

Při sociologických výzkumech prováděných pomocí dotazníku respondenti (popřípadě tazatelé) označují jednu či více zvolených odpovědí. Jiným možným způsobem označení odpovědí, který se dosud málo používá, je stanovení jejich pořadí podle závažnosti, oblíbenosti apod. Přitom chceme vědět, zda mezi uspořádáním jednotlivých odpovědí respondentů je závislost tzn., zda se respondenti v uspořádání až na náhodně kolísání shodují. Jsou-li respondenti (nebo skupiny respondentů) nejméně tři, můžeme shodu mezi uspořádáním odpovědí posoudit na základě Kendallova koeficientu shody W.

*

Obecně je možno výše uvedenou úlohu formulovat takto: známe uspořádání n objektů podle k různých vlastností (kde $k \geq 3$) nebo k různými soudci. Za objekty můžeme považovat odpovědi na určitou otázku. Objekty mohou být i lidé (např. v pedagogickém výzkumu se může jednat o uspořádání žáků podle určitých vlastností, jindy nás může zajímat uspořádání umělců, sportovců apod. podle jejich oblíbenosti atd.). Úlohu soudců, kteří objekty uspořádají, převezmou respondenti. Chceme vědět, zda je shoda mezi uspořádáním objektů 1, 2, 3, ..., n , které provedlo k respondentů. Jako nulovou hypotézu zvolíme tvrzení, že se respondenti v uspořádání neshodují, tj. že mezi k uspořádáním n objektů není závislost.

Odpovědi, které respondenti uspořádali, označíme obecně A, B, C, ..., Z. Pořadí, které respondenti stanovili pro jednotlivé odpovědi, určíme takto:

$$\begin{array}{ccccccc} u_1, & u_2, & u_3, & \dots, & u_n, \\ v_1, & v_2, & v_3, & \dots, & v_n, \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ z_1, & z_2, & z_3, & \dots, & z_n. \end{array}$$

Na základě uspořádání odpovědí sestavíme tabulku 1.

Sečteme-li pořadí jednotlivých odpovědí, dostaneme součty pořadí, které jsme v tabulce 1 uvedli jako $R_1, R_2, R_3, \dots, R_n$.

Tabulka 1

Respondent	Odpovědi				
	A	B	C	...	Z
1	u_1	u_2	u_3	...	u_n
2	v_1	v_2	v_3	...	v_n
.	.	.	.	...	.
.	.	.	.	...	.
.	.	.	.	...	.
k	z_1	z_2	z_3	...	z_n
Součet pořadí	R_1	R_2	R_3	...	R_n

Za předpokladu, že by odpovědi seřadili dva respondenti, mohli bychom pro test nulové hypotézy použít Spearmanův nebo Kendallův koeficient pořadové korelace. V úvodu jsme si však stanovili podmínku, že respondenti budou alespoň tři. Za této podmínky můžeme úlohu řešit dvojím způsobem:

1. Vypočítáme Spearmanův nebo Kendallův koeficient pořadové korelace mezi dvojicemi respondentů a potom stanovíme jejich průměr a výsledek testujeme. Na pořadí dvojice nezáleží, jejich celkový počet je tedy $\binom{k}{2}$. Celý postup by byl značně zdoluhavý, zvláště při větším počtu respondentů. (Např. pro 5 respondentů bychom museli stanovit průměr $z\left(\frac{5}{2}\right) = 10$ koeficientů pořadové korelace stanovených pro jednotlivé dvojice, pro 20 respondentů bychom museli vypočítat průměr $z\left(\frac{20}{2}\right) = 190$ koeficientů.)

2. Vypočítáme Kendallův koeficient shody W. Jeho výpočet je mnohem jednodušší než postup uvedený v bodě 1.

Označíme-li průměrnou hodnotu Spearmanova koeficientu pořadové korelace, vypočteného z jednotlivých dvojic respondentů podle postupu uvedeného v bodě 1,

jako $\bar{e}$, budou mezi $\bar{e}$ a W platit tyto vztahy:

$$\bar{e} = \frac{kW - 1}{k - 1}, \quad (1)$$

$$W = \frac{\bar{e}(k - 1) + 1}{k}. \quad (2)$$

Dále se budeme zabývat pouze výpočtem a testováním Kendallova koeficientu shody W .

Pro součet pořadí $\sum_{j=1}^n R_j$ je lhostejné, jaké pořadí odpovědi jednotliví respondenti zvolí. Z hlediska součtu pořadí můžeme proto odpovědi jednotlivých respondentů seřadit do k aritmetických posloupností a sestavit tuto rovnici:

$$\sum_{j=1}^n R_j = k \frac{n(n + 1)}{2}. \quad (3)$$

Ze součtu pořadí vypočítáme průměrné pořadí:

$$\bar{R} = \frac{1}{n} \sum_{j=1}^n R_j = \frac{k(n + 1)}{2}. \quad (4)$$

Potom vypočítáme součet čtverců rozdílů jednotlivých součtů pořadí a průměrného pořadí a označíme ho S .

$$S = \sum_{j=1}^n (R_j - \bar{R})^2 = \sum_{j=1}^n R_j^2 - 2 \sum_{j=1}^n R_j \bar{R} + \sum_{j=1}^n \bar{R}^2 = \sum_{j=1}^n R_j^2 - n\bar{R}^2. \quad (5)$$

V tabulce 2 je uvedeno schéma uspořádání odpovědí při úplné shodě jejich pořadí u všech respondentů.

Tabulka 2

Respondenti	Odpovědi				
	A	B	C	...	Z
1	1	2	3		n
2	1	2	3		n
.	.	.	.		.
.	.	.	.		.
.	.	.	.		.
k	1	2	3		n
Součet pořadí	k	2k	3k		nk

Z tabulky 2 je zřejmé, že při úplné shodě uspořádání odpovědí jsou součty pořadí:

$$R_1 = k, R_2 = 2k, R_3 = 3k, \dots, R_n = nk. \quad (6)$$

Pro $j = 1, 2, 3, \dots, n$ platí, že

$$R_j = jk. \quad (7)$$

Platí-li (7), tj. je-li úplná shoda uspořádání, nabývá S maximální hodnoty. Označíme ho $S_{\max}$ a vypočítáme z rovnice:

$$S_{\max} = \frac{1}{12} k^2 n (n^2 - 1). \quad (8)$$

Kendallův koeficient shody W je podílem S vypočteného ze zvoleného uspořádání a $S_{\max}$:

$$W = \frac{S}{S_{\max}} = \frac{\sum_{j=1}^n (R_j - \bar{R})^2}{\frac{1}{12} k^2 n (n^2 - 1)} = \frac{12 S}{k^2 n (n^2 - 1)}. \quad (9)$$

W může nabývat hodnot od 0 do 1. W bude rovno 0, bude-li platit pro $j = 1, 2, 3, \dots, n$, že $R_j = \bar{R}$. Při úplné shodě uspořádání bude v čitateli i ve jmenovateli zlomku ve vztahu (9) $S_{\max}$ a W bude rovno 1.

