

3. CARACTERISTICI STATISTICE ALE UNEI SERII DE DATE

3.1. INTRODUCERE

Statistica matematică, mai precis metodele furnizate de aceasta s-au implementat puternic în metodologia de lucru a diferite domenii. Apelul la metodele specifice statisticii matematice s-a făcut în principal din două motive:

- ❑ existența variabilității naturale a fenomenelor, proceselor, caracteristicilor etc. aflate sub observație;
- ❑ necesitatea luării unor decizii asupra acestor fenomene, procese, caracteristici etc.

Statistica matematică, sintetizând o informație, de cele mai multe ori parțială, asupra procesului investigat, poate furniza, cu riscuri controlate de operator, baza metodologică pentru adoptarea anumitor decizii, chiar în condiții specificate de incertitudine. În acest scop, statistica și-a dezvoltat direcțiile principale: principii, modele și metode.

La baza statisticii există două concepte de bază: populația și eșantionul.

Populația statistică este obiectul de studiu și poate fi reprezentat de mulțimea produselor ce pot rezulta dintr-un proces tehnologic, de mulțimea valorilor pe care le poate lua o caracteristică de calitate a unui produs etc. Standardele definesc populația statistică drept o mulțime de obiecte sau fenomene, calitativ omogene.

Eșantionul reprezintă acea parte a populației asupra căreia experimentatorul aplică metode statistice propriu-zise, pentru a obține concluzii pe care le extrapolează asupra întregii populații.

Această operație de extrapolare se numește inferență statistică. Inferența statistică este acea ramură a metodelor științifice de investigare a unei populații, care cu margini specificate de incertitudine exprimată în termeni probabilisti, face trecerea de la observații la concluzii privind populația, [20].

Inferența are două limite:

- ❑ informația pe care se bazează decizia este de natură aleatoare, fiind constituită din observații supuse întâmplării;
- ❑ se recunoaște explicit nesiguranța concluziilor.

În cazurile practice, populațiile luate în studiu nu se pot investiga integral, ci doar prin intermediul unuia sau mai multor eșantioane, ceea ce implică apariția unor riscuri în luarea deciziei finale asupra procesului investigat. Pentru a putea aplica anumite metode cantitative de extragere a informației dorite, este necesar ca populația statistică analizată să poată fi reprezentată matematic într-un mod convenabil, adică trebuie să se poată găsi un model statistic al populației studiate, care să încorporeze cât mai realist caracteristicile esențiale ale populației și care să nu fie prea complicat la manevrarea analitică. Soluția acestei probleme o constituie conceptul de variabilă aleatoare – a cărei valoare se atribuie în funcție de diferite circumstanțe sau evenimente ce se produc într-un experiment. Ex.: timpul de bună funcționare a unui produs, duritatea unor organe de mașini, rezistența la rupere a unor epruvete, o anumită cotă reprezentativă a unei piese etc. Această variabilă aleatoare caracterizează de fapt populația respectivă. Comportarea variabilei aleatoare este descrisă, din punct de vedere matematic, de funcția de repartiție asociată. Funcția de repartiție are o expresie specifică depinzând de tipul variabilei aleatoare: discretă (numărul de impulsuri/unitatea de timp) sau continuă (cota unei piese).

Repartiția de probabilitate a unei variabile aleatoare discrete se exprimă, de regulă, sub forma unui tablou, în care prima linie conține toate valorile posibile, iar a doua linie conține probabilitățile cu care ia aceste valori:

$$X = \begin{bmatrix} x_1 & x_2 & \dots & x_n \\ p_1 & p_2 & \dots & p_n \end{bmatrix}, \quad (3.1)$$

cu condiția $p_i \in [0,1]$ și $\sum_{i=1}^n p_i = 1$, p_i fiind probabilitatea ca X să ia valoarea x_i .

Se adoptă notația: $p_i = P(X = x_i)$. De obicei variabilele aleatoare discrete se asociază experimentelor ce constau din numărare.

Spre deosebire de variabilele discrete, cele continue pot lua orice valori într-un anumit interval. În consecință, nu se vor mai asocia probabilități punctuale ca în cazul anterior, ci se va defini o funcție pozitivă, $f(x)$, numită densitate de probabilitate, astfel încât aria domeniului cuprins între graficul ei și axa Ox este egală cu 1. Probabilitatea ca variabila aleatoare să ia valori într-un interval $(x;$

$x+h$) este $\int_x^{x+h} f(x)dx$ este probabilitatea ca variabila considerată să ia valori în intervalul $(x; x+dx)$ al dreptei reale (fig. 3.1).

În timpul unui proces de măsurare pot interveni trei tipuri de erori:

- erori aberante;
- erori sistematice;
- erori aleatoare.

Erorile aberante apar ca urmare a încălcării principiilor generale de măsurare sau ca rezultat al neatenției experimentatorului. Rezultatele afectate de erori aberante diferă esențial de restul valorilor și în consecință se pot elimina.

Erorile sistematice sunt determinate de diferiți factori, cum ar fi: reglarea necorespunzătoare a mijlocului de măsurare, variația condițiilor de mediu (temperatură, presiune, umiditate etc.). Erorile sistematice pot fi depistate și înlăturate pe baza unor calcule ce pleacă de la principiile fizice ce stau la baza măsurării respective.

Erorile aleatoare apar datorită unei mulțimi de factori, a căror influență individuală este neglijabilă și nu există posibilitatea înlăturării acestor influențe. Studiul influențelor erorilor aleatoare se bazează pe cunoașterea legilor de repartiție a acestor erori.

Dacă o măsurare se repetă de n ori, iar rezultatele se împart în clase (intervale) de lățime Δx , se poate calcula frecvența relativă cu care rezultatele apar în fiecare clasă:

$$f_i = \frac{n_i}{n}, \quad (3.2)$$

unde n_i este numărul de rezultate aflat în clasa i .

Reprezentarea grafică a frecvenței relative este o histogramă (Fig. 3.1), care pentru $n \rightarrow \infty$ și $\Delta x \rightarrow 0$ devine o funcție continuă, numită funcție densitate de repartiție sau densitate de probabilitate, deoarece prin integrarea ei pe un interval se obține probabilitatea cu care variabila ia valori în acest interval.

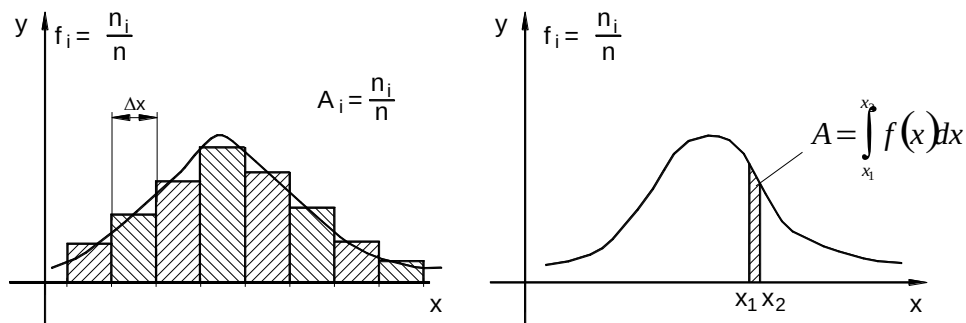


Fig. 3.1 Histograma și densitatea de repartiție

În mod natural se impune condiția:

$$\int_{-\infty}^{\infty} f(x) dx = 1, \quad (3.3)$$

atunci când x ia valori pe toată dreapta reală, adică probabilitatea ca variabila să ia valori reale este 1. Relativ la o variabilă aleatoare prezintă interes probabilitățile unor evenimente de tipul $(X \in I)$ - se citește evenimentul X ia valori în intervalul I . Pentru a calcula aceste probabilități, notate $P(X \in I)$ este suficient să cunoaștem funcția de repartiție a variabilei aleatoare. Funcția care asociază oricărui număr real a , probabilitatea ca x să ia valori mai mici sau egale cu a se numește funcție de repartiție a lui X :

$$F(a) = P(x \leq a), \quad (3.4)$$

Legătura dintre densitate și repartiție este dată de relația:

$$F'(x) = f(x). \quad (3.5)$$

Funcția de repartiție are următoarele proprietăți, [17, 18, 20]:

- valorile funcției de repartiție $F(x)$ aparțin intervalului $[0, 1]$, conform (3.3);
- pentru două numere reale x_1 și x_2 , cu $x_1 < x_2$, funcția de repartiție este nedescrescătoare, adică $F(x_1) \leq F(x_2)$;
- Dacă densitatea este o funcție continuă atunci:

$$P(x_1 \leq x < x_2) = P(x_1 < x < x_2) = P(x_1 < x \leq x_2) = P(x_1 \leq x \leq x_2). \quad (3.6)$$

Dacă variabila aleatoare este discretă funcția ei de repartiție are expresia:

$$F(x) = \begin{cases} 0 & x \leq x_1 \\ p_1 & x_1 \leq x \leq x_2 \\ p_1 + p_2 & x_2 < x \leq x_3 \\ \dots & \dots \\ p_1 + p_2 + \dots + p_{n-1} & x_{n-1} < x \leq x_n \\ 1 & x \geq x_n \end{cases} \quad (3.7)$$

Pe baza proprietăților funcției de repartiție se deduc relațiile:

$$\begin{aligned} P(x_1 \leq x \leq x_2) &= F(x_2) - F(x_1), \\ P(x \geq x_1) &= 1 - F(x_1), \end{aligned} \quad (3.8)$$

relații utile în practica industrială:

- probabilitatea ca o variabilă să ia valori între două limite date,
- probabilitatea ca o variabilă să fie mai mare decât o valoare dată sau într-o interpretare și mai particulară, când variabila x reprezintă numărul de ore de bună funcționare a unui dispozitiv, aparat atunci $P(x \geq x_1)$ este probabilitatea ca acel dispozitiv, aparat să funcționeze minimum x_1 ore.

În figura 3.2 sunt prezentate funcția densitate de probabilitate și funcția de repartiție, fiind marcate ariile ce reprezintă probabilitatea ca:

- variabila ia valori mai mici decât a ;
- variabila ia valori mai mari decât b ;
- variabila ia valori cuprinse în intervalul $[a, b]$.

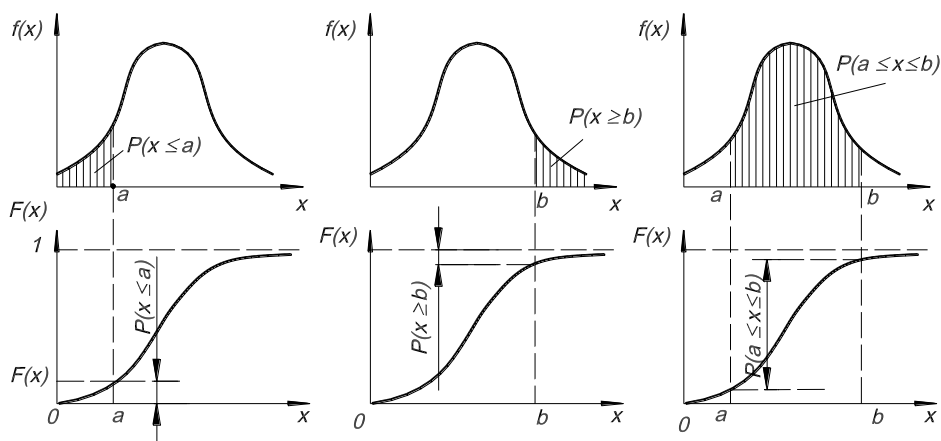


Fig. 3.2 Funcția densitate de repartiție și funcția de probabilitate

Probabilitățile reprezintă un concept esențial în înțelegerea analizei statistice, motiv pentru care trebuie bine înțelese. Principalele aspecte ce vor fi subliniate se referă la: atribuirea probabilităților și reguli ale acestora.

Probabilitatea este o măsură numerică ce cuantifică șansa unui eveniment de a se produce într-un experiment. Se măsoară pe o scală de la 0 la 1, 0 corespunzând evenimentului imposibil, iar 1 evenimentului sigur, adică evenimentul ce s-ar produce ori de câte ori s-ar repeta experimentul în aceleași condiții. În plus, suma probabilităților tuturor evenimentelor posibile trebuie să fie egală cu 1 ($\sum P_i = 1$), în cazul când experimentul are un număr finit de evenimente ce se exclud reciproc.

Există două modalități de atribuire a probabilităților unor evenimente, în funcție de situație:

- metoda frecvențelor relative;
- metoda subiectivă.

Metoda frecvențelor relative se bazează pe raționamente, deoarece cu ajutorul logicii se pot determina probabilitățile de apariție ale unor evenimente. Într-un experiment cu un număr finit de realizări mutual exclusive și echiprobabile, probabilitatea unui eveniment este numărul cazurilor favorabile/numărul cazurilor posibile. De ex.: când se dă cu banul, există două situații posibile, "cap" sau "stema", evenimente ce sunt echiprobabile (au aceeași șansă de apariție), de aceea $P_{\text{cap}} = P_{\text{stema}} = \frac{1}{2} = 0.5$; sau se bazează pe efectuarea unor observații (evidențe empirice), situație când:

$$P = \frac{n_i}{N}, \quad (3.9)$$

unde n_i este numărul de evenimente favorabile, iar N numărul total de evenimente. Ex.: dacă se înregistrează numărul de ore de funcționare a unui bec și se constată că din 1000 de becuri testate, 100 au funcționat mai puțin de 500 ore, atunci probabilitatea ca un bec ales aleator să aibă o durată de viață mai mică de 500 h este:

$$P = \frac{100}{1000} = 0,1.$$

Această metodă nu poate fi aplicată în orice situație, iar pentru evenimente viitoare se utilizează metoda subiectivă, conform căreia se atribuie probabilități pe baza propriilor considerente. Evident, în acest caz diferite persoane vor atribui probabilități diferite unui același eveniment (probabilitatea de câștigare a unui concurs).

Există numeroase situații când trebuie determinată probabilitatea unor evenimente legate. De exemplu, pentru două evenimente, A și B , ne interesează dacă se vor produce ambele sau dacă va avea loc unul dintre ele. Pentru a putea răspunde la acest tip de întrebare trebuie introduse două reguli fundamentale ale probabilităților:

- regula de adunare a probabilităților;
- regula de înmulțire a probabilităților.

Regula de adunare se aplică în situația când există două evenimente și se dorește determinarea probabilității de apariție a cel puțin unuia din cele două evenimente. Există două variante ale adunării, în funcție de tipul evenimentelor; dacă ele se exclud reciproc sau nu. Două evenimente se exclud reciproc, dacă nu se pot produce simultan (cap sau stemă pentru o singură monedă).

Dacă evenimentele se exclud reciproc, probabilitatea de apariție a cel puțin unui eveniment din cele două este:

$$P(A \cup B) = P(A) + P(B) \quad (3.10)$$

Dacă evenimentele nu se exclud reciproc, atunci "A sau B" înseamnă se produce A, se produce B sau se produc ambele. În acest caz, din suma anterioară trebuie scăzută probabilitatea de producere simultană, pentru a evita însumarea ei dublă (vezi fig. 3. 3):

$$P(A \cup B) = P(A) + P(B) - P(A \cap B). \quad (3.11)$$

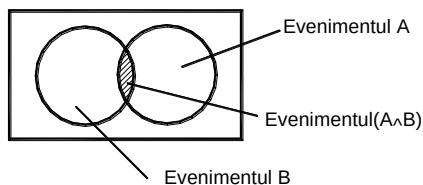


Fig. 3.3 Evenimentul $P(A \cup B)$

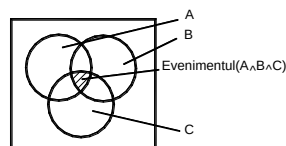


Fig. 3.4 Evenimentul $P(A \cap B \cap C)$

Această regulă poate fi generalizată pentru un număr de mai multe evenimente. De exemplu, pentru trei evenimente, ce se exclud reciproc, se obține:

$$P(A \cup B \cup C) = P(A) + P(B) + P(C), \quad (3.12)$$

iar pentru evenimente ce nu se exclud reciproc:

$$P(A \cup B \cup C) = P(A) + P(B) + P(C) - P(A \cap B) - P(B \cap C) - P(A \cap C) + P(A \cap B \cap C) \quad (3.13)$$

Relația între evenimentele implicate în (3.13) este ilustrată în fig. 3.4, unde se observă că zona hașurată este adunată de trei ori (în $P(A)$, $P(B)$ și $P(C)$), apoi scăzută de trei ori (în $P(A \cap B)$, $P(B \cap C)$ și $P(A \cap C)$), deci trebuie inclusă din nou la sfârșit.

Regula de înmulțire se folosește la găsirea probabilității de producere simultană a două sau mai multe evenimente ($P(A \cap B)$, $P(A \cap B \cap C)$ etc.). Răspunsul la astfel de probleme depinde de tipul evenimentelor, dacă sunt independente sau nu.

