

4. ANALIZA EXPLORATORIE A UNEI SERII DE DATE

4.1. INTRODUCERE

Analiza exploratorie a datelor (EDA – Exploratory Data Analysis) este o tehnică de abordare a prelucrării datelor de dată mai recentă, ce constă într-o colecție de tehnici preponderent grafice, ce permit evidențierea unor structuri în date. John Tukey a lansat această abordare în 1977.

Conform acestei metode, datele trebuie explorate fără a presupune apriori anumite modele statistice, distribuții de erori, relații între variabile etc. Scopul principal este evidențierea caracteristicilor datelor pentru ca analistul să înțeleagă cât mai bine procesul, să-l poată analiza și modela. După depistarea caracteristicilor setului de date, modelarea datelor se poate face în condiții mult mai bune.

Într-o analiză statistică clasică, etapele analizei sunt:

Problemă $\Rightarrow$ Date $\Rightarrow$ Model $\Rightarrow$ Analiză $\Rightarrow$ Concluzii.

În abordarea exploratorie a datelor, modelarea se face doar după depistarea principalelor caracteristici ale setului de date, moment în care modelul are șanse mult mai mari să fie unul corect. Etapele analizei exploratorii sunt:

Problemă $\Rightarrow$ Date $\Rightarrow$ Analiză $\Rightarrow$ Model $\Rightarrow$ Concluzii.

Tehnicile EDA și cele clasice nu sunt mutual exclusive și se recomandă utilizarea lor complementară. În faza inițială a analizei se abordează tehnici EDA, care evidențiază caracteristicile setului de date, pe baza acestora se stabilește modelul corespunzător al datelor, iar validarea modelului se face cu metode cantitative: testarea ipotezelor, ANOVA etc.

Reducerea datelor experimentale la indicatori cantitativi, chiar dacă aceștia se bazează pe toate valorile cuprinse în setul de date, fără o reprezentare grafică prealabilă poate conduce la erori. Un exemplu clasic îl constituie următorul set de date [10]:

X: 10.00; 8.00; 13.00; 9.00; 11.00; 14.00; 6.00; 4.00; 12.00; 7.00; 5.00.

Y: 8.04; 6.95; 7.58; 8.81; 8.33; 9.96; 7.24; 4.26; 10.84; 4.82; 5.68.

În urma efectuării calculelor se obțin următoarele valori:

Volumul setului de date: $N = 11$

Media lui x : $\bar{x} = 9.0$

Media lui y : $\bar{y} = 7.5$

Coeficientul de corelație liniară: $r = 0.816$

Dreapta de regresie a lui y în raport cu x : $y = x/2 + 3$

Deviația standard a erorilor: $\sigma_{\varepsilon} = 1.237$

Aceste informații cantitative, au importanța lor, dar nu oferă o caracterizare completă a datelor. Reprezentarea datelor raportate la dreapta de regresie se poate observa în fig. 4.1.

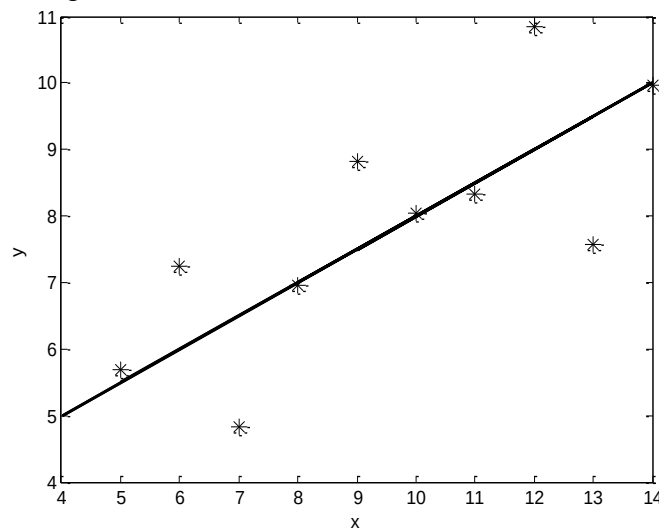


Fig. 4.1. Reprezentarea datelor și a dreptei de regresie corespunzătoare

Din fig. 4.1 se observă că:

- Datele prezintă o dependență liniară;
- Nu este justificată alegerea unui model mai complicat;
- Nu apar valori aberante în setul de date;
- Variabilitatea lui y în raport cu x este aproximativ constantă, deci nu este cazul să se utilizeze ponderi în modelare.

Acest mod de abordare a prelucrării datelor, unde reprezentările grafice se îmbină cu determinarea unor indicatori cantitativi este abordarea corectă, pentru că se evită eventualele interpretări greșite.

Pentru a ilustra pierderea de informații, în cazul când nu se fac reprezentări grafice asociate cu calculele, se prezintă următorul exemplu [10] constând din patru seturi de date. Primul set este cel prezentat anterior, iar celelalte trei sunt

conținute în tabelul 4.1. Se observă că valorile șirului X sunt identice în primele trei cazuri, iar X4 conține o aceeași valoare, 8.00, cu excepția unui singur element egal cu 19.

Tabel 4.1. Seturi de date

X2	10.00;	8.00	13.00	9.00	11.00	14.00	6.00	4.00	12.00	7.00	5.00
Y2	9.14;	8.14	8.74	8.77	9.26	8.10	6.13	3.10	9.13	7.26	4.74
X3	10.00;	8.00	13.00	9.00	11.00	14.00	6.00	4.00	12.00	7.00	5.00
Y3	7.46;	6.77	12.74	7.11	7.81	8.84	6.08	5.39	8.15	6.42	5.73
X4	8.00	8.00	8.00	8.00	8.00	8.00	8.00	19.0	8.00	8.00	8.00
Y4	6.58;	5.76	7.71	8.84	8.47	7.04	5.25	12.5	5.56	7.91	6.89

Tabel 4.2. Indicatori cantitativi

	Set 1	Set 2	Set 3	Set 4
N	11	11	11	11
$\bar{x}$	9.00	9.00	9.00	9.00
$\bar{y}$	7.50	7.50	7.50	7.50
$Y=bx+a$	$Y=0.5x+3$	$Y=0.5x+3$	$Y=0.5x+3$	$Y=0.5x+3$
r	0.816	0.816	0.816	0.817
σ_{ε}	1.237	1.237	1.236	1.236

După cum se observă în tabelul 4.2, cele patru seturi de date au indicatorii cantitativi de valori foarte apropiate. Dar în urma reprezentării grafice a acestor seturi de date (fig. 4.2) se observă că există diferențe mari în ceea ce privește distribuția lor în jurul dreptei de regresie.

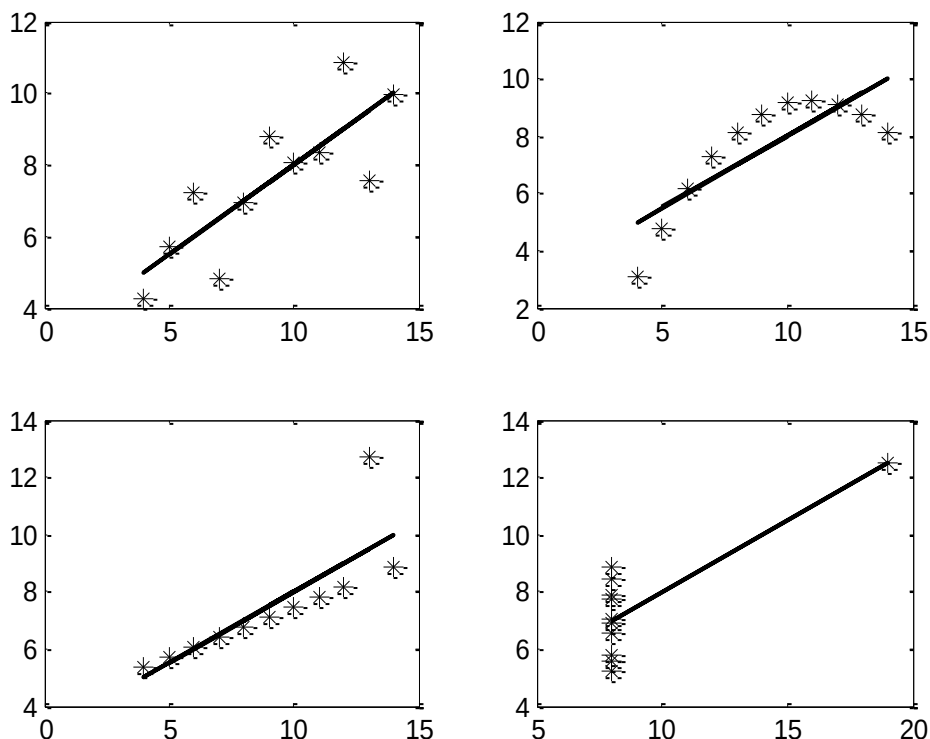


Fig. 4.2. Seturi de date având aceeași dreaptă de regresie

Analizând reprezentările grafice din figura 4.2 se pot trage următoarele concluzii:

- Primul set de date are un comportament liniar, fără valori aberante, la care modelul dreptei de regresie este corespunzător;
- Al doilea set de date are un comportament pătratic, fără valori aberante, modelul ar trebui să fie un polinom de gradul doi;
- Al treilea set de date prezintă o valoare aberantă, care ar trebui eliminată; în acest context ecuația dreptei de regresie se modifică, dreapta de regresie care ar modela corespunzător setul de date ar avea o pantă mai mică;
- Ultimul set de date este rezultatul unui experiment prost organizat, deoarece conține o singură valoare mult distanțată de restul valorilor.

Din acest exemplu de prelucrare a datelor se poate observa importanța analizei exploratorii, care amână alegerea modelului până la cunoașterea comportamentului datelor. În cazul seturilor 2 – 4, eroarea apare din faptul că se presupune un model liniar, fără a face o analiză preliminară a datelor. Importanța determinărilor cantitative nu trebuie subestimată, calcularea unor

indicatori: medie, dispersie, coeficient de corelație etc. permit o prezentare sintetică a datelor, ceea ce oferă un mare avantaj, dar calculele trebuie asociate cu reprezentări grafice.

La baza analizării unei serii de date stau câteva ipoteze fundamentale:

- Caracterul aleator al datelor;
- Datele au o distribuție de probabilitate;
- Localizarea constantă;
- Variabilitatea constantă.

Previziunea este un țel important în inginerie. Dacă ipotezele sunt valabile, se pot face previziuni asupra unui proces, inclusiv afirmații legate de evoluția anterioară, se spune în acest caz că procesul este "în control statistic". Dacă cele patru ipoteze nu sunt valabile, procesul este în derivă (în raport cu locația, variabilitatea sau distribuția), imprevizibil și necontrolabil. Orice caracterizare a unui astfel de proces va conduce la concluzii eronate.

În cazul unei serii de date, cel mai frecvent prin analiza statistică se urmărește înlocuirea seriei de date cu o valoare, la care se asociază un interval de incertitudine. În acest context, modelul matematic asociat unei serii de date este:

$$y_i = a + \varepsilon_i, \quad (4.1)$$

unde a este o constantă, iar ε_i este eroarea aleatoare asociată valorii i a șirului. În funcție de indicatorul statistic utilizat pentru determinarea constantei a , se stabilește mărimea intervalului de incertitudine asociat. Aceasta este abordarea ce se efectuează pentru a se aprecia, de exemplu, dacă un lot de repere îndeplinesc condițiile de calitate, respectiv dacă se încadrează în limitele impuse de intervalele de toleranță.

Pentru ca modelul matematic asociat să fie corect este necesar să fie îndeplinite toate cele patru ipoteze fundamentale. Testarea ipotezelor asigură valabilitatea concluziilor. Verificarea ipotezelor se face cu ajutorul unor tehnici grafice. Tehnicile utilizate sunt:

- graficul secvențial al punctelor $Y_i(i)$
- graficul punctelor succesive $Y_i(Y_{i-1})$
- histograma
- graficul probabilității normale $Y_{\text{exp}}(Y_{\text{estimat normal}})$ – valorile experimentale se reprezintă în raport cu valorile estimate în cazul când datele au o distribuție normală.

