

Teste neparametrice

Conf. dr. habil. Eduard Rothenstein

Testele parametrice funcționează în ipoteza în care datele selectate urmează o repartiție normală sau volumul acestora este suficient de mare, pentru ca aproximarea cu repartiția normală să fie validă. Apare astfel întrebarea dacă mai putem afla informații despre repartiția datelor sau despre parametrii variabilei în cazul în care volumul datelor este redus sau avem dubii în ceea ce privește normalitatea lor.

Testele neparametrice ar putea oferi un răspuns pozitiv la această întrebare. Acestea sunt teste statistice în cadrul cărora nu se fac presupuneri asupra formei repartiției. Ele nu verifică valorile parametrilor tradiționali, de aceea mai sunt cunoscute și sub titulatura de *metode fără parametri* (en., parameter-free methods) sau *metode fără repartiție* (en., distribution-free methods). Testele neparametrice pot fi utilizate atunci când sunt dubii asupra normalității datelor statistice.

Se pot construi teste neparametrice corespunzătoare fiecărui test parametric studiat mai sus, însă aceste teste neparametrice sunt, în general, grupate în următoarele categorii:

- teste pentru diferența dintre grupuri (pentru selecții independente). Este cazul comparării mediilor a două selecții ce provin din populații independente. De regulă, dacă ipotezele acestuia sunt îndeplinite, atunci se utilizează *testul t*. Variante neparametrice ale acestui test sunt: *testul Wald-Wolfowitz* sau *testul Mann-Whitney*.
- teste pentru diferența dintre variabile (pentru selecții dependente). Utilizat la compararea a două variabile ce caracterizează populația din care s-a luat selecția. Teste neparametrice utilizate: *testul semnelor*, *testul Wilcoxon* (signed-rank).
- teste pentru relații între variabile. Pentru a determina corelația între două variabile, de regulă se utilizează coeficientul de corelație al lui Pearson. Există variante neparametrice ale testului bazat pe coeficientul de corelație Pearson, e.g., coeficientul *R* (Spearman), coeficientul *t* (Kendall) sau coeficientul *G* (Goodman și Kruskal).

Avantajul testelor neparametrice este că sunt mai robuste, adică folosesc mai puține ipoteze decât testele parametrice. Testele neparametrice nu au nevoie de o repartiție a priori cunoscută a datelor observate sau de un volum mare de date. Totuși, efectul lipsei unor ipoteze restrictive face ca puterea unui test neparametric să fie (în general) mai mică decât a testului parametric corespunzător (care ar fi folosit dacă ipotezele sale sunt satisfăcute). Astfel, în cazul unui test neparametric sunt șanse mai mici ca ipoteza nulă să fie respinsă atunci când ea este, în realitate, falsă. Acest fapt înseamnă că valoarea P_v este mai mare în cazul unui test neparametric decât în cazul testului parametric corespunzător, calculată pentru același set de date. Testele neparametrice pot fi singurele opțiuni pentru analiza datelor statistice în următoarele cazuri: datele sunt ordinale, datele sunt fără valori numerice, datele conțin valori aberante extreme sau în cazul în care datele sunt rezultatul unor măsurători imprecise. Dacă s-ar dori analiza acestor date folosind teste parametrice, vor fi necesare ipoteze restrictive asupra datelor, cum ar fi ipoteza de normalitate.

În general, dacă atât metodele parametrice cât și cele nonparametrice sunt aplicabile unei anumite probleme, ar trebui să utilizăm procedura parametrică mai eficientă. Cu toate acestea, presupunerile pentru metoda parametrică pot fi dificile sau imposibil de justificat. De exemplu, datele pot fi conectate prin intermediul rangurilor. Aceste situații apar frecvent în practică. De exemplu, un complet de judecători poate fi utilizat pentru a evalua 10 tipuri diferite ale unei băuturi răcoritoare pentru o calitate globală, cu formularea „cea mai bună” rangul 1, formularea „cel mai bun următor” a fost atribuită cu rangul 2 și așa mai departe. Este puțin probabil ca datele de tip rang să satisfacă condiția normalității. Multe metode nonparametrice implică analiza rangurilor și, în consecință, sunt ideale pentru acest tip de probleme.

1 Testul semnelor

Testul semnelor se mai numește și testul medianei. Este un test neparametric bazat pe semnele anumitor valori și nu pe valorile în sine. Testul semnelor este util atunci când avem date ordinale (grupate pe categorii ordonate), fără a ști valorile numerice ale diferențelor dintre categorii. Dacă valorile numerice sunt cunoscute, atunci se poate folosi un test mai puternic, e.g., *testul rangurilor* cu semn al lui Wilcoxon. Este unul dintre cele mai simple teste statistice neparametrice. Pentru ca acest test să poate fi utilizat, trebuie ca datele statistice observate să fie alese aleator și independent din populația considerată. Acest test verifică valoarea centrală a setului de date observate și nu impune nicio ipoteză referitoare la repartiția datelor. La *testul t* clasic, valoarea centrală

testată este media (în condițiile normalității datelor sau pentru un volum suficient de mare de date), iar la testul semnelor se testează valoarea mediană a observațiilor. Dacă setul de date este simetric (așa cum este cazul datelor empirice pentru o repartiție normală), atunci valoarea mediană este egală cu media. În acest caz, testul semnelor poate da informații despre media datelor observate, deși este un test mai puțin precis decât testul t .

Condițiile testului: Datele $x_1, \dots, x_n$ sunt observații aleatoare și independente asupra unei caracteristici continue X a unei populații.

Ipoteza nulă:

$$(H_0) : Me = Me^* \text{ (valoarea mediană a datelor este o valoare dată, } Me^*),$$

la nivelul de semnificație α . În funcție de ipoteza alternativă, putem avea un test unilateral sau un test bilateral.

Test unilateral stânga:

$$(H_1)_s : Me < Me^*.$$

$$\text{Statistica test este } S^* = S_{<} = \sum_{i=1}^n 1_{\{x_i < Me^*\}}.$$

Test bilateral:

$$(H_1) : Me \neq Me^*.$$

$$\text{Statistica test este } S^* = S_{\neq} = \max \{S_{<}, S_{>}\}.$$

unde $S_{<}$ este numărul datelor mai mici decât Me^* .

Pentru testul unilateral dreapta, ipoteza alternativă este $(H_1)_d : Me > Me^*$, iar statistica test este $S^* = S_{>} = \sum_{i=1}^n 1_{\{x_i > Me^*\}}$, adică numărul datelor mai mari decât Me^* .

Dacă ipoteza nulă este adevărată și mediana este Me , atunci S este o variabilă binomială $S \sim B(n, 0.5)$. Pe baza acestor statistici se calculează nivelul de semnificație observat, P_v , care reprezintă probabilitatea de a obține un rezultat cel puțin la fel de extrem ca și cel observat, dacă ipoteza nulă este adevărată. Vom avea:

$$\text{cazul unilateral: } P_v = \mathbb{P}(S \geq S^*); \quad \text{cazul bilateral: } P_v = 2\mathbb{P}(S \geq S^*),$$

unde $S \sim B(n, 0.5)$. Dacă valoarea P_v este mai mare decât α , atunci acceptăm ipoteza nulă (nu avem motive să o respingem). Altfel, acceptăm ipoteza alternativă.

Egalitate în Semnul testelor. Deoarece populația pentru care caracteristica investigată X este considerată a fi continuă, $\mathbb{P}(X_i = Me^*) = 0$. Cu toate acestea, în practică, datorită modului în care sunt culese datele empirice, aceasta valoare mediană poate fi chiar atinsă. Atunci când se întâmplă acest lucru, aceste măsurători sunt eliminate și se aplică Testul semnelor pentru datele rămase.

Observația 1.1 Dacă volumul observațiilor este mare (e.g. $n \geq 10$) și $S^* \sim B(n, 0.5)$ atunci, conform Teoremei Limită Centrală, repartiția binomială este bine aproximată prin intermediul unei repartiții normale. Prin urmare, repartiția test folosită este $S^* \sim \mathcal{N}(n/2, \sqrt{n}/2)$. În acest caz, testul pentru mediană se poate face pe baza statisticii

$$Z_0 = \frac{S^* - n/2}{\sqrt{n}/2}, \quad \text{cu valoarea sa calculată în datele empirice obținute } z_0,$$

Decizia finală se ia astfel: respingem ipoteza nulă (H_0) dacă

$$z_0 < -z_{1-\alpha} \text{ (pentru } (H_1)_s), \quad z_0 > z_{1-\alpha} \text{ (pentru } (H_1)_d), \quad |z_0| > z_{1-\frac{\alpha}{2}} \text{ (pentru } (H_1)),$$

unde $z_{1-\alpha}$ și $z_{1-\frac{\alpha}{2}}$ sunt cuantilele (tabelate) ale repartiției normale standard.

Exemplul 1.1 Montgomery, Peck și Vining (2001) prezintă o analiză asupra unui motor rachetă legând informațiile unui propulsor de aprindere de cele ale unui propulsor de susținere, în interiorul unei carcase metalice. Rezultatele testării a 20 de motoare selectate la întâmplare sunt prezentate în tabelul următor. Am dori să testăm ipoteza conform căreia forța de forfecare medie între cele două tipuri de motoare este de 2000 psi, folosind un prag de semnificație $\alpha = 0.05$.