Na rozdíl od koeficientu korelace nenabývá W hodnot od -1 do 0, což odpovídá tomu, že není možné interpretovat záporné hodnoty koeficientu.

Výpočet W ukážeme na příkladu.

Příklad 1.

Při sociologické sondě uvedlo podle určitého hlediska 15 respondentů pořadí u čtyř odpovědí, které označíme A, B, C, D. Uspořádání odpovědí jednotlivými respondenty je uvedeno v tabulce 3 na str. 173.

Součty pořadí jsou uvedeny v tabulce 3. Vypočítáme průměrné pořadí:

$$\bar{R} = \frac{15 \cdot 5}{2} = 37,5.$$

Po dosazení do vzorce obdržíme:

$$S = 25^2 + 28^2 + 41^2 + 56^2 - 4(37,5)^2 = 601.$$

Hodnota Kendallova koeficientu shody W je:

$$W = \frac{12 \cdot 601}{15^2 \cdot 4 \cdot (16 - 1)} = 0,534.$$

Tabulka 3

Respondent	Odpověď			
	A	B	C	D
1	2	1	3	4
2	1	2	3	4
3	1	3	2	4
4	2	1	4	3
5	1	2	3	4
6	1	2	4	3
7	4	1	2	3
8	1	2	3	4
9	1	4	2	3
10	3	1	2	4
11	1	2	3	4
12	3	1	2	4
13	1	3	2	4
14	1	2	3	4
15	2	1	3	4
Součet pořadí	25	28	41	56

Další postup závisí na velikosti k a n :

1. Pro $k = 3, 4, 5, 6, 8, 10, 15$ a 20 a pro $3 \leq n \leq 7$ najdeme kritické hodnoty S v příložené tabulce I. Pro test tedy stačí vypočítat hodnotu S . Pro $k = 7, 9, 11, 12, 13, 14$ můžeme kritické hodnoty zjistit lineární interpolací. Pro $n = 3$ a $k = 9, 12, 14, 16$ a 18 jsou kritické hodnoty S uvedeny přímo v tabulce I.

$S_{\max}$ je pro dané hodnoty k a n konstantní, takže v tabulce stačilo uvést S .

Kritické hodnoty S jsou v tabulce I uvedeny pro hladiny významnosti $\alpha = 0,05$ a $\alpha = 0,01$, označíme je $S_{0,05}$ a $S_{0,01}$. Nulovou hypotézu zamítáme na 5 %, popřípadě 1 % hladině významnosti, je-li $S \geq S_{0,05}$ respektive $S \geq S_{0,01}$.

Test ukážeme na příkladu 1.

Pokračování příkladu 1.

Chceme vědět, zda je mezi uspořádáním čtyř odpovědí shoda. Pro test zvolíme 1 % hladinu významnosti. Platí, že $S = 601 > S_{0,01} = 269,8$ a proto zamítáme nulovou hypotézu a činíme závěr, že se respondenti v uspořádání odpovědí signifikantně shodují.

2. Pro větší hodnoty k a n , než jaké jsou uvedeny v příložené tabulce I, je možno použít *aproximace χ^2 — distribucí*. Podmínkou však je, že $n \geq 8$.

Postupujeme tak, že nejdříve vypočítáme testové kritérium χ_z^2 podle vztahu:

$$\chi_z^2 = k(n-1)W = \frac{12S}{kn(n+1)} \quad (10)$$

Potom zjistíme počet stupňů volnosti z rovnice:

$$v = k - 1 \quad (11)$$

Testové kritérium χ_z^2 srovnáme s kritickou hodnotou $\chi^2_{\alpha}(v)$, kterou nalezneme v běžně dostupné literatuře. Nulovou hypotézu zamítáme na zvolené hladině významnosti α , je-li $\chi_z^2 \geq \chi^2_{\alpha}(v)$. Celý postup ukážeme na příkladu.

Příklad 2.

Při sociologickém výzkumu stanovilo 100 respondentů pořadí osmi umělců podle stupně obliby. Podle vztahu (5) jsme vypočetali $S = 81\,600$. Zajímá nás, zda se respondenti v uspořádání umělců shodují. Pro test nulové hypotézy jsme zvolili 1 % hladinu významnosti.

Vypočítáme χ_z^2 .

$$\chi_z^2 = \frac{12 \cdot 81\,600}{800 \cdot 9} = 136,0.$$

Počet stupňů volnosti $v = 99$.

Platí, že $\chi_z^2 = 136,0 > \chi^2_{0,01}(99) = 134,6$. Nulovou hypotézu zamítáme a konstatujeme, že mezi uspořádáním umělců je statisticky významná shoda.

Můžeme-li pro test použít aproximaci χ^2 — distribucí a vyskytnou-li se v uspořádání odpovědi se stejným pořadím, vypočítáme pro ně průměrné pořadí. Předpokládejme, že se průměrná pořadí vyskytla u respondenta $i = 1, 2, 3, \dots, r$, kde r je celkový počet respondentů, kteří stanovili u některých odpovědí stejná pořadí. Počet skupin odpovědí se stejným pořadím u respondenta i označíme t_i . (Počítáme počet skupin odpovědí se stejným pořadím bez ohledu na počet odpovědí ve skupině.)

Pro každého z r respondentů vypočítáme:

$$T_i = \frac{1}{12} \sum (t_i^3 - t_i), \quad (12)$$

kde $i = 1, 2, 3, \dots, r$.

Součet všech T_1 za všechny respondenty, kteří uvedli odpovědi se stejným pořadím je:

$$\sum_{i=1}^r T_i. \quad (13)$$

Testové kritérium vypočítáme z rovnice:

$$\chi_r^2 = \frac{S}{\frac{1}{12} k n (n+1) - \frac{1}{n-1} \sum_{i=1}^r T_i}. \quad (14)$$

Další postup je stejný jako u testového kritéria χ_r^2 .

Podmínka pro použití aproximace χ^2 — distribucí, aby $n \geq 8$, je nevýhodná. Respondenti dokáží seřadit spíše menší počet odpovědí než větší.

3. Jsou-li hodnoty k větší, než jaké jsou uvedeny v příložené tabulce I a je-li $n \leq 7$, použijeme aproximaci pomocí Fisherovy z-distribuce¹⁾. Výpočet je sice složitější než v předchozích dvou případech, ale umožňuje použít Kendallův koeficient shody W i pro malá n a větší počet respondentů, což jsou případy, které se často vyskytnou v praxi sociologických výzkumů.

Nejprve vypočítáme veličinu z podle rovnice:

$$z = \frac{1}{2} \ln \frac{(k-1)W}{1-W}. \quad (15)$$

Na pravé straně rovnice (15) je přirozený logaritmus. V logaritmičeských tabulkách však bývá přirozený logaritmus uveden jen pro některé hodnoty. Je proto výhodné převést přirozený logaritmus na dekadický. Můžeme tak učinit podle rovnice:

$$\ln a = \log a \cdot \ln 10, \quad (16)$$

kde $\ln 10$ je modul přirozených logaritmů, který označíme M .

$$M = 2,302\,585\,029 \dots \quad (17)$$

Pro naše výpočty stačí používat M s přesností na 6 desetinných míst.

