

MIHAELA MATCOVSCHI

OCTAVIAN PĂSTRĂVANU

APLICAȚII ALE REȚELELOR NEURALE ÎN AUTOMATICĂ

Editura POLITEHNIUM

2008

Editura POLITEHNIUM

a Universității Tehnice „Gh. Asachi” din Iași
Bd. Dimitrie Mangeron nr. 67
RO-700050, Iași, România
Tel/Fax: +40 – 0232 – 231 343

Editura POLITEHNIUM (fostă „Gh. Asachi”) este recunoscută de
Consiliul Național al Cercetării Științifice
din Învățământul Superior (CNCSIS)

Director editură:

Prof. univ. dr. ing. Mihail VOICU
Membru corespondent al Academiei Române

Referenți științifici:

Prof. univ. dr. ing. Corneliu LAZĂR
Prof. univ. dr. ing. Doru PĂNESCU

Redactor:

Ing. Elena MATCU-ZBRANCA

Răspunderea pentru tot ceea ce conține această carte aparține
în întregime autorului (autorilor) ei.

Descrierea CIP a Bibliotecii Naționale a României
MATCOVSCHI, MIHAELA HANAKO**Aplicații ale rețelelor neurale în automatică /**

Mihaela Matcovschi, Octavian Păstrăvanu. – Iași:
Politehniun, 2008.

Bibliogr.

ISBN: 978-973-621-221-5

I. Păstrăvanu, Octavian

681.5

© **Mihaela MATCOVSCHI, Octavian PĂSTRĂVANU**
Universitatea Tehnică “Gh. Asachi” Iași
Facultatea de Automatică și Calculatoare
Bd. D. Mangeron 53A, 700050 Iași
email: {mhanako, opastrav}@delta.ac.tuiasi.ro

Prefață

În pofida volumului mare de informații (cărți, articole, documentații de pachete software etc) existent asupra rețelelor neurale artificiale, realizarea unui material didactic eficient la nivelul primului ciclu de instruire al învățământului universitar tehnic constituie o sarcină dificilă. Dificultatea provine din condițiile restrictive pe care trebuie să le satisfacă un atare material pentru a-și îndeplini cu adevărat rolul său formativ. Altfel spus, o carte despre aplicații ale rețelelor neurale în automatică care să capteze interesul studenților trebuie (i) să pună la dispoziție o descriere riguroasă a structurilor complexe pe care le reprezintă rețelele neurale cu utilizare reală în practică, dar fără a apela la construcțiile matematice laborioase specifice acestor probleme; (ii) să selecteze arhitecturi de rețele și tipuri de aplicații pentru care se pot realiza experimente ilustrative pe calculator, dar fără a implica elaborarea de programe sofisticate care să deplaseze centrul de greutate al experimentelor pe partea de programare, distorsionând sensul intuitiv al experimentelor.

Autorii volumului de față au conștientizat delicata lor misiune pe tot parcursul elaborării textului. În organizarea cărții și a fiecărui capitol în parte o importanță deosebită a avut-o experiența de predare a cunoștințelor respective studenților specializării automatică, de la Facultatea de automatică și calculatoare a Universității Tehnice “Gh. Asachi” din Iași. Într-un sens mai general dar expresiv, am măsurat și croit după talia clientului cu speranța că vestimentația rezultată va fi purtată cu plăcere. Ne-am străduit să răspundem apetitului pentru tehnologii “miraculoase” care este specific vârstei cititorilor noștri, dar, în aceeași măsură, am căutat să inoculăm ideea fundamentului teoretic solid ca singură premisă pentru dezvoltarea unor aplicații funcționale. În ansamblu, am încercat să realizăm un material de nivel elementar spre mediu, care să ofere o orientare profesionistă în acumularea cunoștințelor, extrăgând cititorul din sfera lecturilor facile de popularizare a științei și antrenându-l ca viitor specialist.

Volumul include câteva teme de studiu care au fost considerate drept esențiale pentru introducerea cititorului în universul amplu al rețelelor neurale artificiale. Parcurgerea acestor teme asigură înțelegerea proprietăților de clasificare și aproximare pe care le posedă anumite tipuri de rețele. Pe această bază sunt prezentate ulterior aplicațiile cele mai importante din automatică și anume identificare și control neliniar cu modele de tip rețea neurală dinamică. Deși concepută pentru profilul de instruire ingineria sistemelor, cartea poate fi utilizată și de cititori cu formație profesională înrudită. Primele capitole dedicate proprietăților rețelelor neurale artificiale sunt accesibile oricărui cititor care posedă cunoștințele de bază ale ingineriei electrice.

Pentru a răspunde intențiilor noastre educaționale, volumul include două tipuri de probleme care permit aprofundarea și fixarea cunoștințelor. Acestea sunt organizate în două capitole distincte și anume: probleme care se rezolvă pe cale analitică și, respectiv, probleme care se rezolvă prin procedee numerice. Problemele cu soluții analitice vizează rețele cu arhitectură simplă, care pot fi abordate printr-un volum redus de calcule. Problemele rezolvabile cu ajutorul calculatorului se referă la arhitecturi mai mari, care necesită un volum mai ridicat de calcule, dar sunt în măsură de a ilustra potențialul aplicativ real al rețelelor neurale. Problemele din prima categorie sunt acompaniate de răspunsuri și indicații de rezolvare, iar pentru problemele din categoria a doua sunt furnizate informații despre întrebuințarea rutinelor din pachetul Neural Network Toolbox for MATLAB elaborat de compania The MathWorks Inc. Programele MATLAB realizate sunt accesibile pe web, la adresa http://81.180.214.164/web_ARNA/index.html.

Cititorii interesați în dobândirea de cunoștințe suplimentare (detalii, demonstrații) ce depășesc nivelul propus pentru această carte, sunt invitați să consulte materialul bibliografic recomandat în lista de referințe, care include monografiile cotate (la momentul elaborării cărții) cu un impact științific pe plan internațional foarte ridicat. De asemenea cititorilor le sunt puse la dispoziție două anexe care furnizează informații sumare despre două arii conexe rețelelor neurale, și anume o trecere în revistă a metodelor de optimizare utilizabile pentru antrenarea rețelelor neurale și a tehnicilor de clasificare.

Autorii

Cuprins

Capitolul 1. <i>Introducere</i>	1
Capitolul 2. <i>Arhitecturi de rețele neurale feed-forward – Principii generale de organizare</i>	5
2.1. Modelul neuronului artificial	6
2.1.1. Neuronul artificial cu o singură intrare	6
2.1.2. Neuronul artificial cu mai multe intrări	9
2.1.3. Alte tipuri de neuroni artificiali	10
2.2. Rețele de neuroni artificiali	11
2.2.1. Rețea cu un singur strat	11
2.2.2. Rețea cu două straturi	13
2.2.3. Cazul general - rețea cu mai multe straturi	15
2.3. Instrumente software pentru simularea rețelelor neurale cu structură feedforward oferite de <i>Neural Network Toolbox</i>	15
Capitolul 3. <i>Aplicații ale rețelelor feed-forward cu funcții de activare treaptă în probleme de clasificare</i>	17
3.1. Studiarea transferului intrare-ieșire pentru un neuron	18
3.2. Utilizarea rețelelor cu un singur strat	20
3.2.1. Proprietatea de separare folosită în clasificare	20
3.2.2. Obiectivul antrenării rețelei	20
3.2.3. Algoritmul de antrenare al rețelei	22
3.3. Utilizarea rețelelor cu mai multe straturi	25
3.4. Facilități software oferite de <i>Neural Network Toolbox</i>	26
Capitolul 4. <i>Aplicații ale rețelelor feedforward cu funcții de activare liniare în probleme de aproximare</i>	27
4.1. Transferul liniar intrare-ieșire și proprietatea de aproximare	27

