab. č. 8 – vypočtené hodnoty k modelovému příkladu 3

| Pohlaví | Utracené peníze  |                  |                  | Marginální četnosti ( $n_{.}$ ) |
|---------|------------------|------------------|------------------|---------------------------------|
|         | 0 – 200 Kč       | 201 – 300 Kč     | 301 a více Kč    |                                 |
| Chlapci | 25 (31,08; 1,19) | 45 (38,32; 1,16) | 16 (16,60; 0,02) | 86                              |
| Dívky   | 48 (41,92; 0,88) | 45 (51,68; 0,86) | 23 (22,40; 0,02) | 116                             |

|                               |    |    |    |                              |
|-------------------------------|----|----|----|------------------------------|
| Marginální četnosti ( $n_s$ ) | 73 | 90 | 39 | Celková četnost<br>$n = 202$ |
|-------------------------------|----|----|----|------------------------------|

$$\chi^2 = \sum \frac{(P - O)^2}{O} = 1,19 + 1,16 + 0,02 + 0,88 + 0,86 + 0,02 = 4,13$$

Zjistili jsme, že hodnota námi vypočteného testového kritéria je rovna 4,13. Je však nutné ještě zjistit počet stupňů volnosti, který vypočteme z jednoduchého vzorce, kde  $r$  je počet řádků kontingenční tabulky a  $s$  počet sloupců.

$$f = (r - 1) \cdot (s - 1)$$

(Chráška, 2007, s. 78) [3]

Počet stupňů volnosti pro náš modelový příklad je tedy:

$$f = (r - 1) \cdot (s - 1) = (2 - 1) \cdot (3 - 1) = 2$$

Máme tedy 2 stupně volnosti. Nyní již máme všechny potřebné výpočty k tomu, abychom mohli říci, zda existují či neexistují statisticky významné rozdíly. Vypočtená hodnota testová kritéria 4,13, je při porovnání s kritickou hodnotou 5,991 pro 2 stupně volnosti a zvolenou hladinu významnosti 0,05 nižší. Nelze odmítnout nulovou hypotézu a konstatujeme, že neexistují statisticky významné rozdíly mezi zkoumanými proměnnými. Naše věcná hypotéza se tedy nepotvrdila.

Důležitou poznámku na závěr této subkapitoly je informace, kdy nelze použít testu nezávislosti chí-kvadrát: „*ve více než 20 % polí kontingenční tabulky jsou očekávané četnosti menší než 5 a v případě, že v některém poli je očekávaná četnost menší než 1.*“ (Chráška, 2007, s. 78)

Uvedený proces je již složitější a s narůstajícím počtem polí tabulky je ruční výpočet zdouhavější a především, čím více něco ručně počítáme, tím stoupá i šance na naši početní chybu.

### 3.3.2.3 Test nezávislosti chí-kvadrát pro čtyřpolní tabulku

Jedná se v podstatě o podobu kontingenční tabulky, ale čtyřpolní tabulku využíváme v případě, kdy námi zjišťované proměnné mohou nabývat pouze dvou alternativ (muž; žena; lyžař – nelyžař apod.) (Chráska, 2007).

Někdy se ve vlastních dotaznících objevují otázky, ke kterým jsou nabízeny pouze dvě alternativy (zpravidla ano – ne). Ovšem pouze na několik málo otázek lze takto kategoricky odpovědět. Blíže se k tomuto tématu dostaneme v části věnované sestavování dotazníku, kde budou uvedeny i konkrétní příklady.

Test nezávislosti chí-kvadrát pro čtyřpolní tabulku Vám bude jistě sympatický, protože je zde použit zjednodušený výpočet, který nevyžaduje dílčí určení očekávaných četností (i když lze pochopitelně využít postupu, který je uveden u kontingenční tabulky). Využíváme vzorec, který je uveden níže (Chráska, 2007).

$$\chi^2 = n \cdot \frac{(ad - bc)^2}{(a + b) \cdot (a + c) \cdot (b + d) \cdot (c + d)}$$

Tab. č. 9 – schéma čtyřpolní tabulky

|             | $\alpha$   | non $\alpha$ |            |
|-------------|------------|--------------|------------|
| $\beta$     | <b>a</b>   | <b>b</b>     | <b>a+b</b> |
| non $\beta$ | <b>c</b>   | <b>d</b>     | <b>c+d</b> |
|             | <b>a+c</b> | <b>b+d</b>   | <b>n</b>   |

(Chráska, 2007, s. 83) [4]

Opět zde využijeme modelový příklad uvedený níže. Ověřujeme hypotézu: *Muži doma častěji uklízejí než ženy.*

Tab. č. 10 – modelový příklad 4

|      | Uklízí    | Neuklízí  |           |
|------|-----------|-----------|-----------|
| Muži | <b>25</b> | <b>14</b> | <b>39</b> |

|      |    |    |    |
|------|----|----|----|
| Ženy | 7  | 24 | 31 |
|      | 32 | 38 | 70 |

Jednotlivé četnosti vsadíme do výše uvedeného vzorce č. 4 a vypočtená hodnota je rovna výslednému testovému kritériu.

$$\chi^2 = 70 \cdot \frac{(25 \cdot 24 - 14 \cdot 7)^2}{39 \cdot 32 \cdot 38 \cdot 31} = 70 \cdot \frac{(600 - 98)^2}{1470144} = 70 \cdot \frac{252004}{1470144} = 70 \cdot 0,17 = 11,90$$

Čtyřpolní tabulka má vždy 1 stupeň volnosti. Hodnota testového kritéria nám vyšla 11,90<sup>31</sup>. Při srovnání s hodnotou kritickou pro zvolenou hladinu významnosti 0,05 a 1 stupeň volnosti (3,841) zjišťujeme, že vypočtená hodnota je vyšší než hodnota kritická. Nulová hypotéza by v námi uvedeném příkladu byla: *Muži i ženy uklízejí doma stejně často*, zamítáme tedy nulovou hypotézu. Vzhledem k námi formulované věcné hypotéze a uvedeným četnostem konstatujeme, že se hypotéza potvrdila.

Test nezávislosti pro čtyřpolní tabulku je využitelný tehdy, pokud je celková četnost vyšší než 40 a v žádném poli tabulky není očekávaná četnost menší než pět.

*V případě využití chí-kvadrátu k testování jednostranných hypotéz je nutné dbát především na jeho interpretaci, resp. interpretaci výsledku, a to ve vztahu k věcné hypotéze! Samotné zjištění statisticky významných rozdílů neznamená potvrzení hypotézy.*

*U čtyřpolní tabulky se se zásadním problémem interpretace nesetkáváme, stejně tak u tabulky se dvěma řádky a jedním sloupcem, ale kontingenční tabulky (mimo čtyřpolní) již objektivní riziko přinášejí<sup>32</sup>. Existující riziko bychom měli eliminovat již ve fázi přípravy výzkumu, formulace hypotéz a volbou plánovaných testů.*

### 3.3.2.4 Studentův t-test

Než se pustíte do čtení této kapitoly, seznamte se nejprve se slovníkem doplňujících pojmů a kapitolou, která se věnuje číselnému popisu dat.

<sup>31</sup> Ve statistickém softwaru byla hodnota testového kritéria 11,99.

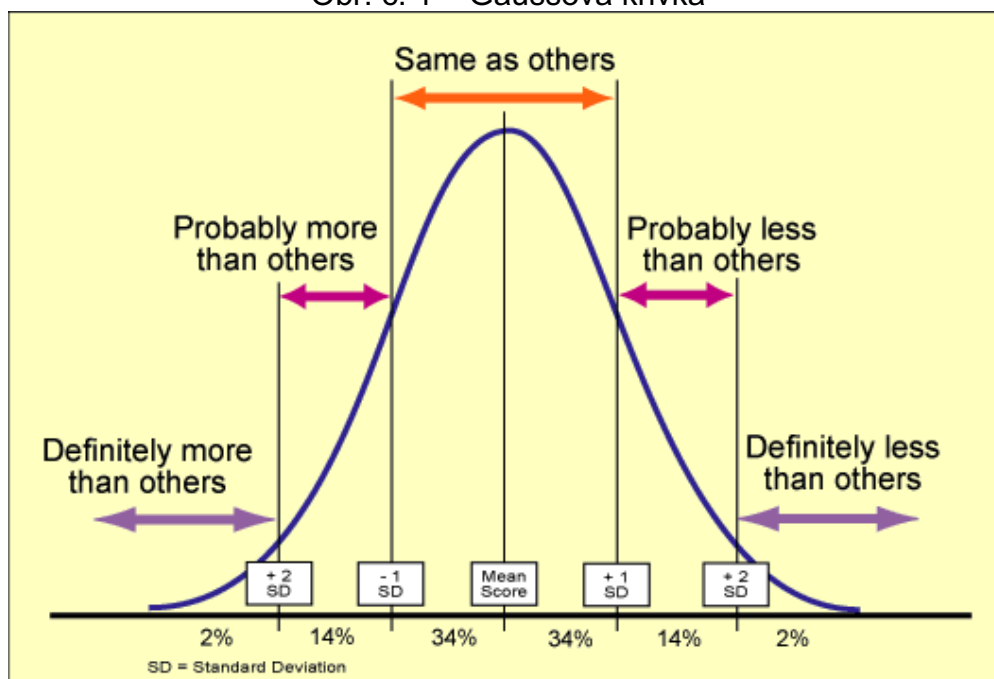
<sup>32</sup> Obecně lze říci, že čím více řádků a sloupců, tím nižší reálnost verifikace věcné hypotézy, která se blíží věštbě z křišťálové koule.

Praktické využití t-testu je sice v BP a DP možné, ale v oblasti sociálních věd je velmi omezené. Zmiňujeme ho zde především z důvodu, aby měli studenti usnadněno studium odborných článků v kvalitních časopisech, kde bylo při výzkumech t-testu využito. Pochopitelně jej zmiňujeme i z důvodu, aby jej mohli sami využít.

Studentův t-test je základním testem pro testování hypotéz u metrických dat. Částečně jej lze využít i v případech, kdy používáme různé typy škál, ale tím přebíráme již určité riziko při oprávnění jeho užití v takových případech a výsledky mohou být zkresleny (Chráska, 2007).

T-test je tzv. testem parametrickým, který vyžaduje splnění určitých podmínek, aby bylo jeho užití správné a výsledky interpretovatelné. Nejdůležitější podmínkou je, že rozdělení metrických dat musí být normální. Tedy musí vzhledově připomínat Gaussovu křivku (viz obr. č. 1) (Hendl, 2009; Chráska, 2007).

Obr. č. 1 – Gaussova křivka



(<http://aweekinthelifeofaredhead.com/a-bell-curve/>)

Na obrázku č. 1 je střední hodnota ( $\mu$  – mean score) a směrodatná odchylka ( $\sigma$  – standard deviation). Pro normální rozdělení platí, že:

- interval  $\mu \pm \sigma$  obsahuje 68,3 % hodnot,
- interval  $\mu \pm 2\sigma$  obsahuje 95,5 % hodnot,



- interval  $\mu \pm 3\sigma$  obsahuje 99,7 % hodnot (Hendl, 2009, s. 146).

Výsledek t-testu nám jako v případě chí-kvadrátu říká, zda existují či neexistují statisticky významné rozdíly, resp. zda odmítáme či nelze zamítnout nulovou hypotézu. Princip je tedy stejný. Vypočítáváme testové kritérium  $t$ , které srovnáváme s hodnotou kritickou (pro zvolenou hladinu významnosti) stanovenou v tabulkách a pro příslušný počet stupňů volnosti  $f$ . Ve vzorcích je  $n_1$  četnost prvního souboru,  $n_2$  četnost druhého souboru,  $s$  výběrová směrodatná odchylka,  $\bar{x}_1$  je průměr prvního souboru,  $\bar{x}_2$  je průměr druhého souboru,  $x_{1i}$  je konkrétní hodnota prvního souboru a  $x_{2i}$  je konkrétní hodnota druhého souboru (Chráska, 2007).

$$t = \frac{\bar{x}_1 - \bar{x}_2}{s} \sqrt{\frac{n_1 n_2}{n_1 + n_2}}$$

(Chráska, 2007, s. 122) [5]

$$s^2 = \frac{1}{n_1 + n_2 - 2} \left[ \sum (x_{1i} - \bar{x}_1)^2 + \sum (x_{2j} - \bar{x}_2)^2 \right]$$

(Chráska, 2007, s. 123)<sup>33</sup> [6]

$$f = n_1 + n_2 - 2$$

(Chráska, 2007, s. 124) [7]

Výše uvedené vzorce využíváme při tzv. dvouvýběrovém t-testu, tedy máme dvě na sobě nezávislé skupiny. Existuje i varianta pro situaci, kdy provádíme opakované měření u jedné skupiny, jedná se o tzv. párový t-test<sup>34</sup>. V tom případě užíváme vzorce.