Pentru a deduce concluzii corecte din analiză este necesară găsirea unui

model corespunzător și a unor estimări corecte ale parametrilor. În cazul când unele ipoteze nu sunt respectate pot apare probleme:

- Neglijarea caracterului aleator conduce la: invalidarea majorității testelor statistice, lipsa de semnificație a incertitudinilor calculate, invalidarea modelului $y_i = a + \varepsilon_i$ etc.
- Un aspect specific al caracterului nealeator este autocorelația, adică legătura dintre y_t și y_{t-k} , unde k este un număr întreg ce definește decalajul pentru autocorelație. Acest lucru este valabil la seriile cronologice. Fenomenul se depistează din graficul de autocorelație.
- Nerespectarea localizării constante. Estimarea uzuală pentru localizare este media. Dacă acest lucru nu este real atunci centrarea poate fi deplasată, estimarea centrării nu mai are semnificație, iar formula incertitudinii mediei ($S_{\bar{y}} = \frac{S}{\sqrt{n}}$ - S fiind abaterea standard) indică valori mult mai mici decât cele reale.
- Aceleași aspecte apar în cazul nerespectării variabilității constante.

Uzual, media este estimatorul localizării. Variabilitatea mediei este legată de ipotezele admise ale distribuției de probabilitate a datelor. Pentru anumite tipuri de distribuții media nu este estimatorul cel mai bun. Se poate utiliza media, mediana, mijlocul amplitudinii datelor, modul. Din acest motiv trebuie stabilită întâi distribuția și după aceea estimatorul.

4.2. TEHNICI GRAFICE UTILIZATE ÎN ANALIZA EXPLORATORIE A DATELOR UNIDIMENSIONALE

4.2.1. GRAFICUL SECVENȚIAL AL PUNCTELOR (RUN SEQUENCE PLOT)

Scop: verifică deplasări ale localizării, variabilității și prezența valorilor aberante.

Reprezentare: $y_i(i)$.

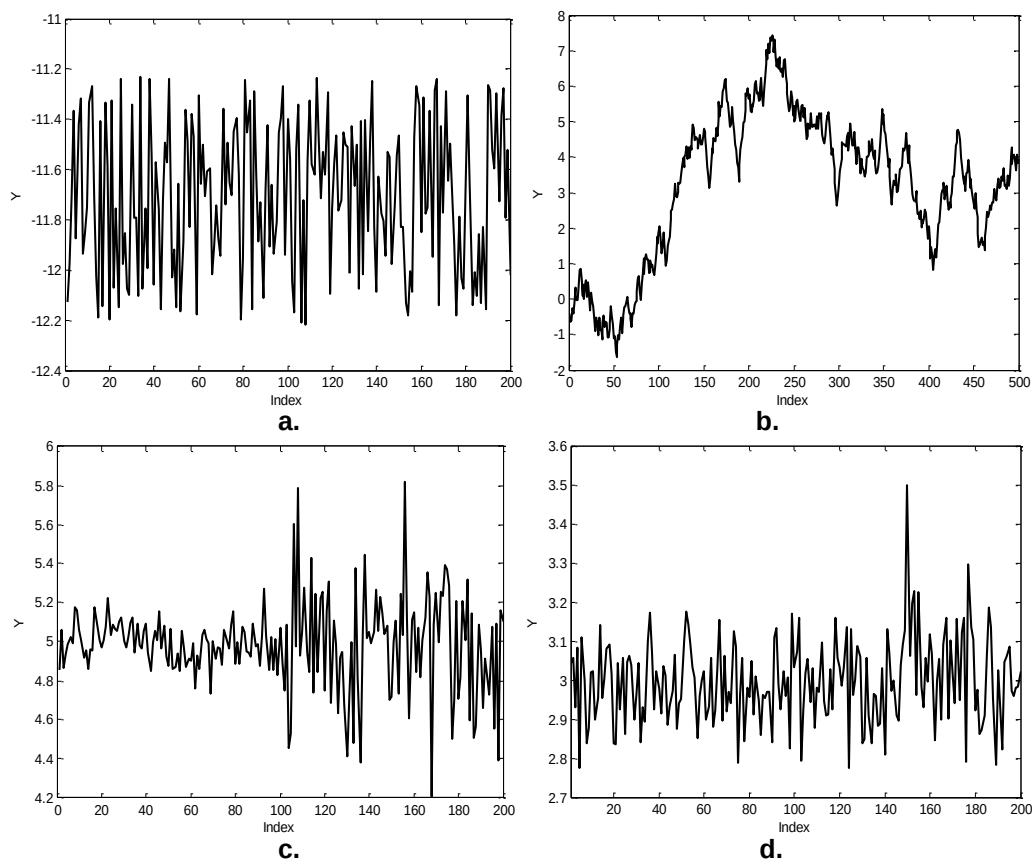


Fig. 4.2. Graficul secvențial al punctelor

În fig. 4.2 se prezintă două exemple de astfel de grafice. În primul caz, fig. 4.2.a, se observă că localizarea setului de date este constantă – dacă s-ar înlocui graficul printr-o dreaptă, aceasta ar fi o dreaptă paralelă cu axa Ox ; variabilitatea șirului este de asemenea constantă – variațiile de o parte și alta în jurul valorii medii este aproximativ aceeași, adică proiecția pe axa Oy în diferite zone ale graficului au aceeași valoare; nu apar valori aberante. În fig. 4.2.b setul de date prezintă modificări ale localizării, în prima jumătate apare o tendință crescătoare, iar în partea a doua una descrescătoare, nu apar modificări de variabilitate sau valori aberante. În fig. 4.2.c se remarcă modificarea de variabilitate care apare în a doua jumătate a setului de date, aproximativ de la jumătate. Localizarea este constantă, iar din punctul de vedere al valorilor aberante, există câteva valori situate la distanță mai mare de valoarea medie, valori ce pot apare datorită creșterii variabilității setului de date. În fig. 4.2.d localizarea setului de date este constantă, variabilitatea de

asemenea, apare însă suspiciunea unei valori aberante, cea situată în jurul valorii 3.5.

Interpretare: graficul trebuie să poată fi aproximat cu o dreaptă paralelă cu axa absciselor, să aibă amplitudinea în direcția axei Oy aproximativ constantă, să nu apară valori situate la distanță mare de restul valorilor. Această interpretare a graficului se poate face și prin metode cantitative, în sensul că se aproximează graficul cu dreapta de regresie (x este numărul de ordine al elementelor, y este seria de date). Panta dreptei de regresie trebuie să fie 0. Se verifică valoarea pantei dacă diferă semnificativ de 0.

4.2.2. GRAFICUL PUNCTELOR SUCCESIVE (LAG PLOT)

Scop: verificarea caracterului aleator al datelor.

Reprezentare: Se reprezintă $Y_i(Y_{i-1})$, adică perechea de puncte (Y_{i-1}, Y_i) .

Interpretare: În cazul datelor aleatoare reprezentarea nu evidențiază nicio structură. Aspectul datelor este asemănător cu cel al unei ținte dintr-un poligon de tragere.

În fig. 4.3 se prezintă câteva exemple de astfel de grafice. În primul caz (fig. 4.3.a) se remarcă aspectul nestructurat al graficului, ceea ce indică un caracter aleator. Nu par să fie valori aberante neexistând puncte izolate. În graficul 4.3.b se observă un model liniar, ceea ce indică un caracter puternic nealeator. În acest caz un model de autoregresie pare mai potrivit. Datele se grupează de-a lungul unei drepte. Pentru o anumită valoare a unui punct, de ex. $y_{i-1} = 0$, se observă că $y_i \in (-3; 2)$, deci se pot face predicții asupra valorilor succesive ale șirului de date. Se sugerează un model de forma:

$$y_i = a_0 + a_1 \cdot y_{i-1} + \varepsilon_i \quad (4.2)$$

În fig. 4.3.c modelul de autoregresie este prezent din nou, de data aceasta cu o autocorelație mai puternică. În acest caz, cât și în cel precedent nu apar valori aberante. Se remarcă gruparea clară a datelor de-a lungul unei drepte. Fenomenul de autocorelare se poate datora fenomenului studiat, variațiilor condițiilor de mediu, distorsionării datelor de la sistemul de achiziție.

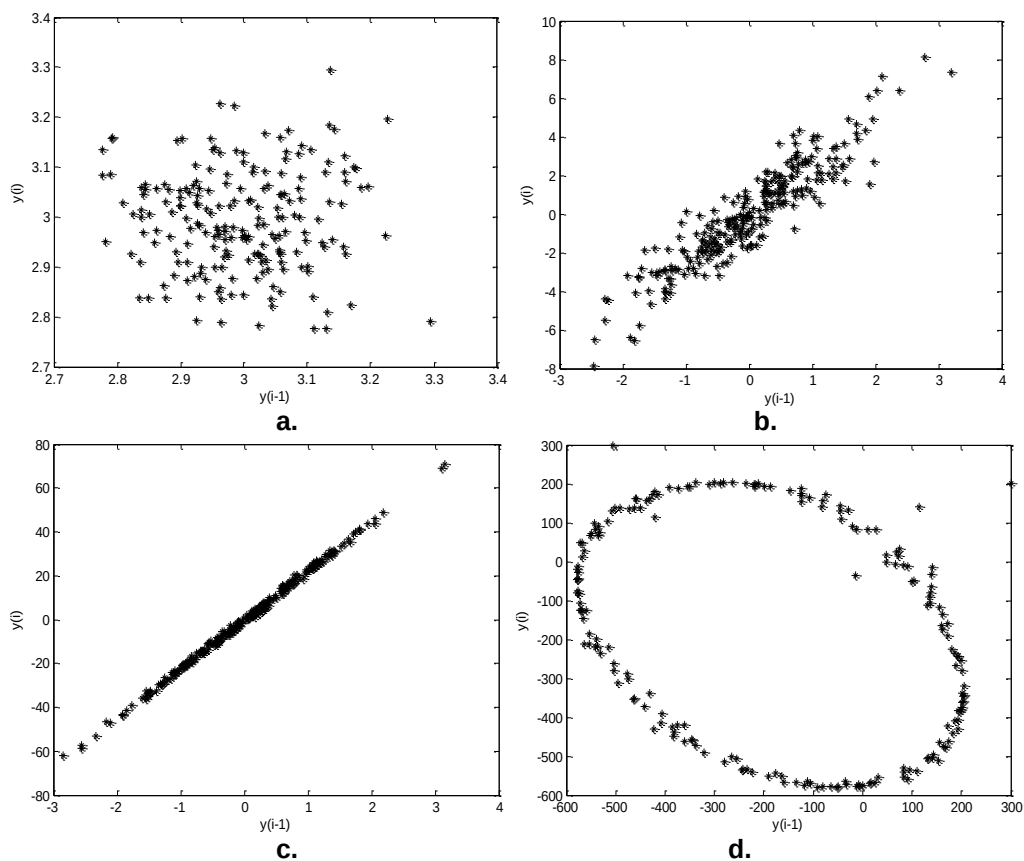


Fig. 4.3. Graficul punctelor succesive

În fig. 4.3.d apare un model de formă eliptică. Acest tip de grupare este caracteristica modelelor sinusoidale. Datele provin dintr-un model periodic de tip armonic; apar și valori aberante. Modelul corect este:

$$y_i = C + \alpha \sin(2\pi\omega t_i + \varphi) + \varepsilon_i. \quad (4.3)$$

iar parametrii modelului se pot estima cu metoda celor mai mici pătrate, [1]. Punctele mai îndepărtate de elipsă indică posibila prezență a valorilor aberante.

În concluzie, dacă datele au caracter aleator, graficul punctelor succesive are un aspect nestructurat, nu se pot face predicții legate de valorile următoare ale setului de date, deci între y_{i-1} și y_i nu există nici o relație.

4.2.3. HISTOGRAMA

Scop: reprezentarea distribuției datelor în intervale de lungime constantă.

Histograma indică: localizarea datelor, variabilitatea acestora, asimetria, prezența valorilor aberante, caracterul uni- sau multimodal al repartiției. Aceste caracteristici furnizează indicații clare referitoare la modelul corespunzător distribuției de probabilitate a datelor. Graficul de probabilitate sau un test de concordanță poate fi utilizat pentru verificarea modelului.

În fig. 4.4 se prezintă mai multe histogramme. În fig. 4.4.a se observă: simetria repartiției, existența unor extremități de anvergură moderată (cozi de lungime moderată), clasică formă de clopot. Acest tip de repartiție apare cel mai frecvent în natură. Dacă histograma este simetrică, cu anvergură moderată la extremități, nu există motive să nu se încerce modelarea setului de date cu o repartiție normală. Pentru confirmarea provenienței datelor din distribuția normală se construiește graficul probabilității normale. În cazul că acest grafic este liniar, atunci modelul este corespunzător.

În fig. 4.4.b histograma indică o repartiție diferită de cea normală având extremitățile fără coadă (short tail).