Două evenimente sunt dependente când apariția unui eveniment influențează probabilitatea de apariție a celuilalt. Regula de înmulțire se aplică diferit, în funcție de tipul evenimentelor:

➤ pentru evenimente dependente:

$$P(A \cap B) = P(A) \times P(B|A), \quad (3.14)$$

unde $P(B|A)$ reprezintă probabilitatea de producere a evenimentului B, știind că A s-a produs:

$$P(B|A) = \frac{P(A \cap B)}{P(A)}. \quad (3.15)$$

➤ pentru evenimente independente:

$$P(A \cap B) = P(A) \times P(B). \quad (3.16)$$

Expresia $P(B|A)$ reprezintă o probabilitate condiționată, adică probabilitatea ca B să fie afectat de apariția lui A. Evident, în cazul a două evenimente independente $P(B|A) = P(B)$. O extindere a acestei idei în contextul mai multor evenimente secvențiale este cunoscută sub numele de regula lui Bayes. În cazul a două evenimente formula are forma (3.15). Pentru n evenimente mutual exclusive $B_1, B_2, \dots, B_n$, când A a apărut:

$$P(B_i|A) = \frac{P(B_i)P(A|B_i)}{P(B_1)P(A|B_1) + P(B_2)P(A|B_2) + \dots + P(B_n)P(A|B_n)} \quad (3.17)$$

3.2. STATISTICA DESCRIPTIVĂ

Statistica descriptivă se ocupă, în principal, cu două probleme:

- ❑ prezentarea datelor sub formă tabelară și vizualizarea lor sau a unor caracteristici prin tehnici grafice;
- ❑ utilizarea unor indicatori numerici pentru caracterizarea datelor.

Pentru a putea organiza un volum mare de date și a le prezenta sub formă grafică sau tabelară este necesar să se detecteze eventualele tendințe de distribuire a lor pe axa reală prin intermediul unor mijloace de analiză. Cele mai utilizate metode grafice de reprezentare sunt: tabele de distribuție de frecvențe, pie-chart-uri, histogramme, poligoane de frecvență, curbe de frecvență cumulată etc

3.2.1. ORGANIZAREA DATELOR

Pentru a putea reprezenta tabelele de frecvență este necesar ca datele să fie grupate în câteva clase sau intervale semnificative (nici prea multe, nici prea puține) și să se contorizeze numărul de valori din fiecare clasă, iar apoi să se calculeze frecvențele relative, adică numărul datelor din fiecare clasă raportat la numărul total. În general, se recomandă ca numărul claselor să fie între 5 – 15 (în funcție de mărimea eșantionului avut la dispoziție) și cel mai adesea mărimea claselor este egală. De asemenea, trebuie ca fiecare valoare existentă în șir să fie introdusă o singură dată (se stabilește o convenție pentru

limitele intervalelor).

În tabelul 3.1 se prezintă un set de date experimentale organizat cu ajutorul tabelului de frecvență.

Tabelul 2.1. Tabel de frecvență

Date	10.824	10.827	10.775	10.826	10.771	10.802	10.816	10.816	10.829
	10.814	10.815	10.801	10.813	10.821	10.812	10.860	10.850	10.865
	10.810	10.854	10.847	10.844	10.796	10.858	10.797	10.843	10.839
Date ordonate	10.771	10.775	10.796	10.797	10.801	10.802	10.810	10.812	10.813
	10.814	10.815	10.816	10.816	10.821	10.824	10.826	10.827	10.829
	10.839	10.843	10.844	10.847	10.850	10.854	10.858	10.860	10.865
Clase				Frecvența absolută			Frecvența relativă [%]		
10.770 – 10.7867				2			7.40		
10.7867 - 10.8023				4			14.81		
10.8023 - 10.8180				7			25.92		
10.8180 - 10.8337				5			18.51		
10.8337 - 10.8493				4			14.81		
10.8493 - 10.8650				5			18.51		

După construirea tabelului de frecvență, în general, pasul următor îl constituie reprezentarea sub o anumită formă grafică.

Un pie-chart este o reprezentare grafică sub forma unui cerc, unde frecvențele relative sunt utilizate pentru divizarea cercului în sectoare de cerc corespunzătoare fiecărei categorii de variabile. Acest tip de reprezentare se pretează în cazul când datele sunt grupate pe diverse categorii. Ex.: Totalul studenților din Facultatea de Mecanică clasificat în studenți integraliști, studenți ce au acumulat între 50 – 60 credite, 40 – 50 credite, mai puțin de 40 credite.

Cel mai utilizat mod de reprezentare îl constituie histogramele. În fig. 3.5 se prezintă histograma aferentă datelor din tabelul 3.1. În cazul când intervalele nu sunt echidistante trebuie ajustată înălțimea dreptunghiurilor astfel încât aria fiecărui dreptunghi să fie proporțională cu frecvența corespunzătoare.

Funcțiile din Matlab asociate cu histograma sunt:

- **N=hist(x)** – returnează un vector linie, ce conține numărul de date din vectorul x, aflate în 10 intervale echidistante;

- ❑ **N=hist(x,m)** – returnează un vector linie, ce conține numărul de date din vectorul x, aflate în m intervale echidistante;
- ❑ **N=histc(x,limite)** – contorizează valorile lui x aflate între elementele vectorului limite, ce trebuie să fie un vector monoton crescător. N va avea dimensiunea cu o unitate mai mică decât vectorul limite, iar N(k) va număra valorile x(i) dacă $\text{limite}(k) \leq x(i) < \text{limite}(k+1)$. Ultimul interval va contabiliza și valorile egale cu limita superioară;
- ❑ **hist(x)** – trasează histograma utilizând 10 intervale sau clase;
- ❑ **hist(x,m)** - trasează histograma utilizând m clase.

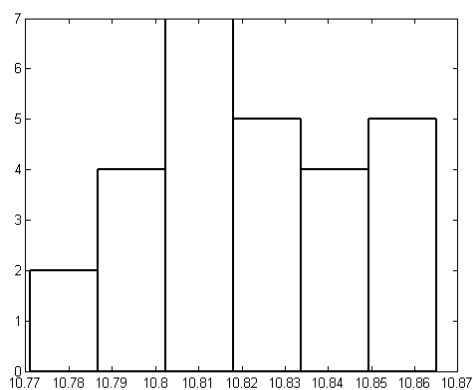


Fig. 3.5. Histograma datelor

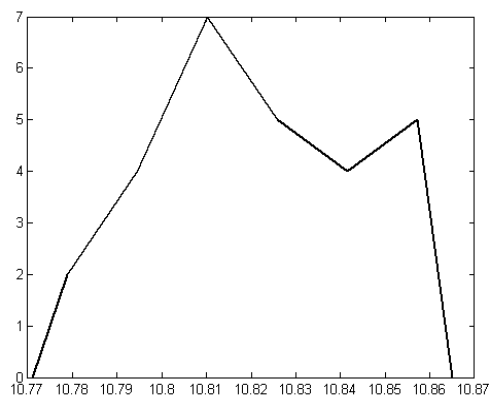


Fig. 3.6. Poligonul de frecvențe

Poligoanele de frecvență se utilizează mai ales în cazul când se dorește compararea a două distribuții în aceeași reprezentare grafică. Ele se obțin prin reprezentarea frecvențelor în dreptul mijlocului fiecărei clase și unirea acestor puncte prin linii drepte. În plus, la extremități, primul, respectiv ultimul punct corespunde valorii minime, respectiv maxime din șirul de date. În fig. 3.6. este prezentat poligonul de frecvențe al datelor cuprinse în tabelul 3.1.

O altă variantă de reprezentare este curba frecvențelor cumulate, ce permite aflarea răspunsului la întrebări de tipul: câte observații sunt egale sau mai mici decât limita superioară a fiecărei clase, respectiv câte sunt mai mari sau egale cu limita inferioară a fiecărei clase. După cum indică și numele, curba frecvențelor cumulate se obține prin reprezentarea frecvențelor cumulate ale claselor. Cumularea se poate face de la valoarea minimă spre maximă sau invers. În primul caz se obține o curbă ascendentă, ce oferă informații privind

numărul de observații mai mic cel mult egal cu limita superioară a fiecărei clase, iar în cazul al doilea, se obține o curbă descendentă, pe baza căreia se poate determina numărul de observații mai mare cel mult egal cu limita inferioară a fiecărui interval. În fig. 3.7 se prezintă curba frecvențelor cumulate corespunzătoare datelor prezentate în tabelul 3.1.

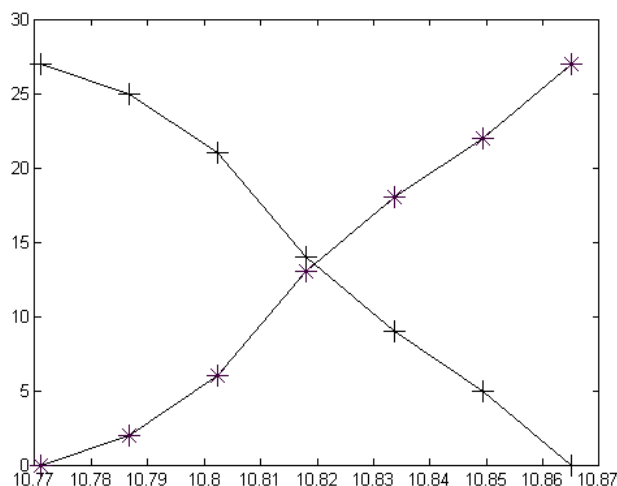


Fig. 3.7. Curba frecvențelor însumate

3.2.2. INDICATORI NUMERICI DE CARACTERIZARE A DATELOR

Există două caracteristici esențiale ce se investighează în cazul seturilor de date:

- centrarea sau localizarea;
- concentrarea sau împrăștierea.

Fie o serie de date $X: x_1, x_2, \dots, x_n$ obținută în urma unui proces de măsurare. Tendența centrală a seriei poate fi caracterizată prin: medie, mediană, modul.

Cel mai utilizat indicator de centrare este media. Media unei serii de date este media aritmetică a valorilor sale:

$$\bar{x} = \frac{1}{n}(x_1 + x_2 + \dots + x_n) = \frac{1}{n} \sum_{i=1}^n x_i. \quad (3.18)$$

Media reprezintă centrul de greutate al seriei de date. În Matlab, media seriei

de date, stocate în vectorul x , se calculează cu funcția **mean** (x) .

Mediana este valoarea centrală a unui set de date ordonate crescător. Ea se calculează cu formula:

$$\begin{aligned} x_{me} &= x_{\frac{n+1}{2}}, \quad n = 2k + 1 \\ x_{me} &= \frac{1}{2} \left(x_{\frac{n}{2}} + x_{\frac{n}{2}+1} \right), \quad n = 2k \end{aligned} \quad (3.19)$$

Prima formulă se aplică pentru un număr impar de elemente ale seriei de date, iar a doua când volumul seriei este un număr par. Mediana caracterizează mai bine valoarea centrală unui set de date în situația când setul de date este asimetric, respectiv apare o concentrare de date la una din extremitățile seriei. În Matlab mediana seriei de date stocate în vectorul x se calculează cu funcția **median** (x) .

Modulul este observația cu cea mai mare frecvență de apariție. Există seturi de date cu un singur modul (unimodale) sau cu mai multe (multimodale). Modulul se calculează cu formula:

$$x_{mo} = \bar{x} + 3(x_{me} - \bar{x}). \quad (3.20)$$

În cazul unor seturi de date simetrice, media, mediana și modulul coincid. La repartițiile asimetrice se poate întâlni cazul $\bar{x} < x_{me} < x_{mo}$ sau $x_{mo} < x_{me} < \bar{x}$ (fig. 3. 8).

Pentru caracterizarea variabilității datelor se utilizează o serie de indicatori de concentrare sau împrăștiere. Printre cei mai uzuali se numără: intervalul de variație, deviația standard, varianța, coeficientul de variație etc.

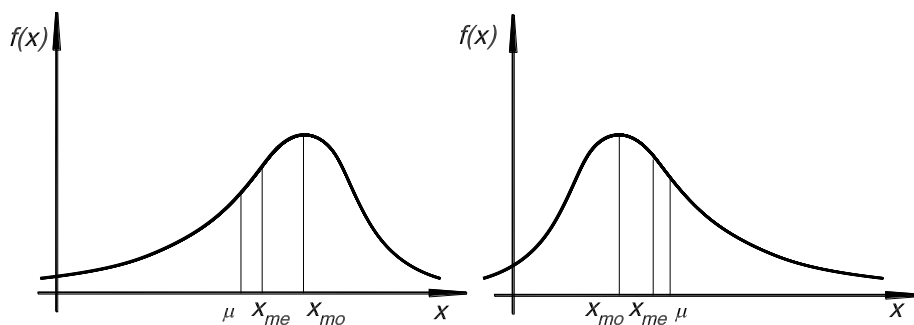


Fig. 3.8 Repartiții asimetrice

Intervalul de variație este diferența dintre valoarea maximă și minimă a seriei

de date. Deoarece acesta se bazează doar pe două valori aparținătoare seriei de date, în multe situații se apelează la p-cvantele.

P-cvanta ($p \in (0,1)$) seriei de date este un număr real, q_p , cu proprietatea că 100p% valori ale seriei sunt mai mici sau egale decât q_p , iar $(1 - p)100\%$ valori sunt mai mari sau egale decât q_p . Numărul q_p poate să aparțină seriei sau nu.

Pentru determinarea p-cvantei unei serii de date se procedează în felul următor [14]:

- se ordonează crescător setul de date X' : $x_1', x_2', \dots, x_n'$;
- se calculează partea întreagă a numărului $n \cdot p + 0.5$, care se notează i ;
- p-cvanta este:

$$q(p) = x_i' + \left(x_{i+1}' - x_{i-1}' \right) \{np + 0.5\}, \quad (3.21)$$

unde $\{np + 0.5\}$ este partea fracționară a numărului $np + 0.5$.

În cazul când $p = 0.5$, $q_{0.5}$ este chiar mediana seriei de date, deoarece $i = [n \cdot 0.5 + 0.5] = [(n+1)/2]$, care pentru n egal cu un număr impar este $(n+1)/2$, deci 0.5-cvanta este elementul șirului ordonat $x_{(n+1)/2}$, iar în cazul când n este

un număr par, 0.5-cvanta devine $\frac{1}{2} \left(x_{\frac{n}{2}} + x_{\frac{n+1}{2}} \right)$, deci se obțin chiar expresiile medianei.

În cazul unor seturi de date de volum mai mare se utilizează procentilele. Procentila 100p, $p \in (0,1)$, este valoarea reală x , cu proprietatea că cel mult 100p% din valorile seriei de date sunt mai mici decât x și cel mult $100(1 - p)\%$ sunt mai mari. Dacă $p(n+1)$ este o valoare întreagă, procentila 100p este valoarea de rang $p(n+1)$ din seria ordonată, iar în cazul când $p(n+1)$ nu este număr întreg, procentila seriei este valoarea reală obținută prin interpolare din valorile pozițiilor adiacente [14].

În Matlab, există funcția `prctile(x,p)` ce calculează procentila serie de date x , iar $p \in (0,1)$. Pentru seria de date prezentată în tabelul 2.1, procentila 0.3 este 10.8126, iar procentila 0.5, adică mediana, este 10.8210.

Printre cei mai uzuali indicatori de concentrare sau împrăștiere a valorilor seriei de date în jurul mediei se numără dispersia sau varianța unei serii de date. Având o serie de date x : $x_1, x_2, \dots, x_n$ cu media $\bar{x}$, dispersia este:

$$S^2(x) = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2. \quad (3.22)$$

Pentru o serie statistică:

$$Y = \begin{pmatrix} x_1 & x_2 & \dots & x_k \\ n_1 & n_2 & \dots & n_k \end{pmatrix} \quad \sum_{j=1}^k n_j = N \quad (3.23)$$

n_i indică de câte ori apare x_i în seria de date, $\bar{x}$ este media seriei, dispersia este:

$$S^2(Y) = \frac{1}{N-1} \sum_{j=1}^k (x_j - \bar{x})^2 n_j. \quad (3.24)$$

În Matlab, varianța unei serii de date se poate calcula cu funcțiile **var(x)**, ce calculează dispersia seriei de date cu formula (3.22), respectiv **var(x,1)** ce aplică formula:

$$S^2(x) = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2. \quad (3.25)$$

Având o serie de date cu dispersia $S^2(x)$, valoarea $S = \sqrt{S^2(x)}$ se numește abatere standard sau abatere medie pătratică a seriei x . În Matlab, abaterea standard se calculează cu funcțiile **std(x)**, respectiv **std(x,1)**.

Un indicator al concentrării relative este coeficientul de variație. Având o serie de date x , cu media $\bar{x}$ și abaterea standard S , acesta are expresia:

$$CV(x) = \frac{100S}{\bar{x}}. \quad (3.26)$$

Cu cât coeficientul de variație este mai apropiat de 0, seria este mai omogenă și media $\bar{x}$ mai reprezentativă. Dacă valoarea coeficientului de variație tinde spre 100, împrăștierea seriei este mare și media este mai puțin reprezentativă.