<sup>33</sup> Ve vzorcích č. 5 a 6 je  $t$  (hodnota testového kritéria),  $n_1$  a  $n_2$  jsou četnosti v jednotlivých skupinách,  $\bar{x}_1$  a  $\bar{x}_2$  jsou průměry v jednotlivých skupinách,  $s$  je směrodatná odchylka a  $x_{1i}$  a  $x_{2i}$  jsou jednotlivé naměřené hodnoty v příslušných skupinách (Chráska, 2007).

<sup>34</sup> Třetí možností je tzv. jednovýběrový t-test, kdy porovnáváme výběrový soubor se základním souborem. Musíme však znát data základního souboru.

$$t = \frac{\bar{d} * \sqrt{n * (n - 1)}}{\sqrt{\sum (d - \bar{d})^2}}$$

[8]

$$\sum (d - \bar{d})^2 = \sum d^2 - \bar{d} \sum d$$

[9]

$$f = n - 1$$

(Chráska, 2007, s. 130-131)<sup>35</sup> [10]

Výpočet samotného t-testu již neprovádíme ručně, resp. můžeme, ale je to práce pro matematické masochisty. Při představě např. 250 respondentů v jedné a 185 respondentů ve skupině druhé je výpočet záležitostí dlouhého večera a především vysokého rizika chyby ve výpočtu. Výše uvedené vzorce jsme uvedli pouze pro představu a pochopení zákonitostí, které hodnoty výsledek ovlivňují. Výpočet si ukážeme v krátkém obrazovém náhledu, kdy byl využit program Microsoft Excel (v aktuální verzi – dále jen Excel).

Smyslem a cílem této části není počítat nějaký relativně reálný příklad, ale vysvětlit studentům ovládání Excelu jako významného pomocníka při t-testu.

Jako první krok je nutné si do programu Excel doinstalovat doplněk s názvem Analýza dat (analytické nástroje). Místo obrázků budou využita videa. U některých i s mluveným komentářem.

Video č. 1 – instalace analytických doplňků ([klikni zde](#))

Nyní si ve vybraném listu Excelu nastavte sloupec B jako soubor č. 1 a sloupec C jako soubor č. 2. V našem modelovém příkladu předpokládáme, že data mají normální rozdělení. Zadaná data již máme, nyní je nutné určit, zda v obou výběrech

---

<sup>35</sup> Ve vzorcích k párovému t-testu je  $d$  difference mezi hodnotami u jednoho páru (např. pokud běžím 100 m při prvním pokusu za 12 s a při druhém za 11 s je  $d = -1$ ; v opačném případě  $+1$ ),  $n$  je počet párů a  $\bar{d}$  je průměrná difference (Chráska, 2007, s. 130).

jsou shodné či rozdílné rozptyly<sup>36</sup>. To zjistíme tzv. F-testem. Klikněte na záložku data a poté vpravo na ikonu s názvem *Analýza dat* a zobrazí se následující tabulka (viz obr. č. 4) a vyberte možnost *Dvouvýběrový F-test pro rozptyl*.

Video č. 2 – F-test ([klikni zde](#))

V Excelu, ale i při ručním výpočtu je vyžadováno, aby jako první (v čitateli či sloupci) byla zadána vyšší hodnota rozptylu. Při ručním výpočtu to poznáte sami. Při využití Excelu máte 50% šanci, že to „trefíte“ (pokud jste si rozptyl nevypočítali předem). Zda jste to učinili správně, popíšeme na následujícím obrázku (Chráska, 2007).

Pokud oba soubory zaměníte, nemá to vliv na *p-hodnotu*. Má to však vliv na testové kritérium F a hodnotu kritickou!

Obr. č. 2 – analýza F-testu

The screenshot shows the 'F-Test Two-Sample for Variances' dialog box in Excel. The 'Variable 1' is 'Mean' with a value of 7,85 and 'Variable 2' is 'Variance' with a value of 7,45. The 'Observations' are 40 for both, resulting in 'df' of 39. The calculated 'F' value is 1,65480957. The 'P(F<=f) one-tail' is 0,06001836. The 'F Critical one-tail' is 1,70446507. A red box highlights the 'Variance' values, and a green box highlights the 'F Critical one-tail' value. A red text box notes that the first sample has a larger variance and the values were entered correctly. A blue text box identifies the 'F' value as the test criterion. A green text box identifies the 'F Critical one-tail' value as the critical value for the given significance level of 0,05.

|    | A                               | B          | C          | D | E |
|----|---------------------------------|------------|------------|---|---|
| 1  | F-Test Two-Sample for Variances |            |            |   |   |
| 2  |                                 |            |            |   |   |
| 3  |                                 | Variable 1 | Variable 2 |   |   |
| 4  | Mean                            | 7,85       | 7,45       |   |   |
| 5  | Variance                        | 45,5666667 | 27,5358974 |   |   |
| 6  | Observations                    | 40         | 40         |   |   |
| 7  | df                              | 39         | 39         |   |   |
| 8  | F                               | 1,65480957 |            |   |   |
| 9  | P(F<=f) one-tail                | 0,06001836 |            |   |   |
| 10 | F Critical one-tail             | 1,70446507 |            |   |   |
| 11 |                                 |            |            |   |   |
| 12 |                                 |            |            |   |   |
| 13 |                                 |            |            |   |   |

<sup>36</sup> Rozptyl je mírou variability metrických dat. Využívá se především v inferenční statistice v rámci testovacích statistik. Rozptyl je: „průměrná kvadratická odchylka měření od aritmetického průměru.“ (Hendl, 2009, s. 102).

Na obr. č. 2 vidíme všechny potřebné údaje pro to, abychom mohli rozhodnout, zda jsou rozptyly shodné či nikoliv. Vidíme, že hodnota testového kritéria je nižší než hodnota kritická, a tudíž můžeme tvrdit, že rozptyly jsou v obou souborech shodné. Pokud by bylo vypočtené testové kritérium vyšší než hodnota kritická, byly by rozptyly rozdílné.

Program Excel nám umožňuje vypočítat t-test i v případě nerovnosti rozptylů. V takovém případě v nabídce vybereme možnost – *dvouvýběrový t-test s nerovností rozptylů*.

Na obrázku č. 2 je uvedena ještě jedna hodnota, a to tzv. *p-hodnota*. V odborných časopisech se častěji setkáte s tzv. p-hodnotami, než s hodnotami testovými a kritickými. P-hodnota nám však poskytuje stejný údaj, s jehož pomocí dojdeme ke stejnému závěru (např. při t-testu, tak chí-kvadrátu, f-testu). Pokud je *p-hodnota* **vyšší**, než námi zvolená hladina statistické významnosti, není mezi srovnávanými daty žádný rozdíl. Pokud je však *p-hodnota* nižší než námi zvolená hladina významnosti, tak rozdíly existují.

V našem příkladu jsme si nastavili hladinu významnosti 0,05, p-hodnota nám vyšla 0,06. Vidíme, že *p-hodnota* (0,06) je vyšší než zvolená hladina významnosti (0,05), a tudíž není mezi srovnávanými daty žádný rozdíl.

**POZOR tedy jakou hodnotu pro interpretaci výsledků a prezentaci závěrů využíváte!!! Zda testové kritérium a jeho konfrontaci s kritickou hodnotou, a nebo p-hodnotu a její srovnání s hladinou významnosti.**

Nyní již tedy víme, že funkci, kterou musíme k našemu modelovému příkladu použít, je dvouvýběrový test se shodností rozptylů (v případě opačného výsledku vybíráme přirozeně test odpovídající, který se nachází ve stejném seznamu).

Znovu klikneme na *Analýzu dat* a z nabídky vybereme *dvouvýběrový t-test s rovností rozptylů*.

## Video č. 3 – dvouvýběrový t-test s rovností rozptylů [\(klikni zde\)](#)

Obr. č. 3 – vyhodnocení t-testu

|    | A                                           | B          | C          | D | E | F                                                                                                  |
|----|---------------------------------------------|------------|------------|---|---|----------------------------------------------------------------------------------------------------|
| 1  | t-Test: Two-Sample Assuming Equal Variances |            |            |   |   |                                                                                                    |
| 2  |                                             |            |            |   |   |                                                                                                    |
| 3  |                                             | Variable 1 | Variable 2 |   |   |                                                                                                    |
| 4  | Mean                                        | 7,85       | 7,45       |   |   |                                                                                                    |
| 5  | Variance                                    | 45,5666667 | 27,5358974 |   |   |                                                                                                    |
| 6  | Observation                                 | 40         | 40         |   |   |                                                                                                    |
| 7  | Pooled Variance                             | 36,5512821 |            |   |   |                                                                                                    |
| 8  | Hypothesized                                | 0          |            |   |   |                                                                                                    |
| 9  | df                                          | 78         |            |   |   |                                                                                                    |
| 10 | t Stat                                      | 0,2958855  |            |   |   | Hodnota testového kritéria t. Hodnota může vyjít i záporně. Vždy hodnotíme její absolutní hodnotu. |
| 11 | P(T<=t) one-tail                            | 0,38405192 |            |   |   | p-hodnota pro jednostranný t-test                                                                  |
| 12 | t Critical one-tail                         | 1,66462464 |            |   |   | Kritická hodnota pro jednostranný t-test                                                           |
| 13 | P(T<=t) two-tail                            | 0,76810383 |            |   |   | p-hodnota pro oboustranný t-test                                                                   |
| 14 | t Critical two-tail                         | 1,99084707 |            |   |   | Kritická hodnota pro oboustranný t-test                                                            |
| 15 |                                             |            |            |   |   |                                                                                                    |

Využití t-testu si v oblasti školství lze představit při komparaci výsledků testů.

Princip interpretace výsledků je shodný s tím, jak interpretujeme výsledky např. u testu dobré shody chí-kvadrát (viz výše). V příkladu (viz obr. č. 3) zjišťujeme, že hodnota testového kritéria  $t$  je nižší, než hodnota kritická pro oboustranný test (i jednostranný). Nebyly zjištěny statisticky významné rozdíly na hladině významnosti

0,05 (viz video). Kterou hodnotu využijeme k porovnání je dáno tím, jak je formulována naše hypotéza – oboustranně či jednostranně.

### 3.3.2.5 Pearsonův koeficient korelace

Pokud jsme ruční výpočet t-testu označili za matematický masochismus, tak ruční výpočet Pearsonova koeficientu korelace je již vyslovená zvrácenost. Využíváme jej pro analýzu určitých metrických dat, kdy hledáme vztah mezi dvěma soubory dat. Výsledek Pearsonova koeficientu korelace nabývá hodnot od -1 do +1. Čím více se přibližuje k oběma krajním hodnotám, tím je vztah těsnější. Záporný výsledek značí, že samotný vztah je záporný. Kladný výsledek značí pozitivní vztah.

Kladný výsledek značí, že nižším hodnotám jedné proměnné odpovídají i nižší hodnoty druhé proměnné a zároveň vyšším hodnotám jedné proměnné odpovídají vyšší hodnoty druhé proměnné (Hendl, 2009; Chráska, 2007).

Záporný výsledek značí, že vyšším hodnotám jedné proměnné odpovídají nižší hodnoty druhé proměnné a naopak (Hendl, 2009; Chráska, 2007).

Budeme to opět demonstrovat na příkladu.

Máme zjistit, jaký vztah je mezi výsledky žáků v testech z matematiky a českého jazyka.