4.2. Utilizarea rețelelor cu un singur strat	31
4.2.1. Obiectivul antrenării rețelei	31
4.2.2. Algoritmul de antrenare Widrow-Hoff	31
4.3. Facilități software oferite de <i>Neural Network Toolbox</i>	34
 Capitolul 5. <i>Aplicații ale rețelelor feedforward multistrat</i>	
în probleme de aproximare și clasificare	36
5.1. Proprietatea de aproximare a rețelelor feed-forward cu două straturi	36
5.2. Utilizarea rețelelor feedforward cu două straturi	39
5.2.1. Obiectivul antrenării rețelei	39
5.2.2. Algoritmul de antrenare prin propagare inversă	40
5.2.3. Algoritmul de antrenare prin propagare inversă privit ca problemă standard de optimizare	45
5.2.4. Strategii de creștere a performanțelor antrenării prin propagare inversă	49
A. Propagarea inversă cu moment	50
B. Propagarea inversă cu rată de învățare adaptivă	50
C. Propagarea inversă cu moment și rată de învățare adaptivă	51
5.2.5. Alți algoritmi de antrenare pentru rețele multistrat	51
5.3. Utilizarea rețelelor feedforward cu trei straturi	51
5.4. Utilizarea rețelelor feedforward în probleme de clasificare	52
5.5. Facilități software oferite de <i>Neural Network Toolbox</i>	53
 Capitolul 6. <i>Aplicații ale rețelelor feedforward multistrat</i>	
în identificarea și controlul sistemelor dinamice	55
6.1. Identificarea sistemelor dinamice utilizând rețele neurale	55
6.1.1. Structura modelelor neurale	56
6.1.2. Estimarea modelelor neurale	59
6.2. Controlul predictiv al sistemelor dinamice neliniare utilizând modele neurale	60
6.3. Blocuri Simulink pentru aplicații ale rețelelor neurale în identificarea și conducerea sistemelor	63
6.4. Studii de caz	69
6.4.1. Exemple ilustrative pentru construcția modelelor	69
A. Construcția unui model liniar	69
B. Construcția unui model neliniar	72
6.4.2. Exemplu de utilizare a controlului predictiv bazat pe model neural	74
Capitolul 7. <i>Probleme cu soluții analitice</i>	76
Capitolul 8. <i>Probleme cu soluții numerice</i>	97

<i>Anexa A. Metode de optimizare implementate ca algoritmi de antrenare supervizată a rețelelor neurale</i>	133
A.1. Introducere	133
A.1.1. Condiții de minim local	134
A.1.2. Rezolvarea ecuației $A\mathbf{x} = \mathbf{b}$	138
A.1.3. Estimatorul celor mai mici pătrate	140
A.1.4. Forma generală a metodelor iterative de optimizare	141
A.1.5. Metode de căutare liniară	143
A.2. Metode de optimizare bazate pe gradient	146
A.2.1. Metoda steepest-descent (de descreștere maximă)	146
A.2.2. Metoda lui Newton	149
A.2.3. Metoda Gauss-Newton	150
A.2.4. Metode de tip Newton modificate	152
Metoda lui Newton cu pas adaptiv	152
Metoda Levenberg-Marquardt	153
A.2.5. Metode quasi-Newton	154
A.3. Metode bazate pe direcții conjugate	155
A.3.1. Metoda direcțiilor conjugate	155
A.3.2. Metoda gradientului conjugat	158
<i>Anexa B. Problematica clasificării formelor</i>	163
B.1. Tipuri de clasificatori	163
B.2. Clasificarea pe baza funcțiilor de decizie	165
B.2.1. Separarea în două clase	165
B.2.2. Separarea în m clase	166
B.3. Clasificarea pe baza distanței	168
B.3.1. Clase reprezentate printr-un singur prototip	168
B.3.2. Clase reprezentate prin mai multe prototipuri	170
<i>Bibliografie</i>	173

Capitolul 1

INTRODUCERE

Studierea unor mecanisme biologice naturale, în urma abstractizării și formalizării matematice, a condus la diferite metode de procesare a informației. Aceste tehnici de procesare, combinate cu concepte din alte domenii au fost dezvoltate pentru a permite simularea pe calculatoare analogice sau numerice. Astfel, prin simulare au fost emulate o serie fenomene naturale, dintre care marea majoritate se desfășoară la nivelul celulelor nervoase ale creierului uman. În această conjunctură a luat ființă un nou domeniu științific de sine stătător denumit *Inteligența Artificială*, care înglobează mai multe direcții de cercetare. Una dintre aceste direcții o constituie *Rețelele neurale artificiale*.

În tentativa de “deconspirare” a comportamentelor ființelor umane sau ale altor organisme vii, s-a constatat o *abordare euristică*, unde nu funcționează algoritmi expliciți (situație care în literatura științifică anglo-saxonă este caracterizată prin expresia consacrată “rule-of-thumb”). Cel mai adesea este greu de formulat reguli precise conform cărora oamenii cunosc anumite fapte, prelucrează date, trag concluzii sau iau decizii. Experimentele au arătat că majoritatea persoanelor care iau anumite decizii nu pot explica, în sens algoritmic, mecanismul prin care au ajuns la respectivele decizii. Prin astfel de investigații, Inteligența Artificială a reușit să pună în evidență deosebirile fundamentale dintre două tipuri de procesare a informației și anume *de tip secvențial* și *de tip contextual*. Procesarea de tip secvențial corespunde algoritmilor implementabili și executabili pe mașini de calcul cu arhitectura clasică Von Neumann. Procesarea de tip contextual are în vedere distribuirea informației către mai multe elemente de calcul care conlucrează între ele, fiecare operând însă în baza unor principii locale. Astfel s-a constatat că în cazul de corupere / alterare a unor informații, procesarea de tip secvențial eșuează, în timp ce procesarea de tip contextual poate furniza rezultate utile.

Până în prezent nu a fost formulată încă o definiție a rețelelor neurale artificiale care să se bucure de o acceptare unanimă în literatura de specialitate. Aceasta se explică prin faptul că termenii tehnici la care s-ar putea face apel ar fi ori prea generali și vagi, ori, dimpotrivă, prea specializați pentru cadrul impus de o definiție. Totuși, pentru a asigura rigoarea prezentării, în capitolele ce urmează vom defini *tipuri particulare de rețele neurale artificiale*, întrucât cititorul va acumula treptat cunoștințele care să permită formalizarea tratării. De asemenea, pentru a simplifica limbajul, vom utiliza exprimarea “rețele neurale”, suprimând atributul “artificiale”, deoarece în textul nostru cu orientare tehnică nu există pericolul apariției unor confuzii prin referirea la rețele neurale biologice.

Deși constituie un participant recent la dezvoltarea științei, Rețelele neurale acoperă un teritoriu larg în procesarea informației, cu multiple conotații de interdisciplinaritate. În termeni largi, o rețea neurală *învață să reprezinte* o mulțime de *trăsături* ale mediului informațional în care funcționează. De asemenea poate învăța o serie de răspunsuri care sunt livrate în corelației cu anumiți stimuli relevanți pentru mediul în care funcționează. De asemenea, multe rețele prezintă proprietăți de *generalizare* care permit elaborarea de răspunsuri adecvate pentru stimuli de intrare nemaiîntâlniți anterior.

Având în vedere importanța conceptelor amintite în paragraful anterior, drept bază de plecare pentru studierea rețelelor neurale, considerăm utilă aprofundarea lor, după cum urmează:

- *învățarea*: Este procesul prin care, în potențialul comportamental al unei rețele neurale au loc permanent schimbări, ca rezultat al experienței. Prin acest proces, parametrii unei rețele neurale sunt supuși unei adaptări continue (în sensul modificării valorilor parametrilor) ca urmare a stimulărilor pe care le primește din partea mediului. Tipul de învățare este determinat de maniera în care au loc schimbările parametrilor respectivi. Capacitatea de învățare poate fi privită ca opusul unei funcționări pre-programate, care dispune, de la bun început, de toate informațiile necesare pentru întreaga durată de funcționare.
- *reprezentarea cunoștințelor*: O rețea neurală acumulează cunoștințe despre lumea înconjurătoare. Modul de reprezentare a cunoștințelor diferă de la un tip de arhitectură de rețea neurală la altul. Totuși există anumite principii general valabile, în sensul că informația se memorează prin intermediul unor ponderi care se alocă conexiunilor dintre elementele de bază ale rețelei.
- *mediul informațional*: O rețea neurală are un acces limitat la informații, de o factură mult mai restrictivă decât procedează o ființă umană. O rețea neurală operează cu date de intrare și date de ieșire pe care le conectează sub forma unor aplicații (în general neliniare). Rețelele neurale sunt aplicate în probleme specifice pe domenii, fiecare domeniu având anumite tipuri de date relevante.
- *generalizarea*: Un model bun se pretează generalizării, în sensul că el poate acoperi submulțimi de date din domeniul problemei considerate, chiar dacă acele date nu au fost întâlnite pe parcursul învățării. Proprietatea de generalizare este cu atât mai bună cu cât rețeaua necesită mai puține eșantioane de învățare pentru a furniza răspunsuri corecte la

intrări necunoscute. Fără această proprietate de generalizare, rețelele neurale s-ar limita la o funcționalitate similară tabelelor de căutare.