Observația	Forța	Diferența	Semn	Observația	Forța	Diferența	Semn
1	2158.70	158.70	+	11	2165.20	165.20	+
2	1678.15	-321.85	-	12	2399.55	399.55	+
3	2316.00	316.00	+	13	1779.80	-220.20	-
4	2061.30	61.30	+	14	2336.75	336.75	+
5	2207.50	207.50	+	15	1765.30	-234.70	-
6	1708.30	-291.70	-	16	2053.50	53.50	+
7	1784.70	-215.30	-	17	2414.40	414.40	+
8	2575.10	575.10	+	18	2200.50	200.50	+
9	2357.90	357.90	+	19	2654.20	654.20	+
10	2256.70	256.70	+	20	1753.70	-246.30	-

Formulăm ipotezele

$$(H_0) : Me = 2000 \quad \text{versus} \quad (H_1) : Me \neq 2000,$$

iar pragul de semnificație $\alpha = 0.05$. Avem $S_{<} = 6$, $S_{>} = 14$ și $S_{\neq} = 14$. Probabilitatea critică este, știind că repartiția caracteristicii investigate este $\mathcal{B}(20, 1/2)$:

$$P_v = 2 \cdot \mathbb{P}(S \geq 14) = 2 \sum_{k=14}^{20} C_{20}^k \frac{1}{2^{20}} = 0.1153 > 0.05 = \alpha,$$

deci, în consecință nu putem respinge ipoteza nulă. Aceasta înseamnă că numărul de semne + observate nu este suficient de mare sau de mic pentru a indica că valoarea mediană este semnificativ diferită de valoarea de 2000, la un nivel de semnificație 0.05.

Aplicăm acum procedura de aproximare prin normalizare. Respingem ipoteza nulă dacă $|z_0| > z_{0.025} = 1.96$. În situația considerată,

$$z_0 = \frac{14 - 20/2}{\sqrt{20}/2} = 1.789,$$

adică nu sunt motive suficiente pentru respingerea ipotezei nule, concluzie deja obținută prin abordarea binomială.

Codul MATLAB corespunzător este:

```
x = [2158.70; 2165.20; 1678.15; 2399.55; 2316.00; 1779.80; 2061.30; 2336.75; 2207.50; 1765.30; 1708.30; 2053.50;
      1784.70; 2414.40; 2575.10; 2200.50; 2357.90; 2654.20; 2256.70; 1753.7];
m = 2000;
[p, h] = signtest(x, m)
[p, h, stats] = signtest(x, m)
```

Obținem:

$$\begin{array}{ccc} p & = & h \\ 0.1153 & & 0 \end{array}$$

și

$$\begin{array}{ccc} p & = & h & = & stats & = \\ 0.1153 & & 0 & & & zval : NaN \\ & & & & & sign : 14 \end{array}$$

2 Testul semnelor pentru date perechi

Vom numi date perechi un set de date bivariate (date ce conțin două valori, adică de forma $(x_i, y_i)_{i=1}^n$ ce reprezintă observatii asupra aceleiași caracteristici, între cele două componente existând măcar o legătură. Pentru aceste seturi de valori, ipoteza de independență între seturile de valori $(x_i)_{i=1}^n$ și $(y_i)_{i=1}^n$ nu mai este satisfăcută.

Exemple:

- masele corporale ale unor persoane înainte și după o anumită dietă (se dorește a studia efectul dietei asupra masei corporale);
- notele elevilor la testarea inițială la Matematică și notele aceluiași elevi la teza de Matematică (se urmărește testarea progresului făcut de elevi într-un semestru);
- starea sănătății unor bolnavi înainte și după administrarea unui tratament (se urmărește testarea eficienței tratamentului);
- salariile individuale pentru un număr de perechi soț - soție (se urmărește testarea diferențelor salariale între soți).

Considerăm X și Y două variabile dependente între ele. Pentru a compara mediile celor două variabile nu se poate aplica *testul t* pentru diferența mediilor, deoarece ipoteza de independență dintre X și Y este una de bază pentru aplicabilitatea *testului t*. Vom vedea mai târziu (vezi *testul t* pentru date perechi) cum putem testa dacă mediile sunt egale.

Deocamdată, să ne îndreptăm atenția asupra medianelor variabilelor.

Presupunem că $(x_1, y_1), \dots, (x_n, y_n)$ sunt datele perechi observate asupra variabilelor (X, Y) . În multe aplicații se dorește a se determina cum este X față de Y . Pentru aceasta, se consideră diferențele $d_i = x_i - y_i$.

Condițiile testului: Se presupune că $d_1, \dots, d_n$ sunt independente și provin dintr-o populație continuă, de mediană Me .

Ipoteze:

$$(H_0) : Me = 0, \quad (\text{diferențele dintre valorile perechi au mediana 0})$$

$$(H_1) : Me \neq 0.$$

Se pot considera și teste unilaterale, dacă $(H_1)_s : Me < 0$ sau $(H_1)_d : Me > 0$.

Ipotezele de mai sus pot fi testate folosind testul semnelor descris anterior, dar acest test nu verifică dacă medianele celor două selecții, Me_X și Me_Y , sunt egale.

Exemplul 2.1 Un dezvoltator auto studiază două dispozitive de măsurare pentru un sistem de injecție, pentru a determina dacă ele diferă în performanța medie. Sistemele, instalate pe 12 autoturisme, furnizează datele din următorul tabel. Utilizăm testul semnelor pentru date perechi pentru a determina dacă consumul este aproximativ egal în medie, pentru un nivel de semnificație 0.05.

Mașina	1	2	3	4	5	6	7	8	9	10	11	12
Sistem 1 (x_i)	17.6	19.4	19.5	17.1	15.3	15.9	16.3	18.4	17.3	19.1	17.8	18.2
Sistem 2 (y_i)	16.8	20.0	18.2	16.4	16.0	15.4	16.5	18.0	16.4	20.1	16.7	17.9
Diferențele d_i	0.8	-0.6	1.3	0.7	-0.7	0.5	-0.2	0.4	0.9	-1.0	1.1	0.3
Semn	+	-	+	+	-	+	-	+	+	-	+	+

Ipoteze formulate sunt:

$$(H_0) : Me = 0, \quad \text{versus} \quad (H_1) : Me \neq 0,$$

pentru $\alpha = 0.05$. Statistica folosită este $S^* = S_{\neq} = \max\{S_<, S_>\}$. Tabelul arată că $S_< = 4$, $S_> = 8$, deci $S_{\neq} = 8$. Prin urmare, probabilitatea critică este

$$P_v = 2 \cdot \mathbb{P}(S \geq 8) = 2 \sum_{k=8}^{12} C_{12}^k \frac{1}{2^{12}} > 0.05 = \alpha,$$

deci nu putem respinge ipoteza nulă.

Codul MATLAB corespunzător este:

$$\left| \begin{array}{l} x = [17.6; 19.4; 19.5; 17.1; 15.3; 15.9; 16.3; 18.4; 17.3; 19.1; 17.8; 18]; \\ y = [16.8; 20.0; 18.2; 16.4; 16.0; 15.4; 16.5; 18.0; 16.4; 20.1; 16.7; 17]; \\ [p, h, stats] = signtest(x, y) \end{array} \right|$$

Obținem:

$$\begin{array}{lll} p = & h = & stats = \\ & 0.3877 & 0 \quad \text{zval : NaN} \\ & & \text{sign : 8} \end{array}$$

Dacă caracteristica studiată este repartizată normal, atunci pentru testarea medianei se poate utiliza fie testul de semn fie testul t . Testul t are cea mai mică valoare pentru erori de tipul II, printre toate testele unilaterale cu un nivel de semnificație prestabilit sau printre teste bilaterale cu regiuni critice simetrice. Prin urmare, este superior testului semnelor în cazul caracteristicilor normale. Dacă populația este simetrică și ne-Gaussiană, dar cu medie finită, testul t va avea o eroare β de tipul II mai mică decât testul semnelor (deci o putere $\pi = 1 - \beta$ mai mare). De aceea testul semnelor este considerat mai curând o procedură pentru testarea valorii mediane, decât un test statistic veritabil. Testul Wilcoxon bazat pe ranguri cu semn va fi de preferat și dă rezultate bune în comparație cu testul t pentru caracteristici ce au repartiții simetrice.

3 Testul Wilcoxon bazat pe ranguri cu semn (Signed-Rank Test)

Testează valoarea centrală a unui set de date. Este folosit ca o alternativă pentru testul t pentru medie când ipotezele acestuia nu sunt verificate. Astfel, testul *signed rank* al lui Wilcoxon este utilizat pentru verificarea dacă

un set de date provine dintr-o **distribuție continuă, simetrică**, de o anumită medie (deci și mediană), în cazul în care datele observate nu urmează neapărat o repartiție Gaussiană.

Condițiile testului: Datele $x_1, \dots, x_n$ sunt observații aleatoare și independente asupra unei caracteristici continue X a unei populații, de mediană Me .

