

Online příloha 2. Kódovací schéma pro obsahovou analýzu článků (s komentářem)

Následující popis je zaměřen na kódování při obsahové analýze článků. Pro část I v článku byl analytickou jednotkou článek, pro části II a III pak jednotlivé analýzy (jednotlivé statistické techniky v článcích použité).

Kromě toho byly zachyceny tyto informace o publikovaných článcích:

1. *Časopis* – 1 = Sociologický časopis, 2 = Pedagogika, 3 = Československá psychologie
2. *Rok vydání časopisu* – číselná proměnná, rozsah 1995–2014
3. *Číslo časopisu* – číselná proměnná, rozsah 1–6
4. *Jazyk článku* – 1 = čeština nebo slovenština, 2 = angličtina
5. *Autor/autoři článku* – textová proměnná, výpis všech autorů článku
6. *Institucionální příslušnost prvního autora* – nominální proměnná, 1 = VŠ ČR, 2 = Akademie věd, 3 = jiné pracoviště v ČR, 4 = mimo ČR
7. *Rok sběru dat* – číselná proměnná

I. V oblasti nesprávného používání statistických testů bylo sledováno, zda data nepocházejí:

- a) z censu (kód 1),
- b) ze záměrného výběru (kvótního, dostupného apod.) (kód 2),
- c) z malého výběru (kód 3) a
- d) z výběru s extrémně velkým počtem vybraných jednotek či daty spojenými z různých datových souborů (kód 4, resp. 5).

Bylo použito trojice proměnných, aby bylo možné zaznamenat více problematických rysů použitých dat, nicméně vždy postačila jen jediná. Na rozdíl od dále uvedených charakteristik jde o charakteristiku, která přináleží článku. Problematické bylo, že ne vždy je v článku charakter analyzovaných dat popsán. Na rozdíl od postupu Cubereka a Frömela [2011], kteří v takovém případě považovali data za pocházející ze záměrného výběru, byla informace o datech hledána na webových stránkách příslušných vědeckých projektů. Pokud článek analyzoval dva datové soubory (jde o okrajový jev), pak by byla připsána charakteristika příslušné analýze, nicméně nikdy nenastal případ, že by dva v článku použité soubory byly kódovány odlišně.

II. V oblasti nesprávného užívání statistické významnosti bylo trojicí proměnných sledováno, zda v jednotlivých analýzách v člancích nedochází k těmto problémům:

- a) mechanická práce s klasickou 5% hladinou statistické významnosti (hvězdičky, stepwise, nejlepší modely apod.) a jiné mechanické aplikace (například skrývání malých hodnot faktorových zátěží) (kód 1),
- b) záměna statistické a věcné významnosti (kód 2),
- c) slovní popis „významné, signifikantní“ pro statisticky významné výsledky (kód 5),
- d) ignorování výsledku testů statistické významnosti, resp. interpretace v rozporu s těmito výsledky (kód 6).

Původní kódovací schéma využívalo ještě kódu 3 (opomíjení statisticky nevýznamných výsledků v abstraktu či závěru) a kódu 4 (neužití statistické inference vůbec), nicméně pilotáž ukázala, že kód 4 je zcela okrajový a kód 3 je naopak velice rozšířen a bude smysluplné mu věnovat samostatnou studii (nadto nelze zjistit, kolik statisticky nevýznamných výsledků bylo autorem odstraněno již před publikací textu, a tak by šlo o zkreslené výsledky). Pro úplnost je nutno dodat, že mechanická aplikace statistické významnosti ve formě zvyrazňování výsledků (hvězdičky, tučné písmo, kurzíva apod.) byla hodnocena jako přítomná bez ohledu na to, zda se nacházela přímo v textu, či v doprovodných tabulkách. Pro zaznamenání výskytu postačoval jediný výskyt mechanické aplikace u příslušné statistické analýzy. Různé variace tohoto fenoménu jsou popsány v analytických výsledcích.

Pro obtížnost, respektive častou nemožnost rozlišování případů pod písmeny b) a c) zejména v případě Československé psychologie jsou výsledky analyzovány společně, tj. nerozlišuje se, zda autor zaměňuje věcnou a statistickou významnost, nebo „jen“ užívá výrazů „významný“, „signifikantní“, „důležitý“ pro výsledky, které jsou statisticky významné. U jednotlivých analýz (statistických technik) je vždy zaznamenáno (ve formě mnohonásobné proměnné), jaké formy nesprávného užívání statistické významnosti se vyskytují.