Pro dosazení (16) a (17) do (15) dostáváme:

$$z = \frac{1}{2} \left[\left(\log \frac{(k-1)W}{1-W} \right) \cdot 2,302\,585 \right]. \quad (18)$$

¹⁾ Použití aproximace pomocí Fisherovy z-distribuce je univerzální. Je možné ji použít i pro hodnoty k a n uvedené v tab. I nebo v těch případech, kdy lze použít aproximace χ^2 - distribucí. Vzhledem k tomu, že je výpočet pomocí Fisherovy z-distribuce složitější, doporučujeme použít tam, kde je to možné, tabulku I nebo aproximaci χ^2 - distribucí.

²⁾ Kendallův koeficient shody W vypočítáme podle vztahu (9). Pro malé hodnoty k (např. pro $k = 3$)

Při testování srovnáváme:

- vypočtenou hodnotu z (podle rovnice (18)),
- hodnotu Fisherovy z -distribuce na zvolené hladině významnosti. Tuto hodnotu označíme obecně z_α (v_1, v_2) a vyhledáme ji v příložené tabulce II.

Fisherova z -distribuce je určena stupni volnosti v_1, v_2 . Musíme proto vypočítat:

$$v_1 = n - 1 - \frac{2}{k}, \quad (19)$$

$$v_2 = (k - 1) v_1. \quad (20)$$

Stupně volnosti v_1 a v_2 vypočtené z rovnic (19) a (20) vyjdou jako desetinná čísla. To je nevýhodné, protože Fisherova z -distribuce je tabelována pouze pro celá čísla (viz příloženou tabulku II). S k rostoucím nade všechny meze konverguje v_1 k hodnotě $n - 1$ a v_2 k ∞ . Vypočtený stupeň volnosti v_1 se tedy blíží tím více k hodnotám uvedeným v tabulce II, čím je počet respondentů vyšší.

Např. pro $n = 4$ a pro:

$k = 100$, je $v_1 = 2,98$ a $v_2 = 295,02$,

$k = 200$, je $v_1 = 2,99$ a $v_2 = 595,01$,

$k = 1000$, je $v_1 = 2,998$ a $v_2 = 2995,002$.

Vypočtená hodnota v_2 je již pro $k = 100$ vyšší, než je nejvyšší hodnota uvedená v tabulce II (tj. $v = 120$). Pro hladiny významnosti $\alpha = 0,05$ a $\alpha = 0,01$ a pro

$$v_1 = 1, 2, 3, 4, 5, 6,$$

$$v_2 = 200, 400, 1000, \infty,$$

nalezneme přibližné hodnoty Fisherovy z -distribuce v příložené tabulce III. Je-li to pro výsledek testu důležité, můžeme provést pro hodnoty v_1 a v_2 , nacházející se mezi hodnotami uvedenými v tabulce III, lineární interpolaci. (Postup při lineární interpolaci je uveden dále.)

Výpočet Fisherovy z -distribuce pro přesné hodnoty v_1 a v_2 nebude vždy nutný. Nulovou hypotézu zamítáme na hladině významnosti α , je-li $z \geq z_\alpha(v_1, v_2)$. Mohou nastat dva případy, kdy nebude zapotřebí počítat hodnotu z -distribuce pro přesné hodnoty v_1 a v_2 :

musíme použít pro výpočet W opravu pro spojitost podle vzorce:

$$W = \frac{S - 1}{S_{\max} + 2} = \frac{S - 1}{\frac{1}{12} k^2 n (n^2 - 1) + 2}.$$

Vzhledem k tomu, že pro malá k bude možno většinou použít tabulku I, popřípadě aproximace χ^2 - distribucí, nepřichází oprava pro spojitost v praxi častěji v úvahu.

A. Bude-li vypočtená hodnota $z \geq z_{\alpha}(v_1, v_2)$ pro stupně volnosti, které jsou v tabulce II nejbližší než přesně vypočtené hodnoty v_1 a v_2 , můžeme nulovou hypotézu zamítnout, protože vypočtená hodnota z je vyšší než $z_{\alpha}(v_1, v_2)$ pro přesně stanovené stupně volnosti. (Tento postup vlastně znamená přísnější měřítko.)

B. Naopak, jestliže bude vypočtená hodnota $z \leq z_{\alpha}(v_1, v_2)$ pro stupně volnosti, které jsou v tabulce II nejbližší vyšší než přesně vypočtené hodnoty v_1 a v_2 , nelze nulovou hypotézu zamítnout, protože vypočtená hodnota z je nižší než $z_{\alpha}(v_1, v_2)$ pro přesné hodnoty stupňů volnosti.

Oba případy ukážeme na příkladech.

Příklad 3.

Při výzkumu stanovilo 100 respondentů pořadí u čtyř odpovědí. Podle vztahu (5) jsme vypočítali hodnotu $S = 3500$. Chceme zjistit, zda se respondenti v uspořádání odpovědí shodují. Pro test zvolíme hladinu významnosti 0,1 %.

Vypočítáme W :

$$W = \frac{12 S}{k^2 n (n^2 - 1)} = \frac{42\,000}{600\,000} = 0,07.$$

Dále vypočítáme z :

$$z = \frac{1}{2} \left[\left(\log \frac{6,93}{0,93} \right) \cdot 2,302\,585 \right] = \\ = \frac{1}{2} (0,87\,225 \cdot 2,302\,585) = 1,004\,2.$$

Stupně volnosti jsme pro náš případ uvedli již vpředu, $v_1 = 2,98$ a $v_2 = 295,02$. Z tabulky II zjistíme hodnotu Fisherovy z -distribuce pro stupně volnosti, které jsou tabelovány pro hodnoty nejbližší nižší, než jsou přesně vypočtené, tj. pro $z_{0,001}(2; 120) = 0,9952$.

Platí, že $z = 1,0042 > z_{0,001}(2; 120) = 0,9952$ a proto zamítáme nulovou hypotézu na hladině významnosti $\alpha = 0,1$ %. Můžeme učinit závěr, že mezi uspořádáním odpovědí je statisticky vysoce významná shoda.

Nyní předpokládejme, že $S = 2500$. Ostatní podmínky (tj. počet respondentů, počet otázek, zvolená hladina významnosti) zůstávají stejné. Potom je $W = 0,05$ a $z = 0,8253$. Z tabulky II zjistíme hodnotu Fisherovy z -distribuce pro stupně volnosti, které jsou tabelovány pro hodnoty nejbližší

vyšší, než jsou přesně vypočtené, tj. pro $z_{0,001}(3; \infty) = 0,8453$. Srovnáním zjistíme, že $z = 0,8253 < z_{0,001}(3; \infty) = 0,8453$ a proto nemůžeme nulovou hypotézu zamítnout na 0,1 % hladině významnosti.

Pro menší k a n je $v_1 < 24$ a $v_2 < 120$. Např. pro $k = 20$ a pro $n = 3$ je $v_1 = 1,9$ a $v_2 = 36,1$. Je-li pro test zapotřebí zjistit přesnou hodnotu $z_{\alpha}(v_1, v_2)$, tj. platí-li např. pro $k = 20$ a $n = 3$, že

$$z_{\alpha}(2; 40) < z < z_{\alpha}(1; 30)$$

nezbývá, než provést lineární interpolaci.