Ultimii indicatori menționați sunt coeficientul de asimetrie (γ_1) și coeficientul de exces (γ_2):

$$\gamma_1 = \frac{m_3}{S^3},$$

$$\gamma_2 = \frac{m_4}{S^4}, \quad (3.27)$$

unde m_3 și m_4 sunt momentele centrate de ordinul 3, respectiv 4. Momentul centrat de ordinul k se calculează cu formula:

$$m_k = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^k. \quad (3.28)$$

Asimetria și excesul ajută la identificarea formei de repartiție a datelor experimentale. Etalonul la această comparație este repartiția normală, pentru care $\gamma_1 = 0$ și $\gamma_2 = 3$. Excesul mai poartă denumirea de “kurtosis”. În funcție de valoarea coeficientului de exces (pozitiv, negativ sau nul) se pot trage concluzii asupra alurii curbei densității de repartiție (Fig. 3.9).

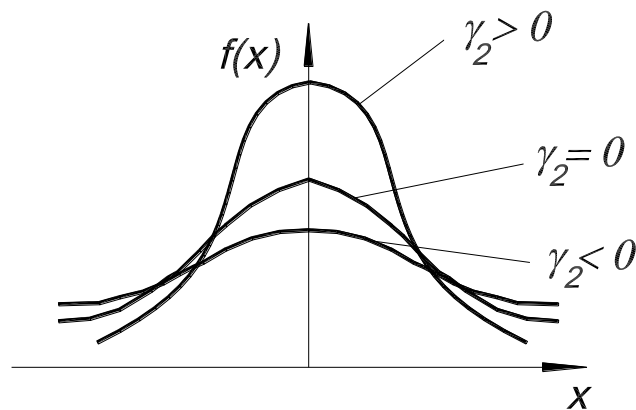


Fig. 3.9 Variația formei de repartiție față de coeficientul de exces

Dacă $\gamma_2 = 0$, repartiția se numește mezokurtică, având o formă apropiată de repartiția normală. Dacă $\gamma_2 > 0$, repartiția este leptokurtică, fiind mai ascuțită decât cea normală, iar dacă $\gamma_2 < 0$, repartiția este platokurtică, fiind mai aplatizată.

Coeficientul de asimetrie caracterizează simetria repartiției. Repartițiile simetrice au $\gamma_1 = 0$. În cazul $\gamma_1 > 0$, repartiția este pozitiv asimetrică, iar în cazul $\gamma_1 < 0$, repartiția este negativ asimetrică.

În Matlab, momentul centrat de ordinul k al seriei de date stocate în vectorul x

se calculează cu funcția **moment (x,k)**, coeficientul de asimetrie se determină cu funcția **skewness (x)**, iar coeficientul de exces cu **kurtosis (x)**.

3.2.3. METODE DE AFIȘARE SAU REPREZENTARE GRAFICĂ ÎN ANALIZA PRIMARĂ A DATELOR

O metodă de afișare și analiză primară a datelor este diagrama tulpină – frunze (stem and leaf). Ea se pretează la serii mici de date ce au același ordin de mărime. A fost introdusă de Tukey în 1977 ca o metodă de afișare a datelor într-o listă structurată. În cazul unui set de date ce conține valori cu minim două cifre se poate construi o astfel de diagramă. Se separă fiecare valoare a șirului de date în două părți : tulpina și frunza. Tulpina reprezintă prima cifră a datelor, iar frunza cea de-a doua, respectiv ultima. În cazul valorii 34, 3 este tulpina, iar 4 este frunza; în cazul valorii 126, tulpina este 12, iar frunza este 6.

Fie seria de date x: 54, 25, 43, 3, 28, 39, 78, 32, 54, 93, 27, 33, 22, 78, 75, 83, 62, 76, 77, 67, 77, 80, 4, 26, 18, 10, 34, 30, 36, 43, 78, 41, 24, 91, 90, 63, 87, 55, 60, 49, 37, 39, 51, 66, 10, 76, 34, 47, 51, 34.

Se constată ca fiecare valoare este de ordinul zecilor sau al unităților. În acest caz, tulpina este cifra zecilor, iar cifra unităților reprezintă frunza. Se ordonează crescător seria. Diagrama tulpină – frunze este constituită din trei coloane. Pe coloana a doua se trec în ordine crescătoare tulpinile, iar în ultima se înregistrează în ordine crescătoare toate frunzele aparținătoare tulpinei respective (în acest caz cifrele unităților). În prima coloană se indică frecvențele cumulate de la prima clasă până la clasa medianei, respectiv de la ultima clasă până la clasa medianei. Frecvența clasei medianei se introduce în paranteză. Diagrama setului de date se prezintă în fig. 3. 10.

Frecvențe cumulate	Tulpina	Frunze
2	0	34
6	1	0058
12	2	245678
22	3	0234446799
(5)	4	13379
23	5	1445
19	6	02367
14	7	56677888
6	8	037
3	9	013

Fig. 3.10 Diagrama tulpină – frunze

Acest tip de afișare se aseamănă cu histograma, dar pe lângă distribuția datelor sunt indicate și valorile respectivei serii. De asemenea, în urma unei astfel de afișări se poate decide modul de divizare a setului de date în clase, în sensul stabilirii mărimii și numărului acestora.

O metodă de reprezentare grafică este box-plot-ul. El se construiește pe baza a 5 valori asociate seriei de date: minimul, cvantila inferioară $q(0.25)$, mediana, cvantila superioară $q(0.75)$ și maximul. Aceste valori pot fi reprezentate pe o axă verticală sau orizontală. Ambele sunt ilustrate în fig. 3.11.

Se desenează un dreptunghi având o latură egală cu diferența dintre cvantila superioară și cea inferioară. În acest dreptunghi se mai trasează o dreaptă paralelă cu latura menționată, în dreptul valorii mediane. Lungimea laturii dreptunghiului, $q(0.75) - q(0.25)$, este cunoscută sub numele de interval intercvantilic (IQR). Acest interval se multiplică cu un coeficient, de regulă 1.5 și se determină valoarea $q(0.25) - 1.5 \cdot \text{IQR}$ și $q(0.75) + 1.5 \cdot \text{IQR}$. Se trasează două segmente ce unesc mijloacele bazelor dreptunghiului cu aceste valori. Aceste segmente se numesc mustăți (whiskers). datele situate în afara intervalului $[q(0.25) - 1.5 \cdot \text{IQR}; q(0.75) + 1.5 \cdot \text{IQR}]$ sunt considerate aberante. Alegerea coeficientului 1.5 nu este standardizată, ea poate fi impusă de utilizator în funcție de natura datelor și a experienței precedente din analiza datelor de un anumit tip.

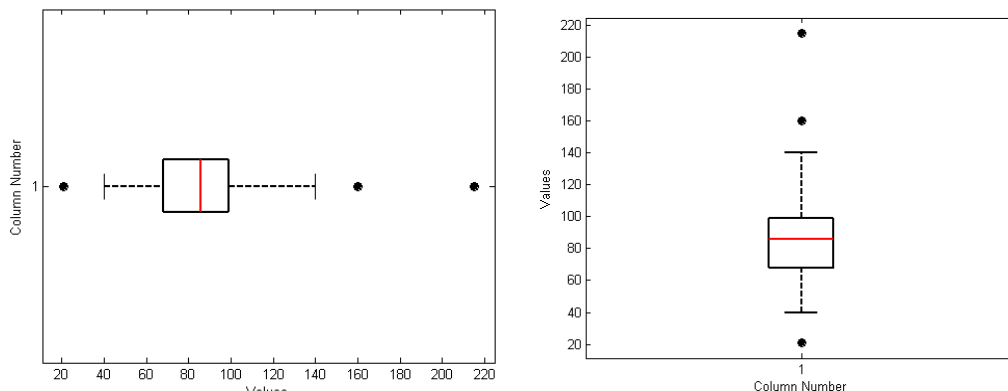


Fig. 3.11 Box-plot

În Matlab, un box-plot se poate trasa cu comanda `boxplot(x,cod-box,symbol,vertic,whisker)`, unde x este setul de date, ce poate fi un vector sau o matrice. În cazul când x este o matrice se va trasa câte un box-plot pentru fiecare coloană. `Cod-box` este o variabilă ce poate lua valoarea 0

sau 1, dacă este 0 se va desena un dreptunghi, dacă este 1 se va desena un dreptunghi din care se decupează un V. **simbol** este caracterul cu care se marchează valorile din exteriorul intervalului $[q(0.25) - 1.5 \cdot IQR; q(0.75) + 1.5 \cdot IQR]$ și este implicit '+'. **vertic** este o variabilă ce poate lua valoarea 0 (trasarea se face pe orizontală) sau 1 (trasarea se face pe verticală). **whisker** definește intervalul, este valoarea cu care se înmulțește IQR și este implicit 1.5. Dacă **whisker** = 0 se vor marca toate valorile aflate în afara intervalului $[q(0.25), q(0.75)]$.

În Matlab, intervalul intercvantilic al unei serii de date se poate calcula cu funcția **iqr(x)**.

Acest tip de reprezentare permite indicarea prezumtivelor valori aberante dintr-o serie de date.

3.3. TIPURI DE REPARTIȚII

Determinarea probabilităților asociate unor evenimente aleatoare poate fi mult simplificată dacă se construiește un model matematic ce descrie cu acuratețe situațiile asociate cu anumite evenimente de interes. Un astfel de model utilizat la determinarea probabilităților de producere a unor evenimente este o distribuție de probabilitate.

Unui experiment aleator în care se determină (prin măsurare sau observare) anumite caracteristici ale unei populații, ce pot fi cuantificate numeric, i se asociază una sau mai multe variabile, numite variabile aleatoare. O variabilă aleatoare ce poate lua valori într-o mulțime numărabilă se numește variabilă discretă, iar în situația când poate lua orice valoare într-un interval real se numește variabilă aleatoare continuă.

Comportarea variabilei aleatoare este descrisă din punct de vedere matematic de distribuția de probabilitate asociată. Distribuția de probabilitate a unei variabile aleatoare discrete se exprimă sub forma unei matrice:

$$X = \begin{bmatrix} x_1 & x_2 & \dots & x_n \\ p_1 & p_2 & \dots & p_n \end{bmatrix}, \quad (3.29)$$

pe prima linie fiind înregistrate valorile variabilelor, iar pe a doua probabilitățile de realizare, cu condițiile: $p_i \in [0,1], \sum_{i=1}^n p_i = 1$.

Pornind de la evenimente elementare ($X = x_i$) ale unui experiment se pot

determina probabilități ale evenimentelor de tipul: $(X \leq a), (X \geq a), (a \leq X \leq b), (a < X < b)$ etc. Aceste probabilități se obțin prin însumarea probabilităților evenimentelor elementare corespunzătoare, deoarece acestea se exclud reciproc.

Exemplu: Fie x variabila aleatoare ce indică numărul de restanțe înregistrate la sfârșitul unui semestru de studenții unui an de studiu:

$$x = \begin{bmatrix} 0 & 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ 0.05 & 0.3 & 0.25 & 0.15 & 0.1 & 0.05 & 0.05 & 0.05 \end{bmatrix}.$$

- Să se determine probabilitatea ca un student să aibă maximum 3 restanțe.
- Dacă un student are minim 3 restanțe, care este probabilitatea ca el să aibă 4 restanțe.

Evenimentul a cărui probabilitate se cere în cazul a) este:

$$(X \leq 3) = (X = 0) \cup (X = 1) \cup (X = 2) \cup (X = 3).$$

Deci:

$$P(X \leq 3) = P(X = 0) + P(X = 1) + P(X = 2) + P(X = 3) = 0.05 + 0.3 + 0.25 + 0.15 = 0.75.$$

În cazul b) se notează $A = (X = 4)$, evenimentul ca studentul să aibă 4 restanțe și $B = (X \geq 3)$, evenimentul ca studentul să aibă minim 3 restanțe. În problemă se cere probabilitatea evenimentului A condiționată de B :

$$P(A|B) = \frac{P(A \cap B)}{P(B)}.$$

$$P(B) = P(X=3) + P(X=4) + P(X=5) + P(X=6) + P(X=7) = 0.15 + 0.1 + 0.05 + 0.05 + 0.05 = 0.4$$

$$P(A \cap B) = P(A) = 0.1$$

$$P(A|B) = 0.1/0.4 = 0.25$$

În cazul variabilelor aleatoare discrete funcția de repartiție definită prin $F(X) = P(X \leq x)$ – probabilitatea ca variabila aleatoare X să ia valori mai mici sau egale cu x - este de forma:

$$F(X) = \begin{cases} 0 & X \leq x_1 \\ p_1 & x_1 < X \leq x_2 \\ p_1 + p_2 & x_2 < X \leq x_3 \\ \dots & \dots \\ p_1 + p_2 + \dots + p_{n-1} & x_{n-1} < X \leq x_n \\ 1 & X \geq x_n \end{cases} \quad (3.30)$$

După cum se extrag informații dintr-o serie de date prin asocierea unor indicatori numerici, în același mod se definesc caracteristici pentru variabile aleatoare.

Valoarea medie a unei variabile aleatoare discretă, X , de valori x_i și probabilitățile $p_i = P(X=x_i)$, $i=1, n$ este numărul notat $M(X)$:

$$M(X) = \sum_{i=1}^n x_i p_i \quad (3.31)$$

Fie o variabilă aleatoare discretă X , ce are media m . Dispersia sau varianța variabilei este:

$$\sigma^2(X) = M((X - m)^2), \quad (3.32)$$

iar abaterea standard $\sqrt{\sigma^2(X)}$ și se notează $\sigma(X)$.

Calcularea acestor indicatori este exemplificată în exemplul următor [14]. Fie distribuția vârstei populației României în 1990 și 2000. Se raportează vârsta pe categorii, în intervale, reprezentate de mijloacele acestora. Pentru fiecare interval se indică procentul populației ce are vârsta în acel interval:

Interval	Mijloc interval	1990	2000
sub 5 ani	3	7.6	6.4
5-13	9	12.8	11.6
14-17	16	5.3	5.2
18-24	21	12.8	9.0
25-34	30	17.3	12.5
35-44	40	15.1	12.2
45-64	55	18.6	22.5
65-84	75	9.3	16.0
≥ 85	92	1.2	4.6

Interpretând procente ca probabilități, distribuția vârstei în 1990 este:

$$X = \begin{bmatrix} 3 & 9 & 16 & 21 & 30 & 40 & 55 & 75 & 92 \\ 0.076 & 0.128 & 0.053 & 0.108 & 0.173 & 0.151 & 0.186 & 0.093 & 0.012 \end{bmatrix}.$$

Vârsta medie a populației României în 1990 era:

$M(X) = 3(0.076) + 9(0.128) + \dots + 92(0.012) = 34.455$, iar abaterea standard:

$$\sigma(X) = \sqrt{(3 - 34.455)^2 + \dots + (92 - 34.455)^2} = 21.87.$$

Considerând distribuția de probabilitate a vârstei populației în 2000, se obține media $M(Y) = 41.2$ și abaterea standard $\sigma(Y) = 25.4$. Comparând valorile corespunzătoare se observă că media vârstei crește, dar și variabilitatea crește.

3.3.1. DISTRIBUȚIA BINOMIALĂ

Distribuția binomială caracterizează variabile discrete, ce se asociază unui proces Bernoulli. Procesul Bernoulli constă din n încercări ale unui experiment, ce are rezultate ce se exclud reciproc (mutual exclusive): succes sau eșec. Încercările sunt independente, adică rezultatul uneia nu influențează rezultatul celeilalte, iar pentru fiecare încercare probabilitatea succesului este aceeași, p .

Într-un experiment Bernoulli prezintă interes numărul de succese și eșecuri din n încercări. Dacă se notează X variabila aleatoare discretă ce reprezintă numărul de reușite din n încercări, atunci X va putea lua valorile discrete $0, 1, \dots, n$. Probabilitățile asociate fiecărei valori i vor avea o distribuție de frecvențe de tip binomial.

Probabilitatea de a înregistra k succese din n încercări, cu probabilitatea de succes p și de eșec $1 - p$ este [5,14]:

$$P(X = k) = C_n^k p^k (1 - p)^{n-k}, \quad (3.33)$$

unde $C_n^k = \frac{n!}{k!(n-k)!}$, reprezintă combinații de n luate câte k .

Exemple de experimente ce pot fi modelate cu o repartiție binomială:

- Un medicament are probabilitatea de 90% de a vindeca o anumită maladie. Medicamentul se administrează la 100 de pacienți, iar în final aceștia sunt vindecați sau nu. Dacă X este numărul de pacienți

vindecați, X este o variabilă aleatoare binomială cu parametrii (100; 0.9); adică $n = 100$, $p = 0.9$.