Video č. 4 – Pearsonův koeficient [\(klikni zde\)](#)

Ve videu č. 4 jsou zaznamenány jednotlivé hodnoty, které pochopitelně musí být ve stejných jednotkách (v tomto případě získaných bodech) Bodové hodnoty v jednotlivých řádcích odpovídají vždy jednomu žákovi.

Výpočet požadovaného koeficientu provedeme jednoduše v Excelu. Vyberte si libovolnou buňku, kde chcete, aby se Vám zobrazil výsledek a následně postupujte dle příslušného videa.

Pak jen stačí kliknout na OK a máme výsledek Pearsonova koeficientu korelace. Výsledek korelačního koeficientu  $r_p$  je roven 0,909. Vidíme, že vztah je pozitivní. Tedy, že jedinci, kteří dosáhli lepšího výsledku v matematice, dosáhli i lepšího výsledku v českém jazyce.

Podle tabulky, kterou uvádí Chráska (2007, s. 105), se jedná o střední sílu pozitivního vztahu.

Tab. č. 11 – interpretace Pearsonova koeficientu korelace

| Koeficient korelace  | interpretace                           |
|----------------------|----------------------------------------|
| $r = 1$              | naprostá závislost (funkční závislost) |
| $1,00 > r \geq 0,90$ | velmi vysoká závislost                 |
| $0,90 > r \geq 0,70$ | vysoká závislost                       |
| $0,70 > r \geq 0,40$ | střední (značná) závislost             |
| $0,40 > r \geq 0,20$ | nízká závislost                        |
| $0,20 > r \geq 0,00$ | velmi slabá závislost                  |
| $r = 0$              | naprostá nezávislost                   |

Pásma, která udává Chráska, lze analogicky využít i pro záporné hodnoty. Samotný korelační koeficient je náchylný k extrémním hodnotám. Jeho výsledek neznamena, že prokazuje příčinnou souvislost, nejsou rozlišeny závislé a nezávislé proměnné a vyjadřuje pouze sílu lineárního vztahu (Hendl, 2009).

Pearsonův koeficient korelace je zde zmiňován především z toho důvodu, abyste mu lépe porozuměli při čtení výzkumů v odborných časopisech, monografiích apod.

Pro ty, kteří by si rádi vyzkoušeli ruční výpočet, je níže uveden vzorec, kde  $n$  je četnost dvojic hodnot  $x_i$  a  $y_i$  (Chráska, 2007, s. 114):

$$r_p = \frac{n \sum x_i y_i - \sum x_i \cdot \sum y_i}{\sqrt{[n \sum x_i^2 - (\sum x_i)^2] \cdot [n \sum y_i^2 - (\sum y_i)^2]}}$$

[11]

### 3.3.3 Chyba I. a II. druhu

Při testování hypotéz je nutné si uvědomit, že se jedná o pravděpodobnost a ne o 100% tvrzení. U dvou druhů chyb má zásadní význam zvolená hladina významnosti (viz výše).

*„Chyba prvního druhu spočívá v tom, že neoprávněně odmítneme nulovou hypotézu, ač je správná.“* (Chráška, 2007, s. 70)

Např. zvolíme hladinu významnosti 0,05, na této hladině zjistíme statisticky významné rozdíly. Pokud bychom však hodnotu testového kritéria srovnávali s kritickou hodnotou pro hladinu významnosti 0,01, tak statisticky významné rozdíly zjištěny nebyly.

Chyba druhého druhu je svým způsobem opačným příkladem. Zvolíme si hladinu významnosti na 0,001, kde nezjistíme žádné statisticky významné rozdíly, které by však zjištěny byly na hladině významnosti 0,01. Neoprávněně tedy neodmítáme nulovou hypotézu, i když je platná (Chráška, 2007).

Výše uvedené informace nemají být návodem k „čarování“ s hladinou významnosti v průběhu testování hypotéz.

## 3.4 Metody sběru dat v kvantitativním výzkumu

Metod sběru dat v kvantitativním výzkumu je relativně mnoho. My se v této subkapitole budeme zabývat pouze dotazníkem, který je nejrozšířenější metodou užívanou studenty v bakalářských pracích.

### 3.4.1 Dotazník

*„Dotazník je psaný soubor otázek. Jedná se o metodický nástroj výzkumu zjišťování informací o osobních znalostech, postojích k aktuální skutečnosti a hodnotových preferencích.“* (Bartošová, Skutil in Skutil a kol., 2011, s. 80)

Tím, že je nejčastější, je zároveň dotazník i velmi podceňován. Jedním z důvodů může být i to, že se s ním zpravidla každý z nás již setkal, a to minimálně v pozici respondenta. Jednotlivé otázky však musí být řádně promyšleny, formulovány apod. V dotazníku nelze otázky pokládat a řadit za sebe v pořadí, jak nás zrovna napadají.



Pro potřeby BP, ale i DP si zpravidla sestavujeme dotazníky vlastní, nevyužíváme dotazníků standardizovaných, i když ani to se nevylučuje. O standardizovaných dotaznících můžeme uvažovat např. u témat spojených se syndromem vyhoření u učitelů.

Punch (2008b, s. 47) k výběru, zda zvolit vlastní či již sestavený dotazník, poznamenává: *„Již bylo provedeno mnoho výzkumných šetření, v mnoha zemích. Proto existuje hodně výborných dotazníků. Hlavním problémem je jejich vyhledání. Také existují jejich sbírky a pravidelné studium vhodných časopisů nebo výzkumných zpráv nás nasměřuje ke vhodným měřicím nástrojům. [...] Každý odborník si v průběhu doby vytváří pomocí tohoto studia sbírku dotazníků. Pro studenta však může vyhledání vhodného dotazníku znamenat dost práce.“*

Než se tedy pustíme do sestavování vlastního dotazníku, měli bychom se podívat, jak to dělali jiní. Taková drobná „rešerše“ může být pro nás velice inspirativní, ovšem nesmíme se inspirovat příliš (musíme dodržet citační etiku).

Je nutno pamatovat, pakliže se rozhodneme využít již existující dotazník, že některé dotazníky jsou zpoplatněny, vyžadují určité vzdělání apod. Nevylučuje se pochopitelně kombinace dotazníku (otázek) vlastní konstrukce a dotazníků existujících (Punch, 2008b).

Všechny odborné publikace se shodují v tom, že zvolený měřicí nástroj musí mít vysokou reliabilitu a validitu<sup>37</sup>. V tom se pochopitelně shodujeme s autory, ale BP s sebou přinášejí specifická témata, lokality apod., kdy jsou daná data většinou využitelná a platná pouze pro danou skupinu respondentů. Studenti a studentky by se neměli bát pouštět do vlastních výzkumných šetření s vlastními dotazníky. Ovšem poté musí umět získaná data správně prezentovat.

### **Výhody dotazníku:**

- „snadná a rychlá administrace,
- lze oslovit větší počet respondentů, a tím získat značné množství údajů,
- je možné získat informace, které nejsme schopni získat jinou technikou,

---

<sup>37</sup> Viz slovník doplňujících pojmů

- *údaje lze většinou plně kvantifikovat.*“ (Bartošová, Skutil in Skutil a kol., 2011, s. 80)

### **Nevýhody dotazníku:**

- *„musíme počítat se subjektivitou výpovědí,*
- *je možné, že se respondent otázce vyhne<sup>38</sup>,*
- *respondentovi vždy nemusí vyhovovat daná forma dotazování,*
- *nemožnost dovysvětlení otázky v případě, kdy sami nebudeme dotazník administrovat,*
- *přesnost vymezených otázek a variant odpovědí striktně omezuje prostor pro odpovědi respondenta, takže někdy je nucen zvolit variantu, kterou by jinak nezvolil,*
- *možnost zkreslení odpovědí žádoucím směrem.*“ (Bartošová, Skutil in Skutil a kol., 2011, s. 80-81)

#### **3.4.1.1 Otázky v dotazníku**

Odvíjí se od formulovaných výzkumných otázek. Do dotazníku volíme otázky podle toho, co chceme zjistit. (Punch, 2008b).

Nesmíme zapomínat ani pro koho, pro jakou cílovou skupinu respondentů je formulujeme. Upravujeme vhodně styl, kterým píšeme, odbornost, popř. můžeme zapojit i specifické slangové výrazy, pokud je to vyžadováno, resp. žádoucí. Lépe se tím k respondentům přiblížíme, a to i přesto, že jsme vybrali relativně „neosobní“ dotazník. Někdy pracujeme s pojmy, které jsou sice známé, ale zároveň víme, že je mohou respondenti chápat rozdílně. Je tedy někdy vhodné, aby byly dané pojmy v dotazníku striktně vysvětleny (Chráška, 2007).

Otázky vždy musí být formulovány tak, aby z nich explicitně vyplýval vztah k našemu výzkumu. Např. pokud bychom položili otázku: *Kolik hodin denně trávíš u počítače?* V první řadě si musíme uvědomit, zda jsou respondenti kompetentní to racionálně posoudit. Dejme tomu, že ano. Respondent napíše 9 hodin. Cílovou skupinou byli studenti středních škol v MSK a zabývali jsme se oblastí volného času

---

<sup>38</sup> V takovém případě je nutné dotazník vyřadit.

cílové skupiny. Je odpověď 9 hodin šokující? Ano i ne, resp. vždy záleží na tom, co na počítači dělá. Z dané otázky a odpovědi nevyplývá, že respondent tráví devět hodin svého volného času na počítači. Dozvíme se pouze, že stráví devět hodin u počítače. Existuje celá řada podobných příkladů.

***Na to vše je nutné pamatovat při sestavování dotazníku a podtrhuje to důležitost realizace předvýzkumu.***

Rozlišujeme několik základních typů otázek: uzavřené, polouzavřené, otevřené, testové (těm se nebudeme podrobněji věnovat), škálovací (Bartošová, Skutil in Skutil a kol., 2011)<sup>39</sup>.

Níže budou uvedeny příklady k jednotlivým typům i tipy, na co si dát pozor. Zvolené příklady a ukázky jsou pouze exemplární, nemají žádnou hlubší logiku. Byly zvoleny z čisté fantazie autora a též jeho zkušeností.

- **Otázky uzavřené:** od respondenta vyžadujeme odpověď, kterou mu striktně nabízíme<sup>40</sup>. Jedná se o relativně jednoduše a rychle vyhodnotitelný dotazník, a to v případě, že v něm používáme pouze uzavřené otázky. Např. *Bylo by žádoucí obnovit trest smrti?* Nabídneme např. možnosti: *ano – ne – nevím (neumím posoudit)*.

Stává se, že student respondentovi nabídne pouze možnosti *ano a ne*<sup>41</sup>. Zapomínáme právě na možnost *nevím*. Respondent se pak neumí rozhodnout, nemá dostatek informací, necítí se kompetentní ke kategorické odpovědi a kvůli tomu může odpověď vynechat. Typickou otázkou, kde se odpověď *nevím* nabízí a proč je tedy nutné ji nabízet, by byla: *„Souhlasíte se současnou zahraniční politikou Toga?“*

U uzavřených otázek, které použijeme v dotazníku, narážíme také na jednu z nevýhod dotazníku (viz výše). Respondenta donutíme k odpovědi, resp.

---

<sup>39</sup> Níže rozepsané teoretické poznámky k jednotlivým typům otázek vycházejí ze zde citované literatury (Bartošová, Skutil in Skutil a kol., 2011). Vzhledem k tomu, že autor využívá vlastních postřehů, rozhodl se zde citovat tímto způsobem. Aby text nebyl zbytečně rušen.

<sup>40</sup> Pokud nabízíme pouze dvě možné odpovědi, jedná se o otázky *dichotomické*. Pokud nabízíme více odpovědí, tak hovoříme o otázkách *polytomických*. Existuje i nabídka možnosti vícenásobné odpovědi.

<sup>41</sup> Pochopitelně existují otázky, kdy mohou být nabídnuty pouze možnosti *ano – ne*. Např. *Vlastníte osobní automobil?*

respondentovi nabízíme takové možnosti, z nichž ani jedna mu nevyhovuje. Předejít tomu můžeme nabídkou možnosti *jiná, jiné* apod. To by nám měl odhalit předvýzkum.