Postup při lineární interpolaci ukážeme na příkladech.

Příklad 4.

Chceme zjistit hodnotu $z_{0,05}(1,9; 36,1)$. V tabulce II nalezneme tyto hodnoty:

Tabulka 4

v_2	v_1	
	1	2
30	0,7141	0,5994
40	0,7037	0,5866

Další postup rozdělíme na 3 kroky:

1. krok: vypočítáme $z_{0,05}(1; 36,1)$,
2. krok: vypočítáme $z_{0,05}(2; 36,1)$,
3. krok: vypočítáme $z_{0,05}(1,9; 36,1)$.

Výpočet neznámé hodnoty na základě lineární interpolace uskutečnime nejdříve obecně. Hodnoty v_2 , které jsou uvedeny v tabulce pro dané v_1 , označíme jako d_1 a d_2 , hodnoty Fisherovy z -distribuce pro d_1 a d_2 označíme jako a_1 a a_2 . (Viz tabulku 5.) Hledanou hodnotu v_2 , která leží mezi hodnotami d_1 a d_2 , označíme d , hodnotu z -distribuce pro d nazveme a . Pro $z_{0,05}(1; 36,1)$ můžeme sestavit tabulku 5.

Tabulka 5

v_2	$v_1 = 1$
$d_1 = 30$	$a_1 = 0,7141$
$d = 36,1$	a
$d_2 = 40$	$a_2 = 0,7037$

V tabulce 6 jsou uvedeny názvy potřebných rozdílů.

Tabulka 6

Rozdíl	Název
$d_2 - d_1$	tabulkový krok
$a_2 - a_1$	tabulková difference
$d - d_1$	naše difference
$a - a_1$	hledaná difference

Pro rozdíly uvedené v předchozí tabulce platí tato úměra:

$$a_2 - a_1 : d_2 - d_1 = a - a_1 : d - d_1. \quad (21)$$

Z úměry (21) vypočítáme a :

$$a = \frac{(a_2 - a_1) \cdot (d - d_1)}{d_2 - d_1} + a_1. \quad (22)$$

Výpočet hledané hodnoty a je pro větší přehlednost uveden v tabulce 7. (Pro druhý a třetí krok musíme do tabulky 5 dosadit příslušné hodnoty v_1 .)

Tabulka 7

Název	Vzorec	1. krok	2. krok	3. krok
		$v_1 = 1; v_2 = 36,1$	$v_1 = 2; v_2 = 36,1$	$v_1 = 1,9; v_2 = 36,1$
Tabulkový krok	$d_2 - d_1$	$40 - 30 = 10$	10	1
Tabulková difference	$a_2 - a_1$	$0,7037 - 0,7141 = -0,0104$	-0,0128	-0,1162
Naše difference	$d - d_1$	$36,1 - 30,0 = 6,1$	6,1	0,9
Tabulková hodnota pro d_1	a_1	0,7141	0,5994	0,7078
Hledaná hodnota	$a = z_{0,05}(v_1; v_2)$	$[(-0,0104 \cdot 6,1) : 10] + 0,7141 = 0,7078$	0,5916	0,603

Příklad 5.

Máme vypočítat pomocí lineární interpolace hodnotu $z_{0,01}(1,9; 36,1)$. V tabulce II nalezneme tyto výchozí hodnoty pro 1 % hladinu významnosti:

Tabulka 8

v_2	v_1	
	1	2
30	0,0116	0,8423
40	0,9949	0,8223

Další výpočet provedeme podobně jako u předchozího příkladu. Je uveden v tabulce 9 na str. 177.

Nyní si ověříme, že výpočet pomocí Fisherovy z-distribuce vede ke stejnému výsledku jako použití přiložené tabulky I. V dalších příkladech uvedeme test pro významné hodnoty S uvedené v tabulce I.

Příklad 6.

Při sociologické sondě uspořádalo 20 respondentů tři odpovědi do pořadí podle jejich závažnosti. Chceme vědět, zda můžeme po-

važovat rozdíly mezi uspořádáním za nahodilé, tj. zda se respondenti v uspořádání odpovědí shodují.

Podle vztahu (5) vypočítáme $S = 119,7$. Pro test zvolíme hladinu významnosti $\alpha = 5 \%$.

Tabulka 9

Název	Vzorec	1. krok	2. krok	3. krok
		$v_1 = 1; v_2 = 36,1$	$v_1 = 2; v_2 = 36,1$	$v_1 = 1,9; v_2 = 36,1$
Tabulkový krok	$d_2 - d_1$	10	10	1
Tab. difference	$a_2 - a_1$	-0,0167	-0,02	-0,17087
Naše difference	$d - d_1$	6,1	6,1	0,9
Tab. hodn. pro d_1	a_1	1,0116	0,8423	1,00097
Hledaná hodnota	$a = z_{0,01}(v_1; v_2)$	1,00097	0,8301	0,847

Dále vypočítáme:

Kendallův koeficient shody $W = 119,7 : 800 = 0,149\ 625$.

$$\text{hodnotu } z \doteq \frac{1}{2} \left[\left(\log \frac{(k-1)W}{1-W} \right) \cdot 2,302\ 585 \right] = \frac{1}{2} \left[(\log 3,343) \cdot 2,302\ 585 \right] = 0,603.$$

$$\text{stupně volnosti: } v_1 = 2 - \frac{1}{10} = 1,9, \\ v_2 = 19 \cdot 1,9 = 36,1.$$

Hodnotu Fisherovy z-distribuce pro hladinu významnosti $\alpha = 0,05$ a pro stupně volnosti $v_1 = 1,9$ a $v_2 = 36,1$ jsme vypočetli v příkladu 4. Zjistili jsme, že $z_{0,05}(1,9; 36,1) \doteq 0,603$.

Platí, že $z \doteq 0,603 = z_{0,05}(1,9; 36,1) \doteq 0,603$. Nulovou hypotézu tedy můžeme zamítnout a učinit závěr, že mezi uspořádáním odpovědí 20 respondentů je statisticky signifikantní shoda. Rozdíly mezi uspořádáním můžeme považovat za nahodilé.

Příklad 7.

Počet respondentů a odpovědí ponechme stejný jako v příkladu 6. Předpokládejme, že $S = 178$. Test uskutečnime na 1 % hladině významnosti.

Dosazením do vzorce vypočítáme:

$$W = 0,2225,$$

$$z \doteq \frac{1}{2} [(\log 5,437) \cdot 2,302\ 585] \doteq 0,847$$

Vzhledem k tomu, že se počet respondentů ani počet odpovědí nezměnil, zůstávají hodnoty pro stupně volnosti stejné jako v předchozím příkladu, tj. $v_1 = 1,9$ a $v_2 = 36,1$. Hodnoty Fisherovy z-distribuce pro zvolenou hladinu významnosti a pro stupně volnosti v_1 a v_2 jsme vypočetli v příkladu 5. Hodnota $z_{0,01}(1,9; 36,1) \doteq 0,847$.