- Institutul Național de Statistică estimează că există șansa ca 20% din adulți să fie simpatizanți al unui anumit partid. 50 de adulți sunt selectați aleator. Dacă X reprezintă numărul adulților simpatizanți ai respectivului partid, X va fi o variabilă aleatoare binomială cu parametrii (50; 0.20).
- Un producător de plăci de calculator are în medie 5% produse defecte. Pentru a monitoriza procesul de producție se extrage un eșantion de 75 de elemente. Dacă acest eșantion prezintă mai mult de 5 unități de produs defecte, procesul de producție se oprește. Numărul de produse defecte se poate modela cu o variabilă aleatoare binomială cu parametrii (75; 0.05)

Aplicație: 10 muncitori lucrează la un atelier și 6 din ei au categoria a 5-a. Care este probabilitatea ca într-un grup format din 3 muncitori să fie toți de categoria a 5-a.

$$P(X = 3) = C_{10}^3 \left(\frac{6}{10}\right)^3 \left(\frac{4}{10}\right)^{10-3} = 120 \cdot 0.126 \cdot 0.00164 = 0.425$$

Există o probabilitate de 42.5% ca toți cei trei muncitori să fie de categoria a cincea.

În numeroase situații practice, trebuie calculată probabilitatea de succes care să fie mai mare cel mult egală cu o valoare dată, respectiv mai mică sau egală decât o valoare dată: $P(x \geq k)$, $P(x \leq k)$. În astfel de situații se însumează probabilitățile corespunzătoare.

Pentru exemplul anterior să se determine probabilitatea ca nu mai mult de doi muncitori selectați să fie de categoria a cincea.

$$\begin{aligned} P(X \leq 2) &= P(X = 0) + P(X = 1) + P(X = 2) = C_{10}^0 \left(\frac{6}{10}\right)^0 \left(\frac{4}{10}\right)^{10} + C_{10}^1 \left(\frac{6}{10}\right)^1 \left(\frac{4}{10}\right)^9 + \\ &+ C_{10}^2 \left(\frac{6}{10}\right)^2 \left(\frac{4}{10}\right)^8 = 0.0001 + 0.0016 + 0.0106 = 0.0123. \end{aligned}$$

În cazul distribuției binomiale media și varianța sunt:

$$M(x) = np, \quad (3.34)$$

$$\sigma^2(x) = np(1-p).$$

În Matlab există funcția `binocdf(a,n,p)` ce returnează funcția de repartiție $F(a)$ a unei variabile aleatoare binomiale de parametri cunoscuți n și p . Pentru exemplul anterior se putea calcula probabilitatea ca nu mai mult de doi muncitori să fie de categoria a cincea, $P(x \leq 2)$, cu `binocdf(2,10,.6)`, ce returnează valoarea 0.0123.

Funcția densitate de probabilitate pentru repartiția binomială este definită de `binopdf(x,n,p)` care returnează valoarea acestei funcții în punctul x . Rezultatul anterior putea fi obținut și prin însumarea funcțiilor densitate de probabilitate pentru valorile: 0, 1 și 2;

```
prob=binocdf(0:2,10,0.6)
```

returnează aceeași valoare 0.0123.

3.3.2. DISTRIBUȚIA POISSON

Este o distribuție ce caracterizează variabile discrete. Ea tratează problema evenimentelor aleatoare ce au loc în unitatea de timp sau spațiu. Ea poate înlocui distribuția binomială în cazul când numărul de evenimente total este foarte mare, iar șansa de realizare favorabilă foarte mică. Principala utilizare a distribuției Poisson este la probleme ce tratează evenimente rare, ce apar într-un interval de timp specificat sau într-o regiune din spațiu. Ex.: numărul de clienți dintr-un magazin/oră, numărul de mașini ce intră într-o parcare/zi, numărul de scurgeri de-a lungul unei conducte de petrol, numărul de defecte/unitatea de suprafață la o tablă etc.

Distribuția Poisson este aplicabilă în anumite condiții:

- evenimentele aleatoare au loc în unitatea de timp sau spațiu;
- numărul de evenimente favorabile trebuie să fie teoretic infinit;
- producerea unui eveniment este independentă de producerea evenimentelor anterioare sau posterioare.

Dacă aceste condiții sunt îndeplinite, variabila aleatoare ce contorizează numărul de apariții a evenimentului într-un interval de timp este o variabilă aleatoare de tip Poisson. În cazul când evenimentele se produc astfel încât în medie apar λ evenimente într-o perioadă de timp sau spațiu, atunci

probabilitatea ca evenimentul să se producă de k ori este:

$$P(X = k) = \frac{e^{-\lambda} \lambda^k}{k!}. \quad (3.35)$$

Pentru calcularea probabilității ca evenimentele să se producă de cel mult k ori, se procedează la fel ca în cazul distribuției binomiale:

$$P(X \leq k) = P(X = 0) + P(X = 1) + \dots + P(X = k)$$

Dacă X este o variabilă aleatoare de tip Poisson, cu parametrul λ , atunci media sa este:

$$M(X) = \lambda, \quad (3.36)$$

iar dispersia:

$$\sigma^2(X) = \lambda. \quad (3.37)$$

Aplicații:

1. Într-un studiu la o spălătorie de mașini s-a constatat că numărul mediu de mașini ce sosesc luni dimineața între ora 8 și 9 este 5. Care este probabilitatea ca într-o anumită zi de luni să sosească exact 5 mașini? Dar care este probabilitatea să vină mai puțin de 3 mașini?

$$P(X = 5) = \frac{e^{-5} 5^5}{5!} = 0.1755.$$

În primul caz probabilitatea este 17.55%.

$$P(X < 3) = P(X = 0) + P(X = 1) + P(X = 2) = 0.0067 + 0.0337 + 0.0842 = 0.1246$$

În cel de-al doilea caz, probabilitatea este de 12.46%.

2. La editarea unei cărți se utilizează corectarea automată oferită de soft-ul de procesare a textului, în plus la editare se face o nouă corectură. Cu toate acestea un anumit număr de greșeli de editare rămân. Se consideră că numărul de erori tipografice per pagină are o repartiție Poisson cu parametrul $\lambda = 0.25$. Se dorește calcularea probabilității ca pe o pagină să apară cel puțin 2 erori.

$$P(X \geq 2) = 1 - P(X = 0) - P(X = 1) = 1 - e^{-0.25} - .25e^{-0.25} = 0.0265$$

Funcția de repartiție Poisson se poate apela în Matlab cu comanda `poisscdf`. Sintaxa este `poisscdf(x,lambda)` - returnează probabilitatea ca variabila să fie mai mică decât o valoare x. Deci în cazul problemei anterioare, rezultatul se poate obține cu:

```
probabilitate=1-poisscdf(1,0.25)
```

3. Într-o intersecție s-a constatat că apar în medie 2 accidente pe săptămână ($\lambda = 2$). Să se determine care este probabilitatea ca în următoarele 2 săptămâni să aibă loc 3 accidente.

În cazul distribuției Poisson numărul de evenimente ce au loc într-un interval depind doar de lungimea intervalului și sunt independente de punctul de început. Numărul de evenimente ce apar în mai multe intervale de timp, k, este egal cu $k\lambda$. În aceste condiții se poate determina probabilitatea de apariție a 3 accidente în 2 săptămâni ca fiind:

```
probabilitate=poisscdf(3,2*2)
```

3.3.3. DISTRIBUȚIA UNIFORMĂ

Variabilele aleatoare discrete se asociază unor experimente ce constau în contorizarea anumitor rezultate. Variabilele continue se asociază, în general, unor experimente ce constau din măsurare, deci au rezultate numere reale, valorile variabilei putând fi oricare într-un interval mărginit sau nu. Ex.: cota unei piese, înălțimea sau greutatea populației, rata șomajului etc.

După cum s-a arătat în paragraful 3.1, dacă se cunoaște funcția de repartiție a variabilei aleatoare continue, X, se poate calcula probabilitatea de apariție a oricărui eveniment: $X < a$, $a < X < b$, $X > b$, $a \leq X \leq b$ etc.

Fie X o variabilă aleatoare continuă, ce are densitatea de probabilitate f. valoarea medie se notează M(X) și se calculează:

$$M(X) = \int_{-\infty}^{\infty} xf(x)dx. \quad (3.38)$$

Dispersia sau varianța unei variabile cu media m este:

$$\sigma^2(X) = M((X - m)^2) = \int_{-\infty}^{\infty} (x - m)^2 f(x)dx, \quad (3.39)$$

iar abaterea standard este:

$$\sigma(X) = \sqrt{\sigma^2(X)}. \quad (3.40)$$

O variabilă aleatoare distribuită uniform pe intervalul (a, b) are funcția densitate de probabilitate de forma:

$$f(X; a, b) = \frac{1}{b-a}; \quad a < X < b \quad (3.41)$$

Parametrii pentru repartiția uniformă sunt capetele intervalului: a și b. Media și varianța unei variabile repartizată uniform sunt:

$$M(X) = \frac{a+b}{2}, \quad (3.42)$$

$$\sigma^2(X) = \frac{(b-a)^2}{12} \quad (3.43)$$

Funcția de probabilitate a unei variabile distribuită uniform este:

$$F(X) = \begin{cases} 0; & X \leq a \\ \frac{x-a}{b-a}; & a < X < b \\ 1; & X \geq b. \end{cases} \quad (3.44)$$

În Matlab există definite funcția densitate de probabilitate, f, (`unifpdf`) și funcția de repartiție, F, (`unifcdf`) pentru repartiția uniformă.

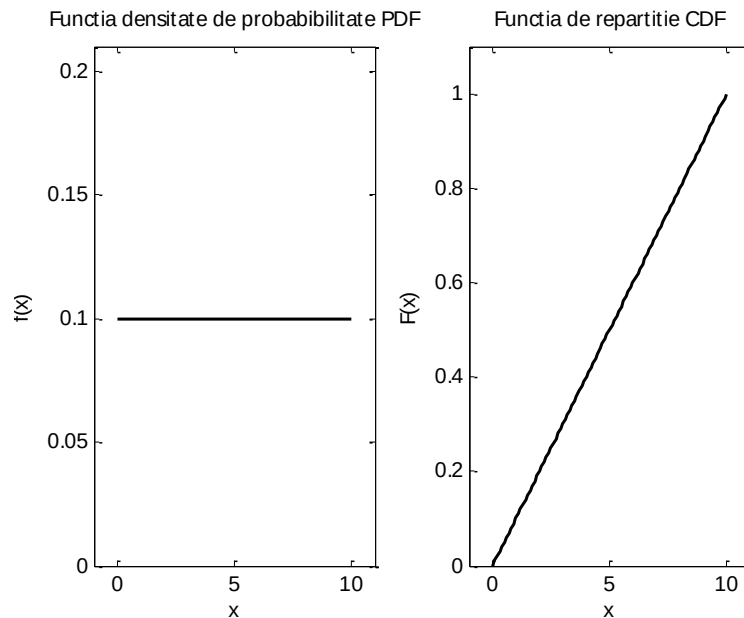


Fig. 3.12 Funcția densitate de probabilitate și funcția de probabilitate pentru o variabilă repartizată uniform în intervalul [0, 10]

În fig. 3.12 se prezintă graficul funcției densitate de probabilitate și funcției de probabilitate pentru o variabilă distribuită uniform în intervalul [0,10]. Reprezentarea grafică a fost obținută cu următoarea secvență de program:

```
x=-1:.1:11;%se divizează intervalul [-1,11] cu pasul h=0.1
pdf=unifpdf(x,0,10);
cdf=unifcdf(x,0,10);
subplot(1,2,1),plot(x,pdf),title('Funcția densitate de
probabilitate PDF')
xlabel('x'),ylabel('f(x)'),axis([-1 11 0 .21])
subplot(1,2,2),plot(x,cdf),title('Funcția de repartiție CDF')
xlabel('x'),ylabel('F(x)'),axis([-1 11 0 1.1])
```

3.3.4. DISTRIBUȚIA NORMALĂ

Distribuția normală este o distribuție fundamentală în statistică, fiind adecvată în modelarea a numeroase fenomene din natură. De asemenea, ea stă la baza inferenței statistice. Funcția de densitate de probabilitate are expresia:

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} \cdot \exp\left(-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right), \quad (3.45)$$

unde μ este media, iar σ abaterea standard. Graficul acestei funcții este un clopot. Cu cât σ este mai mic, cu atât repartiția este mai concentrată (clopotul este mai ascuțit). În figura 3.13 se prezintă câteva exemple de funcții densitate de probabilitate pentru repartiții normale. Se observă că pe măsură ce σ crește înălțimea funcției scade, dar crește anvergura acesteia.

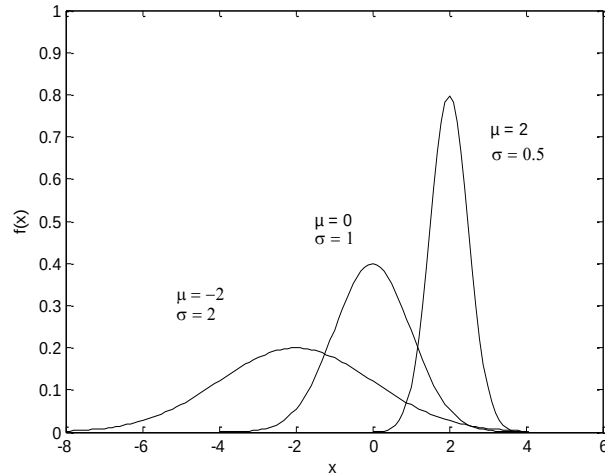


Fig. 3.13 Exemple de funcții densitate de probabilitate pentru distribuții normale

Localizarea și forma repartiției depind de parametrii μ și σ , în sensul că μ determină localizarea repartiției pe axă, iar σ forma curbei (fig. 3. 14).

Densitatea de probabilitate este continuă, are formă de clopot și tinde asimptotic spre 0 pentru $x \rightarrow \pm\infty$. Notăția $X \sim N(\mu, \sigma^2)$ se utilizează pentru a indica că o variabilă aleatoare X este distribuită normal cu media μ și varianța σ^2 .

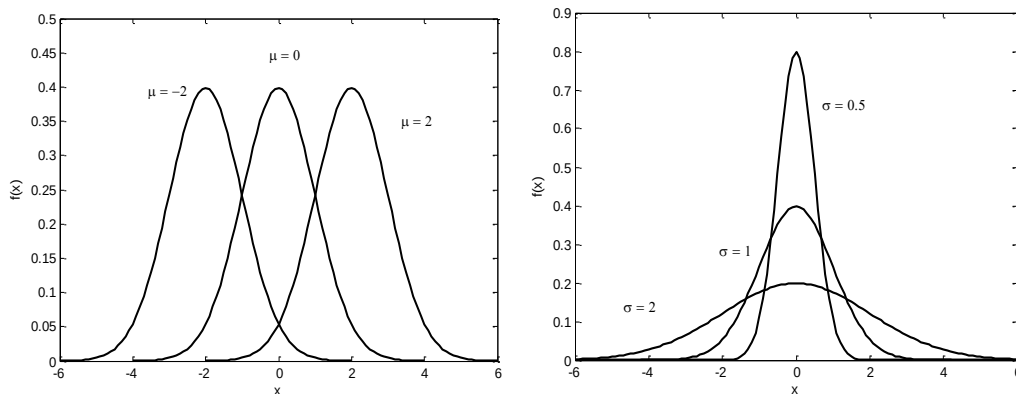


Fig. 3.14. Influența parametrilor asupra repartiției normale

Pentru orice distribuție normală, indiferent de valoarea parametrilor μ și σ , proporția de observații ce aparțin unui interval centrat în μ este aceeași (fig. 3.15):

- 68.26% din valorile lui $x \in [\mu - \sigma; \mu + \sigma]$;
- 95.44% din valorile lui $x \in [\mu - 2\sigma; \mu + 2\sigma]$;
- 99.73% din valorile lui $x \in [\mu - 3\sigma; \mu + 3\sigma]$.

Orice variabilă aleatoare normală poate fi transformată în variabilă aleatoare normală normată prin schimbarea de variabilă: $z = (x - \mu)/\sigma$, deci variabila aleatoare z are o distribuție normală cu media egală cu 0 și dispersia 1. În acest mod se poate calcula aria închisă sub curbă pentru diferite valori ale lui x , arie ce este egală cu probabilitatea ca variabila să aparțină intervalului respectiv (delimitat de limitele de integrare). Aceste valori sunt incluse în tabele statistice (Anexa 1). De remarcat că în tabel sunt înregistrate doar ariile pentru valori pozitive. În cazul când se calculează probabilitatea ca $z \in [-a, a]$, valoarea rezultată din tabel trebuie multiplicată cu 2.

Funcția de repartiție a unei variabile repartizate normal se definește cu relația:

$$\Phi(z) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^z \exp\left(-\frac{y^2}{2}\right) dy \quad (3.46)$$

Funcția de probabilitate se poate calcula utilizând funcția eroare, notată cu erf.