Pentru o repartiție simetrică "corpul repartiției" este centrul acesteia, regiunea unde apar probabilitățile maxime de apariție. Lungimea extremităților indică cât de repede aceste probabilități se apropie de zero. Repartițiile ce au o astfel de caracteristică au un caracter trunchiat, cum este repartiția uniformă – probabilitate constantă pe un domeniu și zero în rest. Setul de date care a generat această histogramă aparține unei repartiții uniforme.

Dacă cozile sunt moderate, probabilitatea ca datele să aparțină unor intervale îndepărtate de medie scade progresiv. Modelul clasic de acest tip este repartiția normală. Dacă coada este mare, ca în fig. 4.4.c, probabilitatea de apariție scade lent, există probabilitate de apariție la distanță mare de corpul repartiției. Modelul clasic pentru un astfel de set de date este repartiția Cauchy.

Estimarea optimă (nedeplasată și consistentă) pentru localizarea centrului repartiției depinde în mare măsură de anvergura extremităților. Alegerea uzuală a N observații și calcularea mediei acestora ca un estimator al localizării este bună în cazul repartiției normale (cozi medii), este o alegere nepotrivită pentru repartiții asemănătoare cu cea uniformă și total nepotrivită, eronată chiar, în cazul datelor provenind din distribuția Cauchy (cu anvergură mare). Media este un estimator corespunzător numai în cazul repartiției normale, [10].

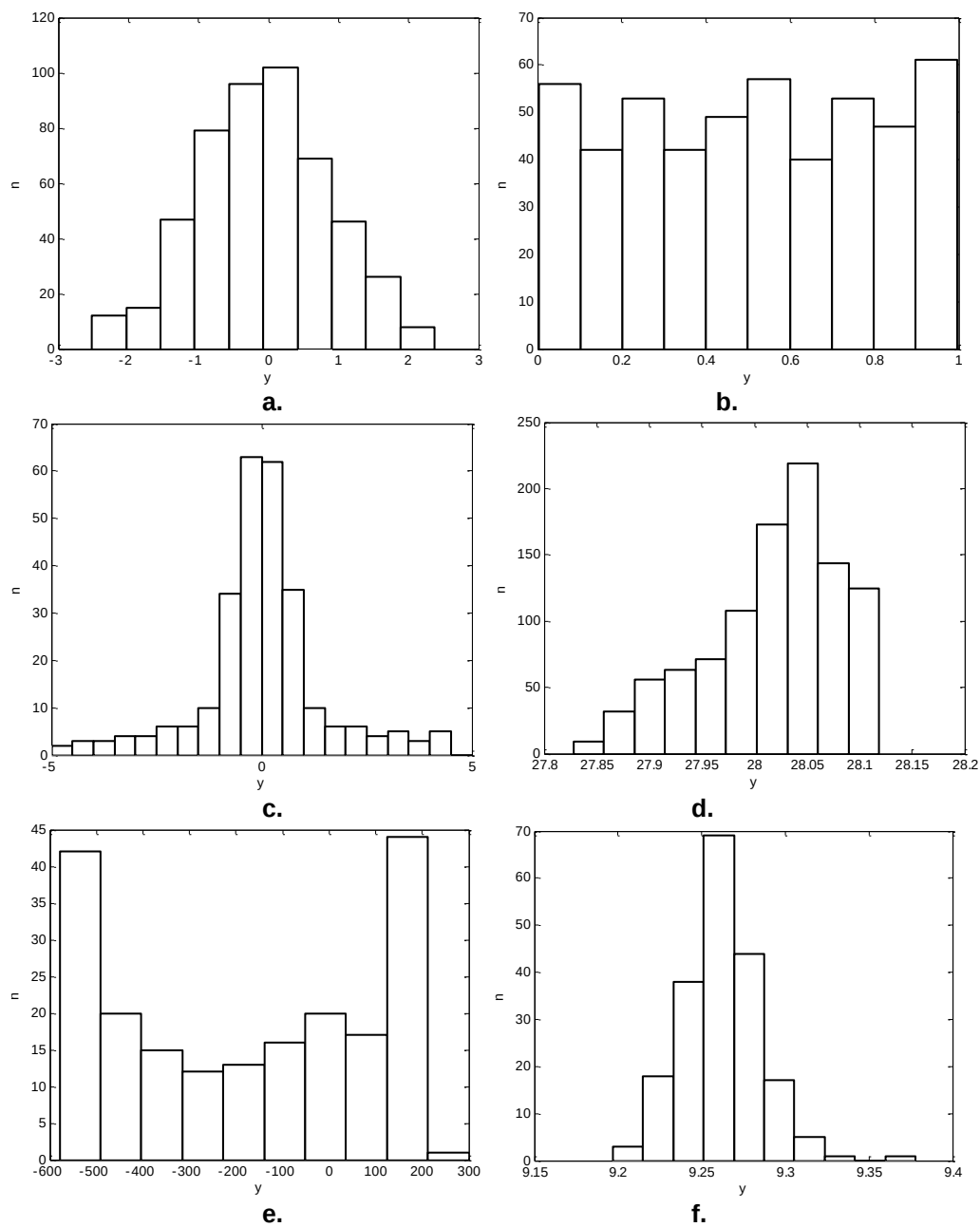


Fig. 4.4. Tipuri de histograme

Pentru repartiția uniformă cel mai bun indicator al localizării este mijlocul amplitudinii (semisuma valorii minime și maxime). Pentru repartiții tip Cauchy,

mediana este cel mai bun estimator al valorii centrale.

În fig. 4.4.d se prezintă o histogramă asimetrică. Pentru repartițiile asimetrice este uzual să aibe o extremitate considerabil mai lungă decât cealaltă. Asimetria dreapta are extremitatea cu anvergură mai mare în dreapta, respectiv asimetria stânga are extremitatea din stânga cu anvergură mai mare. Repartițiile asimetrice ridică probleme de estimare, media nu mai are consistență. Pentru aceste repartiții nu mai apare “centrul” în sensul uzual al cuvântului. Dintre indicatorii statistici, nici modulul nu prezintă semnificație deosebită.

La repartițiile simetrice media, mediana și modulul sunt identice, la cele asimetrice ele sunt distincte. În practică, repartițiile asimetrice se caracterizează cel mai frecvent prin medie, în unele cazuri prin mediană și cel mai rar prin modul. Cel mai corect este să se caracterizeze prin cel puțin doi indicatori, preferabil toți trei. Asimetria poate apare datorită limitării inferioare sau superioare a datelor. Limitarea inferioară generează o asimetrie dreapta, iar limitarea superioară o asimetrie stânga. Asimetria poate fi cauzată de efectul de pornire a unor procese. În aplicațiile de fiabilitate, unele procese pot avea un număr ridicat de căderi inițiale.

Dacă histograma indică asimetrie dreapta sau stânga se recomandă următorii pași:

- rezumarea datelor prin calcularea indicatorilor: medie, mediană și modul;
- să se determine cea mai bună aproximare din familia repartițiilor asimetrice, cum ar fi Weibull, gamma, exponențială etc.
- să se ia în considerare o schimbare de variabilă pentru normalizare.

Pentru multe repartiții este uzual ca răspunsul să fie grupat în jurul unei singure valori (modulul) și să se distribuie cu frecvențe mai scăzute spre extremități. Repartiția normală este exemplul clasic de repartiție unimodală. Modelul de histogramă din fig. 4.4.e prezintă clar aspectul de repartiție bimodală. Pentru a găsi explicații trebuie continuată analiza datelor:

- se construiește graficul secvențial al punctelor pentru a verifica eventualele tendințe ce apar;
- se construiește graficul punctelor succesive pentru a verifica caracterul armonic.

Histograma din fig. 4.4.f indică prezența valorilor aberante. Acestea trebuie investigate pentru a se găsi cauza apariției acestora. Eliminarea valorilor

aberrante prin aplicarea unor algoritmi poate fi înșelătoare. Se recomandă verificarea prezenței valorilor aberrante cu box-plot precum și verificarea cantitativă prin testul Grubbs.

Reprezentarea histogramelor în Matlab se face utilizând frecvența absolută a seriei de date (fig. 4.5).

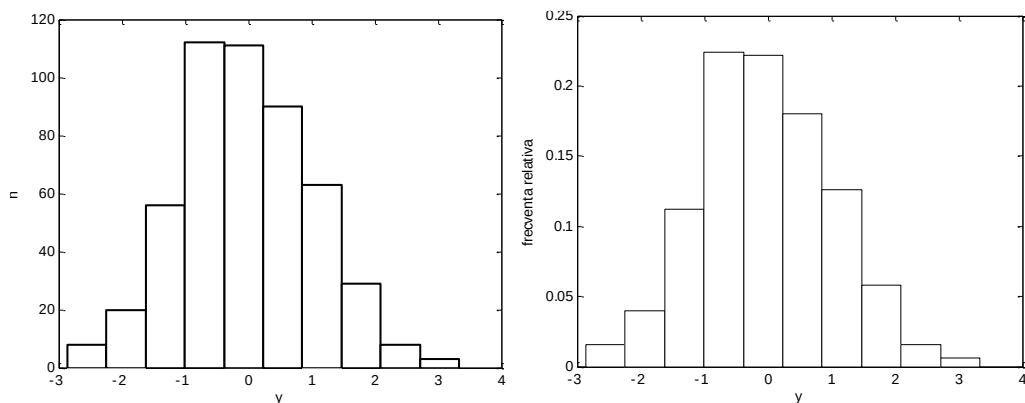


Fig. 4.5 Histograma frecvențelor absolute și relative

Pentru a obține histograma frecvențelor relative, frecvențele absolute ale fiecărui interval trebuie împărțite cu numărul total de valori al șirului. Cele două grafice prezentate în fig. 4.4 se obțin cu următoarea secvență de program:

```
hist(y),xlabel('y'),ylabel('n')
%setare contur cu negru si dreptunghiuri albe
h = findobj(gca,'Type','patch');
set(h,'FaceColor','w','EdgeColor','k')
[n,limite]=hist(y);
figure,bar(limite,n/length(y),1),
h = findobj(gca,'Type','patch');
set(h,'FaceColor','w','EdgeColor','k')
```

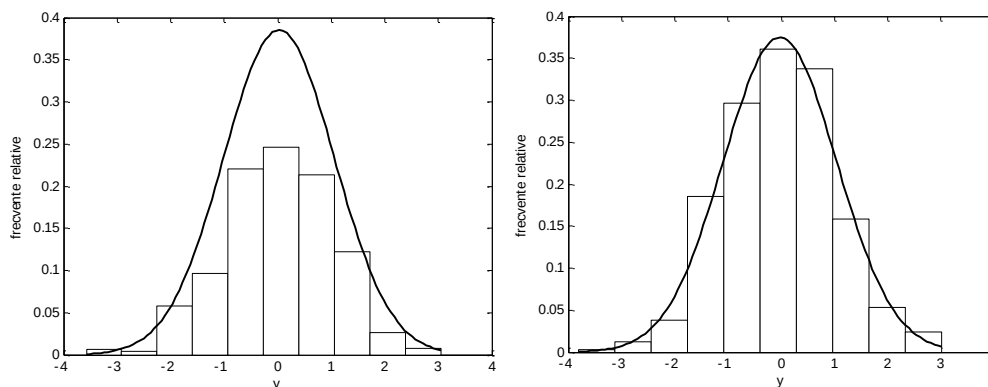


Fig. 4.5 Histograma frecvențelor relative și normalizate împreună cu funcția densității de probabilitate

Pentru ca reprezentarea histogramei să poată fi comparată cu funcția densitate de probabilitate este necesar să se facă normalizarea histogramei, adică suma ariilor dreptunghiurilor trebuie să fie egală cu 1. În fig. 4.5 sunt reprezentate histogramele aceluiasi set de date, având suprapusă funcția densitate de probabilitate normală, estimată pe baza parametrilor setului de date. Se remarcă în primul caz, când nu este făcută normalizarea, că funcția densitate de probabilitate este mult mai înaltă decât histograma, în timp ce în al doilea caz, cele două funcții sunt mult mai apropiate, histograma reprezentând mult mai veridic funcția densitate de probabilitate.