V některých případech též nabídneme respondentovi do jedné možnosti odpovědi dvě alternativy. Např. „*Odkud se dozvídáte informace k dané problematice?*“ Respondentovi nabídneme jako jednu z odpovědí: *od rodičů, a od kamarádů*. Jeho zdrojem však mohou být pouze rodiče či pouze kamarádi. Nevýhodou uzavřených otázek je, že se v některých případech můžeme dostat ke značně zkresleným datům, protože nedokážeme posoudit znalost respondenta o problematice. Např. „*Víš, co je to kyberšikana?*“, respondent odpoví ano. Jedná se o otázku uzavřenou, tudíž nezjistíme, jestli opravdu ví. Pokud bychom se ho zeptali – co to tedy je, mohl by nám odpovědět, že se jedná o druh malé opičky žijící v Kongu.

- **Otázky polouzavřené** – respondentům jsou stejně jako u uzavřených otázek nabídnuty možnosti, ale zároveň mají možnost svou odpověď zdůvodnit. Přirozeně pouze u otázek, kdy je to vhodné. Zpravidla takových, kde se ptáme na určitý názor, resp. k dané problematice. Viz výše příklad s trestem smrti. Z této uzavřené otázky bychom mohli jednoduchým způsobem „udělat“ otázku polouzavřenou. Např. *Bylo by žádoucí obnovit trest smrti? ano – ne – nevím (neumím posoudit). Jestliže odpovíte ano či ne, svou odpověď zdůvodněte.*

Výhodou polouzavřených otázek je, že dáme respondentovi prostor pro vyjádření jeho názoru. Nevýhodu můžeme spatřovat v neochotě respondenta se dlouze rozepisovat. Patrně cítíme, že vyplňování dotazníků není zrovna naší oblíbenou činností, a proto ji chceme mít co nejrychleji za sebou. Jisté úskalí poté spočívá i ve vyhodnocování těchto odpovědí.

- **Otázky otevřené** – respondentovi je ponechán volný prostor pro odpověď. V tomto případě spatřujeme největší riziko ve vyhodnocování získaných odpovědí. Záleží ovšem, na co se respondenta ptáme. Pokud položíme otázku *Kterým zájmovým činností se ve svém volném čase věnujete?*, tak lze předpokládat, že odpovědi nebudou v podobě slohového cvičení, ale spíše v bodech. Pokud položíme otázku, která zjišťuje určitý názor či postoj

respondenta, tak se na slohové cvičení můžeme připravit. Např. *Jaký je Váš názor na trest smrti?*

- **Otázky škálovací** – při užití škálovacích otázek nám jde o zjištění jisté míry vztahu či vlastnosti<sup>42</sup>. Škálovací otázky jsou rozděleny do několika podtypů – Likertovy škály, bipolární škály, numerické, intervalové a sebeposuzovací. V BP či DP zpravidla využíváme škály Likertova typu, škály numerické a intervalové.
  - Škála Likertova typu – respondentům je nabídnuto určité tvrzení a respondenti mají na škále vyjádřit míru souhlasu či nesouhlasu. Např. *T. G. Masaryk byl nejlepší prezident*. Respondenti mají možnost vyjádřit se na škále: *souhlasím – spíše souhlasím – nevím – spíše nesouhlasím – nesouhlasím*. Popř. je daná stupnice reprezentována čísly 1 – 2 – 3 – 4 – 5 (musí však být napsáno, která strana je souhlas a která nesouhlas). Škála by vždy měla být lichá ať již tří-, pěti-, sedmi- nebo devítibodová.
  - Škála numerická – v podstatě se jedná o školní známkování. Setkat se s ním můžeme např. v médiích, kdy lidé vystavují „známku“ politikům. Dávají tak najevo míru spokojenosti s jejich prací apod. Respondenti, kterých se ptáme v BP či DP, mohou vyjadřovat svojí spokojenost s nabízenými službami nějaké instituce apod.<sup>43</sup>
  - Škály intervalové – u nabízených možností existuje reálná možnost subjektivnosti jednotlivých pojmů. Typická intervalová škála může mít následující podobu. Např. u otázky: *Jak často se věnujete strukturovaným pohybovým aktivitám?* *denně – občas – někdy – výjimečně – nikdy*. Pro každého může mít termín *občas* naprosto jiný význam. Jedná se tedy o určité riziko, které mělo být eliminováno v předvýzkumu.

---

<sup>42</sup> Bartošová, Skutil in Skutil a kol., 2011, s. 84

<sup>43</sup> Pokud využíváme numerické škály a škály Likertova typu, je s dosaženými výsledky teoreticky možné operovat jako s klasickými čísly. Je zde však určité riziko zkreslení výsledku. Lze vypočítat aritmetický průměr, směrodatnou odchylku apod. (Chráška, 2007).

Např. *Navštěvuješ pravidelně zámky jako turista?* Pro někoho je pravidelnost jednou za 5 let, ale pravidelně. Opět podobných příkladů nalezneme více.

U posledního typu škálovací otázky jsme naznačili, že některé pojmy mohou být subjektivně vnímány. Buď je můžeme vypustit, nebo je přímo v otázce upřesnit (Bartošová, Skutil in Skutil a kol., 2011).

Ve vlastním dotazníku nemusíme mít jeden typ otázek, můžeme je kombinovat, ale vždy účelně. Jak bylo uvedeno výše, otázky volíme podle našich výzkumných otázek, a to nejen jejich znění, ale i typ.

Jako poslední uvedme ještě jedno pravidlo při formulaci otázek. Nikdy by respondentovi neměla být podsouvána odpověď. Otázka by tedy neměla být sugestivní (Chráška, 2007). Nesmíme tedy formulovat otázku takovým způsobem, aby odpovědi byly takové, jaké si přejeme. Např. jako pracovníci fiktivního nízkoprahového zařízení pro děti a mládež bychom klientům položili otázku: *Jste spokojeni s nabídkou našich služeb stejně jako většina našich klientů?*

Dále by otázky neměly vzbuzovat dojem tzv. sociální žádoucnosti odpovědi. V respondentovi by neměl převládnout názor, že nějaká odpověď je obecně správná. Např. *Domníváte se, že by se měly podporovat organizace, které pečují o handicapované?*

Může se stát, že je naše téma poněkud choulostivé, tudíž na přímo položenou otázku by respondenti nemuseli odpovídat pravdivě nebo by se odpovědi vyhnuli. Otázku tedy pokládáme nepřímou. Neptáme se tedy *Co si myslíte o...*, ale *Co si lidé myslí o ...* (Chráška, 2007).

Nikdy bychom se respondenta neměli ptát na otázky, resp. mu položit otázku, kde vyžadujeme jeho odpověď, ale na kterou pravděpodobně odpovědět neumí. Typickým příkladem jsou otázky začínající slovem *Proč*. Na to by měl zjistit odpověď výzkumník (Chráška, 2007)<sup>44</sup>.

---

<sup>44</sup> V některých případech se však může jednat o výzkumný záměr. Je nutné to zvážit již při plánování výzkumu.

## **Respondent by se nikdy neměl cítit otázkou ohrožen!**

Závěrem ještě jedna důležitá poznámka: Otázky by měly být krátké (Gavora a kol., 2010)<sup>45</sup>.

### **3.4.1.2 Podoba a struktura dotazníku**

Vlastní sestavený dotazník nás do jisté míry reprezentuje, jeho vizuální stránka signalizuje respondentovi míru našeho vlastního zájmu. Nezarovnané okraje, 5 typů písma, 4 velikosti písma, nejednotnost úrovní číslování otázek aj., dávají respondentovi jasnou informaci: *Dotazník jsme sestavili za 5 minut*. Je tedy nutné zaměřit se nejen na obsah, ale i na formální stránku. Zároveň musí být dotazník přehledný a snadno čitelný (Gavora a kol., 2010).

Neexistuje jednotná a závazná podoba formální úpravy dotazníku. Můžeme pouze doporučit zarovnání do bloku, „normální“ font (Times New Roman, Arial, Calibri, Garamond) a přirozeně i rozumnou velikost písma, která může být v závislosti na cílové skupině respondentů.

#### **Struktura dotazníku:**

1. Vstupní část – oslovení respondenta, sdělení účelu dotazníku, k čemu budou sloužit výsledky, zdůraznění anonymity, poděkování za čas, pokyny k vyplnění<sup>46</sup> (Bartošová, Skutil in Skutil a kol., 2011).

Všechny části jsou pochopitelně důležité, ale musíme se věnovat zejména pokynům k vyplnění. Na obr. č. 4 je prezentovaný typický příklad, který může nastat, pokud pokyny zapomene zmínit.

Obr. č. 4 – příklad absence pokynů k vyplnění

---

<sup>45</sup><http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/dotaznik/poziadavky-na-polozky.php?id=i12p6>

<sup>46</sup> Srov. s kap. 3

Pohlaví:

- a) muž  
☒ b) žena

Pohlaví:

- a) muž  
☒ b) žena

Jsou oba respondenti ženami? Možná ano, možná ne. Jelikož jsme zapomněli napsat pokyny k vyplnění, tak respondent/ka vpravo se mohla řídit dle zásady: *nehodící se škrtněte*. Musíme tedy jasně stanovit, jak mají být odpovědi vybrány. Vyhnete se spojení 'správná odpověď'. Všechny odpovědi jsou správné (pokud se nejedná o nějaký didaktický test apod.). Pokud používáme škálu Likertova typu či numerickou škálu, nesmíme zapomenout napsat, co která položka reprezentuje.<sup>47</sup>

U některých otázek nabízíme respondentovi možnost vícenásobné odpovědi. Vždy bychom ho na to měli upozornit. Lépe je upozornit vždy u dané otázky, tedy ne se odkazovat na nějakou poznámku ve vstupní části.

2. Hlavní část – pochopitelně obsahuje samotné otázky. Jejich uspořádání by však nemělo být naprosto libovolné. Doporučuje se, aby byly postupně kladeny otázky od jednodušších přes složitější a opět jednodušší. Na začátek zařadit jakési warm-up otázky. Nedílnou součástí jsou rovněž identifikační údaje. Chráska (2007) je doporučuje zařazovat až na konec hlavní části (Bartošová, Skutil in Skutil *akol.*, 2011).
3. Závěr – v této části především poděkujeme respondentům. Někdy je možné zanechat respondentovi prostor k vyjádření k danému dotazníku (Bartošová, Skutil in Skutil *akol.*, 2011).

#### 3.4.1.3 Rady a tipy ke konstrukci a využití vlastního dotazníku

1. Elementárně dobrý dotazník nevznikne za hodinu, za den nebo za týden.
2. Nepokládejte otázky, které nepotřebujete vědět, nemají vztah k vašim výzkumným otázkám a hypotézám (Gavora a kol., 2010).

---

<sup>47</sup> Z vlastní zkušenosti je vhodné to vždy napsat i u jednotlivých otázek.



3. I když považujete otázku za takovou, na kterou by mohla být zajímavá odpověď, tak to není motivem pro její začlenění. Musí mít vždy vztah k vašemu výzkumu (hypotézám, výzkumným otázkám) (Disman, 2008).
4. Nepokládejte otázky, na které sami můžete lehce zjistit odpověď (Chráška, 2007).
5. Pokud je to možné, tak administrujte dotazník sami.
6. Délka vyplňování dotazníku by neměla překročit 30 minut (Gavora a kol., 2010).
7. Pracujte takovým tempem, abyste se nedostali do časové tísně. „Kvapík“ se projeví na celkové kvalitě. Výzkum nejsou závody. Udělejte si časový harmonogram realizace výzkumu.
8. Ve své BP či DP nesestavujte dotazník pro dotazník, ale pro využití výsledků v praxi (i když zpravidla v určité lokalitě či instituci). Důkladně jej tedy promyslete.
9. Nepodceňujte dotazník.

### **3.5    Prezentace a popis získaných dat**

Po všech úskalích, které jsme úspěšně překonali u všech etap kvantitativního výzkumu, přichází fáze poslední a tou je prezentace získaných dat. V BP se prezentace dat bude odlišovat od klasické výzkumné zprávy, resp. bude zjednodušená. Poslední kapitola těchto skript je zaměřena na prezentaci kvantitativních dat.