Platí, že $z \doteq 0,847 = z_{0,01}(1,9; 36,1) \doteq 0,847$. Nulovou hypotézu zamítáme na 1 % hladině významnosti a konstatujeme, že shoda mezi uspořádaním odpovědí je statisticky význačná.

Poznámka. Hodnota S uvedená v příkladu činí 178, je tedy o 1 vyšší než kritická hodnota na hladině významnosti 0,01, která je uvedena v tabulce I. Kdybychom použili v příkladu $S = 177$, bylo by $W = 0,221\ 25$ a $z \doteq 0,843$, tj. z by bylo o 0,004 menší než $z_{0,01}(1,9; 36,1)$. Rozdíl 0,004 je poměrně malý, zejména vezmeme-li v úvahu, že jsme $z_{0,01}(1,9; 36,1)$ zjišťovali z tabulek trojnásobnou lineární interpolací.

Použití v praxi

Bude-li počet respondentů malý, vystačíme většinou s ručním zpracováním (tak tomu bude převážně při sociologických pilotážích a při předvýzkumech). Pro test použijeme buď tabulku I nebo při větším počtu objektů provedeme test pomocí χ^2 — distribuce.

Při sociologických výzkumech může být počet respondentů poměrně velký. K testu bude v takovém případě vhodné použít samostatný počítač. Pokud je nám známo, nebyl

u нас досуд vypracován program pro test. Uskutečnime-li test s použitím Fisherovy z-distribuce, můžeme uložit do paměti počítače hodnoty uvedené v tabulkách II a III nebo vypracovat program pro přesný výpočet distribuční funkce pro libovolnou velikost stupňů volnosti.

Při vypracování programu pro samočinný počítač bychom měli přihlédnout k tomu, že shodu v uspořádání odpovědí můžeme očekávat spíše u relativně stejnorodých skupin respondentů. Program by proto měl být vypracován tak, aby samočinný počítač provedl test na zvolených hladinách významnosti nejprve u stejnorodějších skupin respondentů vybraných podle zvoleného třídění a potom by podle instrukce test opakoval u větších a méně stejnorodých skupin. Skupiny respondentů musíme volit tak, aby odpovídaly hypotézám pro příslušný výzkum.

Literatura

- [1] M. C. Kendall, Sc. D.: *Rank Correlation Methods*, Third Edition, Charles Griffin and Co, London.
- [2] Sidney Siegel: *Nonparametric Statistics for the Behavioral Sciences*, Mc Graw-Hill, New York 1956.
- [3] V. Břicháček — O. Hampejsová: *Neparametrické techniky statistického hodnocení v psychologickém výzkumu* (II. část), Psychologické štúdie SAV IV/1962.

Резюме

Свобода В.: Применение W-коэффициента согласия Кендалла в социологических исследованиях

В социологических исследованиях до сих пор мало применяется прием, на основе которого можно получить упорядочение n объектов в зависимости от K разных свойств или K судьями, где $K \geq 3$. Объектами могут являться напр. ответы на вопросы, способ упорядочения которых определяется в зависимости от значимости, от интереса и т. п. Или мы можем интересоваться способом упорядочения лиц в зависимости от их определенных свойств (напр. учащихся, артистов, спортсменов итд.). Роль судей могут выполнять респонденты или группы респондентов.

На основе статистического теста можно решить, возможно ли на избранном уровне значимости отказаться от нулевой гипотезы, содержащей утверждение, что респонденты не согласны относительно способа упорядочения объектов. Для тестирования гипотезы удобно применить W-коэффициент согласия Кендалла.

Прием, применяемый для теста, зависит от величины K и n . Если количество объектов и респондентов невелико, то можно сделать тест, применяя критические ценности S , указанные в таблице I. В случае, когда количество упорядоченных объектов составляет по край-

ней мере 8, то возможно применить аппроксимацию χ^2 -дистрибуций. В случае, когда нельзя применить ни таблицу I, ни аппроксимацию χ^2 -дистрибуций, выполним тест на основе аппроксимации при помощи z-дистрибуции Фишера. Величины z-дистрибуции Фишера приводятся в таблицах II и III.

Случаи, когда в социологических исследованиях надо будет применить аппроксимацию при помощи z-дистрибуции Фишера, могут быть довольно частыми. Это будет напр. всегда, если количество объектов меньше восьми и количество респондентов больше.

Аппроксимацию при помощи z-дистрибуций Фишера можно применить во всех случаях (для малых величин следовало бы сделать исправление для непрерывности). Однако расчет при помощи z-дистрибуции Фишера сложнее, поэтому выгодно при разработке данных без помощи вычислительной машины применить таблицу I. или аппроксимацию χ^2 -дистрибуций.

При наличии большого количества респондентов будет целесообразным пользоваться электронно-вычислительной машиной, которая в случае надобности может иметь программу для точного расчета z-дистрибуции Фишера для любой величины степеней свободы. Программу для электронной вычислительной машины следовало бы разработать таким образом, чтобы было возможно, во-первых, выполнить тест в более однородных группах респондентов, от которых можно скорее ожидать согласие со способом упорядочения ответов. Затем, на основе инструкции, тест бы повторялся в больших и менее однородных группах. Группы респондентов следует избирать таким образом, чтобы они соответствовали гипотезам для данного исследования.

Summary

Svoboda V.: The Application of Kendall Coefficient of Concordance W in Sociological Research

In sociological research, the method enabling us to obtain an arrangement of n objects according to k various qualities or by k judges, where $k \geq 3$, has so far been applied very seldom. Objects may be represented e. g. by answers to questions the order of which is established in terms of relevance, popularity, etc. At other times, we may be interested in the organization of people (e. g. of pupils, artists, sportsmen, etc.) according to particular qualities. The role of judges may be performed by the respondents or by groups of respondents.

On the basis of a statistical test we may decide whether it is possible, at the chosen level of significance, to reject the zero-hypothesis containing the statement that the respondents do not agree in arranging the objects. For testing the hypothesis, Kendall Coefficient of Concordance W may adequately be applied.

The procedure used in the test depends on the size of k and n . If the number of both objects and respondents is small, we may carry out the test using the critical values S

shown in Table I. If the number of arranged objects is at least eight, the approximation by the χ^2 -distribution can be used. If it is impossible to apply Table I and the approximation by the χ^2 -distribution, the test may be realized on the basis of an approximation by means of Fisher's z-distribution. The values of Fisher's z-distribution are given in Tables II and III.

The cases where approximations by means of Fisher's z-distribution have to be applied in sociological research may be relatively frequent. This will always be the case e. g. where the number of objects is lower than eight, while the number of respondents is higher.

The approximation by means of Fisher's z-distribution may be applied in all the cases (for small values, a correction should be made for the connection). However, the computation with the aid of Fisher's z-distribution

being more complex, it is convenient to use, in the elaboration of data by hand, Table I or the approximation by the z-distribution — wherever this is possible.