Construirea histogramei normalizate se face cu ajutorul următoarei secvențe de program:

```
% histograma normalizata
miu=mean(y);v=var(y);
xp=linspace(min(y),max(y));
yp=normpdf(xp,miu,v);
[n,limite]=hist(y);
hh=limite(2)-limite(1);
figure,bar(limite,n/(length(y)*hh),1),
h = findobj(gca,'Type','patch');
set(h,'FaceColor','w','EdgeColor','k'),hold on
plot(xp,yp)
```

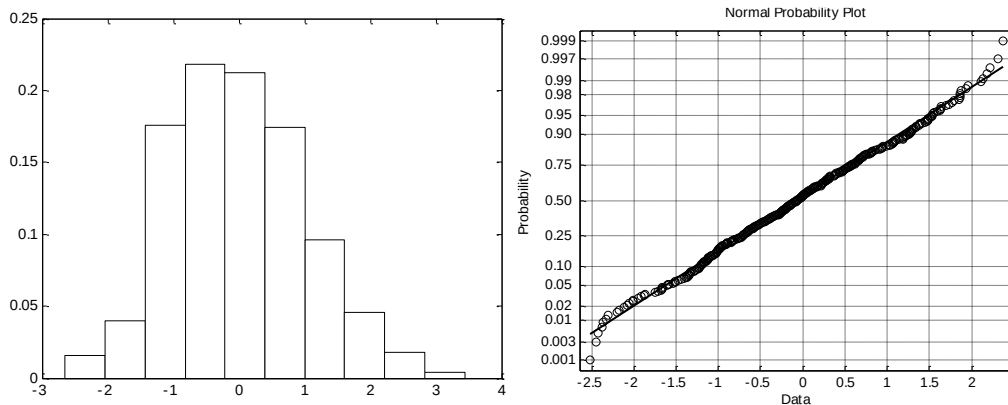
4.2.4. GRAFICUL PROBABILITĂȚII NORMALE (NORMAL PROBABILITY PLOT)

Scop: verificarea provenienței datelor dintr-o distribuție normală.

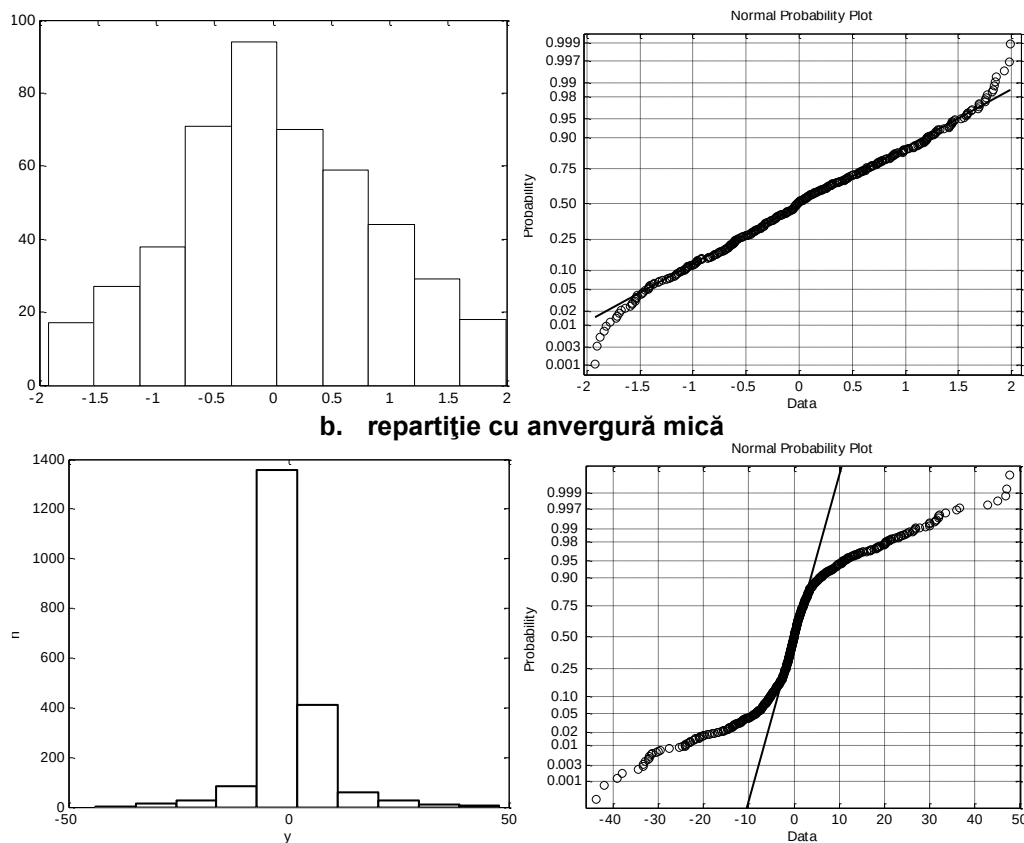
Interpretare: Datele se reprezintă față de repartiția normală teoretică într-un astfel de mod încât punctele ar trebui să se situeze pe o dreaptă. Îndepărtarea de dreaptă indică abateri de la repartiția normală. Graficul probabilității normale este un caz particular al graficului de probabilitate.

Reprezentare: Graficul se generează prin reprezentarea pe axa orizontală a setului de date ordonat, iar pe axa verticală a cvantilei corespunzătoare din repartiția normală. Similar se procedează și pentru alte repartiții, putându-se verifica apartenența la distribuția respectivă. În Matlab acest tip de grafic se construiește cu comanda: `normplot(y)`, `y` fiind setul de date.

Un avantaj al acestei metode este faptul că panta și intersecția în origine a dreptei generate este o estimare a indicatorilor de variabilitate, respectiv de localizare. Acest lucru nu este important la repartiția normală, dar devine semnificativ în cazul altor repartiții. Coeficientul de corelație al punctelor reprezentate pe grafic se poate compara cu valorile critice pentru a testa proveniența dintr-o populație repartizată normal.



a. anvergură medie



b. repartitie cu anvergură mică

c. repartitie cu anvergură mare

Fig. 4.6. Graficul probabilității normale pentru diferite tipuri de anverguri ale extremităților

Metoda furnizează răspunsuri la următoarele întrebări:

- sunt datele repartizate normal;
- care este natura îndepărtării de la normalitate (asimetria, extremități de anvergură mică sau prea mare)

Această metodă grafică oferă răspuns la ipoteza apartenenței la o anumită repartitie. Majoritatea modelelor statistice sunt de forma:

$$y_i = a + \varepsilon_i,$$

unde valoarea deterministă, a , se determină prin estimare, iar valoarea aleatoare este eroarea. Această componentă se ipotezează că este repartizată normal, cu localizarea și variabilitatea constantă. Aceasta este aplicația cea mai frecventă a metodei. Se aproximează un model și se generează graficul probabilității normale pentru erori. Dacă aceste erori nu sunt repartizate normal, înseamnă că una din ipotezele de bază a fost încălcată.

În fig. 4.6 sunt prezentate graficele probabilității normale pentru diferite anverguri ale extremităților asociate cu histogramamele aferente. Se observă alinierea la dreaptă în cazul a., setul de date fiind generat din distribuția normală. În cazul când anvergura extremităților este mică, graficul probabilității normale are o formă de „S”, situație ce se poate observa în fig. 4.6.b (forma de „S” este mai atenuată) și în 4.7, unde această formă este foarte clar evidențiată. Acest aspect de „S” apare și în cazul când setul de date provine dintr-o repartiție cu anvergură mare (fig. 4.6.c – set de date provenind dintr-o repartiție Cauchy). Diferența de amplitudine se manifestă prin aspectul punctelor de la extremitățile graficului:

- la anvergură mică: punctele de început sunt situate sub dreapta corespunzătoare repartiției normale, iar punctele de sfârșit sunt situate deasupra drepte;
- la anvergură mare: punctele de început sunt situate deasupra, iar cele de sfârșit sunt sub dreaptă.

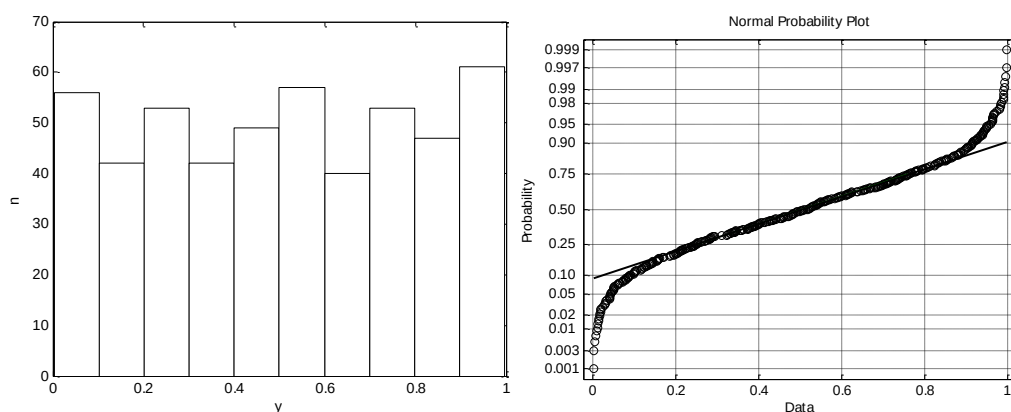


Fig. 4.7 Graficul probabilității normale pentru repartiția uniformă

În fig. 4.8 sunt prezentate graficele probabilității normale pentru două seturi de date ce provin din repartiții asimetrice. Caracteristica după care se poate identifica asimetria este aspectul pătratic al graficului. În cazul unei asimetrii dreapta (anvergura repartiției este amplă în partea dreaptă a axei), fig. 4.8.a, curba probabilității normale are punctele de început și sfârșit situate în partea inferioară drepte corespunzătoare repartiției normale. Dacă setul de date prezintă o asimetrie stânga, ca în fig. 4.8.b, aspectul pătratic se păstrează, dar punctele de început și sfârșit ale curbei sunt situate în partea superioară drepte.

În cazul când la analiza unui set de date se observă un grafic al probabilității normale care indică asimetrie, este necesară modelarea setului de date printr-

o distribuție asimetrică, cum ar fi distribuția exponențială, distribuția Weibull etc.

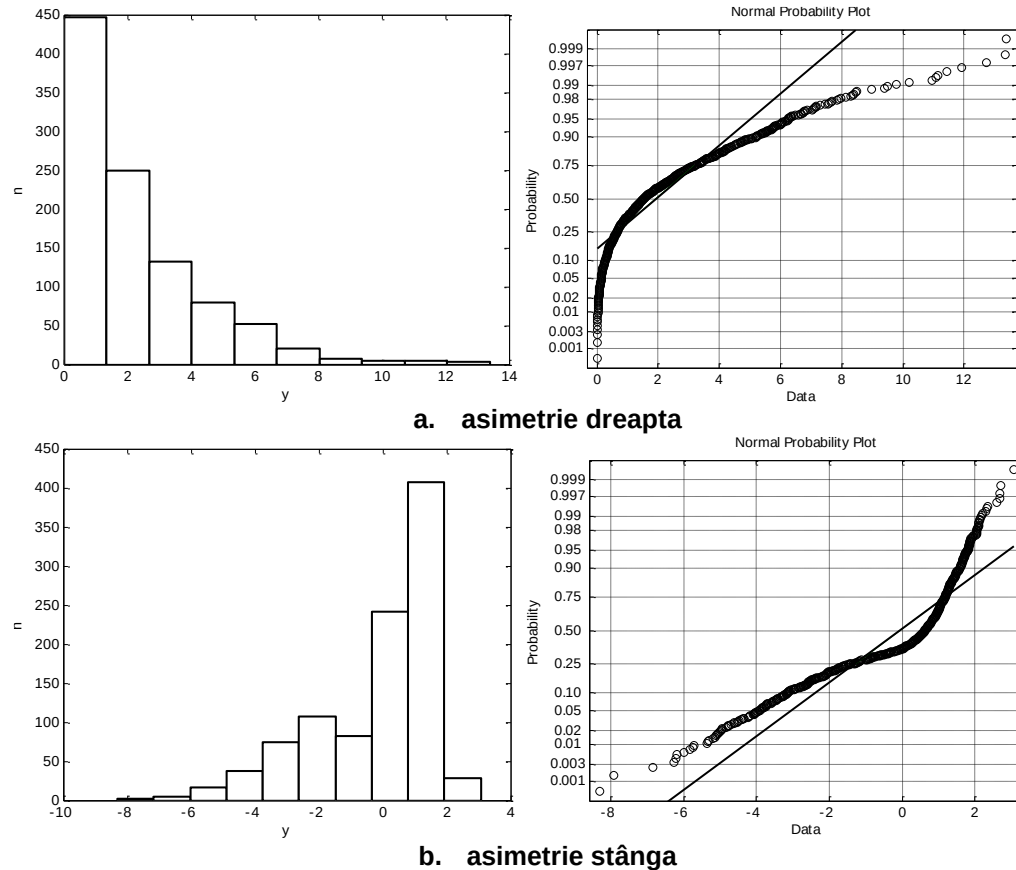


Fig. 4.8 Graficul probabilității normale pentru repartiții asimetrice

4.2.5. GRAFICUL CELOR PATRU

Scop: verificarea ipotezelor statistice fundamentale: caracterul aleator, proveniența dintr-o anumită distribuție de probabilitate, localizarea și variabilitatea constantă.