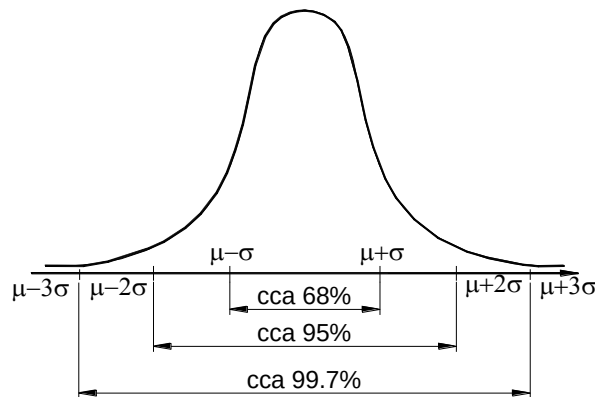


Fig. 3.14. Repartiția normală

În Matlab sunt definite funcțiile:

- `normpdf` ce returnează valorile funcției densitate de probabilitate cu sintaxa:
`y = normpdf(x,miu,sigma)`

unde `miu` este valoarea mediei, iar `sigma` este deviația standard; implicit `miu` este 0, iar `sigma` 1;

- `normcdf`, ce returnează valoarea funcției de repartiție. Sintaxa este `p=normpdf(x,miu,sigma)`
- `normspec` este o funcție ce permite calcularea probabilității ca o variabilă repartizată normal cu media μ și deviația standard σ să aibă valori între anumite limite. Funcția reprezintă grafic densitatea de probabilitate restricționată la intervalul definit de cele două limite. Limitele se indică sub forma unui vector. În cazul când nu există limită inferioară, primul element al vectorului se introduce $-\infty$, iar în cazul inexistenței limitei superioare, al doilea element al vectorului va fi ∞ . Acest lucru se realizează în Matlab cu `-Inf`, respectiv `Inf`. Un exemplu se prezintă în continuare. Valorile implicite pentru μ și σ sunt 0 și 1.
- `norminv` este o funcție ce returnează inversa funcției de repartiție în punctul P. Sintaxa este `norminv(P,miu,sigma)`, unde valorile implicite sunt 0 (pentru μ) și 1 (pentru σ).

Pentru cazul unei variabile distribuită normal, având media 4 și deviația standard egală cu 1.5, probabilitatea ca variabila să aparțină intervalului [2, 7] se poate calcula în felul următor:

```
% se seteaza limitele intervalului in vectorul limite
limite=[2,7];
miu=4;sigma=1.5;
prob=normspec(limite,miu,sigma)
```

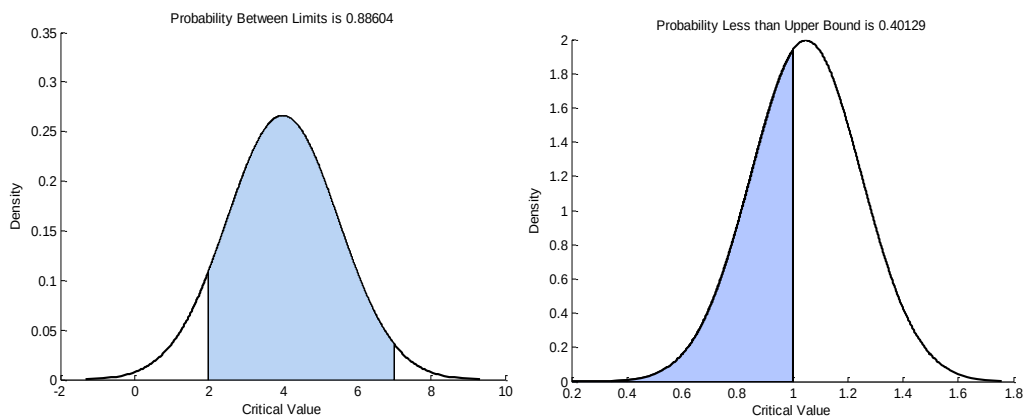


Fig. 3.15 Exemple pentru funcția normspec

Într-o firmă producătoare de alimente, produsele sunt ambalate în cutii de 1 kg. În realitate masa cutiilor este normal distribuită, cu medie 1.05 kg și abatere standard de 0.2. Să se determine probabilitatea să existe cutii sub greutatea

specificată.

În acest caz trebuie determinată probabilitatea ca variabila masă să ia valori în intervalul $(-\infty, 10)$.

```
limite=(-Inf, 1);  
miu=1.05;sigma=.2;  
prob=normspec(limite,miu,sigma);
```

3.3.5. DISTRIBUȚIA T

O altă distribuție cu numeroase aplicații în statistică este distribuția t sau Student. În cazul unei populații distribuită normal, dacă se extrag eșantioane și se calculează media acestora, variabila t calculată cu relația:

$$t = \frac{\bar{x} - \mu}{S / \sqrt{n}}, \quad (3.48)$$

are repartiția t.

Distribuția t este simetrică, are media 0 și o varianță mai mare decât 1 [5]. De fapt, varianța distribuției crește pe măsură ce n, volumul eșantionului scade. Distribuția admite ca parametru numărul gradelor de libertate, $v = n - 1$. Pe măsură ce numărul gradelor de libertate crește, distribuția se apropie de cea normală.

La fel ca în cazul repartiției normale standard, există tabele pentru distribuția t ce permit extragerea valorilor t_α corespunzătoare unei anumite probabilități α . Fiecare rând al tabelului corespunde unui anumit număr de grade de libertate. Probabilitățile sunt calculate pentru o singură ramură a distribuției, $P(t \geq t_\alpha)$. t_α este α -cvantila repartiției Student, adică numărul cu proprietatea că $P(t \leq t_\alpha) = \alpha$.

Dacă se compară din tabele, distribuția normală standard și distribuția t, se constată că pentru valori ale lui n mai mari cele două distribuții sunt aproape identice. Pe baza acestui considerent, în cazurile practice, când $n \geq 30$ se înlocuiește distribuția t cu cea normală standard, iar distribuția t se utilizează mai ales în cazurile când $n < 30$. Mai trebuie menționat faptul că distribuția t necesită ca populația din care s-a extras eșantionul să nu difere mult de repartiția normală. În cazul eșantioanelor mici extrase din populații asimetrice, nici distribuția z și nici t nu pot fi utilizate.

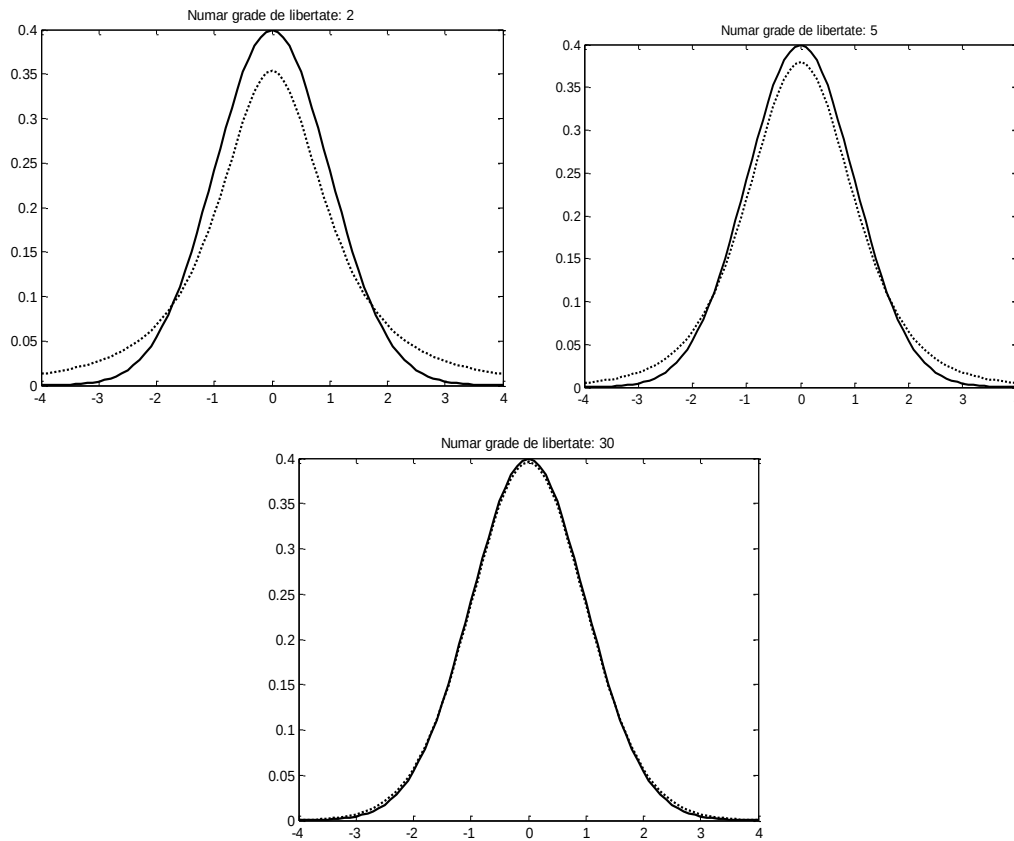


Fig. 3.16 Aproximarea dintre repartiția t și repartiția normală pe măsură ce numărul gradelor de libertate crește

În concluzie se poate afirma că distribuția t se utilizează în următoarele condiții:

- deviația standard a eșantionului, S , este utilizată pentru estimarea lui σ ;
- volumul eșantionului este mic, $n < 30$;
- populația este distribuită aproximativ normal.

În Matlab sunt definite funcțiile `tpdf`, respectiv `tcdf` pentru determinarea densității de probabilitate și a probabilității unei variabile repartizată conform repartiției t. Sintaxa este: $y = \text{tpdf}(x, v)$, respectiv $P = \text{tcdf}(x, v)$, unde v este numărul gradelor de libertate. Este definită și funcția `tinverse(P, v)`, inversa funcției de repartiție, ce returnează cvantila repartiției t cu v grade de libertate în punctul P .

3.3.6. DISTRIBUȚIA EXPONENȚIALĂ

O variabilă aleatoare ce are densitatea de probabilitate:

$$f(x; \lambda) = \lambda e^{-\lambda x}; x \geq 0; \lambda > 0. \quad (3.49)$$

Se numește variabilă aleatoare exponențial distribuită de parametru λ . O variabilă aleatoare exponențial distribuită modelează durata de viață a unui produs sau dispozitiv sau lungimea intervalelor de timp între producerea a două evenimente consecutive contorizate de o variabilă aleatoare Poisson. Ex.: intervalelor de timp între sosirea a doi clienți la o bancă, timpul dintre apelurile telefonice, timpul până la căderea unei piese etc.

Media și varianța unei populații repartizată exponențial sunt:

$$M(x) = \frac{1}{\lambda} \quad (3.50)$$

$$\sigma^2(x) = \frac{1}{\lambda^2}. \quad (3.51)$$

Funcția de probabilitate este de forma:

$$F(x) = \begin{cases} 0; & x < 0 \\ 1 - e^{-\lambda x}; & x \geq 0 \end{cases} \quad (3.52)$$

Repartiția exponențială prezintă o proprietate interesantă și anume intervalul de timp scurs până la apariția unui eveniment nu depinde de timpul anterior [10]. Acest lucru se poate scrie:

$$P(X > s + t | X > s) = P(X > t) \quad (3.53)$$

Cu alte cuvinte, acest lucru înseamnă că probabilitatea ca o piesă sau dispozitiv să fie în stare de funcționare după de $s + t$ unități de timp, dacă ea a funcționat bine deja de timpul s , este egală cu probabilitatea ca piesa să funcționeze timpul t .

Când această distribuție se utilizează la reprezentarea intervalelor de timp între evenimente, parametrul λ reprezintă numărul mediu de evenimente produs în unitatea de timp, adică frecvența acestora.

Exemplu: Intervalul de timp între sosirea mașinilor într-o intersecție este în medie de 12 secunde. Se cere probabilitatea ca intervalul de timp să fie cel mult egal cu 10 secunde. În acest caz numărul mediu al mașinilor este egal cu $1/12$, iar probabilitatea cerută este egală cu:

$$P(x \leq 10) = 1 - e^{-(1/12)10} = 0.5654$$

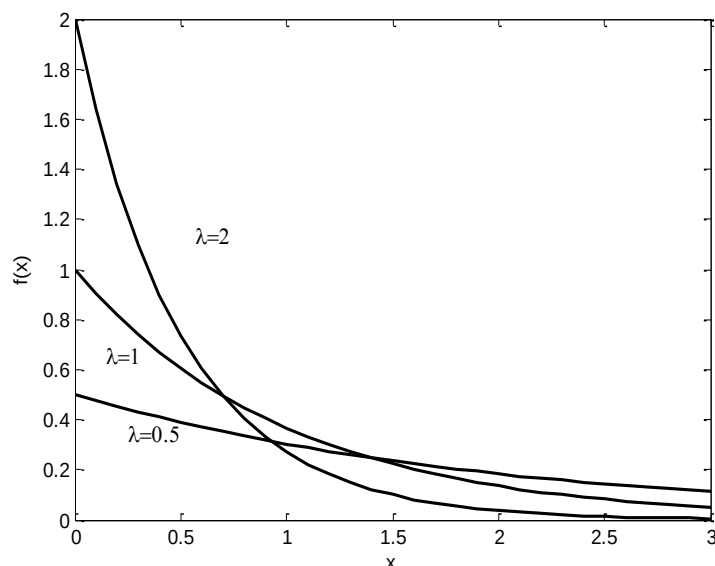


Fig. 3.17 Repartiția exponențială pentru diferite valori ale parametrului λ

În Matlab funcția de repartiție pentru distribuția exponențială este **expcdf(x,1/lambda)** cu care se poate calcula probabilitatea ca X să ia valori mai mici sau egale cu x. Trebuie menționat faptul că în Matlab funcția de probabilitate pentru repartiția exponențială se definește cu relația:

$$f(x; \mu) = \frac{1}{\mu} e^{-\frac{x}{\mu}}; x \geq 0; \mu > 0. \quad (3.54)$$

În acest context rezultatul aplicației anterioare se poate obține prin simpla apelare a funcției **expcdf(10,12)**, ce va returna 0.5654. De asemenea este definită funcția **exppdf(x, 1/lambda)**, ce implementează calculul densității de probabilitate, respectiv **expinv(p,1/lambda)** inversa funcției de repartiție.

3.3.7. DISTRIBUȚIA GAMMA

Funcția densitate de probabilitate a repartiției gamma are expresia:

$$f(x; \lambda, t) = \frac{\lambda e^{-\lambda x} (\lambda x)^{t-1}}{\Gamma(t)}; x \geq 0, \quad (3.55)$$

unde t este parametru de formă, iar λ este parametru de localizare. Funcția gamma este definită de:

$$\Gamma(t) = \int_0^{\infty} e^{-y} y^{t-1} dy, \quad (3.56)$$

relație care în cazul valorilor întregi ale parametrului t devine:

$$\Gamma(t) = (t - 1)! \quad (3.57)$$

De remarcat că pentru $t = 1$ distribuția gamma devine exponențială. Pentru valori pozitive ale lui t , repartiția gamma se poate utiliza la modelarea intervalului de timp scurs până la producerea a t evenimente, dacă numărul evenimentelor rare are distribuție Poisson.

Media și varianța distribuției gamma sunt date de relațiile:

$$M(x) = \frac{t}{\lambda} \quad (3.58)$$

$$\sigma^2(x) = \frac{t}{\lambda^2}. \quad (3.59)$$

Funcția de repartiție a acestei distribuții are expresia [10]:

$$F(x; \lambda, t) = \begin{cases} 0; & x \leq 0 \\ \frac{1}{\Gamma(t)} \int_0^{\lambda x} y^{t-1} e^{-y} dy; & x > 0 \end{cases} \quad (3.60)$$

Relația (3.60) se poate evalua în Matlab cu funcția `gammacdf(lambda*x, t)`.

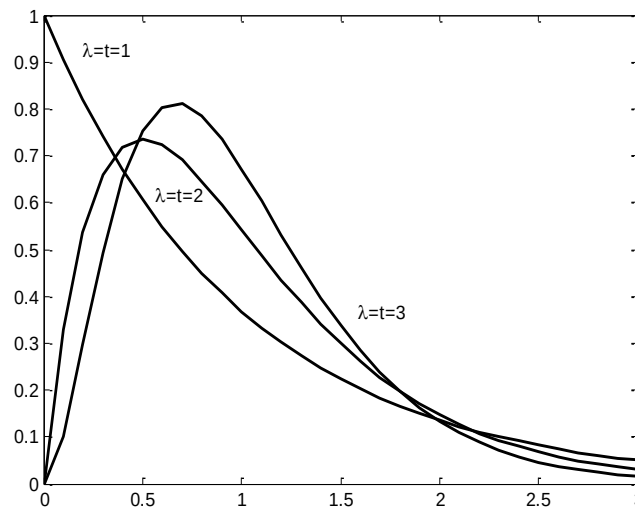


Fig. 3.18 Exemple de funcții densitate de probabilitate gamma

Alura densității de probabilitate gamma pentru diverse valori ale parametrilor se poate observa în fig. 3.18.

Graficul prezentat în figură se obține cu secvența de program:

```
x=0:.1:3;
y1=gampdf(x,1,1);
y2=gampdf(x,2,1/2);
y3=gampdf(x,3,1/3);
plot(x,y1,x,y2,x,y3)
```

La fel ca în cazurile anterioare, în Matlab există definite funcțiile gampdf, respectiv gamcdf, ce pot fi utilizate pentru calcularea densității de probabilitate și funcției de repartiție a distribuției gamma.