In case of a great number of respondents it will be expedient, for the purposes of the test, to make use of a computer eventually provided with the program designed for an accurate computation of Fisher's z-distribution for an arbitrary size of the degrees of freedom. The computer program should be elaborated so as to enable the researchers to carry out the test first of all with more homogeneous groups of respondents, where conformity in arranging the answers may be expected with greater probability. The test would then be repeated with larger and less homogeneous groups according to instruction. The groups of respondents must be chosen so as to correspond to the hypotheses valid in the respective research.

Tabulka I

Kritické hodnoty S (pro koeficient shody W)

Od Friedmana (1940) se svolením autora a vydavatele Annals of Mathematical Statistics

k	n					Součtové hodnoty pro $n = 3$	
	3	4	5	6	7	k	S
Hodnoty pro hladinu významnosti 0,05							
3			64,4	103,9	157,3	9	54,0
4		49,5	88,4	143,3	217,0	12	71,9
5		62,6	112,3	182,4	276,2	14	83,8
6		75,7	136,1	221,4	335,2	16	95,8
8	48,1	101,7	183,7	299,0	453,1	18	107,7
10	60,0	127,8	231,2	376,7	571,0		
15	89,8	192,9	349,8	570,5	864,9		
20	119,7	258,0	468,5	764,4	1 158,7		
Hodnoty pro hladinu významnosti 0,01							
3			75,6	122,8	185,6	9	75,9
4		61,4	109,3	176,2	265,0	12	103,5
5		80,5	142,8	229,4	343,8	14	121,9
6		99,5	176,1	282,4	422,6	16	140,2
8	66,8	137,4	242,7	388,3	579,9	18	158,6
10	85,1	175,3	309,1	494,0	737,0		
15	131,0	269,8	475,2	758,2	1 129,5		
20	177,0	364,2	641,2	1 022,2	1 521,9		

Převzato z knihy: M. C. Kendall, Sc. D.: Rank Correlation Methods, Third Edition, Charles Griffin and Co, London.

Tabulka II
z-rozložení*)
0,1 procentní body (Colcord a Deming)

$\begin{smallmatrix} v_1 \\ \backslash \\ v_2 \end{smallmatrix}$	1	2	3	4	5	6	8	12	24	∞
1	6,5462	6,5612	6,5966	6,6201	6,6323	6,6405	6,6508	6,6611	6,6715	6,6819
2	3,4531	3,4534	3,4535	3,4535	3,4535	3,4535	3,4536	3,4536	3,4536	3,4536
3	2,5604	2,5003	2,4748	2,4603	2,4511	2,4446	2,4361	2,4272	2,4179	2,4081
4	2,1529	2,0574	2,0143	1,9892	1,9728	1,9612	1,9459	1,9294	1,9118	1,8927
5	1,9255	1,8002	1,7513	1,7184	1,6964	1,6808	1,6596	1,6370	1,6123	1,5845
6	1,7849	1,6479	1,5828	1,5433	1,5177	1,4986	1,4730	1,4449	1,4134	1,3783
7	1,6874	1,5384	1,4662	1,4221	1,3927	1,3711	1,3417	1,3090	1,2721	1,2296
8	1,6177	1,4587	1,3809	1,3332	1,3008	1,2770	1,2443	1,2077	1,1662	1,1169
9	1,5646	1,3982	1,3160	1,2653	1,2304	1,2047	1,1694	1,1293	1,0830	1,0279
10	1,5232	1,3509	1,2650	1,2116	1,1748	1,1475	1,1098	1,0668	1,0165	,9557
11	1,4900	1,3128	1,2238	1,1683	1,1297	1,1012	1,0614	1,0157	,9619	,8957
12	1,4627	1,2814	1,1900	1,1326	1,0926	1,0628	1,0213	,9733	,9162	,8450
13	1,4400	1,2553	1,1616	1,1026	1,0614	1,0306	,9875	,9374	,8774	,8014
14	1,4208	1,2332	1,1376	1,0772	1,0348	1,0031	,9586	,9066	,8439	,7635
15	1,4043	1,2141	1,1169	1,0553	1,0119	,9795	,9336	,8800	,8147	,7301
16	1,3900	1,1976	1,0989	1,0362	,9920	,9588	,9119	,8567	,7891	,7005
17	1,3775	1,1832	1,0832	1,0195	,9745	,9407	,8927	,8361	,7664	,6740
18	1,3665	1,1704	1,0693	1,0047	,9590	,9246	,8757	,8178	,7462	,6502
19	1,3567	1,1591	1,0569	,9915	,9442	,9103	,8605	,8014	,7277	,6285
20	1,3480	1,1489	1,0458	,9798	,9329	,8974	,8469	,7867	,7115	,6086
21	1,3401	1,1398	1,0358	,9691	,9217	,8858	,8346	,7735	,6964	,5904
22	1,3329	1,1315	1,0268	,9595	,9116	,8753	,8234	,7612	,6828	,5738
23	1,3264	1,1240	1,0186	,9507	,9024	,8657	,8132	,7501	,6704	,5583
24	1,3205	1,1171	1,0111	,9427	,8939	,8569	,8038	,7400	,6589	,5440
25	1,3151	1,1108	1,0041	,9354	,8862	,8489	,7953	,7306	,6483	,5307
26	1,3101	1,1050	,9978	,9286	,8791	,8415	,7873	,7220	,6385	,5183
27	1,3055	1,0997	,9920	,9223	,8725	,8346	,7800	,7140	,6294	,5066
28	1,3013	1,0947	,9866	,9165	,8664	,8282	,7732	,7066	,6209	,4957
29	1,2973	1,0903	,9815	,9112	,8607	,8223	,7679	,6997	,6129	,4853
30	1,2936	1,0859	,9768	,9061	,8554	,8168	,7610	,6932	,6056	,4756
40	1,2674	1,0552	,9435	,8701	,8174	,7771	,7184	,6463	,5513	,4016
60	1,2413	1,0248	,9100	,8345	,7798	,7377	,6760	,5992	,4955	,3198
120	1,2158	,9952	,8773	,7994	,7426	,6986	,6338	,5519	,4380	,2199
∞	1,1910	,9663	,8453	,7648	,7059	,6599	,5917	,5044	,3786	0

Pro vysoké hodnoty v_1 a v_2 , je přibližně $z(0,1\%) = \frac{3,0902}{\sqrt{h-2,1}} - 1,925 \left(\frac{1}{v_1} - \frac{1}{v_2} \right)$, kde $\frac{2}{h} = \frac{1}{v_1} + \frac{1}{v_2}$.
V ostatních částech tabulky je interpolace přibližně lineární, je-li brána v úvahu reciprocita v_1 a v_2 .

*) Převzato z knihy: Fisher - Yates: Statistical Tables, London 1949 (Third Edition).