Ipoteze statistice:

Teste unilaterale	Test bilateral
$(H_0) : Me = Me^*,$	$(H_0) : Me = Me^*,$
$(H_1)_s : Me < Me^* \text{ (sau } (H_1)_d : Me > Me^* \text{)}.$	$(H_1) : Me \neq Me^*.$

Pentru a efectua testul, procedăm astfel: dacă admitem ipoteza nulă, atunci $Me = Me^*$. Ordonăm următoarele valori în ordine crescătoare: $|x_1 - Me^*|, \dots, |x_n - Me^*|$. Determinăm rangurile asociate acestor valori, iar statistica test va fi

$$S_+ = \sum_{i=1}^n \text{rang}(|x_i - Me^*| : x_1 - Me^* > 0),$$

iar valoarea statisticii, evaluată în datele empirice o notăm cu s_+ .

Regiunile care duc la respingerea ipotezei nule sunt:

$$s_+ \geq c_1, \text{ pentru testul unilateral dreapta,} \quad s_+ \leq \frac{n(n+1)}{2} - c \text{ sau } s_+ \geq c,$$

$$s_+ \leq c_2 = \frac{n(n+1)}{2} - c_1, \text{ pentru testul unilateral stânga.} \quad \text{pentru testul bilateral.}$$

Valorile critice c, c_1 și c_2 sunt date tabelate pentru testele Wilcoxon bilaterale și unilaterale, cu diverse valori pentru pragul de semnificație α . Aceste valori critice verifică

$$\mathbb{P}(S_+ \geq c_1) \simeq \alpha \quad \text{și} \quad \mathbb{P}(S_+ \geq c) \simeq \alpha/2, \quad \text{atunci când ipoteza } (H_0) \text{ este adevărată.}$$

Exemplul 3.1 Un producător de cereale ambalate dorește să verifice dacă un utilaj funcționează corespunzător. Acesta trebuie să umple pungi cu o cantitate medie de 460g. Pentru o selecție aleatoare de 15 pungi, gramajele măsurate sunt:

454.4 470.8 447.5 453.2 462.6 445.0 455.9 458.2 461.6 457.3 452.0 464.3 459.2 453.5 465.8

Se presupune că abaterile de la valoarea mediană pot fi în egală măsură la dreapta sau la stânga, datorită simetriei distribuției. Formulăm ipotezele:

$$(H_0) : Me = 460, \quad \text{versus} \quad (H_1) : Me \neq 460.$$

Tabloul datelor, pregătit pentru utilizarea testului Wilcoxon este:

Magnitudine	0.8	1.6	1.8	2.6	2.7	4.1	4.3	5.6	5.8	6.5	6.8	8.0	10.8	12.5	15.0
Rang	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
Semn	-	+	-	+	-	-	+	-	+	-	-	-	+	-	-

Avem $s_+ = 2 + 4 + 7 + 9 + 13 = 35$. Prin urmare, valoarea critică tabelată c este 25, iar $n(n+1)/2 - c = 95$ și

$$\mathbb{P}(S_+ \geq 95) = \mathbb{P}(S_+ \leq 25) = 0.024, \quad \text{atunci când ipoteza } (H_0) \text{ este adevărată.}$$

Deci, pentru nivelul de semnificație $\alpha = 0.05$, regiunea de respingere este $(25, 95)^c$. Cum $s_+ = 35$, acceptăm ipoteza nulă.

Exemplul 3.2 Codul MATLAB corespunzător este:

```
x = [454.4; 470.8; 447.5; 453.2; 462.6; 445.0; 455.9; 458.2; 461.6; 457.3; 452.0; 464.3; 459.2; 453.5; 465.8];
m = 460;
[p, h, stats] = signrank(x, m, talphat, 0.05, tmethod, texact)
```

Obținem:

$$p = 0.1688 \quad h = 0 \quad stats = \text{signedrank} : 35$$

4 Testul t pentru date perechi

Acesta este un test parametric. Îl menționăm aici doar pentru a face diferența între acest test și alte teste neparametrice ce pot fi utilizate pentru datele perechi. Testul poate fi aplicat pentru perechi de date pentru care diferențele între valorile perechi sunt normale.

Testele parametrice arată cum putem testa dacă mediile a două variabile independente X și Y sunt egale pe baza observațiilor făcute asupra acestor variabile, $\{x_i\}_{i=1}^m$ și $\{y_j\}_{j=1}^n$, unde m și n nu sunt neapărat egale. Există însă situații în care variabilele X și Y nu sunt independente între ele. Spre exemplu, observațiile făcute asupra aceluiași grup de indivizi înainte și după un tratament. În astfel de situații, testul t pentru diferența mediilor studiat anterior nu se mai poate aplica. Presupunem că X și Y sunt două variabile (posibil corelate) și că $(x_1, y_1), \dots, (x_n, y_n)$ sunt datele perechi observate. Notăm mediile teoretice ale acestor variabile prin: $\mu_X = \mathbb{E}(X)$ și $\mu_Y = \mathbb{E}(Y)$. În multe aplicații se dorește a se determina cum este X față de Y . Pentru fiecare pereche, considerăm $d_i = x_i - y_i$. Presupunem că variabilele corespunzătoare diferențelor, $\{D_i\}_{i=1}^n$ sunt normale, de medie μ_D și deviație standard σ_D . Evident, avem că $\mu_D = \mu_X - \mu_Y$, însă σ_D nu mai este neapărat egal cu $\sigma_X^2 + \sigma_Y^2$, egalitatea având loc doar în cazul independenței dintre variabilele X și Y .

Condițiile testului: diferențele d_i sunt aleatoare și repartiția din care au provenit este una normală.

Ipoteze statistice:

Teste unilaterale:

$$(H_0) : \mu_D = \mu_0,$$

$$(H_1)_s : \mu_D < \mu_0 \text{ (sau } (H_1)_d : \mu_D > \mu_0).$$

Test bilateral:

$$(H_0) : \mu_D = \mu_0,$$

$$(H_1) : \mu_D \neq \mu_0.$$

Pentru setul de date $\{d_i\}_{i=1}^n$, notăm cu

$$\bar{d} = \frac{1}{n} \sum_{i=1}^n d_i \quad \text{și} \quad s_D = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (d_i - \bar{d})^2}.$$

Statistica test este

$$t = \frac{\bar{d} - \mu_0}{s_D / \sqrt{n}}.$$

Regiunile care duc la respingerea ipotezei nule sunt:

$$t \leq -t_{1-\alpha, n-1} \text{ pentru testul unilateral stânga,}$$

$$|t| \geq t_{\frac{\alpha}{2}, n-1} \text{ pentru testul bilateral.}$$

$$t \geq t_{1-\alpha, n-1} \text{ pentru testul unilateral dreapta.}$$

De asemenea, testul poate fi efectuat pe baza unei valori P_v , care poate fi calculată în fiecare caz.

5 Testul Wilcoxon pentru date perechi

Este varianta neparametrică a testului anterior. Acest test este utilizat când ipoteza de normalitate a diferențelor nu este verificată.

Condițiile testului: repartiția diferențelor d_i este una continuă și simetrică. În cazul în care observațiile pentru X și Y sunt continue și diferă doar prin valorile medii, atunci repartiția diferențelor va fi continuă și simetrică. Nu este necesar ca repartițiile lui X și Y să fie simetrice.

Acest test verifică ipoteza nulă că valoarea mediană $Me_D = Me_X - Me_Y$ a diferențelor este una dată.

Ipoteze statistice:

Teste unilaterale

$$(H_0) : Me_D = Me^*,$$

$$(H_1)_s : Me_D < Me^* \text{ (sau } (H_1)_d : Me_D > Me^*).$$

Test bilateral

$$(H_0) : Me_D = Me^*,$$

$$(H_1) : Me_D \neq Me^*.$$

Pentru a testa această ipoteza pentru mediana Me_D se continuă cu etapele testului Wilcoxon Signed-Rank Test discutat anterior.

Exemplul 5.1 Acum aproximativ 100 de ani s-a făcut un experiment pentru a vedea dacă medicamentele ar putea ajuta insomnie severă (*The Action of Optical Isomers, II: Hyoscines, J. Physiol., 1905:501–510*). Au fost selectați 10 pacienți care

au avut probleme cu somnul, și fiecare pacient a încercat mai multe medicamente. Aici vom compara grupul de control (fără medicație) și cel ce a luat levo-hyoscină. Oferă medicația o îmbunătățire în timpul mediu de somn? Ipotezele formulate sunt:

$$(H_0) : Me_D = 0, \quad \text{versus} \quad (H_1) : Me_D < 0.$$

Tabelul datelor este următorul:

Pacient	1	2	3	4	5	6	7	8	9	10
Control $(x_i)_i$	0.6	1.1	2.5	2.8	2.9	3.0	3.2	4.7	5.5	6.2
Medicament $(y_i)_i$	2.5	5.7	8.0	4.4	6.3	3.8	7.6	5.8	5.6	6.1
Diferența	-1.9	-4.6	-5.5	-1.6	-3.4	-0.8	-4.4	-1.1	-0.1	0.1
Rangul cu semn	-6	-9	-10	-5	-7	-3	-8	-4	-1.5	1.5

Deoarece se observă o egalitate la ultimele două poziții, cele două ranguri inferioare primesc ca valoare media aritmetică a rangurilor 1 și 2. Pentru un nivel de semnificație $\alpha = 0.05$, ipoteza nulă va fi respinsă dacă valoarea măsurată a statisticii este $s_+ \leq 10 \cdot 11/2 - 44 = 11$. Cum $s_+ = 1.5$, aceasta se găsește în regiunea critică. Prin urmare, medicația oferă o durată de somn semnificativ mai mare în medie, deci acceptăm ipoteza alternativă.