Kromě nesprávného užívání statistické významnosti bylo též u každé analýzy zaznamenáno, *jaká statistická technika je používána*. Využit byl následující seznam, resp. kódy:

- 1 – korelace (bez rozlišení toho, jaký korelační koeficient se počítá)
- 2 – regrese lineární
- 3 – regrese nelineární
- 4 – explorační faktorová analýza
- 5 – konfirmační faktorová analýza
- 6 – strukturní modelování
- 7 – mnohorozměrné škálování

- 8 – procenta
- 9 – průměry
- 10 – analýza rozptylu (bez rozlišení, zda je jedno či vícefaktorová, nebo vícerozměrná) včetně neparametrických období
- 11 – t-test (a případně neparametrické období)
- 12 – loglineární modely (bez logitových modelů, srov. kód 15)
- 13 – kontingenční tabulky
- 14 – diskriminační analýza
- 15 – logitové modely
- 16 – korespondenční analýza
- 17 – shluková (seskupovací) analýza
- 18 – analýza událostí
- 19 – víceúrovňové modely
- 20 – analýza latentních tříd
- 29 – jiná technika

III. V oblasti užívání měř věcné významnosti a věcné interpretace bylo (opětovně na úrovni jednotlivých analýz) sledováno následující:

1. Použité míry věcné významnosti (max. 3):

- 1 – korelační koeficient (bez rozlišení typu, respektive vzorce, s výjimkou ICC, viz kód 13)
- 2 – R^2 (vč. pseudo R^2 pro nelineární regrese)
- 3 – standardizované koeficienty
- 4 – Eta, resp. Eta^2
- 5 – Cohenovo d
- 6 – procenta
- 7 – pravděpodobnosti, šance
- 8 – Wiksovo lambda
- 9 – procento vysvětleného rozptylu
- 10 – kontingenční koeficienty (např. Cramerovo V)
- 11 – ukazatele vhodnosti modelů SEM (např. RMSEA, AGFI)
- 12 – míra úspěšnosti klasifikace
- 13 – vnitrotřídní korelační koeficient (ICC)
- 14 – jedinečnost
- 15 – adjustovaná residua

Dále bylo zaznamenáno ve formě dichotomické proměnné (kódované 0, 1), *zda je provedena interpretace minimálně jedné použité míry věcné významnosti* (1 = ano).

Poslední zaznamenanou charakteristikou bylo, *zda je provedena věcná interpretace analýzy*. Bylo opět využito jednoduché dichotomické proměnné (kódované 0, 1). Kódu 1 (přítomnost věcné interpretace) bylo využito, pokud autor článku u příslušné analýzy kromě zhodnocení statistické významnosti interpretoval (byť částečně) věcné výsledky, například tedy uvedl, že posuzovaný vliv je pozitivní či negativní, velký nebo malý. Oproti studii Bernarda, Chakhaiy, Leopold [2017] jde tedy o minimalistický přístup, proto nelze výsledky obou studií přímo srovnávat.

Literatura citovaná v přílohách

- Bernard, F., L. Chakhaiy, L. Leopold. 2017. „Sing Me a Song with Social Significance: The (Mis)use of Statistical Significance Testing in European Sociological Research.“ *European Sociological Review* 33 (1): 1–15, <https://doi.org/10.1093/esr/jcx044>.
- Cuberek, R., K. Frömel. 2011. „K problematice výzkumného výběru a testování nulové hypotézy.“ *Československá psychologie* 55 (5): 468–477.
- Matějů, P. 1989. „Metoda strukturního modelování. Přehled základních problémů.“ *Sociologický časopis* 25 (4): 399–418.
- Souček M. 2012. „Analýza vědy a výzkumu na základě dat z databáze RIV.“ *Ikaros* [online] 16 (12). Dostupné z: <http://ikaros.cz/node/14018>.
- Urbánek, T. 2007. „K prezentaci výsledků statistických analýz – 1. část.“ *Československá psychologie* 51 (6): 601–609.
- Urbánek, T. 2008. „K prezentaci výsledků statistických analýz – 2. část.“ *Československá psychologie* 52 (1): 70–79.
- Ziliak, S. T., D. M. McCloskey. 2008. *The Cult of Statistical Significance (How the Standard Error Costs Us Jobs, Justice, and Lives)*. The University of Michigan Press.