Tabulka II
z-rozdělení*)
1 procentní body

Pokračování 1

$\begin{smallmatrix} v_1 \\ \backslash \\ v_2 \end{smallmatrix}$	1	2	3	4	5	6	8	12	24	∞
1	4,1535	4,2585	4,2974	4,3175	4,3297	4,3379	4,3482	4,3585	4,3689	4,3794
2	2,2950	2,2976	2,2984	2,2988	2,2991	2,2992	2,2994	2,2997	2,2999	2,3001
3	1,7649	1,7140	1,6915	1,6786	1,6703	1,6645	1,6569	1,6489	1,6404	1,6314
4	1,5270	1,4452	1,4075	1,3856	1,3711	1,3609	1,3473	1,3327	1,3170	1,3000
5	1,3943	1,2929	1,2449	1,2164	1,1974	1,1838	1,1656	1,1457	1,1239	1,0997
6	1,3103	1,1955	1,1401	1,1068	1,0843	1,0680	1,0460	1,0218	,9948	,9643
7	1,2526	1,1281	1,0672	1,0300	1,0048	,9864	,9614	,9335	,9020	,8658
8	1,2106	1,0787	1,0135	,9734	,9459	,9259	,8983	,8673	,8319	,7904
9	1,1786	1,0411	,9724	,9299	,9006	,8791	,8494	,8157	,7769	,7305
10	1,1535	1,0114	,9399	,8954	,8646	,8419	,8104	,7744	,7324	,6816
11	1,1333	,9874	,9136	,8674	,8354	,8116	,7785	,7405	,6958	,6408
12	1,1166	,9677	,8919	,8443	,8111	,7864	,7520	,7122	,6649	,6061
13	1,1027	,9511	,8737	,8248	,7907	,7652	,7295	,6882	,6386	,5761
14	1,0909	,9370	,8581	,8082	,7732	,7471	,7103	,6675	,6159	,5500
15	1,0807	,9249	,8448	,7939	,7582	,7314	,6937	,6496	,5961	,5269
16	1,0719	,9144	,8331	,7814	,7450	,7177	,6791	,6339	,5786	,5064
17	1,0641	,9051	,8229	,7705	,7335	,7057	,6663	,6199	,5630	,4879
18	1,0572	,8970	,8138	,7607	,7232	,6950	,6549	,6075	,5491	,4712
19	1,0511	,8897	,8067	,7521	,7140	,6854	,6447	,5964	,5366	,4560
20	1,0457	,8831	,7985	,7443	,7058	,6768	,6355	,5864	,5253	,4421
21	1,0408	,8772	,7920	,7372	,6984	,6690	,6272	,5773	,5150	,4294
22	1,0363	,8719	,7860	,7309	,6916	,6620	,6196	,5691	,5056	,4176
23	1,0322	,8670	,7806	,7251	,6855	,6555	,6127	,5615	,4969	,4068
24	1,0285	,8626	,7757	,7197	,6799	,6496	,6064	,5545	,4890	,3967
25	1,0251	,8585	,7712	,7148	,6747	,6442	,6006	,5481	,4816	,3872
26	1,0220	,8548	,7670	,7103	,6699	,6392	,5952	,5422	,4748	,3784
27	1,0191	,8513	,7631	,7062	,6655	,6346	,5902	,5367	,4685	,3701
28	1,0164	,8481	,7595	,7023	,6614	,6303	,5856	,5316	,4626	,3624
29	1,0139	,8451	,7562	,6987	,6576	,6263	,5813	,5269	,4570	,3550
30	1,0116	,8423	,7531	,6954	,6540	,6226	,5773	,5224	,4519	,3481
40	,9949	,8223	,7307	,6712	,6283	,5956	,5481	,4901	,4138	,2952
60	,9784	,8025	,7086	,6472	,6028	,5687	,5189	,4574	,3746	,2352
120	,9622	,7829	,6867	,6234	,5774	,5419	,4897	,4243	,3339	,1612
∞	,9462	,7636	,6651	,5999	,5522	,5152	,4604	,3908	,2913	0

Pro vysoké hodnoty v_1 a v_2 , je přibližně $z(1\%) = \frac{2,3263}{\sqrt{h-1,4}} - 1,235 \left(\frac{1}{v_1} - \frac{1}{v_2} \right)$, kde $\frac{2}{h} = \frac{1}{v_1} + \frac{1}{v_2}$.

V ostatních částech tabulky je interpolace přibližně lineární, je-li brána v úvahu reciprocita v_1 a v_2 .

*) Převzato z knihy: Fisher - Yates: Statistical Tables, London 1949 (Third Edition).

Tabulka II
z-rozdělení*)

5 procentní body

Pokračování 2

$v_1 \backslash v_2$	1	2	3	4	5	6	8	12	24	∞
1	2,5421	2,6479	2,6870	2,7071	2,7194	2,7276	2,7380	2,7484	2,7588	2,7693
2	1,4592	1,4722	1,4765	1,4787	1,4800	1,4808	1,4819	1,4830	1,4840	1,4851
3	1,1577	1,1284	1,1137	1,1051	1,0994	1,0953	1,0899	1,0842	1,0781	1,0716
4	1,0212	,9600	,9429	,9272	,9168	,9093	,8993	,8885	,8767	,8639
5	,9441	,8777	,8441	,8236	,8097	,7997	,7862	,7714	,7550	,7368
6	,8948	,8188	,7798	,7558	,7394	,7274	,7112	,6931	,6729	,6499
7	,8606	,7777	,7347	,7080	,6896	,6761	,6576	,6369	,6134	,5862
8	,8355	,7475	,7014	,6725	,6525	,6378	,6175	,5945	,5682	,5371
9	,8163	,7242	,6757	,6450	,6238	,6080	,5862	,5613	,5324	,4979
10	,8012	,7058	,6553	,6232	,6009	,5843	,5611	,5346	,5035	,4657
11	,7889	,6909	,6387	,6055	,5822	,5648	,5406	,5126	,4795	,4387
12	,7788	,6786	,6250	,5907	,5666	,5487	,5234	,4941	,4592	,4156
13	,7703	,6682	,6134	,5783	,5535	,5350	,5089	,4785	,4419	,3957
14	,7630	,6594	,6036	,5677	,5423	,5233	,4964	,4649	,4269	,3782
15	,7568	,6518	,5950	,5585	,5326	,5131	,4855	,4532	,4138	,3628
16	,7514	,6451	,5876	,5505	,5241	,5042	,4760	,4428	,4022	,3490
17	,7466	,6393	,5811	,5434	,5166	,4964	,4676	,4337	,3919	,3366
18	,7424	,6341	,5753	,5371	,5099	,4894	,4602	,4255	,3827	,3253
19	,7386	,6295	,5701	,5315	,5040	,4832	,4535	,4182	,3743	,3151
20	,7352	,6254	,5654	,5265	,4986	,4776	,4474	,4116	,3668	,3057
21	,7322	,6216	,5612	,5219	,4938	,4725	,4420	,4055	,3599	,2971
22	,7294	,6182	,5574	,5178	,4894	,4679	,4370	,4001	,3536	,2892
23	,7269	,6151	,5540	,5140	,4854	,4636	,4325	,3950	,3478	,2818
24	,7246	,6123	,5508	,5106	,4817	,4598	,4283	,3904	,3425	,2749
25	,7225	,6097	,5478	,5074	,4783	,4562	,4244	,3862	,3376	,2685
26	,7205	,6073	,5451	,5045	,4752	,4529	,4209	,3823	,3330	,2625
27	,7187	,6051	,5427	,5017	,4723	,4499	,4176	,3786	,3287	,2569
28	,7171	,6030	,5403	,4992	,4696	,4471	,4146	,3752	,3248	,2516
29	,7155	,6011	,5382	,4969	,4671	,4444	,4117	,3720	,3211	,2466
30	,7141	,5994	,5362	,4947	,4648	,4420	,4090	,3691	,3176	,2419
40	,7037	,5866	,5217	,4789	,4479	,4242	,3897	,3475	,2920	,2057
60	,6933	,5738	,5073	,4632	,4311	,4064	,3702	,3255	,2654	,1644
120	,6830	,5611	,4930	,4475	,4143	,3885	,3506	,3032	,2376	,1131
∞	,6729	,5486	,4787	,4319	,3974	,3706	,3309	,2804	,2085	0