În ciuda similitudinii dintre proprietățile enumerate mai sus pentru rețelele neurale și acțiunile (cu denumiri identice) pe care le poate desfășura creierul uman, atragem atenția asupra faptului că suportul care asigură funcționarea rețelelor neurale este de factură matematică, cu diverse implementări, în timp ce suportul biologic este încă în mare măsură necunoscut. În înțelegerea corectă a rolului și limitelor de utilizare a rețelelor neurale, trebuie evitată orice confuzie, de tip falsă popularizare a științei, care ar crea impresia că rețelele neurale pot fi (oricând și oriunde) un substitut al creierului uman. Purtând denumirea de rețele neurale pentru a sugera că funcționarea lor este de inspirație biologică, acestea sunt instrumente de procesare a informației, cu avantaje și dezavantaje ca orice altă ustensilă inventată de om în scopul cunoașterii. Deși poate părea desuet, considerăm că cititorul trebuie să fie pe deplin avertizat că de la stadiul actual de dezvoltare a Inteligenței Artificiale, în general și a Rețelelor neurale, în particular, până la creația tehnică (robotul) cu comportament integral humanoid (prezent în diverse ficțiuni), știința și tehnologia mai au de parcurs pași însemnați.

La momentul de față, potențialele beneficii obținute prin utilizarea rețelelor neurale se referă la: abilitatea de a extrage date / informații fără specificarea unui anumit model; abilitatea de generalizare a unor date / informații necunoscute anterior; dezvoltarea suportului tehnico-științific pentru decizii autonome (neasistate); dezvoltarea sistemelor de interfațare “prietenoase”; dotarea unor produse cu un anumit grad de “inteligență”. Din punct de vedere practic, aplicativ, aceste beneficii trebuie evaluate și prin prisma complexității problemei care se dorește a fi rezolvată făcând apel la teoria și tehnologia rețelelor neurale. Altfel spus, se vor căuta răspunsuri la întrebări de tipul: Există o soluție alternativă pentru utilizarea rețelelor neurale? Prin ce este mai profitabilă soluția bazată pe rețele neurale decât soluția alternativă? Care este gradul de dificultate în implementarea soluției bazate pe rețele neurale? Există personalul cu calificarea necesară pentru exploatarea unei aplicații bazată pe rețele neurale? Răspunsuri corecte la aceste întrebări pot evita experiențe nefructuoase care pot conduce (fără reale justificări) la instalarea unei atitudini de adversitate față de utilizarea rețelelor neurale.

Dezvoltarea Inteligenței Artificiale în general și a rețelelor neurale în particular a stimulat și progresul unor domenii științifice învecinate, cum ar fi Procesarea semnalelor, Teoria și ingineria controlului, Vederea computerizată, Analiza și sinteza vorbirii, Transmisiunea informației etc. Prin preluarea și adaptarea unor rezultate și / sau instrumente, la începutul anilor 90, domeniile amintite au promovat noi direcții de cercetare care ulterior s-au dovedit foarte productive.

Ultimul deceniu a reușit să pună în valoare studiile și cercetările experimentale asupra rețelelor neurale prin realizarea unor aplicații de anvergură industrială. În paragraful curent enumerăm câteva dintre acestea, în baza informațiilor furnizate de (Howard Demuth, Mark Beale, Martin Hagan, 2008, Neural Network Toolbox User’s Guide, The MathWorks Inc.):

- Tehnologii aerospațiale: Pilot automat de înaltă performanță pentru aeronave; simularea și detecția erorilor pentru componentele aeronavelor.
- Tehnologii pentru autovehicule: analiza activităților de garanție
- Tehnologii pentru electronică: predicția secvențelor de codare, realizarea layout-ului pentru chip-uri cu circuite integrate, sinteză vocală, recunoaștere de imagini și vorbire, conversie text-vorbire.
- Tehnologii pentru telecomunicații: compresii de date și imagini, traducere automată, servicii de informații automatizate, sisteme de procesare a plăților
- Tehnologii de fabricație: analiză și proiectare de produse chimice; sisteme de control al calității (bere, sudură, hârtie, chipuri, pulberi și particule); planificare și management de proiecte
- Tehnici de conducere și diagnoză pentru: procese de fabricație, roboți, manipolatoare.
- Tehnici și servicii medicale: analiză automatizată de celule, analiză automatizată a diagramelor EEG și EKG; proiectare de proteze; asistarea unor teste ATI
- Operații și servicii financiar-bancare: evaluări de proprietăți; consiliere împrumuturi; analiza liniilor de credit, investigații ipotecare; urmărirea activității de pe cărțile de credit; analiză și predicție financiară; citirea automată a unor documente; evaluarea solicitărilor de credite; predicția piețelor
- Operații și servicii de asigurări: evaluarea solicitărilor pentru polițe; optimizarea produselor asigurate
- Servicii de transport: planificarea și routarea vehiculelor
- Televiziune și cinematografie: animații și efecte speciale:
- Tehnologii militare: ghidarea diferitelor tipuri de armament; urmărirea țintelor; clasificarea obiectelor, recunoaștere facială, tipuri noi de senzori; noi strategii de procesare a informațiilor provenite de la sonare și radare, noi strategii de identificare a semnalelor și imaginilor.

Capitolul 2

ARHITECTURI DE REȚELE NEURALE FEEDFORWARD – PRINCIPII GENERALE DE ORGANIZARE

Rețelele neurale feedforward sunt sisteme statice, adică sisteme ce realizează un transfer instantaneu intrare-ieșire. Acest transfer poate fi descris prin compunerea unor funcții dependente de topologia (arhitectura) rețelei. Termenul de „feedforward” (însemnând „cu alimentare înainte” – semnalele propagându-se numai în sensul de la intrare către ieșire) marchează diferența față de *rețelele neurale recurente*, sau cu feedback, care sunt sisteme dinamice, posedând căi de reacție ce permit propagarea semnalelor și de la ieșire către intrare. Deosebirea dintre cele două clase de rețele neurale, introdusă prin terminologie, se regăsește și în descrierea matematică a transferului intrare-ieșire. Rețelele neurale recurente sunt sisteme dinamice și, drept urmare descrierea lor se realizează prin ecuații diferențiale (în cazul rețelelor cu timp continuu) și prin ecuații cu diferențe (pentru rețelele cu timp discret). Elementul de bază în structura unei rețele neurale (feedforward sau recurente) îl constituie *neuronul artificial*. Pentru construirea unei rețele neurale, neuronii artificiali se conectează în unul sau mai multe straturi, definind, astfel, *arhitectura rețelei*.

Rețelele neurale artificiale s-au dezvoltat ca obiect de studiu în domeniul informaticii, matematicii aplicate, calculatoarelor, etc. Ele prezintă proprietăți de factură matematică (riguros dovedite) care seamănă cu o serie de activități curente desfășurate de mintea umană (învățare, recunoaștere, aproximare, adaptare). Aceste activități desfășurate de ființa umană au inspirat ca denumire proprietățile rețelelor neurale artificiale, adică în sens științific aproximarea, învățarea, etc. a unei rețele neurale artificiale înseamnă un concept foarte precis definit matematic.

Punctul de vedere formulat de acest curs: rețelele neurale artificiale ca parte componentă a inteligenței artificiale (inteligență computațională) oferă instrumente cu fundament matematic care aparent copie activitatea creierului uman. În prezent, inteligența artificială nu a deconspirat mecanismele raționamentului uman; ea pune la dispoziție doar unele instrumente de factură matematică care permit abordări euristice.