Pokud bychom zjednodušili obecné struktury závěrečných výzkumných zpráv, a to pro potřeby BP i DP, získali bychom následující strukturu, resp. strukturu praktické části:

- 1. Uvedení do problematiky** – obsahuje zdůvodnění výběru daného tématu; zvolený druh výzkumu (kvalitativní či kvantitativní) a jeho zdůvodnění; formulovaný výzkumný problém, výzkumné otázky či hypotézy (nebo obojí); vhodné je formulovat cíl výzkumu.
- 2. Metodologie výzkumu** – jedná se o popis výzkumu; popis respondentů a druh jejich výběru; kde jsme výzkum realizovali; metody sběru dat; postupy při vyhodnocení; etické aspekty.

3. **Výsledková část (věnováno samotným datům)** – popisujeme získaná data; neočekávané výsledky; verifikujeme hypotézy; odpovědi na výzkumné otázky. Při verifikaci by neměla být využívána věta: „*přijímáme nulovou hypotézu*“ (Hendl, 2009, 636).
4. **Diskuze** – možná je nejtěžší částí celého výzkumu. Jedná se o závěry, které vyplývají ze získaných dat. Výsledky dáváme do souvislosti s teoretickými informacemi v literatuře. Popisujeme faktory, které omezují získaná data, slabé stránky našeho výzkumu. V BP zpravidla z časových a studijních možností není příliš reálné, aby bylo provedeno druhé měření (zejm. v případech, kdy se *p hodnota* u jistých statistických testů pohybuje na hranici statistické významnosti. To souvisí s validitou a reliabilitou dotazníků vlastní konstrukce), je tedy nutné takové skutečnosti v diskuzi zmínit, upozornit na ně a být si toho vědom.
5. **Důsledky a doporučení** – jaké mají naše výsledky dopad pro praxi, pro jejich praktické využití; náměty pro další výzkum; změny zažité praxe aj. (Hendl, 2009; Švaříček, Šedřová a kol., 2007; Zikl in Skutil a kol., 2011).

Existuje celá řada různých statistických výpočtů. O těch, které využijeme, bychom měli mít jasno již před zpracováváním výsledků, i když někdy dochází ke změně, a to z důvodu, že data nesplňují podmínky pro daný test či výpočet. Presentace výsledků by neměla být matematickou olympiádou, kdy chceme předvést naše znalosti a dokázat, že toho dovedeme spoustu spočítat. Vždy by měla být předním hlediskem účelnost a smysluplnost daných výpočtů.

### 3.5.1 Třídění dat

Při prezentaci dat je musíme nějakým způsobem třídit. To bychom však měli mít ujasněno již ve fázi samotné přípravy. Avšak v průběhu realizace sběru dat si můžeme povšimnout některých „zajímavostí“, které nás mohou vést k jinému typu třídění.

Stupňů třídění máme v podstatě neomezeně (pokud by náš dotazník měl neomezený počet otázek), resp. omezeni jsme časem. Čím vyšší stupeň třídění = tím více tabulek, grafů apod.

Jako třídění prvního stupně označujeme takové třídění, kdy jsou respondenti rozděleni na základě jednoho znaku (např. dle pohlaví, věku, vzdělání apod.). Tříděním druhého stupně označujeme takové třídění, kdy jsou respondenti rozděleni podle dvou znaků – např. rozdělíme respondenty na muže a ženy a následně na muže kuřáky, muže nekuřáky a ženy na kuřáčky a nekuřáčky. Tříděním třetího stupně označujeme takové třídění, kdy jsou respondenti rozděleni podle tří znaků – např. rozdělíme respondenty na páry sezdané a nesezdané, podle místa bydliště město – vesnice, a podle výdajů na domácnost. Můžeme takto třídit dále a dále a dále (Chráska, 2007; Hendl, 2009, Punch, 2008b).

### 3.5.2 Číselný popis nominálních (kategoriálních) dat

Tato data představují nejčastější typ dat, které získáváme v rámci výzkumu v BP. Vyhotovují se tabulky četností a grafy. Pamatovat by se však mělo především na přehlednost a účelnost. Často dochází k tomu, že se z výzkumu v BP stane festival barviček v grafech, které jsou vyvedeny přes celou stranu. Někdy však naprosto neúčelně.

Jak vytvořit tabulku četností v Excelu – [klikněte zde](#).

Hendl (2009) hovoří o tom, že je otázkou citu, zda zvolíme tabulky či grafy. Vždy však musí být zvoleny účelně. **Zdůrazňuje, že společné využití tabulky a grafu by mělo být pouze ve výjimečných a důležitých případech.**

Hendl (2009, s. 637) prezentuje principy, kterými bychom se měli řídit při sestavování tabulek (nejen tabulek četností):

- „*Titulek má identifikovat všechny proměnné a kategorie proměnných.*
- *Nadpisy sloupců i řádků mají jasně označit, co řádky nebo sloupce obsahují.*

- *Pokud to je namístě, uvádějí se rozsahy výběrů a skupin.*
- *Pokud to je nutné, jasně se označují statisticky významné výsledky.*
- *Výsledky se uvádějí pouze s požadovaným počtem desetinných míst.*
- *Popis tabulky obsahuje vše, co je potřeba k porozumění (jen výjimečně se odvolává na informace v textu).“*

Citované zásady se pochopitelně vztahují i k tabulkám, které prezentují metrická data. Analogicky lze použít výše uvedené zásady i na grafy.

V BP dochází ke zdvojování či ztrojování informací. Student zkonstruuje tabulky četností obsahující (absolutní i relativní četnost), následně je relativní četnost přenesena i do podoby znázorněné grafem a komentována výčtem v textu pod tabulkou.

Jednotlivé tabulky či grafy mají být vhodně propojeny textem, neměli bychom však opakovat informace v nich uvedených. Tyto informace jsou čtenářům jasné z tabulek či grafů. To platí u všech dat, která jsou vyjádřena „číslem“ (Gavora a kol., 2010).<sup>48</sup>

Nejjednodušším popisem nominálních dat je uvádění absolutní četnosti v tabulce četností a její doplnění četností relativní (uváděnou nejčastěji v procentech). Procenta jsou sice snadno srozumitelné, ale ve své podstatě chudé údaje, které již nelze téměř komentovat.

U číselného popisu nominálních dat zmiňme především tzv. *normovanou nominální varianci*, která popisuje míru variability získaných dat a *normovaný koeficient kontingence*, který udává míru závislosti mezi jevy.

NORMOVANÁ NOMINÁLNÍ VARIANCE popisuje, jak se jednotlivé četnosti rozptylují, jakou mají variabilitu. Tedy rozptýlenost celkové četnosti do jednotlivých možností výběru, které měli respondenti k dispozici. Využití normované nominální variance (NNV) dává možnost srovnání jednotlivých hodnot NNV, aniž by byly zkresleny počtem kategorií<sup>49</sup>.

<sup>48</sup> <http://www.e-metodologia.fedu.uniba.sk/index.php/kapitoly/interpretacia-udajov.php?id=i20>

<sup>49</sup> Ve vzorci č. 12 je  $n$  celkový počet respondentů;  $n_k$  jsou četnosti v jednotlivých kategoriích a  $k$  je počet nabízených kategorií.

$$NNV = \frac{n^2 - \sum n_k^2}{n^2} \cdot \frac{k}{k-1}$$

(Chráška, 2007, s. 58) [12]

Čím více se vypočtená hodnota blíží +1, tím jsou jednotlivé četnosti rovnoměrněji zastoupeny. Čím více se blíží 0, udává nám NNV pravý opak. Výsledku 0 by bylo dosaženo v případě, že všichni respondenti vybrali pouze jednu z nabízených možností. Výsledku +1 by bylo dosaženo v případě, že všechny nabízené možnosti jsou vybrány stejným počtem respondentů.

Praktické využití NNV budeme demonstrovat na následujícím příkladu. NNV může nabývat pouze hodnot od 0 do +1(Chráška, 2007).

Obyvatelé vybraného města v ČR byli v rámci marketingového průzkumu dotazováni, jakou značku auta vlastní a jakou značku jízdního kola vlastní<sup>50</sup>. Celkem bylo dotázáno 100 respondentů.

Tab. č 12. – četnost jednotlivých značek kol

| <b>Značka kola</b> | <b>Četnost</b> |
|--------------------|----------------|
| <b>Olpran</b>      | 35             |
| <b>Merida</b>      | 45             |
| <b>Author</b>      | 20             |
| <b>Σ</b>           | 100            |

Tab. č. 13 – četnost jednotlivých značek aut

| <b>Značka automobilu</b> | <b>Četnost</b> |
|--------------------------|----------------|
| <b>Škoda</b>             | 28             |
| <b>Fiat</b>              | 15             |
| <b>Peugeot</b>           | 16             |
| <b>Mazda</b>             | 20             |
| <b>Toyota</b>            | 21             |
| <b>Σ</b>                 | 100            |

<sup>50</sup> Nepátřejme po logice, smysluplnosti a jiných aspektech, je to jen příklad.

NNV pro tab. č. 12

$$\begin{aligned} NNV &= \frac{100^2 - (35^2 + 45^2 + 20^2)}{100^2} \cdot \frac{3}{3-1} = \frac{10000 - (1225 + 2025 + 400)}{10000} \cdot 1,5 \\ &= \frac{10000 - 3650}{10000} \cdot 1,5 = 0,635 \cdot 1,5 = 0,95 \end{aligned}$$

NNV pro tab. č. 13

$$\begin{aligned} NNV &= \frac{100^2 - (28^2 + 15^2 + 16^2 + 20^2 + 21^2)}{100^2} \cdot \frac{5}{5-1} \\ &= \frac{10000 - (784 + 225 + 256 + 400 + 441)}{10000} \cdot 1,25 = \frac{10000 - 2106}{10000} \cdot 1,25 \\ &= 0,7894 \cdot 1,25 = 0,99 \end{aligned}$$

Co jsme se dozvěděli z uvedených výpočtů? V první řadě jsme zjistili, že NNV pro tab. č. 13 je „nepatrně“ vyšší než NNV pro tab. č. 12. Data v tab. č. 13 jsou tedy rovnoměrněji rozložena než data v tab. č. 12. Vidíme, že ve výběru značky kola existují určité preference určité značky, na rozdíl od výběru auta. Možná si řeknete, k čemu to je dobré. Především je to dobré k lepšímu popisu a charakteristice dat, k určení míry variability a rozptýlenosti jednotlivých četností. V obou tabulkách vidíme, že data „nějak“ rozptýlena jsou, ale nevíme, jak moc, proto NNV počítáme.

NORMOVANÝ (KORIGOVANÝ) KOEFICIENT KONTINGENCE (NKK) nám určuje míru závislosti, vztahu mezi jevy. Počítá se pro kontingenční tabulku a využívá výpočet testového kritéria chí-kvadrát. Jeden výpočet tedy v podstatě využijeme dvakrát. NKK může nabývat hodnot od 0 do +1, přičemž čím vyšší hodnota, tím vyšší závislost mezi proměnnými. Normovaný koeficient kontingence nám opět umožňuje srovnávat jednotlivé výpočty mezi sebou, a to i pokud byly docíleny z kontingenčních tabulek o jiném počtu řádků či sloupců (Hendl, 2009).

$$C_{norm} = \sqrt{\frac{\frac{\chi^2}{n + \chi^2}}{\frac{r-1}{r}}}$$

(Chráška, 2007, s. 86) [13]

Ve vzorci je  $\chi^2$  hodnota testového kritéria chí-kvadrát pro danou kontingenční tabulku,  $n$  je celková četnost v kontingenční tabulce a  $r$  je menší hodnota z počtu řádků či sloupců v kontingenční tabulce (Chráska, 2007).

Opět pro představu uvedme příklad. V našem výzkumu jsme chtěli zjistit, zda existují statisticky významné rozdíly mezi pohlavím a vydělánou mzdou.