Codul MATLAB corespunzător este:

$$\begin{cases} x = [0.6; 1.1; 2.5; 2.8; 2.9; 3.0; 3.2; 4.7; 5.5; 6.2]; \\ y = [2.5; 5.7; 8.0; 4.4; 6.3; 3.8; 7.6; 5.8; 5.6; 6.1]; \\ [p, h, stats] = \text{signrank}(x, y, \alpha, \text{method}, \text{exact}) \end{cases}$$

Obținem:

$$\begin{array}{lll} p = & h = & stats = \\ 0.0059 & 1 & \text{signedrank} : 1.5000 \end{array}$$

6 Testul Wilcoxon bazat pe suma rangurilor (Wilcoxon rank-sum test)

Acest test este varianta neparametrică a testului t pentru compararea mediilor. Este utilizat în cazul în care ipotezele testului t nu sunt satisfăcute (lipsa normalității a cel puțin unui set de date sau volumul datelor nu este suficient de mare). Acest test mai se regăsește sub denumirea de testul Mann-Whitney.

Presupunem că avem două seturi independente de date continue, ale căror valori observate în urma unui sondaj statistic sunt reprezentate de $\{x_i\}_{i=1}^m$ și $\{y_j\}_{j=1}^n$. Notăm cu Me_X și Me_Y medianele teoretice corespunzătoare repartițiilor din care provin aceste date. Se presupune că X și Y au aceeași distribuție, singura diferență posibilă fiind valorile lor medii. La nivelul de semnificație α se dorește a se testa ipoteza nulă:

Teste unilaterale	Test bilateral
$(H_0) : Me_D = Me^*,$	$(H_0) : Me_D = Me^*,$
$(H_1)_s : Me_D < Me^* \text{ (sau } (H_1)_d : Me_D > Me^* \text{)}.$	$(H_1) : Me_D \neq Me^*.$

Pentru a efectua testul, procedăm astfel: dacă admitem ipoteza nulă, atunci $Me_D = Me^*$. Presupunem că $m \leq n$ (dacă nu e adevărat, renotăm selecțiile). Ordonăm următoarele valori în ordine crescătoare:

$$x_1 - Me^*, \dots, x_m - Me^*, y_1 - Me^*, \dots, y_n - Me^*.$$

Statistica test va fi S^* = suma rangurilor asociate cu valorile $(x_i - Me^*)$ din șirul anterior. Cum, pentru orice întreg K , suma numerelor naturale până la el este $K(K+1)/2$, valoarea **minimală** a statisticii test este $m(m+1)/2$, valoare ce este atinsă atunci când toate cantitățile $x_i - Me^*$ sunt situate la stânga diferențelor $y_j - Me^*$. Similar, valoarea maximală posibilă pentru statistica S^* este atinsă atunci când diferențele $y_j - Me^*$ preced toate diferențele $x_i - Me^*$. În consecință, dacă notăm cu s^* valoarea, pentru datele măsurate a statisticii S^* , valoarea sa **maximă** devine

$$s^* = (n+1) + \dots + (n+m) = \frac{(m+n)(m+n+1)}{2} - \frac{n(n+1)}{2} = \frac{m(m+2n+1)}{2}$$

Repartiția statisticii test este simetrică față de mijlocul intervalului dat de minimul și de maximul valorilor sale posibile, iar valoarea de simetrie va fi $(m(m+2n+1) + m(m+1))/4 = m(m+n+1)/2$. Valoarea critică superioară pentru testul statistic va putea fi deci obținută prin intermediul valorii critice inferioare.

Regulile care duc la respingerea ipotezei nule sunt:

$$s^* \geq c_1, \text{ pentru testul unilateral dreapta,} \quad s^* \geq c \text{ sau } s^* \leq m(m+n+1) - c, \\ s^* \leq c_2 = m(m+n+1)/2 - c_1, \text{ pentru testul unilateral stânga.} \quad \text{pentru testul bilateral,}$$

unde c și c_1 sunt date în tabele. Avem

$$\mathbb{P}(S^* \geq c_1) \simeq \alpha, \quad \mathbb{P}(S^* \geq c) \simeq \alpha/2, \quad \text{atunci când ipoteza } (H_0) \text{ este adevărată.}$$

Cum statistica de test are o repartiție discretă, simetrică, este posibil ca, în general, să nu existe o valoare critică care să corespundă cu exactitate pragul dorit de semnificație α . Tabelele furnizează informații pentru $\alpha \in \{0.05, 0.025, 0.01, 0.005\}$ și $3 \leq m \leq n \leq 8$. Pentru valori mai mari ale numărului eșantioanelor trebuie utilizată o aproximare prin repartiția normală.

Exemplul 6.1 Concentrația de fluor (părți per milion) a fost măsurată la un eșantion de animale ce a păscut într-o zonă de pășunat expusă anterior la poluarea cu fluor cât și pentru un eșantion de animale ce a păscut într-o regiune nepoluată. Datele obținute sunt cuprinse în tabelul următor.

Poluat	21.3	18.7	23.0	17.1	16.8	20.9	19.7
Nepoluat	14.2	18.3	17.2	18.4	20.0		

Datele arată că, la un nivel de semnificație $\alpha = 0.01$, valoarea medie a concentrației de fluoriide este consistent mai mare în zona poluată față de cealaltă?

Tabelul ordonat al datelor empirice este:

x	y	y	x	x	x	y	y	x	y	y	y
14.2	16.8	17.1	17.2	18.3	18.4	18.7	19.7	20.0	20.09	21.3	23.0
1	2	3	4	5	6	7	8	9	10	11	12

Avem $m = 5$ și $n = 7$, $\mathbb{P}(S^* \geq 47 : (H_0) \text{ este adevărată}) \simeq 0.01$. Valoarea critică (inferioară) pentru ipoteza alternativă stânga este $c_2 = 5 \cdot (5 + 7 + 1) / 2 - 47 = 18$. Ipoteza nulă va fi respinsă pentru $s_* \leq 18$. În exemplul prezentat, $s_* = 1 + 5 + 4 + 6 + 9 = 25$, deci nu avem motive să respingem ipoteza nulă pentru un prag de semnificație $\alpha = 0.01$.

Codul MATLAB corespunzător este:

```
x = [0.6; 1.1; 2.5; 2.8; 2.9; 3.0; 3.2; 4.7; 5.5; 6.2];
y = [2.5; 5.7; 8.0; 4.4; 6.3; 3.8; 7.6; 5.8; 5.6; 6.1];
[p, h, stats] = signrank(x, y, talpha, 0.05, tmethod, texact)
```

Obținem:

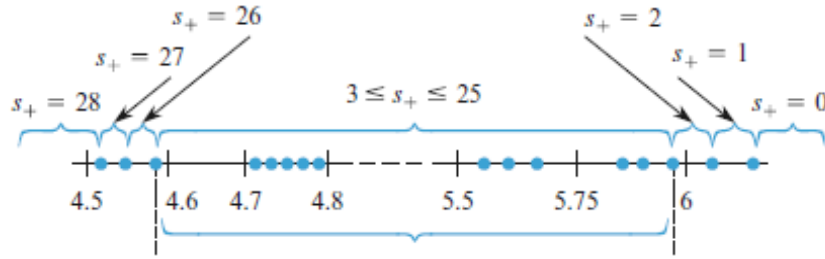
```
p = 0.2677      h = 0      stats = ranksum : 25
```

6.1 Alte teste de tip Wilcoxon

Metoda utilizată de regulă pentru determinarea intervalelor de încredere având ca scop estimarea parametru-lui unui repartiții presupune utilizarea unei statistici (Z, t, χ^2, F) care depinde de parametru și evaluarea unei inegalități de tip probabilistic ce oferă capetele aleatoare ale intervalului ce va acoperi parametrul investigat. O metodă alternativă constă în folosirea legăturii dintre testele statistice și intervalele de încredere. Un interval de încredere cu un nivel de încredere de $100(1 - \alpha)\%$ pentru un parametru λ poate fi obținut prin intermediul unui test statistic cu grad de semnificație α , la care formulăm ipotezele $(H_0) : \lambda = \lambda_0$ versus $(H_1) : \lambda \neq \lambda_0$. Instrumentele folosite ulterior vor fi cele două teste Wilcoxon prezentate deja.

Testul Wilcoxon bazat pe ranguri cu semn pentru intervale

Pentru valorile empirice observate în urma unui sondaj de volum n , $x_1, \dots, x_n$, un interval de încredere având la baza rangurile cu semn și $100(1 - \alpha)\%$ nivel de încredere, este format din toate valorile Me^* pentru care



ipoteza nulă (H_0) : $Me = Me^*$ nu este respinsă pentru pragul de semnificație α . Pentru aceasta, este suficient să exprimăm statista pentru test sub o altă formă:

$$S_+ = \# \left\{ (i, j) : i \leq j, \frac{X_i + X_j}{2} \geq Me^* \right\},$$

cu valoarea măsurată $s_+ = \# \{ (i, j) : i \leq j, (x_i + x_j)/2 \geq Me^* \}$. Echivalența celor două metode pentru calcularea lui s_+ este ușor de justificat. Numărul mediilor obținute este $C_n^2 + n$ (primul termen fiind dat de perechile cu elemente distincte, iar al doilea apare atunci când facem media fiecărei valori empirice cu ea însăși), cantitate egală cu $n(n+1)/2$. Dacă prea multe sau prea puține valori medii sunt mai mari sau egale cu Me^* , atunci respingem ipoteza nulă.