Rețelele neurale artificiale au drept componentă de bază neuronul artificial. Datorită numărului mare de neuroni structura ce descrie rețeaua neurală operează cu un număr mare de semnale de intrare, interne și de ieșire. Din acest motiv, tratarea matematică este specifică sistemelor de dimensiuni mari și operării cu proprietăți generice. Pe tot parcursul cursului, proprietăți de factură matematică le vom analiza pe cazuri simple ce evită complexitatea rețelelor de lucru (rațiuni didactice); rețelele de lucru pentru care se urmăresc performanțe au întotdeauna dimensiuni mari.

Tipurile de rețele neurale artificiale folosite azi în practică au rezultat ca urmare a încercărilor mai vechi de a modela activitățile neuronilor biologici și a creierului. Modelarea neuronului biologic este o problemă foarte complexă și, la un anumit moment al dezvoltării de modele, s-a constatat că ceea ce se vroia a fi modele pentru activitatea nervoasă biologică, puteau fi folosite ca instrumente de cunoaștere matematică. Modelele la nivelul respectiv au devenit de fapt ceea ce azi denumim neuron artificial, rețele neurale artificiale.

Altfel spus, actualmente neuronul artificial nu mai e privit ca un model al neuronului biologic. În prezent cercetările pentru modelarea neuronului biologic și al sistemului nervos țin de biomedicină, bioinformatică.

2.1. Modelul neuronului artificial

Neuronul artificial constituie elementul fundamental din structura unei rețele neurale. În această secțiune, vom prezenta neuronul artificial într-o manieră suficient de generală, care să permită unele particularizări pe parcursul capitolelor ce urmează. Pentru a asigura înțelegerea cât mai deplină a noțiunilor, expunerea se realizează gradat, tratând, inițial, cazul neuronului cu o intrare și, apoi, cazul neuronului cu mai multe intrări.

2.1.1. Neuronul artificial cu o singură intrare

Pentru prezentarea arhitecturii vom utiliza descrieri de tip diagramă bloc (schemă bloc) cu semnificați simbolurilor de la disciplina „Introducere în Automatică”. În figura 2.1.1 se prezintă schema bloc a unui neuron cu o singură intrare (intrare scalară), notată $x \in \mathbb{R}$ și ieșirea scalară $y \in \mathbb{R}$. Constanta $w \in \mathbb{R}$ se numește *pondere* (eng. *weight*), iar constanta $b \in \mathbb{R}$ poartă denumirea de *deplasare* (eng. *bias*). Intrarea cu valoarea precizată “1” arată că, indiferent de semnalul de intrare x , valoarea cu care este alimentat blocul b este tot timpul 1. Referirea comună a constantelor $w, b \in \mathbb{R}$ se face prin termenul de “*parametrii neuronului*”.

În utilizarea neuronului, valorile parametrilor $w, b \in \mathbb{R}$ sunt *ajustabile*. Pe ajustarea adecvată a acestor valori se bazează capacitatea de învățare a neuronului artificial.

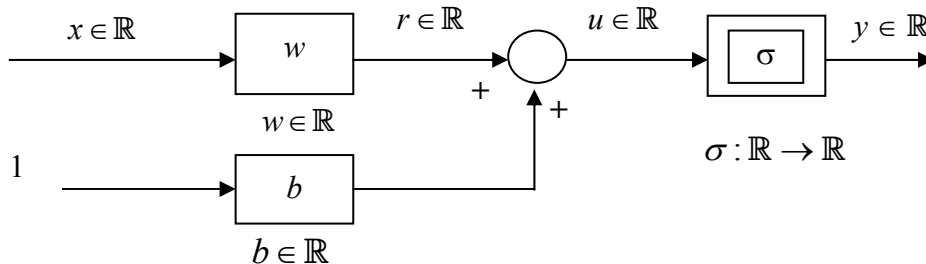


Figura 2.1.1. Schema bloc a unui neuron cu o singură intrare

Funcția $\sigma: \mathbb{R} \rightarrow \mathbb{R}$ se numește *funcția de activare a neuronului* (eng. *activation function*). Pentru funcția de activare σ se folosesc *funcții analitice standard*, liniare sau neliniare. Blocul σ definește *nodul* neuronului, iar ansamblul care generează semnalul u definește partea *liniară*.

Funcția de activare σ aparține unor clase de funcții care au anumite destinații în raport cu aplicațiile practice de utilizare. Clasele de funcții uzuale sunt liniară, treaptă (unipolară sau bipolară) și sigmoidală (unipolară sau bipolară). În tabelul 2.1.1 sunt prezentate analitic și grafic funcțiile de activare cu cea mai frecventă utilizare în practică; de asemenea sunt furnizate *numele funcțiilor* MATLAB aferente, disponibile în **Neural Network Toolbox**.

Revenind la schema bloc din figura 2.1.1, se constată că neuronul artificial cu o singură intrare este *un sistem static* (asigură transferul *instantaneu* al semnalului de intrare către ieșire, spre deosebire de un sistem dinamic, a cărui ieșire depinde atât de intrarea la momentul respectiv cât și de istoria procesului), a cărui funcționare este descrisă prin aplicația:

$$f: \mathbb{R} \rightarrow \mathbb{R}, \quad (2.1.1)$$

$$y = f(x) = \sigma(wx + b). \quad (2.1.2)$$

Așadar f este o funcție realizată prin *compunere*, în următoarea manieră ilustrată sugestiv de schema bloc din figura 2.1.1:

$$y = f(x) = \sigma(u(x)), \quad (2.1.3)$$

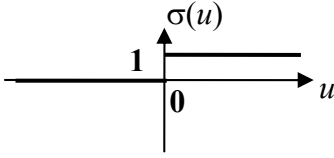
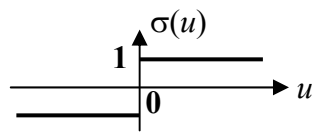
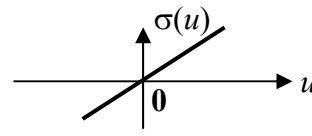
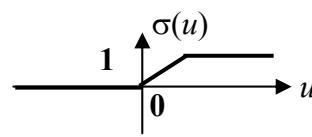
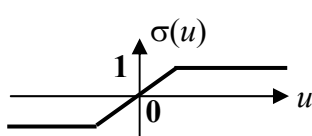
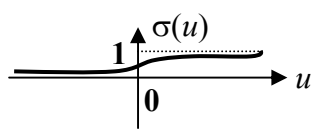
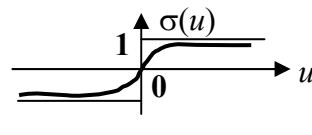
unde:

$$u(x) = r(x) + b, \quad (2.1.4)$$

$$r(x) = wx. \quad (2.1.5)$$

Din expresia lui $y = f(x)$ dată în relația (2.1.2), se observă că *transferul instantaneu* intrare-ieșire $x \rightarrow y$ este *parametrizat* de $w, b \in \mathbb{R}$ sau, cu alte cuvinte, *dependența* lui y de x este *parametrizată* de valorile ponderii $w \in \mathbb{R}$ și deplasării $b \in \mathbb{R}$. Subliniem faptul că, deși conform (2.1.2), ieșirea y se prezintă ca funcție de trei variabile reale (x, w, b), *din punct de vedere sistemic*, variabila x este *singura* cu semnificație de *semnal de intrare*, remarcă pusă în evidență de către schema bloc din fig. 2.1.1. Dintr-un semnal de intrare x nu se poate obține orice formă a semnalului de ieșire y (parametrizat de w și b), cu alte cuvinte, y va păstra informații specifice nodului.