Acest grafic este, de fapt o colecție de 4 tehnici grafice:

- graficul secvențial al punctelor (Run Sequence Plot);
- graficul punctelor succesive (Lag Plot);
- histograma;

- graficul probabilității normale.

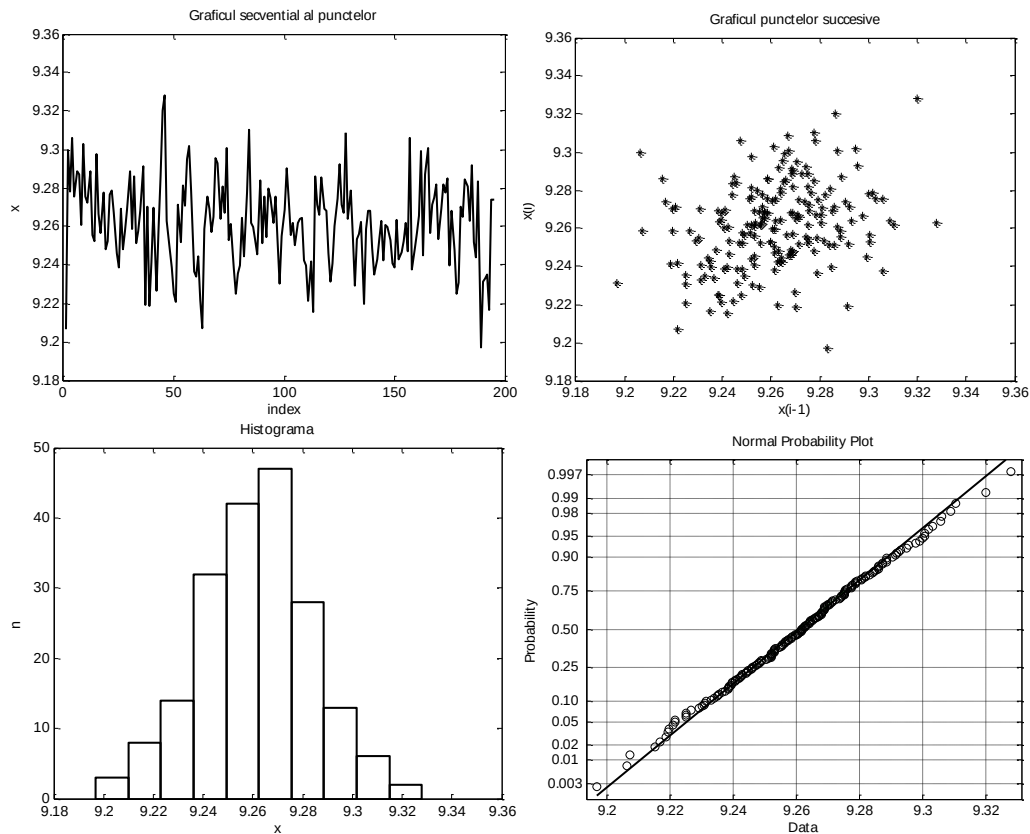


Fig. 4.9. Graficul celor 4 în cazul unui set de date provenind dintr-o repartiție normală

Dacă cele 4 ipoteze fundamentale ale unui proces de măsurare sunt verificate, cele 4 grafice vor avea o alură caracteristică. Dacă una din ipoteze nu este îndeplinită, atunci această anomalie va fi evidențiată de unul sau mai multe grafice.

Deși cele 4 grafice se utilizează la șiruri unidimensionale și serii cronologice, este util să fie extinse și mai departe. Multe modele statistice se modelează prin:

$$y_i = f(x_1, x_2, \dots, x_n) + \varepsilon_i, \quad (4.4)$$

adică cele 4 ipoteze trebuie să fie valabile pentru erorile ε_i , deci metoda se utilizează la validarea modelului prin analiza erorilor.

Se recomandă ca cele patru grafice să fie reprezentate pe aceeași pagină pentru a putea fi interpretate cu ușurință.

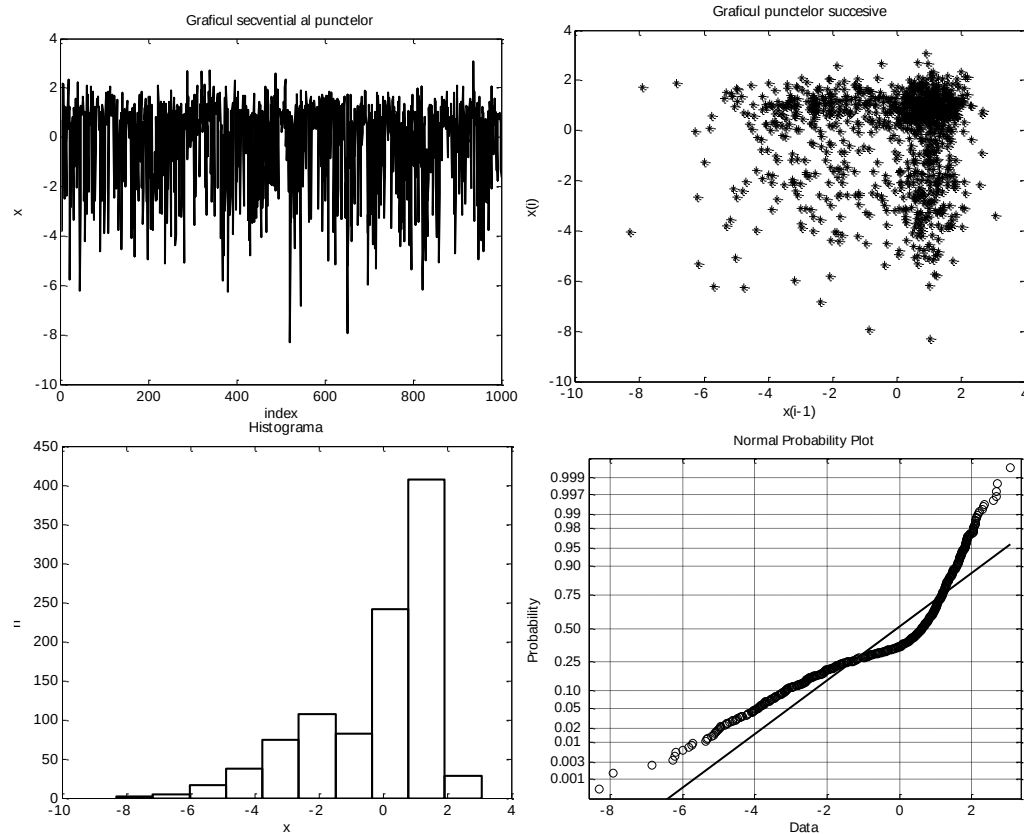


Fig. 4.10. Graficul celor patru pentru un set de date experimentale

Analizând cele patru grafice din fig. 4.9 se poate constata:

- graficul secvențial al punctelor indică variabilitate și localizare constantă, nu are valori aberante;
- graficul punctelor succesive indică existența unui caracter aleator;
- histograma are forma clasică de clopot, deci repartiția de proveniență este simetrică, are anvergură moderată a extremităților, nu are valori aberante;
- graficul probabilității normale are un caracter liniar, indicând faptul că datele provin dintr-o distribuție normală.

În concluzie, nu sunt motive să nu se considere setul de date ca provenind dintr-o repartiție normală, caz în care localizarea setului de date se poate

caracteriza prin media șirului, iar variabilitatea se poate caracteriza pe baza abaterii standard.

În fig. 4.10 este prezentat un alt exemplu pentru graficul celor patru pentru un set de date experimentale. Se poate constata:

- localizarea și variabilitatea setului de date sunt aproximativ constante;
- datele au caracter aleator;
- repartiția datelor prezintă o asimetrie stânga;
- repartiția nu este o repartiție normală.

4.2.6. GRAFICUL DE AUTOCORELARE

Scop: verificarea caracterului aleator.

Caracterul aleator se verifică prin calcularea autocorelației pentru setul de date corespunzătoare la diferite valori ale decalajului. Dacă procesul este aleator, atunci autocorelația trebuie să fie 0 pentru orice decalaj. Dacă setul de date nu este aleator pentru unul sau mai multe decalaje, atunci autocorelația este semnificativ diferită de 0.

Reprezentare: Graficul se construiește prin reprezentarea pe axa verticală a mărimii R_h , unde:

$$R_h = \frac{C_h}{C_0}, \quad (4.5)$$

C_h este autocovarianța:

$$C_h = \frac{1}{N} \sum_{t=1}^{N-h} (Y_t - \bar{Y})(Y_{t+h} - \bar{Y}), \quad (4.6)$$

iar C_0 este dispersia de selecție necentrată:

$$C_0 = \frac{\sum_{t=1}^N (Y_t - \bar{Y})^2}{N}. \quad (4.7)$$

Pe axa orizontală se reprezintă decalajul $h = 1, 2, \dots, n-1$. Se mai trasează și 5 segmente de referință: un segment central la cota 0, iar celelalte 4 segmente delimitează intervalele de încredere de 95% și 99%. Limitele de încredere se recomandă să fie calculate cu formula:

$$\pm \frac{Z_{1-\frac{\alpha}{2}}}{\sqrt{N}},$$

unde N este volumul eșantionului, z cvantila repartiției normale, iar α nivelul de semnificație.

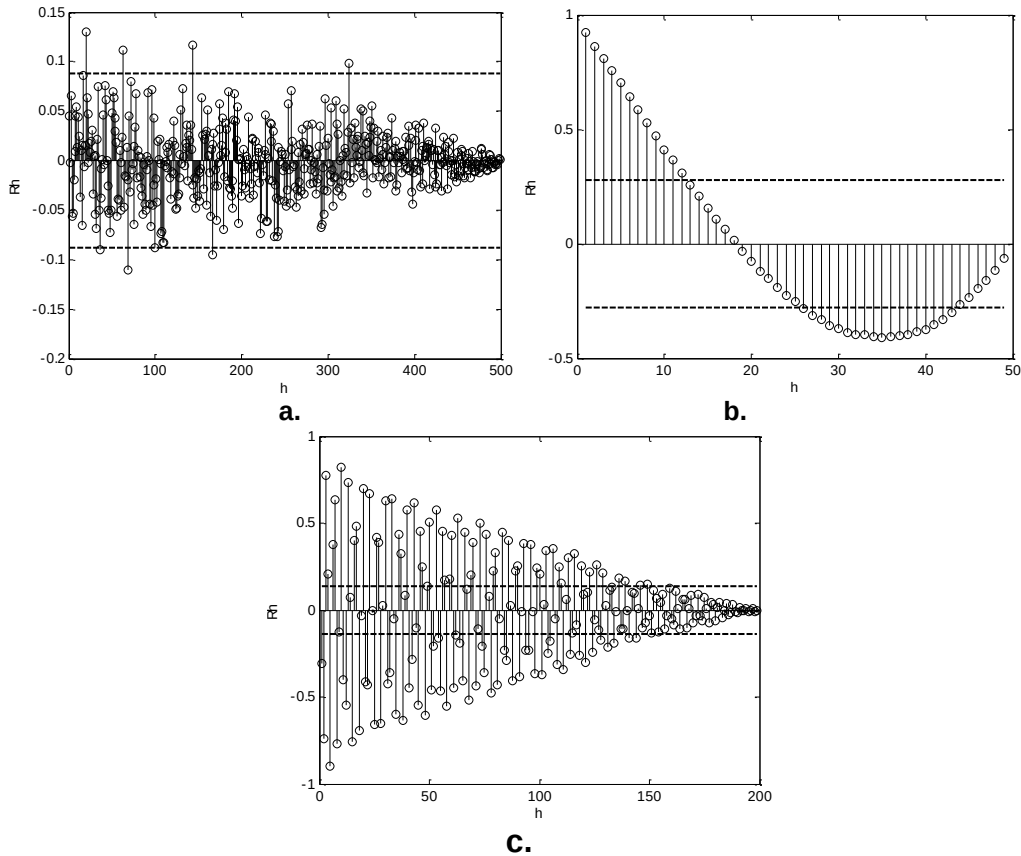


Fig. 4.11. Grafice de autocorelare

Construirea acestui grafic și interpretarea corectă asigură validitatea concluziilor. Caracterul aleator este una din caracteristicile fundamentale ale proceselor de măsurare. Caracterul aleator are o importanță critică din următoarele motive:

- Formula utilizată cel mai frecvent pentru deviația standard a mediei de selecție este $S_y = S / \sqrt{N}$. Aceasta este adevărată numai în condițiile caracterului aleator .
- Pentru date unidimensionale, modelul implicit este $Y = a + \varepsilon$, a fiind o constantă (estimarea valorii centrale) și ε eroarea. Acest model

devine incorect și nu mai este valid în cazul datelor nealeatoare, iar estimarea parametrilor nu mai are semnificație.