Exemplul 6.2 Observațiile următoare reprezintă ratele metabolismului cerebral pentru 7 indivizi ai unei populații: $x_1 = 4.51, x_2 = 4.59, x_3 = 4.90, x_4 = 4.93, x_5 = 6.80, x_6 = 5.08, x_7 = 5.67$. Cele 28 de medii obținute în urma realizării perechilor, sunt, în ordine crescătoare:

4.51	4.55	4.59	4.705	4.72	4.745	4.76	4.795	4.835	4.90
4.915	4.93	4.99	5.005	5.08	5.09	5.13	5.285	5.30	5.375
5.655	5.67	5.695	5.85	5.865	5.94	6.235	6.80		

La pragul 0.0469, (H_0) este acceptată. Datorită caracterului discret al distribuției statisticii S_+ , $\alpha = 0.05$ nu poate fi atins cu exactitate. Regiunea de respingere $\{0, 1, 2, 26, 27, 28\}$ are cea mai apropiată valoare în vecinătatea lui α pe 0.046. Deci, pentru un număr de medii cuprins între 3 și 25 inclusiv, acceptăm ipoteza nulă.

În general, odată ordonate perechile în mod crescător, capetele intervalului Wilcoxon sunt două dintre mediile "extreme". Notăm cea mai mică medie a perechilor cu $\bar{x}_{(1)}, \dots$, iar cea mai mare cu $\bar{x}_{(n(n+1)/2)}$.

Propoziția 6.1 Dacă testul Wilcoxon de semnificație α , bazat pe ranguri cu semn are (H_0) : $Me = Me^*$ versus (H_1) : $Me \neq Me^*$, atunci un interval de încredere cu $100(1 - \alpha)\%$ nivel de încredere pentru Me este

$$(\bar{x}_{(n(n+1)/2-c+1)}, \bar{x}_{(c)}) .$$

Tabelele oferă valorile lui c pentru $n \in \{5, 6, \dots, 25\}$. Pentru volume de selecție mai mare, se folosește ca și statistică de test standardizarea lui S_+ . Valoarea aproximantă a punctului critic c va fi acum

$$c_{approx} = \frac{n(n+1)}{4} + z_{\alpha/2} \sqrt{\frac{n(n+1)(2n+1)}{24}} .$$

Eficiența intervalului Wilcoxon în raport cu intervalul obținut prin testul t este aproximativ aceeași cu eficiența testului Wilcoxon în raport cu testul t . În particular, pentru volume mari de selecție, atunci când populația este repartizată normal, intervalul Wilcoxon va avea tendința de a fi puțin mai lung decât intervalul t . În cazul în care populația este non-Gaussiană (dar cu repartiție simetrică), atunci intervalul Wilcoxon va tinde să fie mult mai scurt decât intervalul obținut prin testul t .

Testul Wilcoxon bazat pe suma rangurilor pentru intervale

Pentru a obține intervalul de încredere asociat valorilor empirice x_i, y_j , cu $1 \leq i \leq m, 1 \leq j \leq n$, exprimăm, din nou statistica de test sub o nouă formă. Cea mai mică valoare pentru S^* este $m(m+1)/2$ și sunt mn diferențe de forma $(X_i - Me^*) - Y_j$. Obținem

$$S^* = \# \{ (i, j) : X_i - Y_j \geq Me^* \} + \frac{m(m+1)}{2} .$$

Neacceptarea ipotezei nule (H_0) : $Me_D = Me^*$ este echivalentă cu acceptarea ipotezei alternative dacă cantitatea $s^* = \# \{ (i, j) : x_i - y_j \geq Me^* \} + m(m+1)/2$ este prea mică sau prea mare. Similar testului precedent, avem următorul rezultat.

Propoziția 6.2 Fie $x_1, x_2, \dots, x_m$ și $y_1, y_2, \dots, y_n$ valori observate pentru două caracteristici continue, ce nu diferă prin forma repartiției. Notând $d_{ij} = x_i - y_j$, pentru toți i, j , ordonăm diferențele valorilor observate, pentru a forma șirul crescător $d_{ij(1)}, d_{ij(2)}, \dots, d_{ij(mn)}$. Intervalul de încredere, cu un nivel de încredere $100(1 - \alpha)\%$ pentru Me_D este

$$(d_{ij(mn-c+1)}, d_{ij(c)}),$$

unde c este valoarea critică a testul Wilcoxon bazat pe suma rangurilor, corespunzătoare unui prag de semnificație α .

Pentru volume de selecție mai mare, se folosește ca și statistică de test standardizarea lui S^* . Valoarea aproximantă a punctului critic c va fi acum

$$c_{approx} = \frac{mn}{2} + z_{\alpha/2} \sqrt{\frac{mn(m+n+1)}{12}},$$

unde $z_{\alpha/2}$ este cuantila de ordin $\alpha/2$ pentru repartiția normală.

7 Testul seriilor pentru caracterul aleator

Testul seriilor (en., *runs test*) este un test neparametric ce verifică ipoteza ca un șir de date bivariante este aleator generat (i.e., datele statistice constituie o selecție aleatoare dintr-un șir infinit de valori).

Condițiile testului: Datele bivariante sunt independente.

Ipotezele testului:

(H_0) : valorile observate provin dintr-un șir aleator,

(H_1) : valorile observate nu provin dintr-un șir aleator.

Dacă o anumită valoare a unui anumit șir de caractere este influențată de poziția sa sau de valorile ce o preced, atunci selecția generată nu poate fi aleatoare.

Definim noțiunea de serie sau faza (en., *run*) ca fiind o succesiune a unui sau mai multe simboluri de același tip, care sunt precedate și urmate de simboluri de alt tip sau niciun simbol. De exemplu:

001111010010 sau MFFFFMFFFFM sau ++--++--++--

Numărul de faze și lungimea lor pot fi folosite în determinarea gradului de stochasticitate a unui șir de simboluri. Prea puține sau prea multe faze, sau de lungimi excesiv de mari sunt rare în serii cu adevărat aleatoare, de aceea ele pot servi drept criterii statistice pentru testarea stochasticității. Aceste criterii sunt adiacente: prea puține faze implică faptul că unele faze sunt prea lungi (se observă o persistență), prea multe faze implică faptul că unele faze sunt prea scurte (se observă o secvență în zigzag). Așadar, ne vom preocupa doar de numărul total de faze.

Fie n_1 și n_2 numărul de simboluri 0, respectiv, 1 din șir. Numărul total de semne este $n = n_1 + n_2$. Fie r_1 și r_2 numărul de faze ce corespund simbolul 0, respectiv, 1 din șir. Numărul total de faze este $r = r_1 + r_2$. Vom nota cu litere mari, R_1, R_2 sau R , variabilele aleatoare ale căror realizări sunt r_1, r_2 , respectiv, r .

Să exemplificăm aceste noțiuni pentru șirurile considerate mai sus. Primul șir de $n = 12$ cifre este constituit din $r = 7$ faze: $r_1 = 4$ faze de 0 și $r_2 = 3$ faze de 1; $n_1 = n_2 = 6$. Al doilea șir este format din $r = 4$ faze, $r_1 = 2$ de M și $r_2 = 2$ de F , iar ultimul șir de $n = 19$ este constituit din $r = 10$ faze, câte $r_1 = r_2 = 5$ din fiecare $+$ sau $-$. Alegem ipoteza nulă:

(H_0) : șirul este aleator (fiecare aranjament de 0 și 1 este echiprobabil),

(H_1) : șirul nu este aleator.

Se pot considera și ipoteze alternative:

$(H_1)_s$: datele au o tendință de a se aduna chiorchine,

$(H_1)_d$: datele au o tendință de a se răsfirea.

Putem găsi repartițiile variabilelor aleatoare R_1, R_2 sau, mai reprezentativ pentru test, R .

Presupunem că există n_1 elemente de 0 și n_2 elemente de 1. Numărul de faze este dat de variabila aleatoare R , a cărei distribuție discretă o vom determina. Pentru început, dat $r \in \mathbb{N}$,

$$\mathbb{P}(R = r) = \frac{\# \text{ permutărilor ce produc } r \text{ faze}}{\# \text{ permutărilor obținute cu } n_1 \text{ de } 0 \text{ și } n_2 \text{ de } 1} = \frac{\# \text{ permutărilor ce produc } r \text{ faze}}{C_{n_1+n_2}^{n_1}}.$$

În ceea ce privește numărătorul, procedăm astfel. Putem ignora cazul $r = 1$ deoarece acesta are loc dacă $n_1 = 0$ sau $n_2 = 0$.