Tabelul 2.1.1. Funcții de activare uzuale

Nume Funcție	Expresie analitică	Reprezentare Grafică	Nume funcție MATLAB
Treaptă unipolară	$\sigma(u) = \begin{cases} 0, & u < 0 \\ 1, & u \geq 0 \end{cases}$		hardlim
Treaptă bipolară	$\sigma(u) = \begin{cases} -1, & u < 0 \\ 1, & u \geq 0 \end{cases}$		hardlims
Funcție liniară	$\sigma(u) = u$		purelin
Liniară cu saturație	$\sigma(u) = \begin{cases} 0, & u < 0 \\ u, & 0 \leq u \leq 1 \\ 1, & u > 1 \end{cases}$		satlin
Liniară cu saturație, simetrică	$\sigma(u) = \begin{cases} 0, & u < -1 \\ u, & -1 \leq u \leq 1 \\ 1, & u > 1 \end{cases}$		satlins
Sigmoid unipolar	$\sigma(u) = \frac{1}{1 + e^{-u}}$		logsig
Sigmoid bipolar	$\sigma(u) = \frac{e^u - e^{-u}}{e^u + e^{-u}}$		tansig

Observație

“1” ca semnal care alimentează blocul b nu este tratat ca o intrare efectivă, întrucât valoarea lui nu se poate schimba (este un semnal constant). Dacă folosim pentru “1” interpretarea de semnal de intrare, atunci putem considera vectorul $\tilde{x} = [x \ 1]^T$ ca un semnal vectorial extins cu a doua componentă constantă, egală cu 1. Notând $\theta = [w \ b]$ (vectorul parametrilor

neuronului), obținem $u(x) = \theta \tilde{x}$. Această convenție de notație prezintă avantajul că, prin înglobarea deplasării b într-un vector, se ajunge la o scriere liniară, mai compactă. ■

2.1.2. Neuronul artificial cu mai multe intrări

Neuronul cu mai multe intrări generalizează modelul prezentat în paragraful anterior, în sensul că *semnalul de intrare este vectorial* (posedă mai multe componente). În figura 2.1.2 se prezintă schema detaliată a unui neuron cu ieșirea $y \in \mathbb{R}$ și cu n intrări, notate $x_i \in \mathbb{R}$, $i = 1, \dots, n$.

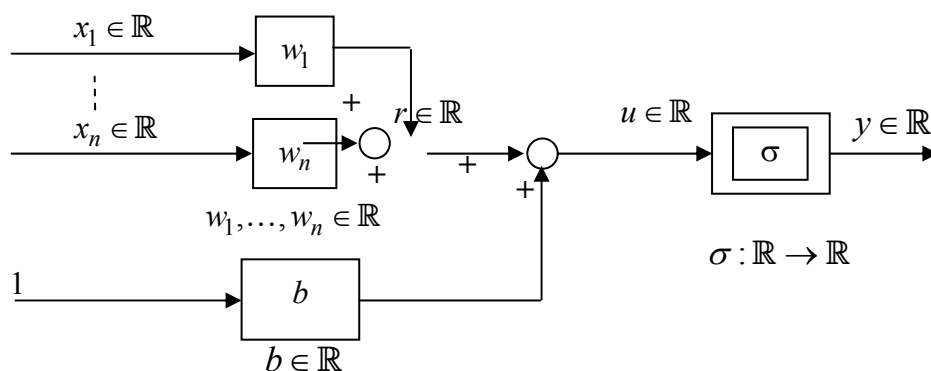


Figura 2.1.2. Schema detaliată a unui neuron cu n intrări

Funcția de activare rămâne neschimbată, dar intrarea ei, u , are expresia:

$$u = w_1 \cdot x_1 + \dots + w_n \cdot x_n + b = \sum_{i=1}^n w_i x_i + b, \quad (2.1.6)$$

care se poate scrie în formă vectorial-matriceală:

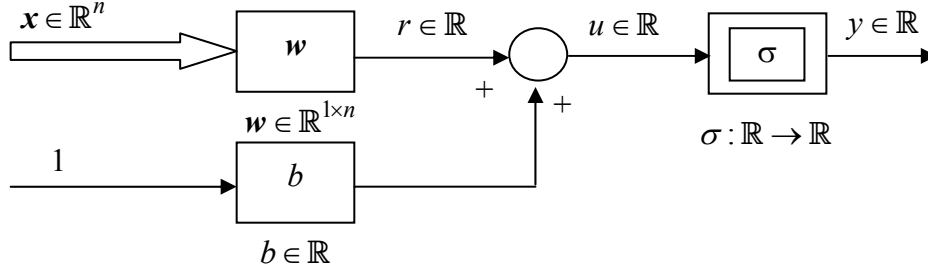
$$u = \mathbf{w} \mathbf{x} + b, \quad (2.1.7)$$

unde $\mathbf{x}$ este *vectorul coloană* conținând intrările, $\mathbf{x} = [x_1 \dots x_n]^T \in \mathbb{R}^n$, iar $\mathbf{w}$ este *vectorul linie* (adică un caz particular de matrice) conținând ponderile, $\mathbf{w} = [w_1 \dots w_n] \in \mathbb{R}^{1 \times n}$.

Această scriere ne arată că reprezentarea grafică compactă a neuronului cu mai multe intrări se poate realiza ca în figura 2.1.3. În respectiva figură, pentru semnalul de intrare $\mathbf{x} \in \mathbb{R}^n$ s-a utilizat convenția de reprezentare grafică cu săgeată dublă (adică $\Rightarrow$), ceea ce permite compactarea reprezentării.

Observație

Precizăm că, începând din acest punct al textului nostru, vom face apel la convenția de reprezentare grafică cu săgeată dublă ($\Rightarrow$) a semnalelor vectoriale fără a mai atrage atenția în mod special asupra semnificației, considerând că cititorul și-a însușit deja respectiva convenție. ■

Figura 2.1.3. Schema bloc a unui neuron cu n intrări

Neuronul cu n intrări este un *sistem static* a cărui funcționare este descrisă prin *aplicația de variabilă vectorială*:

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, \quad (2.1.8)$$

$$y = f(\mathbf{x}) = \sigma(\mathbf{w} \mathbf{x} + b), \quad (2.1.9)$$

ce rezultă din următoarea compunere de funcții (pusă în evidență de schema din figura 2.1.3):

$$y = f(\mathbf{x}) = \sigma(u(\mathbf{x})) \quad (2.1.10)$$

unde:

$$u(\mathbf{x}) = r(\mathbf{x}) + b, \quad (2.1.11)$$

$$r(\mathbf{x}) = \mathbf{w} \mathbf{x}. \quad (2.1.12)$$

Se observă că transferul $\mathbf{x} \rightarrow y$ este *parametrizat* de $n+1$ elemente: n elemente conținute în vectorul linie al ponderilor $\mathbf{w} \in \mathbb{R}^{1 \times n}$ și scalarul $b \in \mathbb{R}$ (deplasarea fiind unică pentru un neuron).

2.1.3. Alte tipuri de neuroni artificiali

Există tipuri de neuroni artificiali care diferă de structura neuronului artificial prezentată în fig. 2.1.2 (sau, echivalent, în fig. 2.13). Un exemplu îl constituie neuronul artificial cu funcție de activare de tip radial (eng. *Radial Basis Function*, abreviat RBF). Fără a furniza detalii referitoare la alte tipuri de neuroni artificiali, considerăm oportună prezența acestui paragraf pentru a asigura o vedere de ansamblu corectă asupra problematicei tratate în capitolul curent. Astfel, deși modul de abordare vizează un grad mare de generalitate, capitolul de față nu își propune să acopere integral prezentarea structurilor neuronilor artificiali și a rețelelor neurale feedforward.

Totodată mai precizăm că elementele de noutate ce vor apărea în viitoarele capitole vor fi grefate cu ușurință pe scheletul informațional construit în capitolul de față, ca o extindere naturală a manierei prin care se poate defini semnalul $u(\mathbf{x})$ aplicat drept intrare în nodul neuronului. Concret, în cazul altor tipuri de neuroni, transferul intrare-ieșire este descris tot printr-o relație de forma (2.1.10), cu deosebirea că expresia analitică a lui $u(\mathbf{x})$ nu mai este dată de (2.1.11) și (2.1.12). Atragem însă atenția asupra faptului că, și pentru aceste tipuri noi de neuroni artificiali, expresia analitică a lui $u(\mathbf{x})$ va rămâne parametrizată de $\mathbf{w}$, b (existența parametrilor fiind cea care asigură proprietățile neuronului și respectiv a unei rețele compuse din mai mulți neuroni).

2.2. Rețele de neuroni artificiali

Prin conectarea în diverse moduri a neuronilor artificiali, se pot construi rețele neurale artificiale. Rețelele neurale cu arhitecturi specifice posedă proprietăți specifice, ceea ce face ca, în aplicații, unui anumit tip de problemă să i se asocieze un anumit tip de rețea neurală. Scopul acestei secțiuni constă în prezentarea generală a principiilor ce stau la baza organizării unei rețele neurale feedforward, urmând ca în capitolele viitoare să se detalieze asocierea dintre tipologia rețelelor și specificitatea aplicațiilor.