În fig. 4.11 sunt prezentate trei exemple de grafice de autocorelare, pe care sunt marcate cu linie întreruptă limitele intervalului de incertitudine corespunzător unui nivel de încredere de 95%. În primul caz se observă că nu apare o autocorelație semnificativă. Datele sunt aleatoare. Toate valorile sunt în intervalul de incertitudine de 95%, nu apare nici un tipar. Lipsa unei structurări a graficului indică caracterul aleator. Câteva valori sunt situate puțin în afara intervalului de încredere de 95%. Pentru o probabilitate de 95% este normal ca o valoare din 20 să fie în afara intervalului de încredere. Nu apare nici o asociativitate între Y_i și Y_{i+1} . Acest lucru este chiar esența caracterului aleator.

În fig. 4.11.b datele provin dintr-un model de autoregresie cu autocorelație puternică. Graficul indică autocorelare mare la decalaj 1 și scade până la valori negative. Descreșterea este aproximativ liniară cu zgomot scăzut (unirea punctelor reprezentate generează o curbă netedă). Un astfel de tipar indică o autocorelare puternică, adică o previziune ridicată. Pentru pasul următor, se indică estimarea parametrilor pentru modelul de autoregresie:

$$Y_i = A_0 + A_1 \cdot Y_{i-1} + \varepsilon_i.$$

Această estimare se face cu metoda celor mai mici pătrate, mai precis pe baza dreaptelor de regresie.

În cel de-al treilea exemplu, fig. 4.11.c, datele provin din model armonic. Apare o secvență alternativă de vârfuri pozitive și negative, care manifestă tendință de scădere.

Construirea graficului de autocorelare se face cu următoarea secvență de program Matlab:

```
load x1000.dat; y=x1000;
n=length(y);
c0=var(y,1);
C(1:n-1)=0;
for i=1:n-1
    temp=0;
    for j=1:n-i
        temp=temp+(y(j)-mean(y))*(y(i+j)-mean(y));
    end
    C(i)=temp/n;
end
```

```

R=C/c0;
h=1:(n-1);
stem(h,R)
hold on
ls(1:(n-1))=1.96/sqrt(n);
li=-ls;
plot(h,ls,'r',h,li,'r')

```

4.2.7. GRAFICUL DE INCERTITUDINI (BOOTSTRAP PLOT)

Scop: determinarea incertitudinii unei statistici.

Din setul de date se extrage un eșantion de volum mai mic sau egal cu volumul datelor inițiale. Eșantionul se obține cu returnare, astfel încât fiecare element poate fi prelevat de mai multe ori sau deloc. Acest proces se repetă de mai multe ori, în general cel puțin de 500 ori. Valorile calculate (mediana, mijlocul intervalului de variație) se constituie într-o estimare a repartiției eșantionului.

Pentru fiecare eșantion se calculează indicatorul dorit, ex. mediana. În urma eșantionărilor repetate se obțin 500 de valori pentru mediană. Pentru determinarea incertitudinii, se ordonează șirul de 500 valori și se selectează valoarea din poziția 25, respectiv 475. În aceste condiții incertitudinea este determinată cu un grad de încredere de 90% ($25/500 = 0.05$, $475/500 = 0.95$, deci valorile determinate reprezintă o estimare a cvantilei). Cele 2 valori reprezintă limita inferioară și cea superioară pentru incertitudine.

Reprezentare:

- pe axa verticală valoarea calculată pentru fiecare eșantion;
- pe axa orizontală – numărul eșantionului.

În general, acest grafic se asociază și cu histograma aferentă, pentru a urmări atât localizarea, cât și împrăștierea statisticii datelor.

În fig. 4.12 se prezintă graficul pentru medie, mediană și centrul intervalului de variație, împreună cu histogramele aferente. Se observă că variabilitatea minimă o are centrul intervalului de variație. În aceste condiții, acest indicator se recomandă pentru estimarea localizării setului de date ce a generat graficul.

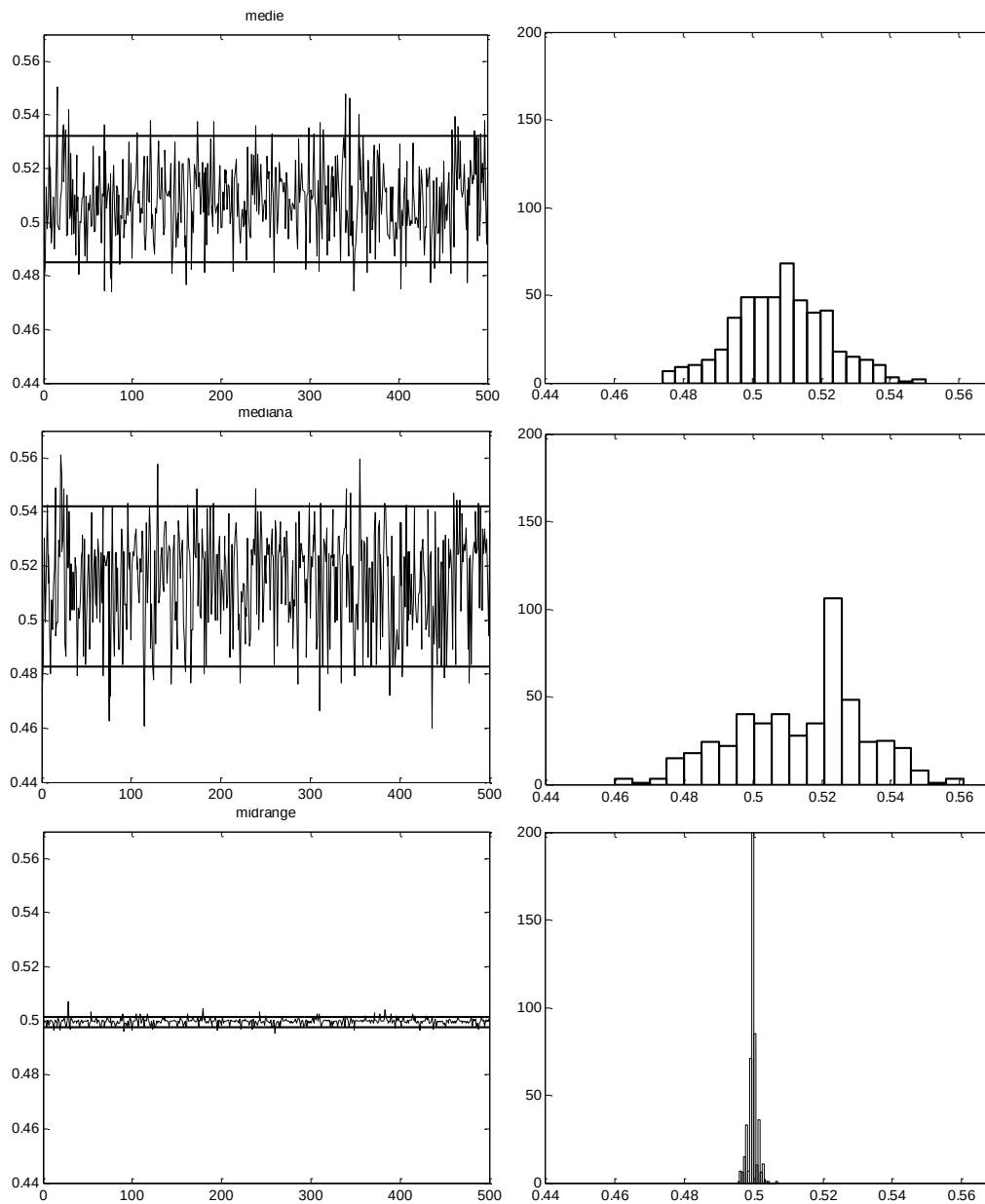


Fig. 4.12. Graficul pentru incertitudini

Construirea acestor grafice se poate face cu următoarea secvență de program:

```
load y1_500.dat
y=y1_500;
n=length(y); z(1:n)=0;
```

```

medie(1:500)=0;mediana(1:500)=0;midrange(1:500)=0;
for i=1:500
    indice=randsample(n,n,true);
    for j=1:n
        z(j)=y(indice(j));
    end
    medie(i)=mean(z);
    mediana(i)=median(z);
    midrange(i)=(min(z)+max(z))/2;
end
x=1:500;
medie_ord=sort(medie);
mediana_ord=sort(mediana);
midrange_ord=sort(midrange);
li(1:500)=medie_ord(25);ls(1:500)=medie_ord(475);
figure,plot(x,li,x,ls,x,medie),title('medie'),axis([0 500 .44
.57])
figure,hist(medie,20),axis([0.44 .57 0 200])
li(1:500)=mediana_ord(25);ls(1:500)=mediana_ord(475);
figure,plot(x,li,x,ls,x,mediana),title('mediana'),axis([0 500
.44 .57])
figure,hist(mediana,20),axis([0.44 .57 0 200])
li(1:500)=midrange_ord(25);ls(1:500)=midrange_ord(475);
figure,plot(x,li,x,ls,x,midrange),title('midrange'),axis([0 500
.44 .57])
figure,hist(midrange,20),axis([0.44 .57 0 200])

```

În program s-au impus scalările graficelor pentru a facilita comparația între cei trei indicatori.

4.2.8. GRAFICUL CVANTILĂ – CVANTILĂ (Q – Q PLOT)

Scop: Verifică dacă două seturi de date provin din aceeași distribuție de probabilitate. De fapt, se verifică dacă cele 2 seturi de date provin din populații având aceeași repartiție. De asemenea, cu ajutorul acestui tip de grafic se poate compara un eșantion cu o distribuție de probabilitate prescrisă, pentru a verifica ipoteza provenienței datelor din distribuția de probabilitate respectivă. În această situație se ajunge la grafice de probabilitate, cum este cazul graficului probabilității normale, prezentat anterior.

Reprezentare: Această metodă grafică reprezintă cvantila primului set de date în raport cu cvantila setului al doilea, conform metodologiei prezentate în 3.2.2. Fie cele două seturi de date: $X: x_1, x_2, \dots, x_n$ și $Y: y_1, y_2, \dots, y_m$ cu $m \leq n$. Se ordonează ambele șiruri crescător: $X' = x'_1, x'_2, \dots, x'_n$ și $Y' = y'_1, y'_2, \dots, y'_n$. În cazul

când eşantioanele au volume egale ($n = m$) se reprezintă puncte având coordonatele egale cu cvantilele corespondente ale celor două seturi de date. (În acest caz coordonatele punctelor sunt chiar elementele şirurilor ordonate $P_i(x'_i, y'_i)$).

Un avantaj major al acestei metode grafice este faptul că nu necesită eşantioane de volum egal. În situaţia când $m < n$, se vor reprezenta puncte y'_i în raport cu cvantila $(i-0.5)/m$ a celui alt şir.

Interpretare: În cazul când datele provin din populaţii cu aceeaşi repartiţie, punctele prezintă doar o abatere uşoară faţă de dreapta de referinţă. Cu cât punctele sunt situate la o distanţă mai mare de linia de referinţă, cu atât este mai evident faptul că seturile provin din repartiţii diferite. Graficul oferă informaţii calitative şi pentru ca rezultatul să aibă credibilitate este necesar ca volumul datelor să fie mare.

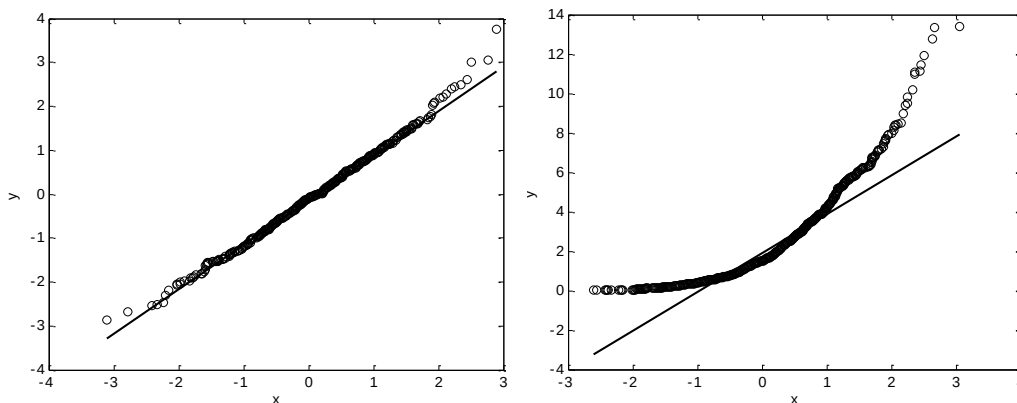


Fig. 4. 13. Grafice cvantila – cvantila

Principalul avantaj al acestei metode constă în faptul că eşantioanele nu trebuie să fie de acelaşi volum.