Tab. č. 14 – intervaly příjmu dle pohlaví

| Pohlaví                                       | Mzda          |                 |                 | Marginální četnosti ( $n_r$ )      |
|-----------------------------------------------|---------------|-----------------|-----------------|------------------------------------|
|                                               | 8000-12000 Kč | 12001- 18000 Kč | 18001 a více Kč |                                    |
| <b>Muž</b>                                    | 25            | 45              | 105             | <b>175</b>                         |
| <b>Žena</b>                                   | 51            | 68              | 52              | <b>171</b>                         |
| <b>Marginální četnosti (<math>n_s</math>)</b> | 76            | 113             | 157             | <b>Celková četnost<br/>n = 346</b> |

Na základě informací (viz výše) jsme vypočítali hodnotu testového kritéria chí-kvadrát, která je rovna 31,43. Při porovnání s hodnotou kritickou pro 2 stupně volnosti a zvolenou hladinou významnosti zjišťujeme, že mezi zaznamenanými proměnnými existují statisticky významné rozdíly.

Do výše uvedeného vzorce dosadíme všechny potřebné údaje, které již máme k dispozici.

$$C_{norm} = \sqrt{\frac{\frac{31,43}{346 + 31,43}}{\frac{2 - 1}{2}}} = 0,41$$

Vypočtená hodnota ukazuje střední míru závislosti. Opět nám tento výpočet slouží ke zpřesnění jiných výpočtů. Vidíme, že statisticky významné rozdíly existují, ale závislost mezi proměnnými je střední. To znamená, že z našeho příkladu vyplývá jistá statistická závislost mezi mzdovým příjmem mužů a žen, ale tato závislost není silná.

Jednotlivé NKK můžeme mezi sebou porovnávat a opět tak odhalit nové dimenze našich dat.

Z jednotlivých výpočtů NNV a NKK lze vytvořit samostatné tabulky pro lepší přehlednost jednotlivých výpočtů, opatřit je příslušným komentářem, a dostat se tím k lepší charakteristice dat.

Pokud pracujete s nominálními (kategoriálními) daty nezapomínejte, že je můžete sčítat a odčítat (resp. četnosti). Pochopitelně pouze v případech, kdy je to logicky správně. Např. zjišťujeme volnočasové aktivity u skupiny respondentů, můžeme je poté rozdělit na pasivní, aktivní anebo si vybrat jiný klíč. Proč je to někdy užitečné, můžeme demonstrovat na konkrétním příkladu.

Bylo zjišťováno, jestli existují nějaké statisticky významné rozdíly v tom, kdy padne ve fotbale gól. Byly brány v úvahu údaje ČMFS (dnes FAČR) za 1. polovinu 1. fotbalové ligy v ČR. Pokud byl fotbalový zápas rozdělen do šesti 15minutových úseků a k jednotlivým úsekům přiřazeny četnosti gólů dle úseku jejich vstřelení, ukázalo se, že žádné statisticky významné rozdíly neexistují. Pokud se však pracovalo s četnostmi gólů v jednotlivých poločasech, tak zjištěny statisticky významné rozdíly byly (Janiš ml., 2011).

Pokud nám to tedy zjištěné odpovědi umožňují a je to účelné, lze jednotlivé četnosti sčítat a odhalit tím určité zákonitosti, které před tím zůstaly skryté. I výše uvedený příklad s výší příjmu by takový postup umožňoval, a to za předpokladu, že by nabízené intervaly byly 4. Mohli bychom tedy respondenty rozdělit pouze do dvou příjmových skupin (sečetli bychom četnosti 2 nejnižších a 2 nejvyšších intervalů). Mohli bychom využít chí-kvadrát pro čtyřpolní tabulku.

### 3.5.3 Číselný popis rozložení dat

Nejčastěji hovoříme o mírách centrální tendence a mírách rozptýlenosti.

Za základní **míry centrální tendence** pokládáme – *aritmetický průměr, medián a modus*. Jejich smyslem je určit, jak jsou na číselné ose rozložena data (Hendl, 2009)

ARITMETICKÝ PRŮMĚR je všeobecně známou mírou a s jeho výpočtem by neměl být problém. Počítáme jej ze vzorce, kde  $n$  je četnost všech hodnot a  $x_i$  je určitá hodnota (Hendl, 2009, s. 99):



$$\bar{x} = \frac{\sum x_i}{n}$$

[14]

Aritmetický průměr má výpovědní hodnotu v případě, že data jsou symetrická, nevyskytují se extrémní hodnoty (Chráška, 2007)

MEDIÁN je hodnota, která řadu hodnot seřazených podle velikosti rozděluje na dvě stejné poloviny. Označujeme je jako  $\tilde{x}$ . Medián je možno počítat, i když jsou v souboru dat extrémní hodnoty (Hendl, 2009).

*Pokud je medián určován ze souboru, kde celkový počet dat je lichý, např. z číselné řady 1, 2, 3, 5, 5, 9, 15, je medián roven číslu 5. Pokud je soubor tvořen sudým počtem hodnot, např. z číselné řady 1, 2, 3, 5, 5, 9, vidíme, že v podstatě střední hodnotu nemáme. Medián tedy dopočítáme jako aritmetický průměr hodnot  $n/2 + n/2 + 1$ . V uvedeném příkladu se sudým počtem hodnot to jsou hodnoty 3 a 5. Medián je tedy roven číslu 4.*

MODUS neboli modální hodnota je taková hodnota, která se ve skupině dat vyskytuje nejčastěji. Využít ji můžeme i u dat nominálních, kdy udává možnost s největší četností. V souboru dat mohou modální hodnoty nabývat více údajů. Modus značíme jako  $\hat{x}$  (Hendl, 2009; Chráška, 2007).

Vztah výše uvedených (aritmetický průměr, medián, modus) budeme opět demonstrovat na příkladu. Pokud někomu předložíme pouze aritmetický průměr, tak ten opravdu není příliš vypovídající a je nutné data dále „upřesnit“.

*Získali jsme hodnoty: 1, 2, 3, 3, 3, 9, 15. Z těchto dat je aritmetický průměr = 5,14; medián = 3 a modus = 3.*

Pokud by bylo přesné symetrické rozdělení dat, byly by jednotlivé hodnoty shodné. Z výše uvedených je patrné, že v souboru dat existuje nějaká extrémní hodnota, která nám „deformuje“ aritmetický průměr. To usuzujeme z doplňujících hodnot mediánu a hodnoty modální.

Typickým praktickým příkladem, kdy extrémní hodnoty deformují průměr, je průměrná mzda v ČR, která je deformována vysokými příjmy úzké skupiny lidí. Podle Českého statistického úřadu byla průměrná mzda ve 3. čtvrtletí 2012 24 514 Kč. Medián je však cca o 4 000 Kč nižší ([http://www.czso.cz/csu/tz.nsf/i/nejrychleji\\_rostou\\_vyssi\\_mzdy20120531](http://www.czso.cz/csu/tz.nsf/i/nejrychleji_rostou_vyssi_mzdy20120531); <http://www.czso.cz/csu/csu.nsf/informace/cpmz120412.doc>)

**Míry rozptýlenosti** nám dále upřesňují data, jejich rozptýlenost na ose. To, že mají dva soubory stejný aritmetický průměr, neznamená, že jsou oba soubory stejné. Níže zmíníme *variační rozpětí (šíře)*, *rozptyl*, *směrodatnou odchylku* a *variační koeficient*.

VARIAČNÍ ROZPĚTÍ je nejjednodušší mírou rozptýlenosti, a to z hlediska výpočtu, ale i výpovědní informace. Ze vzorce je patrné, že odečítáme nejnížší naměřenou hodnotu od nejvyšší naměřené hodnoty. Stejně jako aritmetický průměr je náchylný na extrémní hodnoty (Hendl, 2009, s. 102)<sup>51</sup>.

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

[15]

Variační rozpětí nám lépe charakterizuje aritmetický průměr, kdy ze dvou souborů dat můžeme získat dva stejné aritmetické průměry, ale přesto se data budou výrazně lišit a pouhé uvedení průměrů by mohlo být poněkud zavádějící. Níže to budeme demonstrovat na příkladu.

*Od respondentů jsem zjišťovali počet zájmových aktivit, kterým se věnují. Získali jsme následující data. V první skupině byly od jednotlivých respondentů zjištěny následující četnosti 2, 3, 3, 3, 5, 6, 7, 7. Aritmetický průměr první skupiny činí 4,5. Četnosti ve druhé skupině 1, 3, 4, 4, 5, 5, 5, 9. Aritmetický průměr je roven 4,5. Pokud však doplníme údaje z jednotlivých skupin do vzorce variační rozpětí, tak v první*

---

<sup>51</sup> Ve skriptech jsme zmínili, že je možné u škálových odpovědí využít i výpočty pro metrická data. Variační šíře je však v tomto případě téměř k ničemu.

skupině je  $R = 7 - 2 = 5$ ; ve druhé skupině  $R = 9 - 1 = 8$ . Data ve druhé skupině jsou tedy více „rozptýlena“ kolem aritmetického průměru.

„ROZPTYL je definován jako průměrná kvadratická odchylka měření od aritmetického průměru.“ (Hendl, 2009, s. 102). Zpravidla rozptylu jako hodnoty nevyužíváme v popisné statistice. Vycházíme z následujícího vzorce, kde  $n$  je celková četnost hodnot,  $x_i$  jednotlivé změřené hodnoty a  $\bar{x}$  je aritmetický průměr.

$$\sigma^2 = \frac{\sum (x_i - \bar{x})^2}{n}$$

(Chráška, 2007, s. 52) [16]

Tento vzorec se vztahuje k výpočtu rozptylu pro celý základní soubor, tedy známe data všech prvků v něm. Toho v BP dosahujeme velice výjimečně, ale u specificky zaměřených výzkumů to možné je.

Pokud neznáme data z celého základního souboru, označujeme rozptyl jako  $s^2$ . Vycházíme ze stejného vzorce, který je uveden výše, ale nedělíme celkovou četností  $n$ , ale  $n-1$  (Hendl, 2009, s. 102)

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

[17]

Rozptyl se zpravidla nevyužívá v popisné statistice, ale slouží k další analýze získaných dat (např. ANOVA – analýza rozptylu). Uvádíme jej zde především ve vztahu k směrodatné odchylce.

SMĚRODATNÁ ODCHYLKA nám udává míru rozptýlenosti kolem aritmetického průměru. Je náchylná k extrémním hodnotám (provázanost s aritmetickým průměrem). Je vhodné ji používat, pokud jsou data symetrická. Směrodatná odchylka se vypočítá jako odmocnina rozptylu. Rozeznáváme výběrovou směrodatnou odchylku ( $s$ ) a směrodatnou odchylku ( $\sigma$ ), a to ve vztahu k základnímu či výběrovému souboru (viz výše). V případě většího rozsahu získaných dat (nemusí

se jednat o tisíce) se rozdíly mezi uvedenými typy stávají zanedbatelnými (v setinách a tisícinách) (Chráška, 2007; Hendl, 2009).

U směrodatné odchylky je nutné její hodnotu vždy chápat pouze k danému příkladu. Nemůžeme obecně stanovit, že čím nižší číslo, tím lepší (0 je asi nejlepší), ale její velikost je relativní.

*Např. Měřili jsme časy žáků v běhu na 100 m a směrodatná odchylka byla rovna 10 vteřinám. Jelikož v tomto případě uvažujeme o vteřinách, je směrodatná odchylka poměrně vysoká. Ve druhém jsme měřili časy žáků v běhu na 5 km a směrodatná odchylka byla rovna opět 10 vteřinám. Jelikož jsme opět uvažovali ve vteřinách, je v tomto případě směrodatná odchylka v podstatě zanedbatelná a aritmetický průměr má větší výpovědní hodnotu.*

Možná Vás překvapilo, že v této části skript neukazujeme konkrétní výpočty a příklady. Vzorce jsou relativně jednoduché. Možné příklady pro tuto kapitolu si budeme demonstrovat v příloze č. 2, kde bude ukázán postup jejich výpočtu v Excelu. Ruční počítání v praxi je příliš zdlouhavé a s větším rizikem chyby při výpočtu.