2.2.1. Rețea cu un singur strat

În figura 2.2.1 este ilustrată schema bloc a unei rețele neurale cu un singur strat (eng. *layer*) conținând p neuroni, utilizând convenția de reprezentare grafică cu săgeți duble pentru semnale multivariabile. Rețeaua prezintă n intrări, notate compact $\mathbf{x} \in \mathbb{R}^n$, și p ieșiri, notate compact $\mathbf{y} \in \mathbb{R}^p$:

$$\mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} \in \mathbb{R}^n, \quad \mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_p \end{bmatrix} \in \mathbb{R}^p. \quad (2.2.1)$$

Numărul de ieșiri ale rețelei coincide cu numărul de neuroni: fiecărui neuron îi corespunde o ieșire a stratului. Cei p neuroni din strat au funcții de activare proprii: neuronul i are funcția de activare σ_i .

Semnalele de intrare $x_1, \dots, x_n$ sunt utilizate de toți cei p neuroni, dar cu combinații diferite de ponderi pentru fiecare neuron în parte. Funcția de activare a fiecărui neuron în parte are un singur semnal de intrare (semnal scalar) u_i corespunzător neuronului respectiv. (Se respectă arhitectura studiată a neuronului cu mai multe intrări unde s-a pus în evidență faptul că funcția de activare are un singur semnal de intrare). Pentru compactitatea desenului diagramei, funcțiile de activare a celor p neuroni sunt reprezentate unitar în blocul Σ având drept intrare semnalul multivariabil $\mathbf{u}$ și drept ieșire semnalul multivariabil $\mathbf{y}$. În diagramă este specificat faptul că fiecare componentă a funcției vectoriale Σ acționează doar pe componenta corespunzătoare ei a semnalului de intrare $\mathbf{u}$, adică $\sigma_i(\mathbf{u}) = \sigma_i(u_i) = y_i$, $i = \overline{1, p}$.

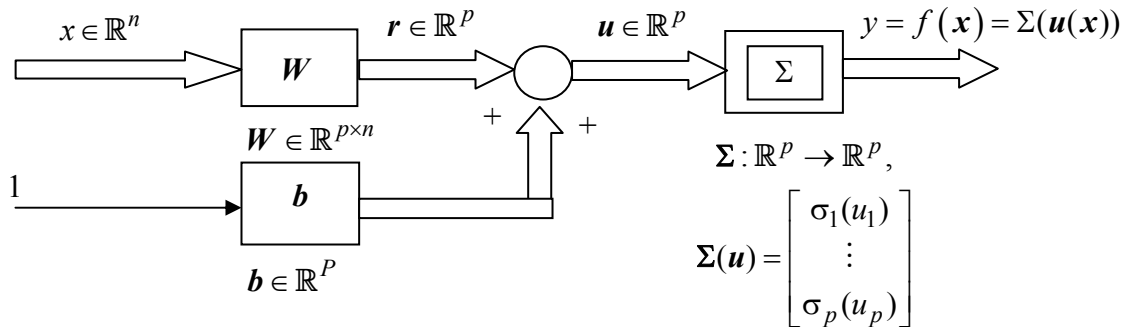


Figura 2.2.1. Schema bloc a unei rețele neurale cu un singur strat de p neuroni

Fiecare din cei p neuroni are în structură $n+1$ parametri (n parametri de tip pondere, care fac legătura cu cele n semnale de intrare $x_1, x_2, \dots, x_n$, și o deplasare). Cele n ponderi ale neuronului i , notate $w_{i,1}, w_{i,2}, \dots, w_{i,n}$, constituie elementele liniei i a matricei $W \in \mathbb{R}^{p \times n}$; deplasarea corespunzătoare aceluiași neuron, notată b_i , reprezintă componenta i a vectorului coloană $b \in \mathbb{R}^p$.

Rețeaua neurală cu un singur strat este un *sistem static* a cărui funcționare este descrisă prin *aplicația vectorială, de variabilă vectorială*:

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^p, \quad (2.2.2)$$

$$y = f(x) = \Sigma(Wx + b), \quad (2.2.3)$$

ce rezultă din următoarea compunere de funcții (pusă în evidență de schema din figura 2.2.1)

$$y = f(x) = \Sigma(u(x)), \quad (2.2.4)$$

unde:

$$u(x) = r(x) + b, \quad (2.2.5)$$

$$r(x) = Wx. \quad (2.2.6)$$

Se observă că transferul $x \rightarrow y$ este *parametrizat* de cele pn elemente ale matricei $W \in \mathbb{R}^{p \times n}$ și cele p elemente ale vectorului $b \in \mathbb{R}^p$.

Observație

În arhitecturile de rețele cu aplicabilitate practică toți neuronii au aceeași funcție de activare de tipul celor precizate în tabelul 2.1.1. Toate componentele funcției Σ sunt identice ca exprimare analitică, adică $y_i = \sigma_i(u_i) = \sigma(u_i)$, $i = \overline{1, p}$. ■

2.2.2. Rețea cu două straturi

În figura 2.2.2 se prezintă schema bloc a unei rețele neurale cu două straturi, conținând p_1 și, respectiv, p_2 neuroni. Rețeaua posedă n_1 intrări, notate compact $x^1 \in \mathbb{R}^{n_1}$, și p_2 ieșiri, notate compact $y^2 \in \mathbb{R}^{p_2}$:

$$x^1 = \begin{bmatrix} x_1^1 \\ \vdots \\ x_{n_1}^1 \end{bmatrix} \in \mathbb{R}^{n_1}, \quad y^2 = \begin{bmatrix} y_1^2 \\ \vdots \\ y_{p_2}^2 \end{bmatrix} \in \mathbb{R}^{p_2}. \quad (2.2.7)$$

Arhitectura acestei rețele rezultă prin înserierea a două rețele cu câte un singur strat fiecare (de tipul celor prezentate în paragraful 2.2.1), în așa manieră încât ieșirile primului strat să constituie intrările celui de-al doilea strat.

Primul strat al rețelei mai este referit cu denumirea de *strat de intrare* (eng. *input layer*) iar al doilea cu denumirea de *strat de ieșire* (eng. *output layer*). Pentru conectarea în serie a celor două rețele se impune ca numărul ieșirilor primului strat (adică p_1) să coincidă

cu numărul intrărilor în cel de al doilea strat (adică n_2) deoarece semnalul vectorial de la ieșirea primului strat, $\mathbf{y}^1 \in \mathbb{R}^{p_1}$, constituie semnalul vectorial de intrare pentru al doilea strat, $\mathbf{x}^2 \in \mathbb{R}^{n_2}$. Astfel, are loc identitatea de semnale:

$$\mathbf{y}^1 (\in \mathbb{R}^{p_1}) \equiv \mathbf{x}^2 (\in \mathbb{R}^{n_2}). \quad (2.2.8)$$

Numărul de intrări ale rețelei este numărul intrărilor primului strat, iar numărul de ieșiri ale rețelei este dat de numărul de neuroni din stratul doi.