3.3.8. DISTRIBUȚIA χ^2

Distribuția χ^2 este un caz particular al distribuției gamma, având $\lambda = 0.5$ și $t = v/2$, cu v o valoare întreagă pozitivă, ce reprezintă numărul gradelor de libertate. Repartiția χ^2 are aplicații în testele de concordanță, cu ajutorul cărora se verifică ipoteza modelării eșantionului cu o anumită repartiție.

Funcția densitate de probabilitate pentru o variabilă aleatoare distribuită conform repartiției χ^2 , cu v grade de libertate este;

$$f(x, v) = \frac{1}{\Gamma(v/2)} \left(\frac{1}{2}\right)^{v/2} x^{v/2-1} e^{-\frac{x}{2}}; \quad x \geq 0. \quad (3.61)$$

Media și varianța distribuției se pot obține pe baza repartiției gamma, având valorile:

$$M(x) = v \quad (3.62)$$

$$\sigma^2(x) = 2v. \quad (3.63)$$

3.4 INFERENȚA STATISTICĂ

Investigările statistice constau în studiul unor caracteristici ale unei populații. În acest scop, din populație (ce poate fi finită sau infinită) se extrag eșantioane de volum finit. În urma investigării eșantionului se obțin informații ce se extrapolează la întreaga populație.

Inferența statistică dezvoltă metode de investigare a unei populații prin sondaj, adică metode și tehnici de analiză a datelor obținute prin eșantioane. Având în

vedere că populația este caracterizată de câțiva indicatori de interes, statistica formulează inferențe, adică predicții, asupra parametrilor populației cu diferite nivele de încredere.

În practică inferența statistică se ocupă cu două mari categorii de probleme:

- estimarea valorilor parametrilor unei populații (media, dispersia);
- testarea ipotezelor referitoare la parametrii populației.

Spre deosebire de statistica descriptivă, în inferența statistică se asociază un model probabilist caracteristicii investigate, în sensul că valorile numerice asociate caracteristicii investigate au o anumită distribuție de probabilitate, ce este definită de o densitate de probabilitate, f. Această densitate de probabilitate depinde de un parametru necunoscut, $\theta \in R^n$, în cazurile cele mai frecvente n fiind 1 sau 2.

Pe baza tehnicilor dezvoltate de inferența statistică se estimează parametrul θ prin:

- o valoare punctuală;
- un interval de valori ce va conține valoarea estimată cu o probabilitate prescrisă;
- o modalitate de testare a ipotezelor $\theta \leq \theta_0, \theta \geq \theta_0$; unde θ_0 este o valoare dată.

3.4.1. EȘANTIONAREA

Dintr-o populație de volum N se pot extrage eșantioane de volum n , $n < N$. Numărul de astfel de eșantioane este mare și valoarea unei anumite statistici, de ex. media, calculată pentru fiecare eșantion va diferi de la un eșantion la altul. Frecvența de distribuție a tuturor acestor eșantioane dă informații primare despre distribuția unui eșantion.

Conceptul de eșantionare al unei populații poate fi ilustrat prin următorul exemplu. Fie o populație în care variabila poate lua orice valoare întreagă între 0 și 4 (incidența apariției a 4 defecte la un produs). Fie populația formată din 5 obiecte, fiecare având 0, 1, 2, 3 sau 4 defecte. Se extrag eșantioane formate din două exemplare. Eșantionarea se poate face cu sau fără returnare. La eșantionarea cu returnare, după extragerea unui element, acesta se introduce din nou în cadrul populației, putând fi ulterior extras.

În tabelul 2.2 se prezintă toate posibilitățile de eșantioane ce apar. Pentru

cazul eșantionării cu înlocuire se obține: $\mu_{\bar{x}} = 2.0$, $\sigma_{\bar{x}}^2 = 1.0$, iar în cazul eșantionării fără înlocuire: $\mu_{\bar{x}} = 2.0$, $\sigma_{\bar{x}}^2 = 0.75$. Deoarece s-au luat în considerare toate eșantioanele posibile se poate utiliza notația cu indicatorii populației (μ și σ) și nu ai eșantionului ($\bar{x}$ și S).

Se remarcă faptul că cele două medii sunt egale ($\mu_{\bar{x}} = 2.0$), dar varianța diferă ($\sigma_{\bar{x}}^2 = 1.0$, respectiv $\sigma_{\bar{x}}^2 = 0.75$).

Tabelul 2.2. Eșantionarea și prelucrarea datelor în cazul unei populații formate din 5 elemente

Eșantionarea cu returnare					Eșantionarea fără returnare				
0, 0	1, 0	2, 0	3, 0	4, 0	0, 1 0, 2 0, 3 0, 4	1, 2 1, 3 1, 4	2, 3 2, 4	3, 4	
0, 1	1, 1	2, 1	3, 1	4, 1					
0, 2	1, 2	2, 2	3, 2	4, 2					
0, 3	1, 3	2, 3	3, 3	4, 3					
0, 4	1, 4	2, 4	3, 4	4, 4					
Mediile fiecărui eșantion									
0.0	0.5	1.0	1.5	2.0	0.5 1.0 1.5 2.0	1.5 2.0 2.5	2.5 3.0	3.5	
0.5	1.0	1.5	2.0	2.5					
1.0	1.5	2.0	2.5	3.0					
1.5	2.0	2.5	3.0	3.5					
2.0	2.5	3.0	3.5	4.0					
Distribuția mediilor									
$\bar{x}$	Nr. eșantioane f		Probabilitatea de apariție $f/\Sigma f$		$\bar{x}$	Nr. eșantioane f		Probabilitatea de apariție $f/\Sigma f$	
0.0	1		1/25		0.5	1		1/10	
0.5	2		2/25		1.0	1		1/10	
1.0	3		3/25		1.5	2		2/10	
1.5	4		4/25		2.0	2		2/10	
2.0	5		5/25		2.5	2		2/10	
2.5	4		4/25		3.0	1		1/10	
3.0	3		3/25		3.5	1		1/10	
3.5	2		2/25						
4.0	1		1/25						
Total	25				10				

Următorul pas este găsirea relațiilor dintre indicatorii eșantionului și cei ai populației. În exemplul prezentat, media populației este:

$$\mu = (0+1+2+3+4)/5 = 2.0,$$

valoare ce este egală cu media determinată în cele două cazuri de eşantionare.

În tabelul 2.3 se prezintă rezultatele obținute în cazul extragerii de eşantioane de 3, respectiv 4 elemente.

Pentru eşantioanele de volum $n = 3$ se obține $\mu_{\bar{x}} = 2.0$, $\sigma_{\bar{x}}^2 = 0.66$, iar pentru cele cu $n = 4$, $\mu_{\bar{x}} = 2.0$, $\sigma_{\bar{x}}^2 = 0.5$. Se remarcă faptul că media în toate situațiile este egală cu media populației. De asemenea, se observă că mediile eşantioanelor se distribuie simetric în jurul valorii medii ideale, μ , și pe măsură ce volumul eşantionului crește frecvența de apariție a valorilor apropiate de media ideală crește. În cazul dispersiei și abaterii standard apar diferențe între parametrii populației și cei ai eşantioanelor. Pe măsură ce volumul eşantionului crește, dispersia și abaterea standard a mediei scade. Datorită variației mediei eşantioanelor în jurul mediei populației datorită eşantionării, o eroare de eşantionare apare ori de câte ori o media unui singur eşantion, $\bar{x}$ este utilizată pentru estimarea mediei populației, μ . Din acest motiv termenul de eroare standard se folosește uneori în locul abaterii standard, când se referă la distribuția mediei ($\sigma_{\bar{x}}$).

Tabelul 2.3. Distribuția mediilor la eşantionarea cu înlocuire

Volumul eşantionului $n = 3$			Volumul eşantionului $n = 4$		
$\bar{x}$	Nr. eşantioane f	Probabilitate de apariție	$\bar{x}$	Nr. eşantioane f	Probabilitate de apariție
0.00	1	0.008	0.00	1	0.0016
0.33	3	0.024	0.25	4	0.0064
0.67	6	0.048	0.50	10	0.0160
1.00	10	0.080	0.75	20	0.0320
1.33	15	0.120	1.00	35	0.0560
1.67	18	0.144	1.25	52	0.0832
2.00	19	0.152	1.50	68	0.1088
2.33	18	0.144	1.75	80	0.1280
2.67	15	0.120	2.00	85	0.1360
3.00	10	0.080	2.25	80	0.1280
3.33	6	0.048	2.50	68	0.1088
3.67	3	0.024	2.75	52	0.0832
4.00	1	0.008	3.00	35	0.0560
			3.25	20	0.0320
			3.50	10	0.0160
			3.75	4	0.0064
			4.00	1	0.0016
Total	125	1.00	Total	625	1.0000

În cazul dispersiei, se observă că pe măsură ce volumul eşantionului creşte varianţa scade. Varianţa populaţiei este:

$$\sigma^2 = \frac{\sum (x_i - \bar{x})^2}{N} = 2.00$$

La eşantionarea cu returnare, relaţia dintre $\sigma_{\bar{x}}^2$ şi σ^2 este: $\sigma_{\bar{x}}^2 = \frac{\sigma^2}{n}$.

La eşantionarea cu returnare, practic populaţia are un volum infinit, deoarece după fiecare extragere, elementele se introduc din nou, volumul populaţiei rămânând acelaşi. La eşantionarea fără returnare, volumul populaţiei scade pe măsură ce se extrag elemente şi, ca urmare, dispersia distribuţiei scade. Din acest motiv această variantă de extragere a eşantioanelor este mult mai eficientă în colectarea informaţiei. Factorul de scădere a dispersiei depinde de mărimea populaţiei şi de volumul eşantionului, factorul de corecţie datorat numărului finit al populaţiei fiind $(N - n)/(N - 1)$. În acest caz, relaţia dintre $\sigma_{\bar{x}}^2$ şi σ^2 este:

$$\sigma_{\bar{x}}^2 = \frac{\sigma^2}{n} \left(\frac{N - n}{N - 1} \right). \quad (3.64)$$

În concluzie, se poate afirma că în cazul eşantionării cu returnare $\sigma_{\bar{x}}^2 = \frac{\sigma^2}{n}$, iar

în cazul eşantionării fără returnare $\sigma_{\bar{x}}^2 = \frac{\sigma^2}{n} \left(\frac{N - n}{N - 1} \right)$.

În cazul când volumul populaţiei este foarte mare şi volumul eşantionului este mic, valoarea factorului de corecţie poate fi ignorată. De obicei, în situaţia când $n \leq 0.05N$, factorul de corecţie poate fi ignorat şi se poate considera [5]:

$$\begin{aligned} \mu_{\bar{x}} &= \mu \\ \sigma_{\bar{x}} &= \frac{\sigma}{\sqrt{n}}. \end{aligned} \quad (3.65)$$

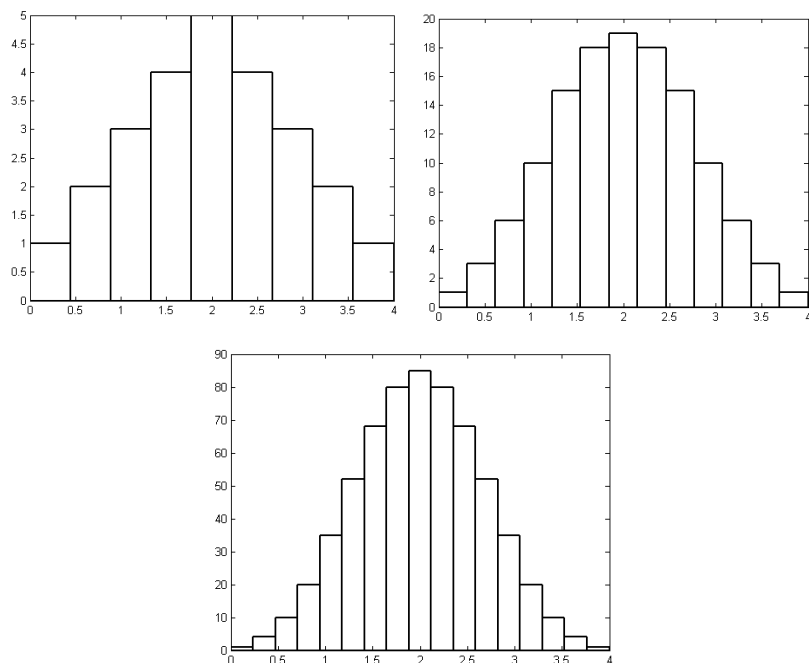


Fig. 3. 19. Histogramele pentru eșantionarea cu 2, 3 și 4 elemente

În fig. 3.19 se prezintă histogramele pentru cele trei serii de date asociate mediilor $\bar{x}$ ale eșantioanelor cu $n = 2$, $n = 3$, respectiv $n = 4$ elemente. Se observă că pe măsură ce n crește, histograma aproximează din ce în ce mai bine clopotul se definește repartiția normală. Această tendință este adevărată pentru toate eșantioanele, indiferent de distribuția populației. S-a demonstrat că pe măsură ce volumul eșantionului crește, distribuția tinde spre una normală, chiar în situația când populația are o altfel de distribuție. Această proprietate este enunțată în teorema limită centrală.

În consecință, proprietățile repartiției normale pot fi aplicate pentru determinarea erorii ce apare datorită eșantionării. În practică, teorema limită centrală se poate aplica la eșantioane de volum $n \geq 30$. În cazul când populația are o distribuție normală, media eșantionului, indiferent de volumul său, va avea o repartiție normală.

3.4.2. ESTIMĂRI

În cadrul acestui paragraf se tratează problema estimării indicatorilor statistici ai populației pe baza eșantionării și determinarea gradului de încredere a

acestei estimări. Estimarea mediei se poate face printr-o estimare punctuală sau cu un interval de încredere.

În cazul estimării punctuale se consideră abaterea standard σ a populației cunoscută. În estimarea punctuală, se utilizează datele dintr-un eșantion pentru a calcula valoarea indicatorului statistic medie, cu alte cuvinte media experimentală, $\bar{x}$, este un estimator punctual pentru μ .

Pentru ca estimarea să fie bună, este necesar să îndeplinească trei condiții [5]:

- să fie nedeplasată, o estimare este nedeplasată dacă media eșantionului (în acest caz $\mu_{\bar{x}}$) este egală cu parametrul estimat (μ);
- este eficientă, adică deviația standard să fie cât mai mică posibil;
- este consistentă, condiție ce este îndeplinită dacă valoarea estimată tinde spre valoarea adevărată pe măsură ce volumul eșantionului crește.

Media unui eșantion îndeplinește aceste condiții, deci poate fi considerată o estimare punctuală. Trebuie reținut faptul că o singură estimare a mediei diferă de media populației, datorită erorii de eșantionare. Din acest motiv, în multe situații se preferă estimarea printr-un interval de încredere.

Estimarea intervalului de încredere pentru media μ a distribuției normale se bazează pe un eșantion de volum n și medie $\bar{x}$. Intervalul de încredere, de nivel de încredere de 95% pentru media μ a populației este:

$$\mu = \bar{x} \pm 1.96\sigma_{\bar{x}} = \bar{x} \pm 1.96 \frac{\sigma}{\sqrt{n}}. \quad (3.66)$$

Se pot utiliza diferite nivele de încredere în estimarea intervalului valorii medii. În cazul general, nivelul de încredere se notează $100(1 - \alpha)\%$, unde α reprezintă suma celor două arii din extremele repartiției normale (vezi fig. 3.20). Pentru o probabilitate de 95%, $\alpha = 0.05$. În cazul general, media cu incertitudinea asociată este:

$$\mu = \bar{x} \pm z_{1-\frac{\alpha}{2}}\sigma_{\bar{x}} = \bar{x} \pm z_{1-\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}, \quad (3.67)$$

unde $100(1 - \alpha)$ este nivelul de încredere, iar $z_{1-\alpha/2}$ este valoarea variabilei Z ce exclude la fiecare extremitate a distribuției o arie de $\alpha/2$, adică $\alpha/2$ -cvantila distribuției normale standard.

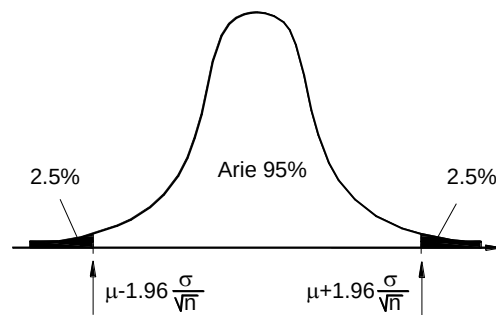


Fig. 3.20. Stabilirea intervalului de încredere

Pentru o probabilitate de 90%, $Z_{\alpha/2} = 1.65$ și pentru $P = 99\%$, $Z = 2.58$.

În mod uzual, σ este o mărime necunoscută, singurele informații certe sunt cele legate de eșantion. Se poate demonstra, [5, 10], că S este o estimare nedeplasată a abaterii standard. În cazul eșantioanelor de volum mare, estimarea este foarte bună, însă în cazul eșantioanelor de volum mic, S poate fi subestimat. Din acest motiv la eșantioanele de volum mic, trebuie lărgit intervalul de încredere. Acest lucru se realizează prin utilizarea repartiției t .