Dacă datele provin din populaţii ale căror repartiţie diferă doar prin localizare, punctele se vor situa pe o dreaptă aproximativ paralelă cu linia de referinţă, dar decalată în sus sau jos faţă de aceasta.

În fig. 4.13 se prezintă două grafice cvantilă – cvantilă, primul generat pe baza a două eşantioane provenite din repartiţia normală, iar cel de-al doilea pentru un eşantion provenit dintr-o repartiţie normală şi unul din repartiţie exponenţială. Se observă alinierea la o dreaptă în primul caz, ceea ce confirmă apartenenţa la o aceeaşi repartiţie a celor două eşantioane. Pentru a interpreta

mai ușor graficul, la grafic se mai atașează o dreaptă obținută pe baza cvantilelor inferioare și superioare ale celor două șiruri. Dreapta din grafice este determinată de două puncte, un punct are coordonatele egale cu cvantila inferioară a șirului x , respectiv y , iar cel de-al doilea punct are coordonatele egale cu cvantilele superioare ale șirurilor. Dreapta se extinde până la valorile minime, respectiv maxime ale datelor. În Matlab, acest grafic se poate construi direct cu comanda `qqplot(x,y)`, unde x și y sunt vectori ce conțin seriile de date.

Dacă există 2 seturi de date este important să se știe dacă ipoteza unei repartiții comune este justificată. Dacă da, localizarea și variabilitatea pot fi estimate pentru ambele loturi. Dacă loturile diferă, este important să se cunoscă clar diferențele.

Graficul cvantilă – cvantilă este similar graficului de probabilitate. Pentru un grafic de probabilitate, cvantilele unui set de date se înlocuiesc cu cele ale repartiției teoretice. În fig. 4.14 se prezintă graficul repartiției uniforme pentru două seturi de date.

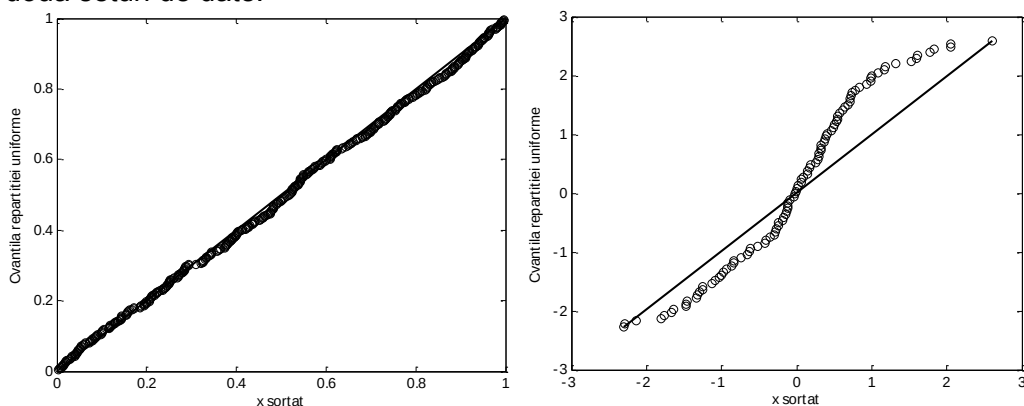


Fig. 4.14. Graficul repartiției uniforme

Se observă în primul grafic din fig. 4.14, că punctele se aliniază pe o dreaptă, ceea ce confirmă proveniența setului de date dintr-o distribuție uniformă. În cel de-al doilea caz, apar abateri semnificative de la dreaptă, dovadă a faptului că punctele nu provin dintr-o distribuție de acest tip. Setul de date utilizat în acest caz a fost generat dintr-o distribuție normală.

Construirea unui astfel de grafic se face cu următoarea secvență de program:

```
load y1_500.dat
x=y1_500;
n=length(x);
```

```

prob=(1:n)-.5)/n;% se calculeaza probabilitatile
corespunzatoare setului de date
qp=unifinv(prob,min(x),max(x));% se determina cvantilele
repartitiei teoretice
plot(sort(x),qp,'o')
% Se ataseaza dreapta de referinta
xa=quantile(x,.25);xb=quantile(x,.75);
ya=0.25;yb=0.75;
x_dr=[xa xb];y_dr=[ya yb];
figure,plot(sort(x),qp,'o',x_dr,y_dr)
m=(yb-ya)/(xb-xa);
ymin=ya+m*(min(x)-xa);
ymax=ya+m*(max(x)-xa);
dreapta_x=[min(x) max(x)];dreapta_y=[min(qp) max(qp)];
figure,plot(sort(x),qp,'o',dreapta_x,dreapta_y)

```

Prin înlocuirea funcției unifinv, ce definește cvantila distribuției uniforme, cu funcția corespunzătoare altei distribuții, se poate construi orice grafic de probabilitate dorit.

4.3. TEHNICI CANTITATIVE ASOCIATE ANALIZEI EXPLORATORII

Analiza exploratorie a datelor se asociază cu o serie de tehnici cantitative. Analiza cu tehnici grafice oferă doar informații calitative, ce trebuie validate cu metode cantitative. Majoritatea tehnicilor cantitative se împart în două mari categorii:

- a – estimări cu intervale de încredere;
- b – testarea ipotezelor.

De exemplu, cel mai uzual indicator al localizării este media. Media teoretică a populației se obține ca valoarea medie a variabilei aleatoare ce reprezintă caracteristica investigată a populației. Ea se estimează prin calcularea mediei unui eșantion și este o estimare punctuală. Intervalul de încredere extinde estimarea punctuală prin includerea unei incertitudini. Prin extragerea diferitelor eșantioane rezultă diferite medii. Se determină un interval pe baza unui nivel de încredere. Acest interval conține valoarea medie teoretică cu o anumită probabilitate.

Prin testarea ipotezei, în loc să se asocieze un interval de încredere, se respinge o anumită afirmație referitoare la parametrul propus pe baza

informației extrase din eșantion.

De exemplu, ipoteza nulă poate fi una din afirmațiile despre parametru:

- media populației este egală cu 10;
- deviația standard este egală cu 5;
- mediile a două populații sunt egale;
- deviațiile standard pentru cinci populații sunt egale.

Respingerea ipotezei nule înseamnă că este falsă. Acceptarea unei ipoteze nu înseamnă că este adevărată, doar că nu avem dovada să credem acest lucru. De aceea testarea ipotezei se face prin două ipoteze, una ce nu inspiră încredere – ipoteza nulă (H_0), și o ipoteză ce se crede adevărată – ipoteza alternativă (H_1).

4.3.1. ESTIMĂRI ALE LOCALIZĂRII

O sarcină fundamentală în numeroase analize statistice este estimarea unui parametru de localizare pentru repartiție, găsirea unei valori tipice sau centrale pentru a descrie cel mai bine datele.

În primul rând trebuie definită o valoare tipică. Pentru șiruri unidimensionale există trei estimări:

- media:

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

este valoarea utilizată cel mai frecvent pentru caracterizarea unei repartiții;

- mediana – valoarea ce împarte setul de date în două părți egale:

$$x_{me} = x_{\frac{n+1}{2}}, \quad n = 2k + 1$$

$$x_{me} = \frac{1}{2} \left(x_{\frac{n}{2}} + x_{\frac{n}{2}+1} \right), \quad n = 2k$$

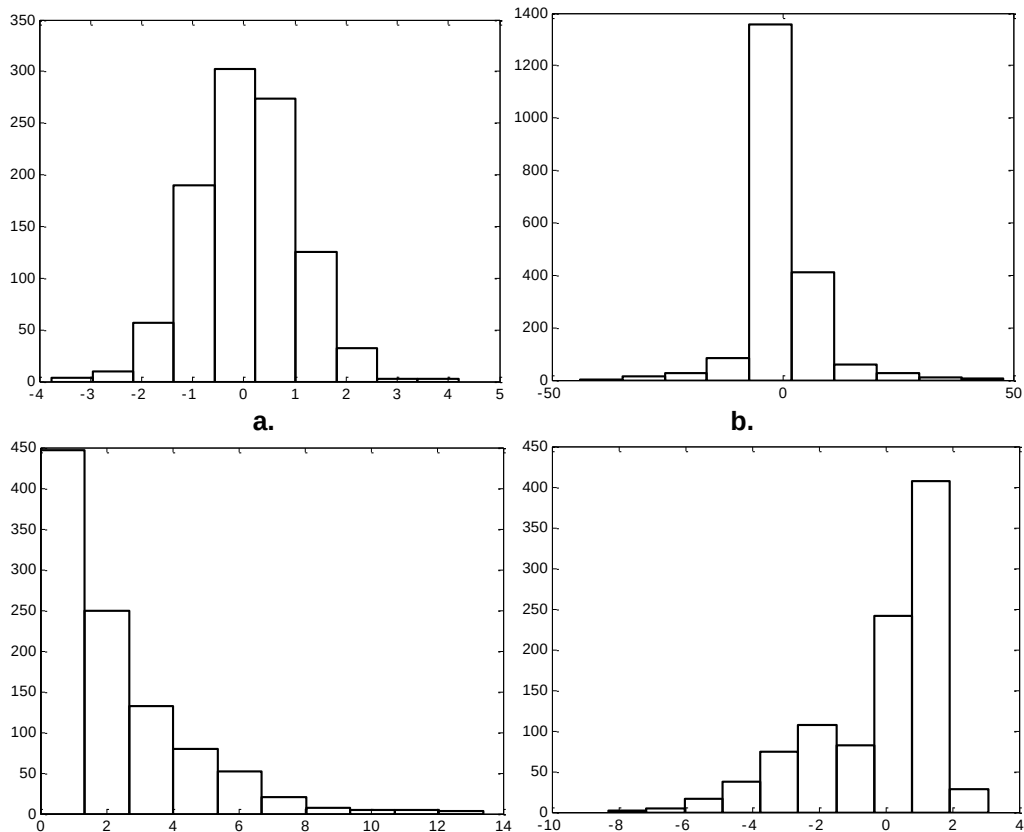
- modulul - valoarea ce apare cu cea mai mare frecvență din toate valorile șirului de date. Nu este neapărat o valoare unică.

Sunt necesare mai multe posibilități de estimare a localizării, deoarece în anumite situații se preferă una din ele datorită repartiției datelor.

Pe baza repartiției normale (având $\mu = 0$, $\sigma = 1$) s-a generat un șir de date a cărui histogramă este prezentată în fig. 4.15.a. Media este 0.0434, mediana

0.0666, iar modulul este -0.0521. Repartiția normală este simetrică, cu cozi moderate și un singur punct de maxim al funcției densitate de probabilitate. Pentru repartiții normale: media, mediana și modulul sunt echivalente. Dacă histograma sau graficul de probabilitate indică faptul că datele pot fi approximate de o repartiție normală, este rezonabil să se utilizeze media pentru estimarea localizării.

Pe baza repartiția Cauchy a fost generat șirul de date reprezentat în fig.4.15.b; are media 0.2965, mediana 0.0021 și modulul -43.9168. Repartiția Cauchy este simetrică, dar are extremități cu anvergură mare și un singur vârf în centru. Această repartiție are proprietatea interesantă că prin colectarea unui număr mai mare de date, nu se generează o estimare mai precisă a mediei (repartiția Cauchy nu are medie, [9]. Acest lucru înseamnă că media nu are semnificație ca măsură a localizării. Mediana este o măsură mai potrivită. Valorile mari din extremități distorsionează media, ceea ce nu se întâmplă cu mediana, ce depinde de indicele elementelor din șir.



c.

d.

Fig. 4.15. Diferite tipuri de repartiții

În general, pentru seturi de date cu cozi extinse, mediana este o estimare recomandată pentru localizare.

Pentru șirul de date provenit dintr-o repartiție exponențială (fig. 4.15.c), media este 2.2191, mediana 1.5204 și modulul 0.0037. Repartiția exponențială este asimetrică. Pentru repartiția asimetrică media este diferită de mediană. Media este deplasată spre extremitatea cu anvergură mai mare. Asimetrie dreaptă înseamnă medie mai mare decât mediana și invers. În fig. 4.15.d este reprezentată o repartiție cu asimetrie stânga. În acest caz media este -0.1574, mediana 0.6234, iar modulul -8.3048. Se poate observa că se respectă condiția $\text{media} < \text{mediana}$.