Rețeaua neurală cu două straturi este un sistem static a cărui funcționare este descrisă prin aplicația vectorială, de variabilă vectorială:

$$\mathbf{f} : \mathbb{R}^{n_1} \rightarrow \mathbb{R}^{p_2}, \quad (2.2.9)$$

$$\mathbf{y}^2 = \mathbf{f}(\mathbf{x}^1) = \Sigma^2 \left(\mathbf{W}^2 \sigma^1(\mathbf{W}^1 \mathbf{x}^1 + \mathbf{b}^1) + \mathbf{b}^2 \right). \quad (2.2.10)$$

Această aplicație rezultă din următoarea compunere de funcții (pusă în evidență de schema bloc din figura 2.2.2):

- La nivelul celui de al doilea strat:

$$\mathbf{y}^2 = \mathbf{f}^2(\mathbf{x}^2) = \Sigma^2 \left(\mathbf{u}^2(\mathbf{x}^2) \right), \quad (2.2.11)$$

unde:

$$\mathbf{u}^2(\mathbf{x}^2) = \mathbf{r}^2(\mathbf{x}^2) + \mathbf{b}^2, \quad (2.2.12)$$

$$\mathbf{r}^2(\mathbf{x}^2) = \mathbf{W}^2 \mathbf{x}^2. \quad (2.2.13)$$

- La nivelul primului strat:

$$\mathbf{x}^2 = \mathbf{y}^1 = \mathbf{f}^1(\mathbf{x}^1) = \Sigma^1 \left(\mathbf{u}^1(\mathbf{x}^1) \right), \quad (2.2.14)$$

unde:

$$\mathbf{u}^1(\mathbf{x}^1) = \mathbf{r}^1(\mathbf{x}^1) + \mathbf{b}^1, \quad (2.2.15)$$

$$\mathbf{r}^1(\mathbf{x}^1) = \mathbf{W}^1 \mathbf{x}^1. \quad (2.2.16)$$

Se observă că transferul $\mathbf{x}^1 \rightarrow \mathbf{y}^2$ realizat de rețeaua cu două straturi de neuroni este parametrizat de elementele matricelor $\mathbf{W}^1 \in \mathbb{R}^{p_1 \times n_1}$, $\mathbf{W}^2 \in \mathbb{R}^{p_2 \times n_2}$ (ponderile celor două straturi) și elementele vectorilor $\mathbf{b}^1 \in \mathbb{R}^{p_1}$, $\mathbf{b}^2 \in \mathbb{R}^{p_2}$ (deplasările celor două straturi).

Observație

Funcțiile de activare de pe primul strat pot fi diferite de funcțiile de activare de pe stratul doi (dar, de regulă, pe același strat funcțiile de activare sunt identice).

În aplicațiile practice arhitecturile uzuale nu depășesc trei straturi.

Ca terminologie rețeaua considerată în diagrama bloc conține, după unele texte, un strat de intrare și unul de ieșire iar, după alte texte, un strat ascuns și un strat de ieșire. ■

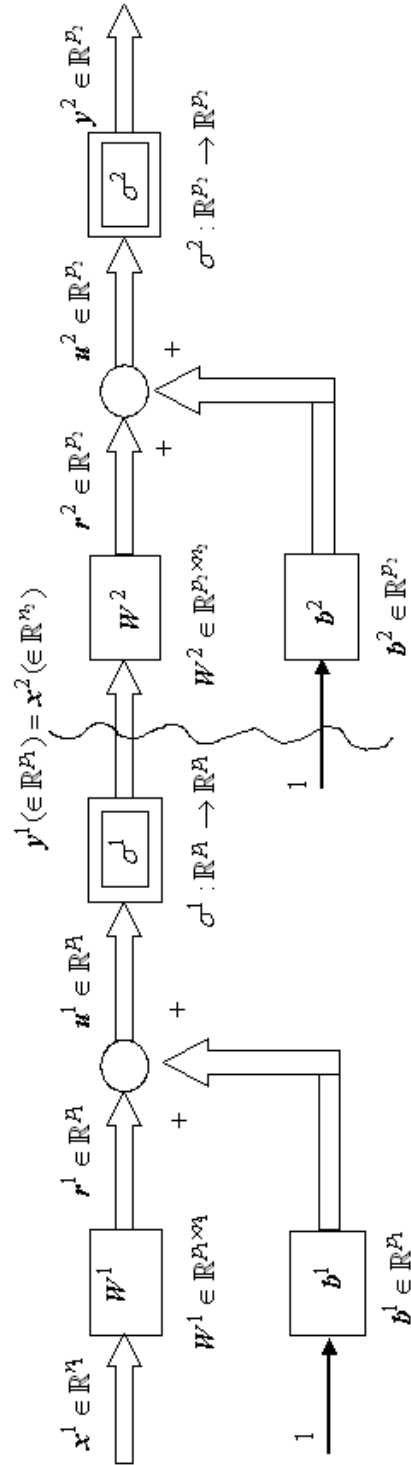


Figura 2.2.2. Schema bloc a unei rețele neurale cu două straturi, conținând p_1 și, respectiv, p_2 neuroni

2.2.3. Cazul general - rețea cu mai multe straturi

Într-o manieră similară celei prezentate în paragraful 2.2.2 se poate construi o rețea cu trei (sau mai multe) straturi de neuroni, conectând trei (sau mai multe) rețele cu un singur strat, astfel încât intrările fiecărui strat (exceptând primul) să coincidă cu ieșirile stratului care îl precede. Cele amintite mai înainte cu privire la terminologia prin care se referă straturile rețelei rămân valabile. Existența mai multor straturi are drept efect atât posibilitatea de a defini transferul intrare-ieșire cu ajutorul a mai mulți parametri, cât și mai multe nivele de neliniarități care se compun în acest transfer. Asupra acestor aspecte se va reveni în detaliu atunci când se va discuta despre capacitatea de aproximare pe care o posedă unele rețele neurale. În practica uzuală se face apel la arhitecturi *cu cel mult trei straturi*.

Straturile de neuroni cuprinse între stratul de intrare și cel de ieșire se numesc *straturi ascunse* (eng. *hidden*). În unele texte, orice strat al rețelei care nu este strat de ieșire este referit drept *strat ascuns* (inclusiv stratul de intrare). Ca urmare, exprimarea fără nici un fel de echivoc este cea care atribuie numere de ordine straturilor, astfel în cazul rețelei cu două straturi, stratul 1 desemnează stratul de intrare și stratul 2 este stratul de ieșire.

2.3. Instrumente software pentru simularea rețelelor neurale cu structură feedforward oferite de *Neural Network Toolbox*

Neural Network Toolbox este unul dintre primele pachete de programe MATLAB lansate de The MathWorks, ceea ce demonstrează importanța arătată domeniului rețelelor neurale de către firma dezvoltatoare a acestui software. Versiunea 4.0 a toolbox-ului dedicat rețelelor neurale a fost lansată în anul 2000, împreună cu versiunea MATLAB Release 12. Odată cu dezvoltarea software-ului de bază (MATLAB), pachetului *Neural Network Toolbox* i-au fost aduse modificări minore, astfel că versiunea 5.0 a fost lansată împreună cu versiunea MATLAB R2006a.

Neural Network Toolbox admite lucrul cu diferite tipuri de rețele neurale. Această flexibilitate se datorează reprezentării orientate obiect a rețelelor, care permite definirea unor arhitecturi diverse și implementarea algoritmilor specifici acestora. Un obiect de tip *network* constă din numeroase proprietăți ce pot fi setate în mod adecvat pentru specificarea arhitecturii și comportării rețelei neurale și poate fi creat cu funcția *network*. Pentru crearea unor rețele neurale cu o structură standard există însă funcții special concepute (de exemplu *newp*, *newlin*, *newlind*, *newff*) ce vor fi studiate în capitolele ce urmează.

Proprietățile unui obiect *network* sunt reunite, în funcție de tipul lor, în mai multe grupe, și anume:

- **architecture:** proprietăți ce permit setarea numărului de intrări și de straturi ale rețelei neurale, precum și modul de conectare a acestora;
- **functions:** pentru definirea funcțiilor pentru operațiile de bază cu rețeaua neurală (inițializare, modul de adaptare a parametrilor, criteriul de performanță, funcția de antrenare, precum și parametri specifici acestora);

- **weight and bias values:** pentru inspectarea sau setarea valorilor ponderilor și deplasărilor neuronilor din rețea;
- **other:** pentru păstrarea altor date dorite de utilizator (de exemplu, comentarii).

În pachetul **Neural Network Toolbox** există și o interfață grafică cu utilizatorul (*Graphical User Interface - GUI*), **nntool**, pentru importarea, crearea, utilizarea și exportarea rețelelor neurale și/sau datelor. De asemenea, **Neural Network Toolbox** pune la dispoziția utilizatorilor un set de blocuri pentru construirea rețelelor neurale în Simulink, precum și funcția **gensim** pentru generarea versiunii Simulink a unei rețele neurale oarecare care a fost creată în MATLAB. Modul de utilizare a funcțiilor este ilustrat prin numeroase programe demonstrative ce pot fi lansate în execuție prin intermediul comenzii **demons**.

Pentru simularea comportării unei rețelei neurale se poate utiliza funcția **sim**.