VARIAČNÍ KOEFICIENT je vyjádřením vztahu mezi směrodatnou odchylkou (i výběrovou) a aritmetickým průměrem. Jeho výhodou je, že můžeme jednotlivé variační koeficienty porovnávat mezi sebou. Lze je uvádět i v procentech. Variační koeficient udává, kolik procent z průměrné hodnoty tvoří směrodatná odchylka. Vycházíme z následujícího vztahu, kde ve vzorci uvádíme  $s$ , tedy výběrovou směrodatnou odchylku (Chráška, 2007; Hendl, 2009, s. 103):

$$VK = \frac{s}{\bar{x}}$$

[18]

Pokud chceme uvádět procentuální vyjádření, násobíme výsledek 100. Obecně lze říci, že čím nižší výsledek, tím lepší vypovídající hodnota aritmetického průměru.

*Např. ze souboru dat jsme vypočetli aritmetický průměr 14,5 a výběrovou směrodatnou odchylku 3,2. Podle výše uvedeného vzorce dostáváme výsledek 0,22.*

*Tedy víme, že směrodatná odchylka tvoří 22 % aritmetického průměru. Porovnání VK můžeme demonstrovat na příkladu, který byl uveden u směrodatné odchylky u běhů. Aritmetický průměr časů v běhu na 100 m byl 13,5 sekund a výběrová směrodatná odchylka 10.  $VK = 0,740 * 100 \% = 74 \%$ . Ve druhém příkladu byl aritmetický průměr časů v běhu na 5 km 1200 sekund a výběrová směrodatná odchylka 10.  $VK = 0,008 * 100 \% = 0,8 \%$ . Při porovnání obou vypočtených variačních koeficientů vidíme, že data získaná z běhu 5 km jsou mnohem méně rozptýlena než data získaná v běhu na 100 m.*

Je nutné si uvědomit, že výše uvedené „výpočty“ slouží především k popisu, zpřesnění a vyšší vypovídající hodnotě získaných dat.

## **Závěr**

Skripta, která jste právě dočetli, jsou pouze elementárním vhledem do předkládané problematiky. Měla by Vám pomoci v orientaci v základních oblastech kvantitativního výzkumu. Jsou psána jednoduchým jazykem, který již příliš zjednodušit nelze. Nezoufejte, pokud jste některé kapitoly nepochopili na první pokus. Vždy je nutné se k jednotlivým informacím vracet.

Jednotlivé informace vycházely ze základních a nejrozšířenějších publikací k problematice výzkumu. Jednotlivé parafráze a citace Vás přirozeně směřují k původním zdrojům, kterými byste si měli vybavit svou studijní knihovnu.

## **Slovník doplňujících pojmů k výzkumu**

Níže uvedené pojmy jsou definovány co nejjednodušeji s ohledem na pochopení jejich podstaty. Vysvětlení některých pojmů slouží především k tomu, abyste se lépe orientovali v odborných člancích apod.

**Analýza** – rozbor. Jedná se o metodu výzkumu, kdy je složitější zkoumaný jev rozložen na jednodušší elementy.

**Analýza rozptylu (ANOVA)** – zjišťuje rozdíly v průměrech více než dvou skupin. V případě dvou skupin využíváme zpravidla t-test.

**Anketa** – jedná o průzkum, který nenaplnuje znaky výzkumu. Typický znak je malý počet respondentů. Nejsou naplněny požadavky validity a reliability. Výsledky mají omezený význam.

**Decil** – typ kvantilu, kdy je soubor dat rozdělen na desetiny.

**Dedukce** – typ úsudku, kdy se od obecného dochází ke konkrétnímu. Závěry vzniklé dedukcí nejsou odhadem (pravděpodobné), ale jsou jisté.

**Evaluace** – jedná se o hodnocení či posouzení jevu, procesu apod.

**Evidence base practice** – praxe založená na důkazech. Jedná se o výběr postupu (apod.) na základě již realizovaných výzkumů.

**Experiment** – je úmyslná manipulace s jednou či více nezávisle proměnnými.

**Ex post facto** – typ výzkumu, kdy jsou zjištěná data retrospektivně analyzována; hledají se nezávisle proměnné, tedy ty proměnné, které ovlivňují výsledky.

**Faktorová analýza** – metoda, kterou se snažíme zjistit proměnné, které ovlivňují realizovaná měření určitého objektu.

**Focus group** – metoda kvalitativního výzkumu. Jedná se o profesionálně řízené diskuze, jejichž cílem je zjistit postoje, názory apod. určité skupiny lidí. Vyžaduje vysokou odbornost tazatele, resp. toho, kdo diskuzi řídí.

**Indukce** – typ úsudku, kdy se na základě dílčích informací dochází k obecnému závěru.

**Korelace** – vzájemný vztah mezi dvěma zkoumanými veličinami, procesy, znaky.

**Kvantil** – číslo, které dělí řadu seřazených hodnot na stejné části. Příkladem kvantilu je např. medián.

**Metaanalýza** – smyslem metaanalýzy je shrnout výsledky dvou či více podobných výzkumů (na dané téma) a analyzovat velikost zjištěných účinků, které byly zjištěny v jednotlivých výzkumech (Hendl, 2009).

**Neparametrický test** – testuje hypotézu, která se nevztahuje k parametru základního souboru. Rozdělení může být nenormální. Jedná se např. o chí-kvadrát.

**Parametrický test** – váže se k hypotézám, které se týkají nějakého parametru základního souboru (aritmetický průměr apod.). Zpravidla vyžaduje normální rozdělení. Jedná se např. o t-test.

**Percentil** – jedná se o typ kvantilu, kdy je soubor dat rozdělen na setiny.

**Případová studie** – intenzivní a detailní studie jednoho či více případů.

**Q metodologie** – jedná se o výzkum, kdy je respondentům představeno určité množství Q-typů (zjednodušeně „kartiček“), na nichž jsou vyjmenovány jisté výroky, vlastnosti, hodnoty apod. Respondenti jim přisuzují míru významu od nejnižší po nejvyšší, přičemž každá úroveň smí obsahovat předem stanovený maximální počet Q-typů. Rozložení úrovní připomíná Gaussovu křivku, kdy na krajích je stanoven nejnižší počet možných Q-typů.

**Q třídění** – viz Q metodologie.

**Regresní analýza** – zjišťuje, jaký je vztah mezi příčinnou (nezávisle proměnná) a následkem (závisle proměnná).

**Reliabilita** – českým ekvivalentem je spolehlivost. Pokud bychom provedli měření opakovaně, získali bychom stejné výsledky.

**Reprezentativnost** – „zmenšenina“ základního souboru poměrově odpovídající jeho determinantům (pohlaví, věk, vzdělání apod.).



**Sekundární analýza** – opětovný rozbor již získaných dat, avšak pro jiný výzkumný záměr. U sesbíraných dat se snažíme dojít k novému závěru.

**Smíšený výzkum** – v sobě integruje kvantitativní a kvalitativní metody, techniky apod. v jednom výzkumu.

**Sociogram** – jedná se o zápis (výstup) vzájemných vztahů ve skupině.

**Sociometrický index** – vyjádření struktury skupiny.

**Sociometrie** – soubor metod, které zjišťují vzájemné vztahy ve skupině (sympatie, antipatie apod.).

**Triangulace** – využití rozdílných zdrojů informací pro náš výzkum. Triangulací se snažíme zvýšit „kvalitu“ našich výsledků. Snižujeme tedy jakési zkreslení výsledků, které by mohlo vyplynout na základě pouze jednoho zvoleného přístupu, teorie apod. Rozlišujeme triangulaci datovou, metodologickou, triangulaci výzkumníků, implicitní, explicitní a triangulaci teoretickou (Hendl, 2008).

**Validita** – českým ekvivalentem je platnost. Naše měření zjišťuje opravdu to, co zjišťovat chceme. Rozlišujeme validitu obsahovou, souběžnou, predikční a konstruktovou (Chrásky, 2007).

Jednotlivé pojmy vycházejí z publikací uvedených v seznamu použité literatury. Tímto se autor skript nezabývá odpovědností při uvádění citací a parafrází, ale některé pojmy jsou natolik rozšířené, že je prakticky nemožné vystopovat původce prvotní myšlenky či definice.

## Použitá literatura

*CODE OF CONDUCT FOR RESEARCH*. [online] [cit. 2012-8-8]. Dostupné z: <http://www.une.edu.au/policies/pdf/codeofconductresearch.pdf>

DISMAN, Miroslav *Jak se vyrábí sociologická znalost*. Praha: Karolinum, 2008. ISBN 978-80-246-0139.

*Etický kodex výzkumných pracovníků v Akademii věd České republiky*. [online] [cit. 2012-8-8]. Dostupné z: [http://www.avcr.cz/o\\_avcr/zakladni\\_informace/dokumenty/eticky\\_kodex.html](http://www.avcr.cz/o_avcr/zakladni_informace/dokumenty/eticky_kodex.html)

GAVORA, Peter, a kol. *Elektronická učebnica pedagogického výskumu*. [online] [cit. 2012-8-8]. Bratislava: Univerzita Komenského, 2010. Dostupné z: <http://www.e-metodologia.fedu.uniba.sk/> ISBN 978–80–223–2951–4.

HENDL, Jan. *Kvalitativní výzkum*. Praha: Portál, 2008. ISBN 978-80-7367-485-4.

HENDL, Jan. *Přehled statistických metod*. Praha: Portál, 2009. ISBN 978-80-7367-482-3.

CHRÁSKA, Miroslav. *Metody pedagogického výzkumu: základy kvantitativního výzkumu*. Praha: Grada, 2007. ISBN 978-80-247-1369-4.

JANIŠ, Kamil ml. Kdy začít sledovat fotbalový zápas kvůli gólům? *Tělesná výchova a sport mládeže*, 2011, roč. 77, č. 1, s. 39-41. ISSN 1210-7689.

MARTIN, Paul, BATESON, Patrick. *Úvod do teorie a metodologie měření chování*. Praha: Portál, 2009. ISBN 978-80-7367-526-4.

MIOVSKÝ, Michal. *Kvalitativní přístup a metody v psychologickém výzkumu*. Praha: Grada, 2006. ISBN 80-247-1362-4.

PRŮCHA, Jan *Přehled pedagogiky*. Praha: Portál, 2000. ISBN 80-7178-399-4.

PUNCH, Keith. *Úspěšný návrh výzkumu*. Praha: Portál, 2008. ISBN 978-80-7367-468-7. (a)

PUNCH, Keith. *Základy kvantitativního šetření*. Praha: Portál, 2008. ISBN 978-80-7367-381-9. (b)

SKUTIL, Martin, a kol. *Základy pedagogicko-psychologického výzkumu pro studenty učitelství*. Praha: Portál, 2011. ISBN 978-80-7367-778-7.

STÖRIG, H., J. *Malé dějiny filozofie*. Kostelní Vydří: Karmelitánské nakladatelství, 200. ISBN 80-7192-500-2.

ŠVAŘÍČEK, Roman, ŠEĐOVÁ, Klára, a kol. *Kvalitativní výzkum v pedagogických vědách*. Praha: Portál, 2007. ISBN 978-80-7367-313-0.

UNIVERSITY COLLEGE CORK. *CODE OF RESEARCH CONDUCT*. [online] [cit. 2012-8-8]. Dostupné z:

[http://www.ucc.ie/research/rio/documents/CodeofGoodConductinResearch\\_000.pdf](http://www.ucc.ie/research/rio/documents/CodeofGoodConductinResearch_000.pdf)

UNIVERSITY OF BIRMINGHAM. *CODE OF PRACTICE FOR RESEARCH*. [online] [cit. 2012-8-8]. Dostupné z:

<http://www.birmingham.ac.uk/Documents/university/legal/research.pdf>

VRÁNA, Stanislav. *Slovník moderní pedagogiky*. Brno: Komenium, 1948.