Presupunem, pentru început că numărul de faze r este **par**, având forma $r = 2m$, $m \in \mathbb{N}$. Considerăm exemplul:

$$00 \ 1111 \ 0 \ 1 \ 00 \ 1$$

Avem

$$\begin{cases} n_1 = n_1(0) = 5, & n_2 = n_2(1) = 6, & n = n_1 + n_2 = 11 \\ r_1 = r_1(0) = 3, & r_2 = r_2(1) = 3, & r = r_1 + r_2 = 6 = 2m, \quad m = 3. \end{cases}$$

- Dacă primul element din secvență este 0, determinăm numărul de secvențe ce conțin exact r ($= 6$) faze. Cum secvența începe cu 0, ea trebuie să se încheie cu 1, pentru ca r să fie par. După cum vedem în exemplul considerat, avem $m = r_1 = 3$ faze de 0 și $m = r_2 = 3$ faze de 1.
- Pentru început, determinăm numărul de posibilități în care putem împărți un șir de n_1 ($= 5$) de 0 în r_1 ($= m = 3$) faze distincte. Ori, aceasta se poate realiza prin plasarea între zerouri a câte unui separator de 1 în $r_1 - 1$ ($= m - 1 = 2$) poziții. Aceasta se poate realiza în $C_{n_1-1}^{m-1} = C_{n_1-1}^{r_1-1}$ moduri.
- Similar, determinăm numărul de posibilități în care putem obține r_2 ($= m = 3$) faze (grupe distincte) având la dispoziție cele n_2 ($= 6$) valori de 1. Acest număr de posibilități este $C_{n_2-1}^{m-1} = C_{n_2-1}^{r_2-1}$.
- Prin urmare, numărul de aranjări distincte ce pornesc cu primul element 0, având fazele intercalate în modul solicitat, va fi deci

$$C_{n_1-1}^{r_1-1} C_{n_2-1}^{r_2-1} = C_{n_1-1}^{m-1} C_{n_2-1}^{m-1} = C_{n_1-1}^{r/2-1} C_{n_2-1}^{r/2-1}.$$

- Analog, în cazul în care r este par, iar primul element din șir este 1, atunci ultimul trebuie să fie neapărat 0. Repetăm raționamentul, dar rolul lui 0 îl joacă acum 1 și reciproc. Obținem, exact ca în etapele anterioare că există tot $C_{n_1-1}^{r/2-1} C_{n_2-1}^{r/2-1}$ posibilități de obținere a fazelor, adică, adunând cu valoarea obținută anterior, avem un total de $2C_{n_1-1}^{m-1} C_{n_2-1}^{m-1}$ moduri de a forma cele $r = 2m = m + m = r_1 + r_2$ faze.
- Concluzionând, dacă r este par,

$$\mathbb{P}(R = r) = \frac{2C_{n_1-1}^{r/2-1} C_{n_2-1}^{r/2-1}}{C_n^m}. \quad (1)$$

Presupunem acum că numărul de faze r este **impar**, având forma $r = 2m + 1$, $m \in \mathbb{N}$. Considerăm exemplul:

$$00 \ 1111 \ 0 \ 11 \ 00 \ 1 \ 0$$

Avem

$$\begin{cases} n_1 = n_1(0) = 6, & n_2 = n_2(1) = 7, & n = n_1 + n_2 = 13 \\ r_1 = r_1(0) = 4, & r_2 = r_2(1) = 3, & r = r_1 + r_2 = 7 = 2m + 1, \quad m = 3. \end{cases}$$

Este evident că, și în acest caz, intercalarea sevențelor va implica faptul că $r_1 = r_2 \pm 1$. Dacă $r_1 = r_2 + 1$, atunci șirul va începe obligatoriu cu un 0, iar dacă $r_1 = r_2 - 1$, atunci primul element trebuie să fie neapărat 1. După cum am văzut în cazul precedent, dacă $r_1 = r_2$ (adică r este par) atunci șirul poate începe, în egală măsură, atât cu 0 cât și cu 1 și de aceea probabilitățile de la numărătorul formulei (1) se dublează.

Pentru cazul când r este impar, utilizăm aceleași raționamente pentru a număra modul de a obține fazele dorite. De data aceasta, secvențele trebuie să înceapă și să se termine cu același simbol. Dacă acesta este 0, cum trebuie inserat un 1 suplimentar cazului în care r era par, obținem $C_{n_1-1}^m$ secvențe posibile în care un 1 singular acționează ca un delimitator între fazele cu 0. Din nou, sunt $C_{n_2-1}^{m-1}$ moduri de a distribui pe 1 și $C_{n_1-1}^m C_{n_2-1}^{m-1}$ moduri de a obține r faze dacă primul element din secvență este 0. Similar, dacă secvența începe cu simbolul 1, găsim $C_{n_1-1}^{m-1} C_{n_2-1}^m$ moduri de a determina fazele. Prin urmare, dacă r este impar,

$$\mathbb{P}(R = r) = \frac{C_{n_1-1}^m C_{n_2-1}^{m-1} + C_{n_1-1}^{m-1} C_{n_2-1}^m}{C_n^m} = \frac{C_{n_1-1}^{(r-1)/2} C_{n_2-1}^{(r-3)/2} + C_{n_1-1}^{(r-3)/2} C_{n_2-1}^{(r-1)/2}}{C_n^m}. \quad (2)$$

Putem determina distribuția vectorului aleator (R_1, R_2) și obținem următorul rezultat (vezi Dickinson Gibbons, Chakraborti [6, Theorem 2.1]).

Teorema 1 Dacă R_1 și R_2 sunt variabilele aleatoare corespunzătoare numărului de faze r_1 , respectiv r_2 , secvența având n_1 simboluri de 0 și n_2 simboluri de 1, atunci repartiția vectorului aleator (R_1, R_2) este $f_{(R_1, R_2)} : \mathbb{R}^2 \rightarrow [0, 1]$,

$$f_{(R_1, R_2)}(r_1, r_2) = \frac{c \cdot C_{n_1-1}^{r_1-1} C_{n_2-1}^{r_2-1}}{C_{n_1+n_2}^{m_1}}, \quad r_1 = 1, 2, \dots, n_1, \quad r_2 = 1, 2, \dots, n_2, \quad r_1 = r_2 \text{ sau } r_1 = r_2 \pm 1,$$

unde $c = 2$ dacă $r_1 = r_2$ și $c = 1$ dacă $r_1 = r_2 \pm 1$.

Rezultatul permite obținerea repartițiilor marginale pentru variabilele aleatoare R_1 și R_2 .

Corolarul 7.1 Distribuția marginală a variabilei aleatoare R_1 este $f_{R_1} : \mathbb{R} \rightarrow [0, 1]$,

$$f_{R_1}(r_1) = \frac{C_{n_1-1}^{r_1-1} C_{n_2+1}^{r_1}}{C_{n_1+n_2}^{m_1}}, \quad r_1 = 1, 2, \dots, n_1.$$

Pentru variabila aleatoare R_2 obținem o formulă similară, prin interschimbarea valorilor n_1 și n_2 .

Demonstrație. Densitatea vectorului aleator (R_1, R_2) arată că $r_2 \in \{r_1, r_1 - 1, r_1 + 1\}$, pentru orice valoare a lui r_1 . Prin urmare,

$$f_{R_1}(r_1) = \sum_{r_2} f_{(R_1, R_2)}(r_1, r_2).$$

Putem deci obține, în urma unor calcule combinatorice,

$$\begin{aligned} C_{n_1+n_2}^{m_1} f_{R_1}(r_1) &= 2C_{n_1-1}^{r_1-1} C_{n_2-1}^{r_1-1} + C_{n_1-1}^{r_1-1} C_{n_2-1}^{r_1-2} + C_{n_1-1}^{r_1-1} C_{n_2-1}^{r_1} = C_{n_1-1}^{r_1-1} (C_{n_2-1}^{r_1-1} + C_{n_2-1}^{r_1-2} + C_{n_2-1}^{r_1-1} + C_{n_2-1}^{r_1}) \\ &= C_{n_1-1}^{r_1-1} (C_{n_2-1}^{r_1-1} + C_{n_2}^{r_1}) = C_{n_1-1}^{r_1-1} C_{n_2+1}^{r_1}, \end{aligned}$$

iar demonstrația este, astfel, încheiată.

Testul exact. Dacă secvența conține simbolurile 0 și 1, atunci numărul minim, respectiv maxim, de faze ce se pot forma sunt:

$$R_{\min} = 2, \quad \text{respectiv} \quad R_{\max} = 2 \min\{n_1, n_2\} + 1.$$

Având în vedere repartiția discretă a statisticii R , dată de formulele (1) și (2), pentru ipoteza alternativă la dreapta "prea multe faze" (datele au o tendință de a se răsfira) putem determina exact valoarea critică superioară pentru numărul total observat de faze:

$$\mathbb{P}(R \geq r) = \sum_{v=r}^{R_{\max}} \mathbb{P}(R = v).$$

Pentru ipoteza alternativă la stânga "prea puține faze" (datele au o tendință de a se aduna ciorchine) putem determina exact valoarea critică inferioară pentru numărul total observat de faze:

$$\mathbb{P}(R \leq r) = \sum_{v=R_{\min}}^r \mathbb{P}(R = v).$$

Pentru testul bilateral, valoarea critică este dată prin:

$$\mathbb{P}(|R - \mathbb{E}(R)| \geq |r - \mathbb{E}(R)|) = \sum_{v=\mathbb{E}(R)-|r-\mathbb{E}(R)|}^{\mathbb{E}(R)-|r-\mathbb{E}(R)|} \mathbb{P}(R = v) + \sum_{v=\mathbb{E}(R)+|r-\mathbb{E}(R)|}^{R_{\max}} \mathbb{P}(R = v).$$

Testul exact este, în mod evident, mai precis decât testul Z și trebuie folosit dacă este posibil. Însă acest test nu poate fi aplicat dacă secvența aleatoare conține mai mult de două tipuri de simboluri distincte.