Similar cu construirea intervalului de încredere pe baza repartiției normale standard $\mu = \bar{x} \pm z_{1-\alpha/2} \cdot \frac{\sigma}{\sqrt{n}}$, în cazul utilizării distribuției t intervalul este:

$$\mu = \bar{x} \pm t(1 - \alpha/2, n-1) \cdot S_{\bar{x}} = \bar{x} \pm t(1 - \alpha/2, n-1) \cdot \frac{S}{\sqrt{n}}, \quad (3.68)$$

unde $100(1-\alpha)$ este nivelul de încredere.

Cvantelele repartiției normale și t se evaluează cu funcțiile Matlab `norminv(P,miu,sigma)`, respectiv `tinv(P,v)`.

Ex. 1: Masa unor persoane ce utilizează un ascensor este distribuită normal. Un eșantion furnizează următoarele valori: 71, 85, 68, 72, 58, 76, 74, 80. Să se estimeze cu o probabilitate de 95% valoarea medie.

$$\mu = \bar{x} \pm t(1 - \alpha/2, n-1) \cdot \frac{S}{\sqrt{n}}$$

$$t(1-\alpha/2, n-1) = t(0.05/2, 8-1) = 2.365$$

$$\bar{x} = 73, S=8.09, \Rightarrow S_{\bar{x}} = \frac{S}{\sqrt{n}} = \frac{2.365}{\sqrt{8}} = 2.86$$

$$\mu = 73 \pm (2.365 \cdot 2.86) = 73 \pm 6.76$$

$$\mu \in [66.24; 79.76].$$

Valoarea cvatilei repartiției t se obține în Matlab cu `tinv([0.025 0.975], 7)`.

Ex. 2: Din 40 de copii ce merg la aceeași grădiniță, au fost selectați 8 și s-a măsurat înălțimea lor în cm.: 110, 112, 85, 117, 100, 98, 104, 90. Știind că înălțimea este o variabilă aleatoare normal distribuită să se estimeze media înălțimii copiilor.

Se observă că volumul eșantionului este mai mare decât 5% din populație ($8/40 = 0.20 > 0.05$), deci va trebui aplicat un factor de corecție datorită populației finite.

$$\mu = \bar{x} \pm t(\alpha/2, n-1) \cdot \frac{S}{\sqrt{n}} \cdot \sqrt{\frac{N-n}{N-1}}$$

$$t(\alpha/2, n-1) = t(0.05/2, 8-1) = 2.365$$

$$\bar{x} = 102, S = 10.994 \Rightarrow S_{\bar{x}} = \frac{S}{\sqrt{n}} = \frac{10.994}{\sqrt{8}} = 3.89$$

$$\mu = 102 \pm (2.365 \cdot 3.89) \sqrt{\frac{40-8}{40-1}} = 102 \pm 8.33$$

$$\mu \in [93.67; 110.33].$$

3.4.3 TESTAREA IPOTEZELOR

Informațiile legate de o anumită populație statistică se obțin prin selectarea unui eșantion, cu ajutorul căruia se estimează parametrii populației respective. Pe baza acestor parametri se poate forma o imagine asupra caracteristicilor analizate. De exemplu, dintr-un lot de piese se extrage un eșantion de n elemente și se efectuează măsurări ale lungimii l , lungime presupusă a fi o

variabilă repartizată normal cu media μ și dispersia σ . Pentru eșantionul extras s-a determinat prin calcul media experimentală $\bar{I}$. Se pune problema dacă $\mu = \bar{I}$, problemă ce constituie o ipoteză statistică, adică o presupunere (supoziție) asupra populației statistice luate în studiu. Chiar dacă ipoteza statistică este avansată pe baza unui eșantion, concluzia referitoare la valoarea parametrului sau la natura repartiției se referă la întreaga populație.

În mod firesc, o presupunere asupra unor elemente ale repartiției ($\mu = \bar{I}$) poate admite o alternativă ($\mu < \bar{I}$). Ipoteza inițială se numește ipoteză nulă și se notează H_0 . Ipoteza alternativă sau concurentă se notează H_1 .

Ipotezele statistice nu sunt echivalente cu ipotezele științifice. În cazul unei ipoteze științifice este suficient un singur exemplu contrar pentru a o infirma, dar o ipoteză statistică poate fi adevărată, chiar dacă într-o anumită situație a fost respinsă ca fiind falsă.

Ipotezele statistice sunt însoțite de două tipuri de erori:

- o eroare apare atunci când se respinge ipoteza inițială, H_0 , în situația în care ea este adevărată. Această eroare este probabilitatea de a respinge ipoteza H_0 , când în realitate ea este adevărată. Eroarea se notează, de obicei, cu α :

$$\alpha = P(\text{respinge } H_0 \mid H_0 \text{ adevărată}). \quad (3.69)$$

- o eroare constă în acceptarea ipotezei H_0 , când în realitate ea este falsă. Această eroare se notează cu β :

$$\beta = P(\text{acceptă } H_0 \mid H_0 \text{ falsă}). \quad (3.70)$$

Aceste erori se mai numesc erori de tip I și II. Este evident că există interesul menținerii valorilor α și β la niveluri cât mai mici.

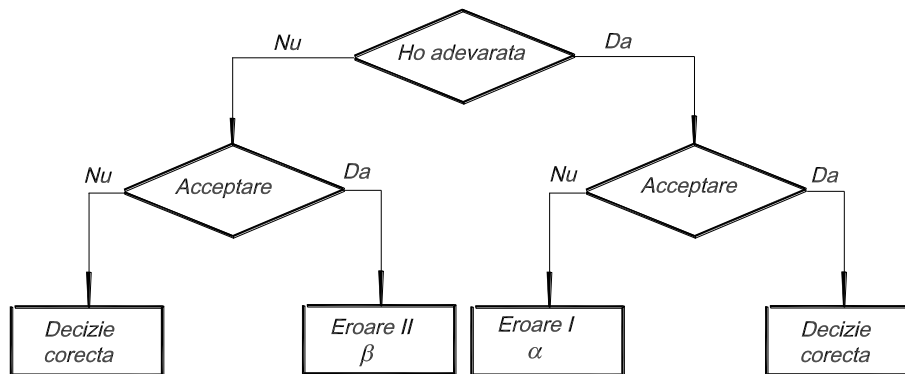


Fig. 3.21 Erori ce apar la verificarea ipotezelor statistice

Metoda efectivă de luare a deciziei asupra acceptării sau respingerii ipotezei H_0 se face cu ajutorul unui test sau criteriu statistic, care este un procedeu furnizat de statistica matematică.

Testul impune luarea uneia dintre cele două decizii – acceptare sau respingere – a ipotezei nule cu un anumit risc. Acest lucru este prezentat în Figura 3.21. Se poate observa că valoarea $1 - \beta$ este probabilitatea de respingere a ipotezei nule când aceasta este falsă. Valoarea $1 - \beta$ poate fi considerată puterea testului respectiv.

Din punct de vedere practic, pentru a verifica ipoteza H_0 , cu alternativa H_1 , sunt necesare, [8]:

- ❑ o statistică (o funcție de măsurările experimentale) a cărei repartiție să fie cunoscută sau cel puțin să fie exprimabilă analitic;
- ❑ o valoare considerată “critică” cu care să se compare valoarea calculată a statisticii;
- ❑ o regulă de decizie prin care să se accepte sau să se respingă H_0 (când se respinge H_0 , automat se acceptă H_1);
- ❑ o valoare a riscului ales α , ce se mai numește nivel de semnificație al testului.

a. Media unei singure populații

În continuarea paragrafului, se tratează problema testării mediei unei populații, în sensul că se presupune cunoscută media populației și se stabilește dacă un eșantion este extras dintr-o populație cu aceeași medie.

Ideea de bază în cazul testării ipotezelor este că se începe cu o presupunere asupra unui parametru, cum ar fi media, după care se utilizează eşantionul pentru testarea ipotezelor. Presupunerea admisă, ipoteza nulă se notează H_0 . De ex.: Se emite ipoteza că μ este μ_0 , adică $H_0: \mu = \mu_0$. În afara ipotezei nule se definesc ipotezele alternative, H_1 . Ipoteza alternativă poate fi una din cele trei situații posibile, $H_1: \mu \neq \mu_0$; $H_1: \mu > \mu_0$; $H_1: \mu < \mu_0$. Procedura de testare constă în compararea statisticii experimentale, în acest caz $\bar{x}$ cu parametrul corespunzător ipotezei nule, μ_0 . Dacă există diferențe între cele două valori, ipoteza nulă se respinge. De obicei apar diferențe, fie datorită întâmplării, fie datorită faptului că H_0 este falsă. În acceptarea și respingerea unei ipoteze statistice intervin cele două tipuri de erori, α și β , deci niciodată nu se poate lua o decizie cu siguranță absolută.

În cazul testării ipotezei $\mu = \mu_0$ cu o probabilitate de 95%, există un risc de 5%. În fig. 3.22 se poate observa că apar zone de acceptare și zone de respingere. În cazul ipotezei alternative $H_1: \mu \neq \mu_0$, zona de respingere, având o arie egală cu α , se împarte în două regiuni distincte, fiecare având o arie egală cu $\alpha/2$ (fig. 2.16 a). Acest tip de test este cu două extremități. În cazul celorlalte ipoteze alternative $\mu > \mu_0$ (fig. 2.16 b) și $\mu < \mu_0$ (fig. 2.16 c) apare o singură zonă de respingere, testele fiind cu o singură extremă.

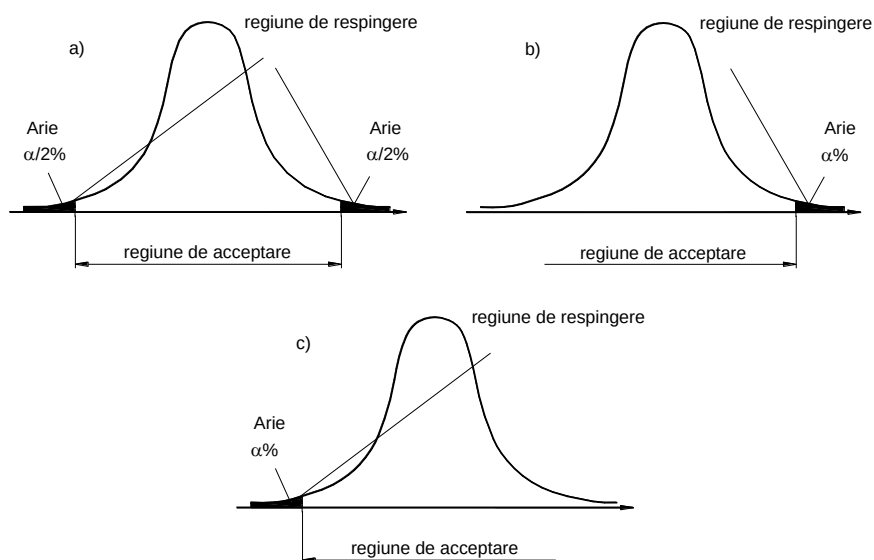


Fig. 3.22 Zone de respingere și acceptare în testarea ipotezelor

În testarea unei ipoteze există totdeauna 5 pași ce trebuie urmăriți:

Pasul 1: Formularea ipotezei nule și a celei alternative. Ipoteza nulă specifică valoarea parametrului ce urmează a fi testat ($\mu = \mu_0$). Forma ipotezei alternative poate diferi. În cazul testelor cu două extremități este importantă doar existența unei diferențe între μ și μ_0 , în timp ce în cazul ipotezelor cu o singură extremitate este important și sensul acestei diferențe. Ipotezele se pot formula:

$$\begin{array}{l} H_0 : \mu = \mu_0 \quad \text{sau} \quad H_0 : \mu \neq \mu_0 \quad \text{sau} \quad H_0 : \mu = \mu_0 \\ H_1 : \mu \neq \mu_0 \quad \text{sau} \quad H_1 : \mu > \mu_0 \quad \text{sau} \quad H_1 : \mu < \mu_0 \end{array}$$

Pasul 2: Specificarea nivelului de semnificație α . Cel mai uzual este să se lucreze cu nivelul de semnificație de 5% sau 1%. Având stabilit nivelul de semnificație, punctul de demarcare a regiunii de acceptare cu cea de respingere reprezintă valoarea critică a testului. Pentru testele cu două extremități există două valori critice. Valorile se numesc critice deoarece sunt valorile cu care se compară parametrii experimentali.

Tabelul 2.4. Valori critice în cazul testării mediei

Tip test	Ipoteze		$\bar{x}$	Z	t
	H_0	H_1			
2 extremități	$\mu = \mu_0$	$\mu \neq \mu_0$	$\mu_0 \pm Z_{\alpha/2} \sigma_{\bar{x}}$	$\pm Z_{\alpha/2}$	$\pm t(\alpha/2, v)$
o extremitate	$\mu = \mu_0$	$\mu > \mu_0$	$\mu_0 + Z_{\alpha} \sigma_{\bar{x}}$	$+ Z_{\alpha}$	$+ t(\alpha, v)$
o extremitate	$\mu = \mu_0$	$\mu < \mu_0$	$\mu_0 - Z_{\alpha} \sigma_{\bar{x}}$	$- Z_{\alpha}$	$- t(\alpha, v)$

Pasul 3: Selectarea statisticii și a valorii critice. Există două posibilități de construire a unui test:

- calcularea valorii critice a parametrului experimental ($\bar{x}$);
- să se lucreze cu variabila normală standard Z sau variabila t.

Valorile critice în fiecare caz de testare a mediei sunt prezentate în tabelul 2.4.

Pasul 4: Determinarea statisticii testului. Dacă testul este condus pe baza valorii mediei, evident că valoarea statisticii va fi chiar media experimentală. Dacă testul se bazează pe Z sau t, valoarea statisticii va fi o transformare a valorii experimentale în forma standard:

$$Z = \frac{\bar{x} - \mu_0}{\sigma_{\bar{x}}} \quad \text{sau} \quad t = \frac{\bar{x} - \mu_0}{S_{\bar{x}}}.$$

Pasul 5: regula de decizie. Odată stabilită valoarea critică se poate concluziona dacă ipoteza nulă se acceptă sau respinge. Este esențial de remarcat că respingerea ipotezei H_0 , implică acceptarea ipotezei H_1 .

În cazul testelor cu două extreme:

- se acceptă H_0 , dacă $\bar{x}$ aparține intervalului de acceptare, determinat de cele două valori critice. $\bar{x} \in [\mu_0 - Z_{\alpha/2} \cdot \sigma_{\bar{x}}; \mu_0 + Z_{\alpha/2} \cdot \sigma_{\bar{x}}]$ și se respinge H_0 dacă acest lucru nu se întâmplă;
- dacă se lucrează cu Z sau t : se acceptă H_0 dacă valoarea Z , respectiv t cade în interiorul domeniului de acceptare, altfel se respinge H_0 .

În cazul testelor cu o extremă:

- cu statistica $\bar{x}$, se acceptă H_0 dacă $\bar{x}$ este mai mare sau mai mic decât valoarea critică, în funcție de ipoteza alternativă, și se respinge H_0 dacă nu se îndeplinește condiția;
- se procedează similar în cazul variabilei Z sau t .

Trebuie reținut că la aplicarea unui test statistic nu se obține niciodată o certitudine. Totdeauna apar erori α și β . Eroarea de tip α se controlează prin specificarea lui α , nivelul de semnificație al testului, specificând riscul asumat de a respinge o ipoteză nulă. Eroarea β nu poate fi controlată simultan cu α , deși ele sunt interdependente. Singura posibilitate de a reduce ambele erori este să se mărească volumul eșantionului. Cele două riscuri pot fi urmărite numai prin intermediul caracteristicii operative [4].

Exemple

1. În urma unui studiu statistic asupra producției unei mașini-unelte s-a constatat că se produc în medie 3.9 repere/h (repere de același tip), cu o deviație standard de 0.6. Urmărind un eșantion de 100 de repere identice pe o altă MU, se constată că producția are o medie de 4.2 repere/h. Să se testeze ipoteza că producția de 4.2 repere diferă de cea de 3.9. Se admite un nivel de semnificație de 0.05.

Deoarece se cunoaște valoarea lui σ , se aplică distribuția normală, cu variabila Z .

Date inițiale: $\mu = 3.9$; $\sigma = 0.6$; $n = 100$

Ipoteze: $H_0: \mu_0 = 3.9$; $H_1: \mu_0 \neq \mu$

Nivel de semnificație $\alpha = 0.05 \Rightarrow Z_{\alpha/2} = Z_{0.025} = \pm 1.96$ (vezi Anexa 1)

Deviația standard a mediei: $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}} = \frac{0.6}{\sqrt{100}} = 0.06$

Valori critice: pentru $\bar{x}$: $\mu_0 \pm Z_{\alpha/2} \sigma_{\bar{x}} = 3.9 \pm 1.96 \cdot 0.06 \Rightarrow 3.78$ și 4.02

pentru Z : $\pm Z_{\alpha/2} \Rightarrow -1.96$ și 1.96

Testul: valoarea $\bar{x} = 4.2$ cade în afara intervalului determinat de valorile critice ($4.02 < 4.2$). Dacă se utilizează variabila Z , statistica este:

$$Z = \frac{\bar{x} - \mu_0}{\sigma_{\bar{x}}} = \frac{4.2 - 3.9}{0.06} = 5,$$

care se găsește în afara intervalului determinat de valorile critice.