Pro vysoké hodnoty v_1 a v_2 , je přibližně z (5 %) = $\frac{1,6449}{\sqrt{h-1}} - 0,7843 \left(\frac{1}{v_1} - \frac{1}{v_2} \right)$, kde $\frac{2}{h} = \frac{1}{v_1} + \frac{1}{v_2}$.

V ostatních částech tabulky je interpolace přibližně lineární, je-li brána v úvahu reciprocita v_1 a v_2 .

*) Převzato z knihy: Fisher - Yates: Statistical Tables, London 1949, (Third Edition).

Tabulka II

z-rozdělení*)

10 procentní body

Dokončení

$\frac{v_1}{v_2}$	1	2	3	4	5	6	8	12	24	∞
1	1,8427	1,9510	1,9907	2,0112	2,0236	2,0320	2,0425	2,0530	2,0636	2,0742
2	1,0716	1,0986	1,1075	1,1120	1,1146	1,1164	1,1186	1,1208	1,1230	1,1252
3	,8558	,8489	,8423	,8379	,8347	,8324	,8293	,8258	,8221	,8179
4	,7570	,7322	,7164	,7064	,6994	,6944	,6875	,6799	,6716	,6623
5	,7006	,6648	,6432	,6293	,6196	,6125	,6029	,5921	,5801	,5665
6	,6643	,6211	,5953	,5786	,5669	,5583	,5465	,5332	,5181	,5007
7	,6390	,5905	,5615	,5427	,5295	,5197	,5061	,4907	,4730	,4523
8	,6203	,5678	,5364	,5160	,5015	,4907	,4757	,4585	,4386	,4148
9	,6060	,5504	,5171	,4953	,4798	,4682	,4520	,4333	,4114	,3849
10	,5947	,5366	,5017	,4788	,4624	,4502	,4330	,4130	,3893	,3602
11	,5855	,5253	,4892	,4653	,4483	,4355	,4173	,3962	,3710	,3395
12	,5779	,5160	,4788	,4541	,4365	,4231	,4043	,3821	,3555	,3219
13	,5715	,5082	,4701	,4447	,4265	,4127	,3932	,3702	,3422	,3066
14	,5661	,5015	,4626	,4366	,4180	,4038	,3836	,3598	,3308	,2931
15	,5614	,4957	,4561	,4296	,4106	,3961	,3754	,3508	,3207	,2813
16	,5572	,4907	,4504	,4235	,4041	,3893	,3681	,3429	,3118	,2706
17	,5537	,4863	,4455	,4181	,3984	,3833	,3617	,3359	,3038	,2611
18	,5505	,4823	,4411	,4134	,3933	,3780	,3560	,3296	,2967	,2524
19	,5476	,4788	,4371	,4091	,3887	,3732	,3508	,3240	,2904	,2445
20	,5451	,4757	,4336	,4052	,3846	,3689	,3462	,3189	,2846	,2373
21	,5427	,4728	,4304	,4017	,3809	,3650	,3420	,3143	,2793	,2307
22	,5407	,4703	,4275	,3986	,3776	,3615	,3382	,3101	,2744	,2245
23	,5388	,4679	,4248	,3957	,3745	,3582	,3347	,3062	,2700	,2188
24	,5370	,4657	,4224	,3931	,3717	,3553	,3315	,3027	,2659	,2135
25	,5354	,4638	,4201	,3906	,3691	,3526	,3286	,2994	,2621	,2086
26	,5339	,4620	,4181	,3884	,3667	,3500	,3258	,2964	,2585	,2039
27	,5326	,4603	,4162	,3863	,3645	,3477	,3233	,2936	,2553	,1996
28	,5313	,4587	,4144	,3844	,3624	,3455	,3210	,2910	,2522	,1955
29	,5301	,4572	,4128	,3826	,3605	,3435	,3188	,2885	,2493	,1916
30	,5290	,4559	,4112	,3809	,3587	,3416	,3167	,2863	,2466	,1880
40	,5211	,4461	,4001	,3688	,3458	,3280	,3019	,2696	,2268	,1599
60	,5133	,4363	,3891	,3567	,3328	,3142	,2868	,2526	,2063	,1279
120	,5054	,4266	,3781	,3446	,3198	,3005	,2717	,2354	,1848	,0081
∞	,4976	,4170	,3671	,3326	,3069	,2866	,2565	,2178	,1622	0

Pro vysoké hodnoty v_1 a v_2 , je přibližně z (10 %) = $\frac{1,2816}{\sqrt{h - 0,8}} - 0,6071 \left(\frac{1}{v_1} - \frac{1}{v_2} \right)$, kde $\frac{2}{h} = \frac{1}{v_1} + \frac{1}{v_2}$.
 V ostatních částech tabulky je interpolace přibližně lineární, je-li brána v úvahu reciprocity v_1 a v_2 .

*) Převzato z knihy: Fisher - Yates: Statistical Tables, London 1949 (Third Edition).

Tabulka III
z-rozdělení
1 procentní body*)

v_2	v_1					
	1	2	3	4	5	6
200	0,978	0,774	0,677	0,613	0,567	0,532
400	0,951	0,769	0,671	0,605	0,559	0,523
1000	0,948	0,765	0,667	0,602	0,555	0,518
∞	0,9462	0,7636	0,6651	0,5999	0,5522	0,5152

5 procentní body**)

v_2	v_1					
	1	2	3	4	5	6
200	0,679	0,555	0,487	0,439	0,407	0,380
400	0,675	0,552	0,481	0,435	0,401	0,376
1000	0,674	0,549	0,479	0,433	0,398	0,371
∞	0,6729	0,5486	0,4787	0,4319	0,3974	0,3706

*) Pro $v_2 = 200, 400, 1000$ mohou být hodnoty uvedené v tabulce, ve srovnání se zaokrouhlenými hodnotami, zatíženy chybou maximálně $\pm 0,001$.

**) Pro $v_1 = 1, 2, 3, 4, 5$ a $v_2 = 200, 400, 1000$ mohou být hodnoty uvedené v tabulce, ve srovnání se zaokrouhlenými hodnotami, zatíženy chybou maximálně $\pm 0,001$. Pro $v_1 = 6$ a $v_2 = 200, 400, 1000$ mohou být hodnoty uvedené v tabulce zatíženy chybou maximálně $\pm 0,002$.