Vom avea nevoie în cele ce urmează de momentele teoretice de ordin k ale variabilei aleatoare R (în special cele de ordinul 1 și 2). Formulele (1) și (2) permit să scriem formula generală a momentului de ordinul k al lui R :

$$\mathbb{E}(R^k) = \sum_r r^k \mathbb{P}(R = r) = \left\{ \sum_{r \text{ par}} 2r^k \frac{C_{n_1-1}^{r/2-1} C_{n_2-1}^{r/2-1}}{C_{n_1+n_2}^{m_1}} + \sum_{r \text{ impar}} r^k \frac{C_{n_1-1}^{(r-1)/2} C_{n_2-1}^{(r-3)/2} + C_{n_1-1}^{(r-3)/2} C_{n_2-1}^{(r-1)/2}}{C_{n_1+n_2}^{m_1}} \right\}.$$

Putem presupune, fără a restrânge generalitatea, că $n_1 \leq n_2$ și atunci $r \in \{2, 3, \dots, 2n_1 + 1\}$. Notăm $r = 2i$ dacă r este par și $r = 2i + 1$ dacă el este impar, caz în care $i \in \{1, 2, \dots, n_1\}$.

Media variabilei R devine:

$$C_{n_1}^{m_1} \mathbb{E}(R) = \sum_{i=1}^{n_1} 4i C_{n_1-1}^{i-1} C_{n_2-1}^{i-1} + \sum_{i=1}^{n_1} (2i+1) C_{n_1-1}^i C_{n_2-1}^{i-1} + \sum_{i=1}^{n_1} (2i+1) C_{n_1-1}^{i-1} C_{n_2-1}^i$$

Pentru evaluarea celor trei sume avem nevoie de următoarele două observații combinatorice, punct în care media devine un simplu exercițiu.

Lema 7.1 *Au loc relațiile:*

$$(a) \sum_{r=0}^c C_m^r C_n^r = C_{m+n}^m, \quad \text{unde } c = \min\{m, n\}.$$

$$(b) \sum_{r=0}^c C_m^r C_n^{r+1} = C_{m+n}^{m+1}, \quad \text{unde } c = \min\{m, n-1\}.$$

Demonstrație. (a) Pornim de la egalitatea evidentă

$$\sum_{i=0}^{m+n} C_{m+n}^i x^i = \sum_{j=0}^m C_m^j x^j \sum_{k=0}^n C_n^k x^k, \quad \text{pentru orice } x \in \mathbb{R}.$$

Presupunem, fără a restrânge generalitatea, că $c = m$ și, egalând coeficienții lui x^m din ambii membri obținem $C_{m+n}^m = \sum_{r=0}^m C_m^{m-r} C_n^r$, iar punctul (a) este demonstrat.

(b) Abordarea este similară punctului (a), egalând coeficienții lui x^{m+1} din ambii membri.

O metodă mai elegantă pentru abordarea momentelor teoretice ale lui R este sugerată de faptul că această variabilă aleatoare poate fi considerată ca fiind suma unor variabile aleatoare elementare. Fie

$$R = 1 + I_2 + I_3 + \dots + I_n,$$

unde am definit

$$I_k = \begin{cases} 1, & \text{dacă elementul de pe poziția } k \neq \text{elementul de pe poziția } k-1, \\ 0, & \text{în caz contrar.} \end{cases}$$

Este evident că, pentru orice k , $I_k \sim \mathcal{B}(1, n_1 n_2 / C_n^2)$ și obținem

$$\begin{cases} \mathbb{E}(I_k) = \mathbb{E}(I_k^2) = \frac{2n_1 n_2}{n(n-1)} & \text{și} & \mathbb{E}(R) = 1 + \sum_{k=2}^n \mathbb{E}(I_k) = 1 + \frac{2n_1 n_2}{n} \\ D^2(R) = D^2\left(\sum_{k=2}^n I_k\right) = (n-1) D^2(I_k) + \sum \sum_{2 \leq j \neq k \leq n} \text{cov}(I_j, I_k) \\ \quad = (n-1) \mathbb{E}(I_k^2) + \sum \sum_{2 \leq j \neq k \leq n} \mathbb{E}(I_j I_k) - (n-1)^2 (\mathbb{E}(I_k))^2 \end{cases}$$

Evaluarea celor $(n-1)(n-2)$ momente $\mathbb{E}(I_j I_k)$ se face astfel:

1. pentru cele $2(n-2)$ situații în care $j = k-1$ sau $j = k+1$:

$$\mathbb{E}(I_j I_k) = \frac{n_1 n_2 (n_1 - 1) + n_2 n_1 (n_2 - 1)}{n(n-1)(n-2)} = \frac{n_1 n_2}{n(n-1)};$$

2. pentru cele $(n-1)(n-1) - 2(n-2) = (n-2)(n-3)$ cazuri rămase, în care $j \neq k$, avem:

$$\mathbb{E}(I_j I_k) = \frac{4n_1 n_2 (n_1 - 1)(n_2 - 1)}{n(n-1)(n-2)(n-3)}$$

Introducând aceste medii în formula dispersiei lui R , obținem:

$$D^2(R) = \frac{2n_1 n_2}{n} + \frac{2(n-2)n_1 n_2}{n(n-1)} + \frac{4n_1 n_2 (n_1 - 1)(n_2 - 1)}{n(n-1)} - \frac{4n_1^2 n_2^2}{n^2} = \frac{2n_1 n_2 (2n_1 n_2 - n_1 - n_2)}{n^2 (n-1)}.$$

Pentru a trage o concluzie asupra modului în care se aplică acest test statistic distingem următoarele două situații.

Cazul I: Când n_1 și n_2 sunt mari (i.e. $n_1 > 12, n_2 > 12$), variabila aleatoare R , corespunzătoare valorii r are o repartiție aproape normală, $R \sim \mathcal{N}(\mu, \sigma^2)$, unde

$$\mu = 2 \frac{n_1 n_2}{n} + 1, \quad \sigma = \sqrt{\frac{2n_1 n_2 (2n_1 n_2 - n)}{n^2 (n-1)}} = \sqrt{\frac{(\mu - 1)(\mu - 2)}{n-1}}.$$

Demonstrarea acestui comportament asimptotic se poate studia în articolul [Wald, A.; Wolfowitz, J., *On a Test Whether Two Samples are from the Same Population*, Annals of Mathematical Statistics, Volume 11, Number 2 (1940), 147-162]. Așadar, statistica

$$Z = \frac{R - \mu}{\sigma} \sim \mathcal{N}(0, 1).$$

poate fi utilizată pentru testarea ipotezei nule (H_0). Pentru testul bilateral, ipoteza nulă este admisă dacă, pentru cuantila de ordin $1-\alpha/2$ a repartiției normale, avem $|(r - \mu) / \sigma| \leq z_{1-\alpha/2}$. Altfel, se respinge ipoteza nulă. Pentru testul unilateral, condiția de respingere a ipotezei nule este

$$\begin{array}{ll} \frac{r - \mu}{\sigma} \leq -z_{1-\alpha}, & \frac{r - \mu}{\sigma} \geq z_{1-\alpha}, \\ \text{pentru test unilateral stânga.} & \text{pentru test unilateral dreapta.} \end{array}$$

Acest test Z asimptotic nu numai că este mai puțin precis decât testul exact, dar este mai puțin precis și decât testul Z asimptotic, cu corelație pentru continuitate. Pentru acesta se utilizează statistica

$$Z_{cc} = \begin{cases} \frac{R - \mu - 0.5}{\sigma}, & \text{dacă } r \geq \mu \\ \frac{R - \mu + 0.5}{\sigma}, & \text{dacă } r < \mu. \end{cases}$$

Cazul II: Când n_1 și n_2 sunt mici ($n_1 \leq 12, n_2 \leq 12$), valorile critice pentru r sunt tabelate. Astfel, pentru testul bilateral, regiunea care asigură acceptarea ipotezei nule este $r_{\alpha/2, L} < r < r_{\alpha/2, U}$.

Pentru testul unilateral stânga, ipoteza nulă va fi respinsă dacă $r < r_{\alpha/2, L}$. Pentru testul unilateral dreapta, ipoteza nulă va fi respinsă dacă $r > r_{\alpha/2, U}$.

Testul pentru caracterul aleator al seriilor poate fi folosit în următoarele situații:

- testarea caracterului aleator a unei selecții de date, prin marcarea cu + a valorilor ce sunt mai mari decât mediana și cu - ale celor ce sunt mai mici decât mediana. Valorile egale cu mediana sunt omise și n este ajustat în consecință.
- testarea ipotezei că două eșantioane sunt observații independente ale aceleiași repartiții (testul Wald-Wolfowitz).
- testarea potrivirii unei funcții cu un set de date, prin marcarea cu + a valorilor ce sunt mai mari decât valoarea funcției și cu - ale celor ce sunt mai mici decât valoarea funcției. Valorile egale cu valoarea funcției sunt omise și n este ajustat în consecință. Acest test nu ține cont de distanțe dintre date și funcție, ci doar de semne, spre deosebire de un test χ^2 .