Concluzia: Indiferent de testul aplicat se poate afirma că se respinge ipoteza nulă cu o probabilitate de 95%, deci media producției în cel de-al doilea caz diferă de prima. Nu se poate afirma că media eșantionului este mai mare, deoarece s-a utilizat un test cu două extreme.

2. Pe baza exemplului 1, să se testeze dacă media eșantionului este mai mare decât media populației cu o probabilitate de 95%.

Ipoteze: $H_0: \mu_0 = 3.9$; $H_1: \mu_0 > 3.9$

Nivel de semnificație $\alpha = 0.05 \Rightarrow Z_{\alpha} = Z_{0.05} = 1.645$

Deviația standard a mediei: $\sigma_{\bar{x}} = 0.06$

Valori critice: pentru $\bar{x}$: $\mu_0 + Z_{\alpha} \sigma_{\bar{x}} = 3.9 + 1.645 \cdot 0.06 = 4.0$

pentru Z : $+Z_{\alpha} = 1.645$

Testul: valoarea $\bar{x} = 4.2 >$ valoare critică. Dacă se utilizează variabila Z :

$$Z = \frac{\bar{x} - \mu_0}{\sigma_{\bar{x}}} = \frac{4.2 - 3.9}{0.06} = 5 > \text{valoare critică}$$

Concluzia: În ambele teste se respinge ipoteza nulă și cu o probabilitate de 95% se acceptă ipoteza alternativă (productivitatea celei de-a doua MU este

mai mare).

3. Un producător de baterii afirmă că bateriile au o durată medie de viață de 55 h. Se testează un eșantion de 40 de baterii și se constată că au o durată de viață de 50 h, iar abaterea standard a mediei este 11.734. Este adevărată afirmația producătorului cu un nivel de semnificație de 1%?

Date inițiale: $\mu_0 = 55$; $\bar{x} = 50$; $S = 11.734$; $n = 40$

Ipoteze: $H_0: \mu_0 = 55$; $H_1: \mu_0 \neq 55$

Nivel de semnificație $\alpha = 0.01$

Deviația standard a mediei: $S_{\bar{x}} = \frac{S}{\sqrt{n}} = \frac{11.734}{\sqrt{40}} = 1.86$

Valori critice: $Z: \pm Z_{\alpha/2} = \pm 2.58$

Testul : $Z = \frac{\bar{x} - \mu_0}{S_{\bar{x}}} = \frac{50 - 55}{1.86} = -2.69$. Statistica calculată > valoarea critică, deci cade în domeniul de respingere.

Concluzie: Cu o probabilitate de 99% se poate susține că afirmația producătorului nu este adevărată.

În caz că se dorește să se specifice dacă durata medie de viață este mai mică decât cea afirmată de producător, trebuie aplicat un test cu o singură extremă.

Ipoteze: $H_0: \mu_0 = 55$; $H_1: \mu_0 < 55$

Nivel de semnificație $\alpha = 0.01$

Deviația standard a mediei: $S_{\bar{x}} = 1.86$

Valori critice: $Z: -Z_{\alpha} = -2.33$

Testul : $Z = \frac{\bar{x} - \mu_0}{S_{\bar{x}}} = \frac{50 - 55}{1.86} = -2.69$. Statistica cade în interiorul domeniului de respingere, deci se acceptă ipoteza alternativă cu o probabilitate de 99%.

b. Media a două populații

În numeroase situații practice există două eșantioane și se pune problema stabilirii dacă între cele două eșantioane apare o diferență statistică

semnificativă. În această situație ipoteza nulă este egalitatea dintre mediile ideale ale celor două populații:

$$H_0: \mu_1 = \mu_2 \text{ (sau } \mu_1 - \mu_2 = 0), \mu_1, \mu_2 \text{ necunoscute}$$

cu alternativele:

$$H_1: \mu_1 \neq \mu_2 \text{ (sau } \mu_1 - \mu_2 \neq 0)$$

$$H_1: \mu_1 > \mu_2 \text{ (sau } \mu_1 - \mu_2 > 0)$$

$$H_1: \mu_1 < \mu_2 \text{ (sau } \mu_1 - \mu_2 < 0)$$

În continuare se tratează separat următoarele cazuri:

b1. eșantioane de volum mare;

b2. eșantioane de volum mic;

b1. Eșantioane independente de volum mare

În acest caz volumul celor două eșantioane trebuie să fie mai mare decât 30, $n_1, n_2 \geq 30$. Datorită acestui fapt, testul Z este mai potrivit, în baza teoremei limită centrale. Testul presupune parcurgerea aceluiași pași prezentați anterior, cu adăugarea unui singur factor, calcularea statisticii Z. Problema care se pune este stabilirea mărimii diferenței dintre cele două medii. Pentru a putea răspunde la această întrebare, trebuie obținute informații în privința distribuției $\bar{x}_1 - \bar{x}_2$, pentru toate valorile posibile $\bar{x}_1$ și $\bar{x}_2$. Având în vedere că volumul eșantioanelor este mai mare decât 30, distribuția $\bar{x}_1 - \bar{x}_2$ poate fi aproximată cu distribuția normală. În plus media distribuției ($\mu_{\bar{x}_1 - \bar{x}_2}$) este 0, când eșantioanele sunt extrase în condițiile ipotezei nule ($\mu_1 = \mu_2$). Se poate demonstra, [5], că abaterea standard este:

$$\sigma_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}, \quad (3.71)$$

σ_1^2, σ_2^2 fiind varianțele celor două populații.

Dacă σ_1^2, σ_2^2 sunt necunoscute, trebuie utilizate valorile experimentale:

$$S_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}. \quad (3.72)$$

Pasul 1: $H_0: \mu_1 - \mu_2 = 0$

$H_1: \mu_1 - \mu_2 \neq 0$ sau $H_1: \mu_1 - \mu_2 > 0$ sau $H_1: \mu_1 - \mu_2 < 0$.

Pasul 2: Stabilirea nivelului de semnificație α .

Pasul 3: Calcularea statisticii Z corespunzătoare diferenței $\mu_1 - \mu_2$:

$$Z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{S_{\bar{x}_1 - \bar{x}_2}}. \quad (3.73)$$

În ipoteza $\mu_1 - \mu_2 = 0$, Z devine:

$$Z = \frac{(\bar{x}_1 - \bar{x}_2)}{S_{\bar{x}_1 - \bar{x}_2}}. \quad (3.74)$$

Pasul 4: Compararea statisticii calculate cu valoarea critică, stabilită pe baza nivelului de semnificație.

Pasul 5: Acceptarea sau respingerea ipotezei nule.

Ex.: În vederea achiziționării unor rulmenți, un beneficiar testează doi producători. Testarea s-a făcut în condiții identice. Datele obținute sunt centralizate în tabelul următor:

	volum	medie [h funcționare]	S
lot 1	$n_1 = 35$	9800	975
lot 2	$n_2 = 40$	10200	120

Să se testeze dacă: a. cele două loturi diferă; b. lotul 1 este inferior lotului 2.

a. Este necesar un test cu două extreme

Ipoteze: $H_0: \mu_1 - \mu_2 = 0$; $H_1: \mu_1 - \mu_2 \neq 0$.

Nivel de semnificație $\alpha = 0.05$

Deviația standard a mediei: $S_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}} = \sqrt{\frac{975^2}{35} + \frac{120^2}{40}} = 28.01$

Valori critice: $\pm Z_{\alpha/2} = \pm 1.96$

Test: $Z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{S_{\bar{x}_1 - \bar{x}_2}} = \frac{9800 - 10200}{28.01} = -14.28$

Concluzie: $|Z| > |Z_{\alpha/2}|$, valoarea calculată cade în zona de respingere, deci H_0 se respinge cu o probabilitate de 95%. Atenție: Nu se obțin informații în privința

superiorității unui anumit producător.

b. Se aplică un test cu o singură extremă.

Ipoteze: $H_0: \mu_1 - \mu_2 = 0$; $H_1: \mu_1 - \mu_2 < 0$.

Nivel de semnificație $\alpha = 0.05$

Deviația standard a mediei: $S_{\bar{x}_1 - \bar{x}_2} = 28.01$

Valori critice: $-Z_{\alpha/2} = -1.645$

Test: $Z = -14.28$

Concluzie: $Z < Z_{\alpha}$, deci H_0 se respinge și se acceptă durata de viață mai ridicată a celui de-al doilea lot.

b2. Eșantioane de volum mic

În situația când σ_1 și σ_2 nu sunt cunoscute, iar $n_1, n_2 < 30$ se recurge la distribuția t. În condițiile ipotezei nule, $H_0: \mu_1 - \mu_2 = 0$, testul statistic t se calculează cu formula:

$$t = \frac{\bar{x}_1 - \bar{x}_2}{S_{\bar{x}_1 - \bar{x}_2}}, \quad (3.75)$$

unde $S_{\bar{x}_1 - \bar{x}_2}$ se calculează cu relația (2.45). Este important de remarcat că statistica t doar aproximează distribuția t. Numărul gradelor de libertate este dat de expresia:

$$\nu = \frac{(S_1^2/n_1 + S_2^2/n_2)^2}{\left(\frac{S_1^2}{n_1}\right)/(n_1-1) + \left(\frac{S_2^2}{n_2}\right)/(n_2-1)}. \quad (3.76)$$

Numărul gradelor de libertate se rotunjește în jos, la cel mai apropiat întreg.

În cazul particular, când se poate considera că varianțele celor două populații sunt egale ($\sigma_1^2 = \sigma_2^2$), se poate obține un test mai puternic. Acest lucru poate fi considerat adevărat în cazul proceselor de lungă durată, la care din determinări anterioare s-a constatat că varianțele sunt egale. În acest caz, având două eșantioane de volum n_1, n_2 , cu varianțele S_1^2, S_2^2 , se calculează

varianța compusă:

$$S_p = \frac{(n_1 - 1)S_1^2 + (n_2 - 1)S_2^2}{n_1 + n_2 - 2}, \quad (3.77)$$

Iar abaterea standard este:

$$S_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{S_p^2}{n_1} + \frac{S_p^2}{n_2}}. \quad (3.78)$$

Statistica t, în cazul ipotezei nule ($H_0: \mu_1 - \mu_2 = 0$) este:

$$t = \frac{\bar{x}_1 - \bar{x}_2}{S_{\bar{x}_1 - \bar{x}_2}}, \quad (3.79)$$

iar numărul gradelor de libertate este $v = n_1 + n_2 - 2$.

Exemple

1. În vederea stabilirii nivelului de pregătire a studenților din anul I, de la două direcții de specializare, s-a testat un număr de studenți pe baza unor chestionare. Ambele eșantioane au fost formate de 10 studenți, iar rezultatele centralizate sunt prezentate în continuare (valoarea înregistrată fiind numărul de puncte realizate).

	volum	medie	S
profil 1	$n_1 = 10$	60.5	8.2
profil 2	$n_2 = 10$	62.5	10.2

Să se stabilească dacă există diferențe semnificative între cele două eșantioane.

Datorită faptului că studenții provin din specializări diferite nu sunt motive să se presupună că varianțele cele două populații ar fi egale. Se va aplica testul t, varianțele fiind necunoscute și volumul eșantioanelor mic.

Ipoteze: $H_0: \mu_1 - \mu_2 = 0$; $H_1: \mu_1 - \mu_2 \neq 0$.

Nivel de semnificație $\alpha = 0.05$

Deviația standard a mediei: $S_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}} = \sqrt{\frac{8.2^2}{10} + \frac{10.2^2}{10}} = 4.14$

Valori critice: $\pm t_{\alpha/2} = \pm 2.11$ cu $v = 17$

$$v = \frac{(S_1^2/n_1 + S_2^2/n_2)^2}{\left(\frac{S_1^2}{n_1}\right)/(n_1-1) + \left(\frac{S_2^2}{n_2}\right)/(n_2-1)} = 17.2 \Rightarrow v = 17$$

Test: $t = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{S_{\bar{x}_1 - \bar{x}_2}} = \frac{60.5 - 62.5}{4.14} = -0.48$

Concluzie: Deoarece $|t| < |t_{\alpha/2}|$ nu se poate respinge ipoteza nulă, deci cu o probabilitate de 95% între cele două specializări nu există diferențe statistice semnificative.

2. Un echipament identic este utilizat în două societăți la fabricarea aceluiași componente. Pe o perioadă de o săptămână se înregistrează numărul defectelor în fiecare societate, obținându-se următoarele rezultate sumarizate:

	volum	medie defecte	S
societate 1	$n_1 = 7$	8.14	3.24
societate 2	$n_2 = 7$	9.71	3.73

Să se stabilească dacă există diferențe semnificative între cele două societăți, respectiv dacă cea de-a doua societate are mai multe defecte decât prima.

Având în vedere că procesul de fabricație este identic se poate presupune că varianțele (necunoscute) sunt egale.

Ipoteze: $H_0: \mu_1 - \mu_2 = 0$; $H_1: \mu_1 - \mu_2 \neq 0$.

Nivel de semnificație $\alpha = 0.05$

Deviația standard a mediei: $S_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{S_p^2}{n_1} + \frac{S_p^2}{n_2}}$ cu

$$S_p = \frac{(n_1-1)S_1^2 + (n_2-1)S_2^2}{n_1 + n_2 - 2} = \frac{(7-1)3.24^2 + (7-1)3.73^2}{7+7-2} = 12.205$$

$$S_{\bar{x}_1 - \bar{x}_2} = 1.867$$

Valori critice: $\pm t_{\alpha/2} = \pm 2.179$ cu $v = n_1 + n_2 - 2 = 12$

$$\text{Test: } t = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{S_{\bar{x}_1 - \bar{x}_2}} = \frac{8.14 - 9.71}{1.867} = -0.84$$

Concluzie: $|t| < |t_{\alpha/2}|$ ($0.84 < 2.179$), deci nu se poate respinge ipoteza nulă; cu o probabilitate de 95% se poate afirma că nu există diferențe semnificative între cele două eșantioane.

Ipoteze: $H_0: \mu_1 - \mu_2 = 0$; $H_1: \mu_1 - \mu_2 < 0$.

Nivel de semnificație $\alpha = 0.05$

Deviația standard a mediei: $S_{\bar{x}_1 - \bar{x}_2} = 1.867$

Valori critice: $\pm t_{\alpha} = -1.782$ cu $v = n_1 + n_2 - 2 = 12$

Test: $t = -0.84$

Concluzie: $t > -t_{\alpha}$, deci nu se poate respinge H_0 , cu o probabilitate de 95% cele două procese de fabricație sunt la fel.

BIBLIOGRAFIE

1. Baron, T., Metode statistice pentru analiza și controlul calității producției, EDP, București, 1979.
2. Cicală, E.F., Metode de prelucrare statistică a datelor experimentale, Ed. Politehnica Timișoara, 1999.
3. Constantinescu, I., ș.a, Prelucrarea datelor experimentale cu calculatoarele numerice, ET, București, 1980.
4. Davidescu, A., Metrologie generală, Ed. Politehnica Timișoara, 2001.
5. Fleming, M., Nellis, J., Principles of Applied Statistics, Routledge Co., London, 1994.
6. Hanselman, D., Littlefield, B., The Student Edition of Matlab®. User's Guide, Prentice Hall, 1995.

7. Hanselman, D., Littlefield, B., Mastering Matlab® 5. A comprehensive tutorial and reference, Prentice Hall, 1998.
8. Iliescu, D.V., Vodă, V., Statistică și toleranțe, ET, București, 1977.
9. J de Leeuw, WEB Statistics Book, <http://www.stat.ucla.edu/textbook>
10. Martinez, W., Martinez, A., Computational Statistics Handbook with Matlab, Chapman&Hall/CRC, Washington DC, 2002.
11. Montgomery, D., Design and Analysis of Experiments, John Willey&Sons, Singapore, 1991.
12. Montgomery, D., Runger, G., Applied Statistics and Probability for Engineers, John Wiley&Sons, New York, 2006.
13. Nichici, A., ș.a., Prelucrarea datelor experimentale, Lito UPT, Timișoara, 1990.
14. Petrișor, E., Probabilități și statistică. Aplicații în economie și inginerie, Ed. Politehnica, Timișoara, 2005.
15. Ross, S., Initiation aux probabilités, Presses polytechniques romandes, Lausanne, 1987.
16. Ruegg, A., Probabilités et statistique, Presses polytechniques romandes, Lausanne, 1985.
17. Tiron, M., Teoria erorilor de măsurare și metoda celor mai mici pătrate, ET, București, 1972.
18. Tiron, M., Prelucrarea statistică și informațională a datelor de măsurare, ET, București, 1976.
19. Tovissi, L., Vodă, V., Metode statistice, Ed. Științifică și Enciclopedică, București, 1982.
20. Wheeler, D., Chambers, D, Understanding Statistical Process Control, SPC Press, Tennessee, 1986.