## Seznam příloh

|                                                                |   |
|----------------------------------------------------------------|---|
| Příloha č. 1 – Kritické hodnoty testového kritéria chí-kvadrát | 1 |
| Příloha č. 2 – Jednostranný a oboustranný test                 | 2 |

## Seznam použitých vzorců

|                                                                     |    |
|---------------------------------------------------------------------|----|
| [1] Test dobré shody chí-kvadrát.....                               | 31 |
| [2] Výpočet očekávané četnosti pro pole v kontingenční tabulce..... | 35 |
| [3] Výpočet stupňů volnosti v kontingenční tabulce.....             | 37 |
| [4] Chí-kvadrát pro čtyřpolní tabulku.....                          | 38 |
| [5] Dvouvýběrový t-test.....                                        | 41 |
| [6] Výběrová směrodatná odchylka pro t-test.....                    | 41 |
| [7] Počet stupňů volnosti pro t-test.....                           | 41 |
| [8] Párový t-test.....                                              | 42 |
| [9] Pomocný výpočet pro párový t-test .....                         | 42 |
| [10] Počet stupňů volnosti v párovém t-testu.....                   | 42 |
| [11] Pearsonův koeficient korelace .....                            | 47 |
| [12] Normovaná nominální variance .....                             | 61 |
| [13] Normovaný koeficient kontingence .....                         | 62 |
| [14] Aritmetický průměr.....                                        | 65 |
| [15] Variační rozpětí .....                                         | 66 |
| [16] Rozptyl základní soubor.....                                   | 67 |
| [17] Rozptyl výběrový soubor.....                                   | 67 |
| [18] Variační koeficient.....                                       | 68 |

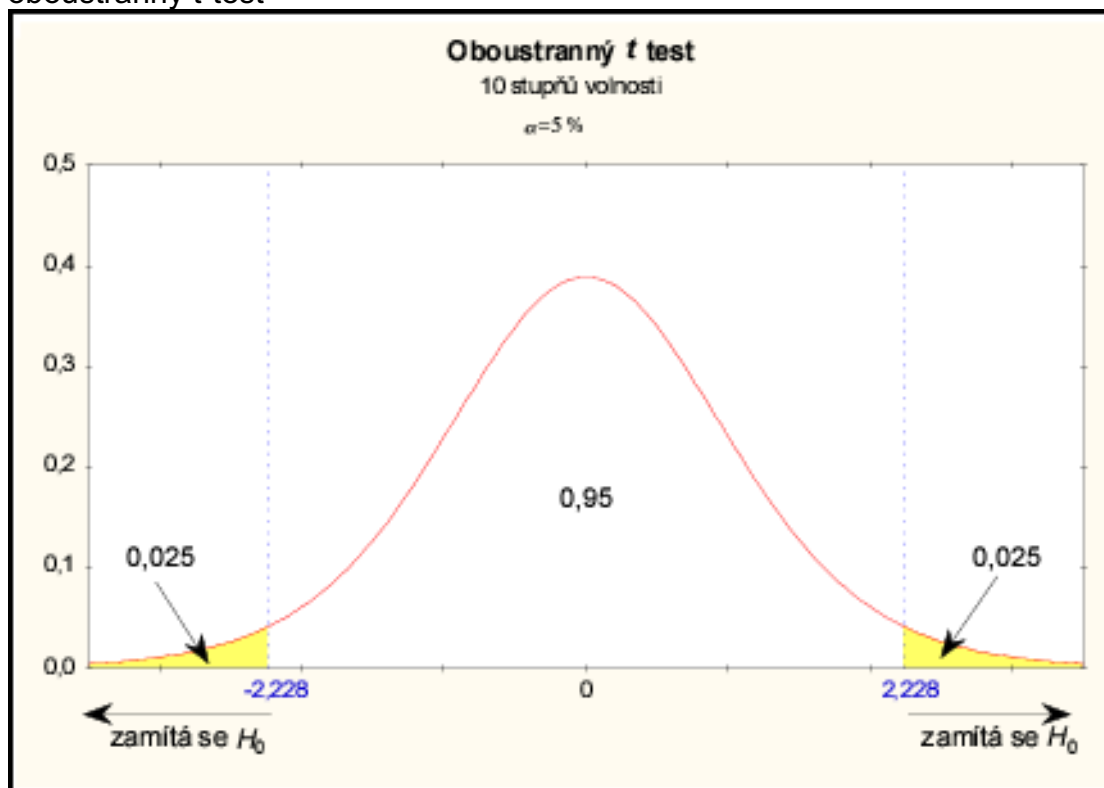
Příloha č. 1 – Kritické hodnoty testového kritéria chí-kvadrát

| Stupeň volnosti | Hladina významnosti |        |
|-----------------|---------------------|--------|
|                 | 0,05                | 0,01   |
| 1               | 3,841               | 6,635  |
| 2               | 5,991               | 9,210  |
| 3               | 7,815               | 11,341 |
| 4               | 9,488               | 13,277 |
| 5               | 11,070              | 15,086 |
| 6               | 12,592              | 16,812 |
| 7               | 14,067              | 18,475 |
| 8               | 15,507              | 20,090 |
| 9               | 16,919              | 21,666 |
| 10              | 18,307              | 23,209 |
| 11              | 19,675              | 24,725 |
| 12              | 21,026              | 26,217 |
| 13              | 22,362              | 27,688 |
| 14              | 23,685              | 29,141 |
| 15              | 24,996              | 30,578 |
| 16              | 26,269              | 32,000 |
| 17              | 27,587              | 33,409 |
| 18              | 28,868              | 34,805 |
| 19              | 30,144              | 36,191 |
| 20              | 31,410              | 37,576 |
| 21              | 32,671              | 38,932 |
| 22              | 33,924              | 40,289 |
| 23              | 35,172              | 41,638 |
| 24              | 36,415              | 42,980 |
| 25              | 37,652              | 44,314 |
| 26              | 38,885              | 45,642 |
| 27              | 40,113              | 46,963 |
| 28              | 41,337              | 48,278 |
| 29              | 42,557              | 49,588 |
| 30              | 43,773              | 50,892 |

(Chráška, 2007, s. 248)

## Příloha č. 2 – jednostranný a oboustranný t-test

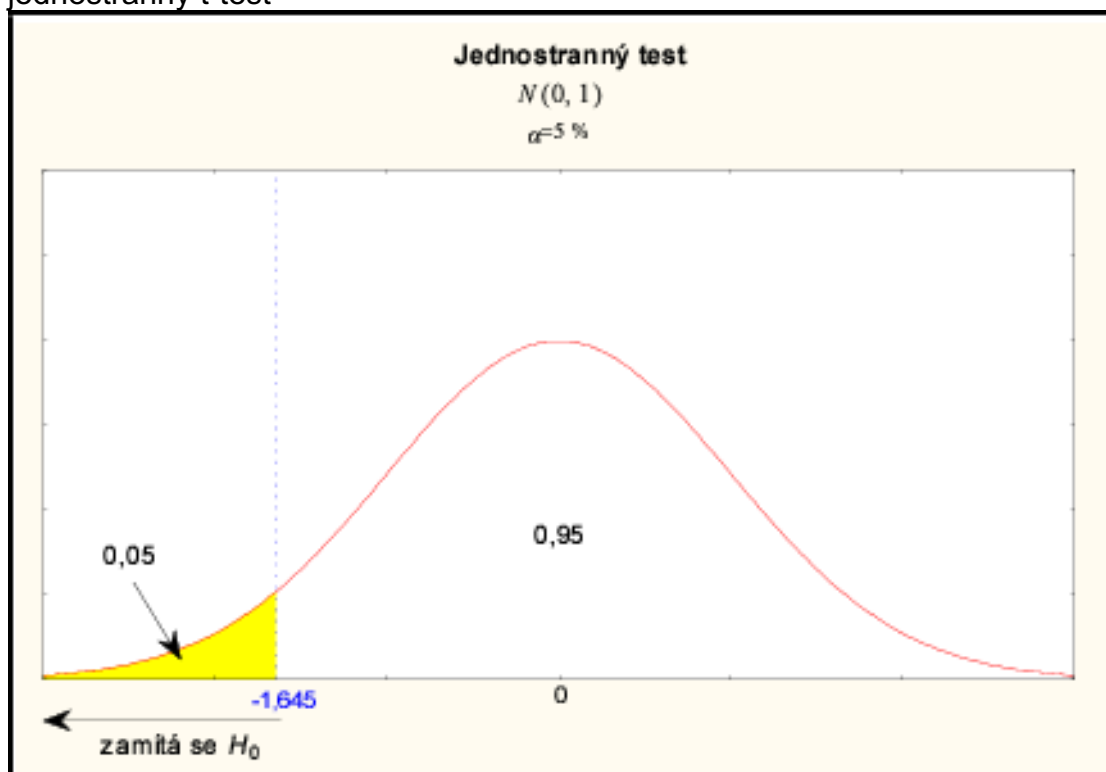
### oboustranný t-test



(<http://ucebnice.euromise.cz/index.php?conn=0&section=biostat1&node=10>)

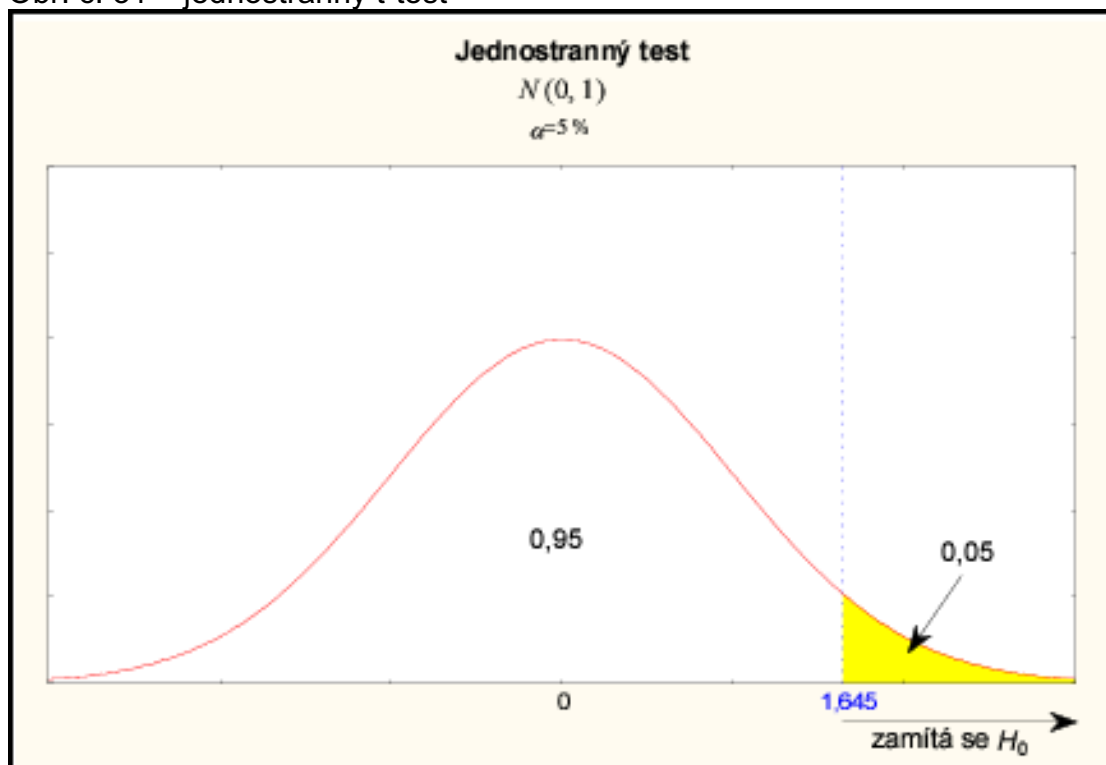
Pro jednostranný test by kritická hodnota pro 10 stupňů volnosti a hladinu významnosti 0,05 byla rovna 1,812. V případě oboustranného testu jasně vidíme, že naše zvolená hladina významnosti je rovnoměrně rozdělena na oba konce 0,05. V podstatě je tedy testové kritérium porovnáváno s kritickou hodnotou pro hladinu významnosti 0,025 (tato pasáž je velice zjednodušená).

jednostranný t-test



(<http://new.euromise.org/czech/tajne/ucebnice/html/html/node9.html>)

Obr. č. 31 – jednostranný t-test





(<http://new.euromise.org/czech/tajne/ucebnice/html/html/node9.html>)