La repartiția asimetrică nu este evident care dintre indicatori caracterizează cel mai bine localizarea repartiției. Se recomandă să se indice toți trei.

Există multe alternative pentru medie și mediană pentru măsurarea localizării. Aceste alternative s-au dezvoltat pentru date ce nu au o repartiție normală. În cazul repartiției normale, media este estimatorul optim. Indicatorii alternativi se adresează repartițiilor cu coadă extinsă.

Indicatori alternativi:

- media intervalului intercvantilic – calculează media între cvantila inferioară și superioară (25% și 75%);
- media scurtată – se calculează pentru datele cuprinse între procentilele 5% - 95%;
- media scurtată modificată – valorile până la cvantila 5% se fac egale cu cea de 5% iar cele mai mari de 95% se modifică la valoarea de 95%;
- centrul intervalului de variație este egal cu media aritmetică dintre minim și maxim, $(\min + \max)/2$.

În caracterizarea localizării unei serii de date se preferă estimarea cu ajutorul unui interval de încredere, deoarece estimarea punctuală este influențată de eșantion. Intervalul de încredere oferă o indicație a cantității de incertitudine din estimarea mediei. Cu cât intervalul este mai îngust, cu atât estimarea este mai precisă. Intervalul de încredere pentru un eșantion din distribuția normală de medie și dispersie necunoscută este:

$$\bar{x} \pm t_{1-\frac{\alpha}{2}}(n-1) \frac{s}{\sqrt{n}}, \quad (4.8)$$

unde $t_{1-\frac{\alpha}{2}}(n-1)$ este 1- $\alpha/2$ cvantila repartiției t cu n-1 grade de libertate și se calculează în Matlab apelând funcția `tinv`. Probabilitatea ca media teoretică să aparțină acestui interval este $(1 - \alpha)$.

Se remarcă faptul că lățimea intervalului de încredere este controlată de doi factori:

- volumul eșantionului, n. Pe măsură ce n crește, intervalul scade în raport cu $\sqrt{n}$ și rezultă că o modalitate de a obține o estimare mai precisă este creșterea volumului eșantionului;
- deviația standard, S. Cu cât S este mai mare, cu atât crește și intervalul de încredere.

Pentru a verifica dacă media unei populații normal distribuite are o valoare prescrisă μ_0 , se testează ipoteza H_0 contra alternativei H_1 :

$$H_0 : \mu = \mu_0$$

$$H_1 : \mu \neq \mu_0$$

pe baza datelor dintr-un eșantion din populație $y_1, y_2, \dots, y_n$ folosind funcția test:

$$T = \frac{\bar{y} - \mu_0}{S / \sqrt{N}}$$

Nivel de semnificație: $\alpha = 0.05$.

Regiune critică: se respinge H_0 dacă $|T| > t_{1-\frac{\alpha}{2}}(n-1)$.

Pentru alți estimatori: mediana, media intervalului intercvantilic determinarea intervalului de încredere este dificilă din punct de vedere matematic. O alternativă este construirea graficului de incertitudini.

4.3.1. ESTIMĂRI ALE VARIABILITĂȚII

O sarcină fundamentală în multe analize statistice este caracterizarea variabilității (împrăștierea) unui set de date. La informațiile legate de variabilitate a datelor, există două componente de bază:

- 1 - cât de împrăștiate sunt valorile datelor în apropierea centrului;
- 2 - cât de împrăștiate sunt datele la extremități.

Diferiți indicatorilor statistici vor oferi informații cu ponderi diferite din acest punct de vedere. Alegerea unui estimator al localizării este determinată de componenta ce prezintă cel mai mare interes. Histograma este un instrument

grafic ce evidențiază ambele componente.

Pentru date unidimensionale, există mai mulți indicatori statistici pentru variabilitate:

1. Dispersia de selecție

$$S^2 = \sum_{i=1}^N (Y_i - \bar{Y})^2 / (N - 1)$$

Dispersia este, de fapt, o aproximare a mediei pătratelor distanțelor de la puncte la valoarea medie. Prin ridicarea la pătrat se atribuie ponderi mai mari valorilor ce sunt situate mai departe de medie. Deci, ea poate fi afectată mult de comportamentul extremităților.

2. Deviația standard

$$S = \sqrt{\frac{1}{N-1} \sum (Y_i - \bar{Y})^2}$$

Deviația standard se exprimă în aceeași unitate de măsură ca și datele asociate eșantionului. Este afectată de comportamentul extremităților.

3. Intervalul de variație (IV) sau amplitudinea setului de date egală cu diferența dintre valoarea maximă și valoarea minimă a șirului. Este un estimator ce se bazează doar pe valorile extreme. Nu oferă nici o informație în privința comportamentului valorilor centrale

4. Media deviației absolute (AAD – average absolute deviation):

$$AAD = \sum_{i=1}^N \frac{|Y_i - \bar{Y}|}{N}. \quad (4.9)$$

Acest indicator este mai puțin afectat de observațiile extreme decât dispersia și deviația standard. Este stabil numeric și se folosește în investigațiile bazate pe calculator.

5. Mediana deviației absolute (MAD – median absolute deviation)

$$MAD = \text{median}(|Y_i - \bar{Y}|), \quad (4.10)$$

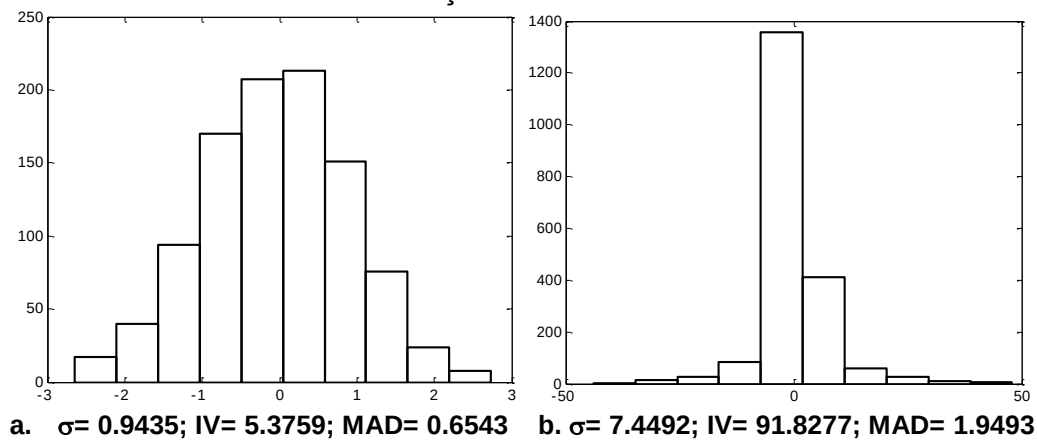
este un indicator ce este și mai puțin afectat de valorile extreme, deoarece ele afectează mult mai puțin valorile medianei.

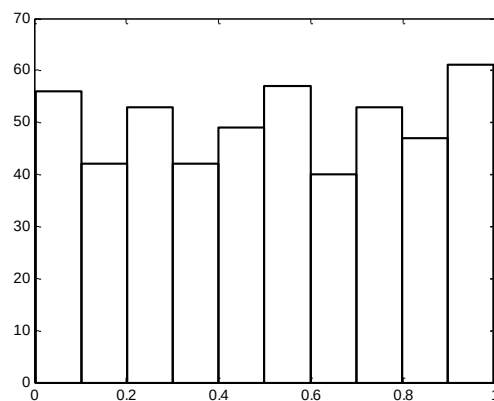
6. Lungimea intervalul intercvantilic: $q_{0.75} - q_{0.25}$, este diferența dintre cvantila superioară și inferioară a setului de date și determină variabilitatea punctelor din zona centrală.

În concluzie, dispersia de selecție, deviația standard, media deviației absolute și mediana deviației absolute măsoară ambele aspecte ale variabilității. Ele diferă prin faptul că AAD și MAD nu oferă ponderi valorilor extreme. Pe de altă parte, intervalul de variabilitate se bazează doar pe două puncte.

Necesitatea folosirii acestor indicatori se ilustrează în exemplele următoare. În fig. 4.16 sunt prezentate histogramele a trei seturi de date asociate cu indicatori ai variabilității: deviația standard, intervalul de variație sau amplitudinea și mediana deviației absolute.

Primul set de date (fig. 4.16.a) provine dintr-o repartiție normală. Repartiția normală este simetrică, cu extremitățile de anvergură moderată, un singur maxim al densității de probabilitate. Există o simetrie față de centru, caz în care mediana deviației absolute este mai mică decât deviația standard. Dacă histograma sau graficul probabilității normale indică faptul că datele sunt bine approximate de o repartiție normală este rezonabil să se utilizeze deviația standard ca indicator al variabilității.





c. $\sigma=0.2943$; $IV=0.9946$; $MAD= 0.2552$

Fig. 4.16. Histograme asociate cu indicatori ai variabilității

În fig. 4.16.b se prezintă histograma unui set de date provenind dintr-o repartiție Cauchy: $\sigma= 7.4492$; $IV= 91.8277$; $MAD= 1.9493$. Comparând repartiția normală cu cea Cauchy se observă că repartiția are un vârf mai pronunțat, scade mai rapid în apropierea centrului și are extremități de anvergură mai mare. Datorită acestor extremități, deviația standard are tendința să crească comparativ cu repartiția normală. Extremitatea de anvergură mai mare se reflectă în valoarea IV. Repartiția are o proprietate interesantă – și anume prin colectarea unui număr mai mare de date nu se obține o medie sau o deviație standard mai precisă. Deci pentru eșantion și populație cei doi indicatori sunt egali. Acest lucru înseamnă că pentru seturi de date provenind din repartiția Cauchy deviația standard nu este o măsură a împrăstierii. Din histogramă se observă că toate datele sunt situate între -5 și 5 . Cu toate acestea, foarte puține valori extreme determină creșterea mare a lui σ și IV. MAD este doar puțin diferită de cea a repartiției normale. În acest caz, MAD este măsura potrivită pentru variabilitate.

Deși repartiția Cauchy este un caz extrem, ea ilustrează importanța extremităților în măsurarea variabilității. Valorile extreme pot distorsiona deviația standard. Ele, însă, nu afectează MAD. Pentru repartiții cu extremități pronunțate, MAD sau intervalul intercvantilic constituie indicatori mai buni pentru variabilitate.

Cel de-al treilea set de date (fig. 4.16.c) provine dintr-o repartiție uniformă: $\sigma=0.2943$; $IV=0.9946$; $MAD= 0.2552$. Repartiția are extremități trunchiate. În acest caz deviația standard și MAD au valori mai apropiate decât în celelalte trei exemple, ce au extremități de diferite anverguri.

Dacă histograma și graficul probabilității normale indică faptul că datele se pot aproxima cu o repartiție normală, are sens utilizarea deviației standard ca indicator al variabilității. În cazul când, datele au alt tip de repartiție, cu extremități de anvergură mare, utilizarea altor indicatori: MAD sau IV este mai potrivită. IV se utilizează în unele aplicații, cum ar fi controlul calității, pentru simplitatea calculelor. În plus, prin compararea IV cu deviația standard avem o indicație legată de împrăștierea datelor la extremități. Acest indicator, IV, trebuie utilizat cu precauție fiind bazat pe doar două valori.

BIBLIOGRAFIE

1. Davidescu, A., Analiza și procesarea datelor cu aplicații în Matlab, Ed. Politehnica Timișoara, 2003.
2. Fleming, M., Nellis, J., Principles of Applied Statistics, Routledge Co., London, 1994.
3. Hanselman, D., Littlefield, B., The Student Edition of Matlab®. User's Guide, Prentice Hall, 1995.
4. Hanselman, D., Littlefield, B., Mastering Matlab® 5. A comprehensive tutorial and reference, Prentice Hall, 1998.
5. J de Leeuw, WEB Statistics Book, <http://www.stat.ucla.edu/textbook>
6. Martinez, W., Martinez, A., Computational Statistics Handbook with Matlab, Chapman&Hall/CRC, Washington DC, 2002.
7. Montgomery, D., Design and Analysis of Experiments, John Wiley&Sons, Singapore, 1991.
8. Montgomery, D., Runger, G., Applied Statistics and Probability for Engineers, John Wiley&Sons, New York, 2006.
9. Petrișor, E., Probabilități și statistică. Aplicații în economie și inginerie, Ed. Politehnica, Timișoara, 2005.
10. Engineering Statistics Handbook :
<http://www.itl.nist.gov/div898/handbook/eda/eda.htm>