Exemplul 7.1 La imbarcarea animalelor în arcă, pentru a evita ruperea punții de urcare, Noe a vrut să vadă dacă șirul animalelor ce urca pe punte este distribuit aleatoriu din punct de vedere al maselor. În acest sens, a decis să marcheze cu - animalele ce au masa mai mică decât a lui și cu + pe cele cu masa superioară. La o secvență de 25 de animale, obține șirul binar:

+ + + - - + + + + - + - - - + + + + - - + + + - -

Pentru nivelul de semnificație considerat $\alpha = 0.05$, se poate accepta ipoteza nulă care afirmă că deviația masei de la medie (considerată cea a lui Noe) este aleatoare (se acceptă ipoteza nulă (H_0)) sau nu (se respinge ipoteza nulă)?

În exemplul considerat, $r = 10$, $n_1 = 15$, $n_2 = 10$, $r_1 = 5$, $r_2 = 5$. Valoarea critică inferioară este $r_{\alpha/2, L} = 11$, iar valoarea critică superioară este $r_{\alpha/2, U} = 21$. Cum $r \leq 11$, respingem ipoteza nulă.

8 Testul Wald-Wolfowitz (Wald-Wolfowitz two-sample runs test)

Testul Wald-Wolfowitz este o alternativă neparametrică a testului t pentru selecții independente. Este utilizat în testarea ipotezei că două eșantioane sunt observații independente ale aceleiași repartiții. Reamintim, testul t pentru două selecții decide dacă două selecții independente provin din două caracteristici ce au aceeași medie. Testul Wald-Wolfowitz poate depista chiar mai multe diferențe dintre cele două repartiții decât poate depista testul t pentru două selecții. Spre exemplu, testul W-W poate depista diferențele dintre mediile sau dintre formele caracteristicilor din care provin cele două seturi de observații. Este eficient pentru un volum al selecției cel puțin moderat, e.g. cel puțin egal cu 10.

Condițiile testului: Datele observate sunt observații aleatoare ale unor caracteristici continue independente. Presupunem că avem două seturi de date, $\{x_i\}_{i=1}^m$ și $\{y_j\}_{j=1}^n$.

Ipotezele testului:

- (H_0) : Cele două seturi de date provin din aceeași repartiție,
 (H_1) : Cele două seturi de date provin din repartiții diferite.

Pentru a testa ipoteza nulă, datele observate se vor scrie împreună, în ordine crescătoare, fiecare observație fiind codată cu 1 sau 2, după cum provine din setul 1 sau 2 de date. Testul Wald-Wolfowitz are la bază ipoteza nulă ca fiecare valoare observată din șirul combinat este extrasă independent dintr-o aceeași repartiție dată. Statistica test este $r =$ numărul de faze (*runs*) observate în șirul obținut prin alipire. Dacă această statistică ar avea o valoare numerică mică, atunci acest fapt indică un anumit trend în datele alipite (datele ce provin din același set tind să se adune în clustere), adică puțin improbabil ca aceste date să fi provenit din aceeași repartiție. Pe de altă parte, un număr mare pentru r este un indiciu că datele sunt observații aleatoare ale unei repartiții, fapt care va duce la acceptarea ipotezei nule. În cazul în care valori ale șirului x coincid cu valori ale șirului y , la codarea lor în șirul alipit se va căuta continuarea fazei deja începute. Decizia se va lua pe baza unor valori tabelate, astfel:

Dacă $r < r_c$, respingem ipoteza nulă, dacă $r > r_c$, acceptăm ipoteza nulă.

Pentru volume mai mari de 20, se poate folosi statistica $R = \frac{r - \mu}{\sigma}$, unde μ este numărul mediu de faze și σ deviația sa standard:

$$\mu = 1 + \frac{2n_1n_2}{n_1 + n_2} \quad \text{și} \quad \sigma = \sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2(n_1 + n_2 - 1)}}.$$

Dacă ipoteza nulă este admisă, atunci statistica R urmează o repartiție normală $\mathcal{N}(0, 1)$. Pentru a lua decizia, procedăm astfel:

Dacă $|R| \geq z_{1-\alpha/2}$, atunci respingem ipoteza nulă. Altfel, o acceptăm.

Exemplul 8.1 Un observator astronomic recepționează, în două seturi de măsurători, semnale radio. În primul set de măsurători înregistrează 12 transmisiuni, iar în cel de al doilea 15 transmisiuni, ale căror secunde se regăsesc în tabelul următor. La un nivel de semnificație $\alpha = 0.05$, să se decidă dacă cele două seturi de date provin de la același emițător, adică provin din aceeași repartiție.

| | | | | | | | | | | | | | | | |
|---------|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|
| Setul 1 | 1.2 | 3.2 | 2.6 | 1.9 | 1.8 | 2.2 | 3.0 | 2.4 | 1.9 | 2.6 | 2.1 | 3.1 | 2.4 | 2.0 | 1.8 |
| Setul 2 | 1.6 | 2.4 | 2.7 | 1.9 | 2.6 | 2.8 | 2.0 | 3.1 | 2.5 | 2.2 | 2.8 | 3.5 | | | |

Ipotezele formulate sunt:

(H_0) datele provin din aceeași repartiție (adică sunt omogene) versus
 (H_1) seturile de date provin din repartiții diferite.

Alipim cele două șiruri de date, le ordonăm crescător și le atașăm codurile seturilor din care provin. Acolo unde sunt valori egale, atașăm codul vecinului din stânga, pentru a realiza continuarea de fază:

| | | | | | | | | | | | | | | | | |
|---|------|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|
| { | Date | 1.2 | 1.6 | 1.8 | 1.8 | 1.9 | 1.9 | 1.9 | 2.0 | 2.0 | 2.1 | 2.2 | 2.2 | 2.4 | 2.4 | 2.4 |
| | Cod | 1 | 2 | 1 | 1 | 1 | 1 | 2 | 2 | 1 | 1 | 1 | 2 | 2 | 1 | 1 |
| { | Date | 2.5 | 2.6 | 2.6 | 2.6 | 2.7 | 2.8 | 2.8 | 3.0 | 3.1 | 3.1 | 3.2 | 3.5 | | | |
| | Cod | 2 | 2 | 1 | 1 | 2 | 2 | 2 | 1 | 1 | 2 | 1 | 2 | | | |

Numărul total de faze este $r = 14$, dintre care $r_1 = 7$ și $r_2 = 7$. Dacă dorim să aplicăm cazul asimptotic,

$$\mu = 14.33, \quad \sigma = 2.515, \quad R_{\text{empiric}} = \frac{14 - 14.33}{2.515} = -0.1312, \quad z_{0.975} = 1.96.$$

Cum $|R_{\text{empiric}}| < z_{0.975}$, nu avem motive pentru respingerea ipotezei nule la un nivel de semnificație $\alpha = 0.05$.

Bibliografie

- [1] Anderson, M., *A characterization of the multivariate normal distribution*, The Annals of Mathematical Statistics, vol. 42, no. 2, 824-827, 1971.
- [2] Benhamou, E.; Melot, V., *Seven proofs of the Pearson Chi-squared independence test and its graphical interpretation*, arXiv:1808.09171v3, 2018.
- [3] Berk, R., *Review 1922 of 'Invariance of Maximum Likelihood Estimators' by Peter W. Zehna*, Mathematical Reviews, 33, 342-343, 1967.

- [4] Devore, J; Berk, K., *Modern Mathematical Statistics with Applications*, 2nd Edition, Springer New York Dordrecht Heidelberg London, 2012.
- [5] Durel, R., *Probability: Theory and Examples*, 5th Edition, Cambridge Series in Statistical and Probabilistic Mathematics, 2014.
- [6] Gibbons Dickinson, J.; Chakraborti, S., *Nonparametric Statistical Inference*, Fourth Edition, Revised and Expanded, Marcel Dekker, INC., New York, Basel, 2003.
- [7] Kendall, M.G., *The Advanced Theory of Statistics*, Volume 1, Distribution Theory, London, Charles Griffin & Company, 1945 (Edition by Stuart, Alan, Ord, Keith, 2010).
- [8] Kendall, M.G.; Stuart, A., *The Advanced Theory of Statistics*, Volume 2, Inference and Relationships, Hafner Publishing Company, 1961 (Edition by Wiley, 2010).
- [9] Klenke, A., *Probability Theory: A Comprehensive Course*, 2nd Edition, Springer, 2014.
- [10] Kolmogorov, A. N., *Sulla Determinazione Empirica di Una Legge di Distribuzione*, Giornale dell'Istituto Italiano degli Attuari, 4. 83-91, 1933.
- [11] Montgomery, D; Runger, G, *Applied Statistics and Probability for Engineers*, 3rd Edition, John Wiley & Sons, Inc, 2003.
- [12] Owen, A, *Lectures on statistics*, Department of Statistics, Stanford University.
- [13] Stoleriu, I., *Statistică aplicată*, note de curs, 2019.
- [14] Wackerly, D.; Mendenhall, W.; Scheaffer, R., *Mathematical Statistics with Applications*, 7th Edition, Thomson Brooks/Cole, 2008.
- [15] Walck, C., *Handbook on Statistical distributions for experimentalists*, Particle Physics Group, University of Stockholm.
- [16] Watson, G.S., *Some recent results in chi-square goodness-of-fit tests*, Biometrics, 15, 440, 1959.