Pentru aprofundarea noțiunilor prezentate în acest capitol se recomandă studierea problemelor cu soluții analitice 7.1, 7.2 și a problemelor cu soluții numerice 8.1 – 8.4.

Capitolul 3

APLICAȚII ALE REȚELELOR FEED-FORWARD CU FUNCȚII DE ACTIVARE TREAPTĂ ÎN PROBLEME DE CLASIFICARE

Neuronul artificial cu funcție de activare treaptă reprezintă forma cea mai simplă a unei rețele neurale artificiale. El a fost introdus în 1943 de Warren McCulloch și Walter Pitts în lucrarea „A Logical Calculus of Ideas Immanent in Nervous Activity”, în încercarea de a modela activitatea unui neuron biologic. Neuronul McCulloch-Pitts acceptă doar intrări binare, furnizând la ieșire un semnal binar: ieșirea unui perceptron este 1 sau 0 după cum neuronul respectiv răspunde sau nu răspunde la semnalele de intrare aplicate. Toate intrările au ponderi egale.

Donald Hebb a revoluționat domeniul rețelelor neuronale artificiale în 1949. În monografia *The Organization of Behavior*, Hebb a propus un mecanism calitativ ce reflectă modul în care se modifică conexiunile sinaptice pentru a reflecta mecanismul de învățare realizat de neuronii interconectați atunci când aceștia sunt influențați de diferiți stimuli din mediu. Acest mecanism poartă numele de *regula lui Hebb* și stă la baza proceselor de învățare (antrenare) a rețelelor neurale artificiale.

Aplicând acest principiu neuronului McCulloch-Pitts, Frank Rosenblatt a dezvoltat primul *perceptron* (1957) capabil să învețe prin modificarea ponderilor corespunzătoare intrărilor sale, modelând astfel capacitatea de (recunoaștere a formelor) clasificare a sistemelor de vedere biologice. Rosenblatt a demonstrat că dacă vectorii prototip utilizați pentru antrenarea unui perceptron fac parte din două clase liniar separabile, atunci algoritmul de antrenare converge într-un număr finit de pași, furnizând un hiperplan ce separă cele două clase. Modelul propus în anul 1962 în cartea *Principles of Neurodynamics* a influențat decisiv dezvoltarea rețelelor neurale.

În prezent, rețelele neurale feed-forward cu *funcții de activare treaptă* poartă denumirea de *rețele perceptron*. Asemenea rețele pot fi utilizate numai în probleme de clasificare liniară și necesită algoritmi de antrenare specifici, care exploatează forma particulară a funcțiilor de activare.

3.1. Studiarea transferului intrare-ieșire pentru un neuron

Considerăm dat un perceptron cu n intrări (figura 3.1.1), având ponderile $w_1, w_2, \dots, w_n$, deplasarea b și funcție de activare de tip treaptă unitate:

$$\sigma(u) = \begin{cases} 0, & u < 0, \\ 1, & u \geq 0. \end{cases} \quad (3.1.1)$$

Cu notațiile introduse în paragraful 2.1.2, și anume $\mathbf{x} = [x_1 \dots x_n]^T \in \mathbb{R}^n$ pentru vectorul coloană corespunzător intrărilor și $\mathbf{w} = [w_1 \dots w_n] \in \mathbb{R}^{1 \times n}$ pentru vectorul linie conținând ponderile perceptronului, intrarea funcției de activare se scrie sub forma

$$u(\mathbf{x}) = \mathbf{w} \mathbf{x} + b. \quad (3.1.2)$$

Ieșirea perceptronului este dată de

$$y(\mathbf{x}) = \begin{cases} 0, & \text{dacă } \mathbf{w} \mathbf{x} + b < 0, \\ 1, & \text{dacă } \mathbf{w} \mathbf{x} + b \geq 0. \end{cases} \quad (3.1.3)$$

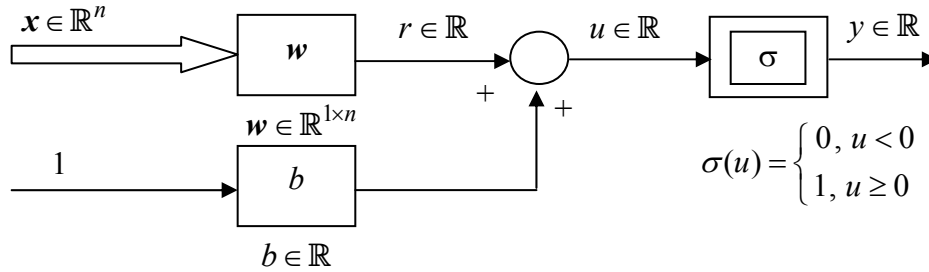


Figura 3.1.1. Schema bloc a unui perceptron cu n intrări

Se observă că spațiul de intrare, $\mathbb{R}^n$ poate fi separat în două clase în funcție de valorile ieșirii perceptronului: $\mathcal{C}_0 = \{\mathbf{x} \in \mathbb{R}^n \mid y(\mathbf{x}) = 0\}$ și $\mathcal{C}_1 = \{\mathbf{x} \in \mathbb{R}^n \mid y(\mathbf{x}) = 1\}$. Suprafața de separație dintre cele două clase este hiperplanul de ecuație $u(\mathbf{x}) = 0$.

Problema inversă este următoarea: având date un număr N de *eșantioane* sau *prototipuri* (eng. *pattern*), $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$, despre care se presupune că aparțin la două clase, notate $\mathcal{C}_0$ și $\mathcal{C}_1$, să se determine parametrii unui perceptron care să realizeze separarea dorită. Dacă această problemă admite soluție (dacă există un hiperplan care să separe punctele ce aparțin unei clase de punctele celeilalte clase) se spune că cele două clase sunt liniar separabile.

În cazuri simple, parametrii unui perceptron pot fi calculați în mod direct astfel încât să se realizeze separarea spațiului de intrare în modul dorit, așa cum se prezintă în continuare.

Exemplul 3.1.1. Se dorește implementarea funcției logice AND utilizând un perceptron cu două intrări și funcție de activare unipolară. Tabela de adevăr și reprezentarea grafică a acestei funcții sunt prezentate în figura 3.1.2.

Se observă că dreapta de ecuație $d(\mathbf{x}) = x_1 + x_2 - 1.5 = 0$ realizează separarea liniară a celor patru puncte ale mulțimii $P = \{(x_1, x_2) | x_1, x_2 \in \{0, 1\}\}$ în funcție de rezultatul operației logice $x_1 \wedge x_2$. Prin identificare se obțin parametrii perceptronului $w_1 = w_2 = 1$ și $b = -1.5$.

Evident că dreapta de separație a celor două clase $\mathcal{C}_0 = \{(x_1, x_2) \in P | x_1 \wedge x_2 = 0\}$ și $\mathcal{C}_1 = \{(x_1, x_2) \in P | x_1 \wedge x_2 = 1\}$ nu este unică, ceea ce arată că implementarea funcției logice AND sub formă de perceptron nu este unică.

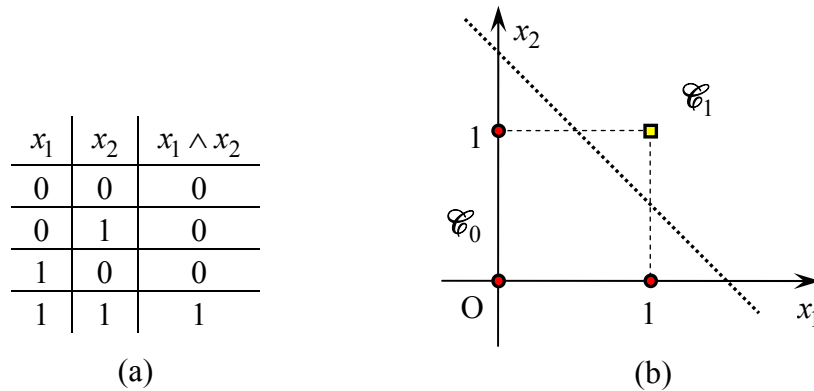


Figura 3.1.2. (a) Tabela de adevăr și (b) reprezentarea grafică a funcției logice AND.

Similar se obțin parametrii unui perceptron cu două intrări și funcție de activare unipolară ce implementează funcția logică OR (figura 3.1.3), cu valorile $w_1 = w_2 = 1$ și $b = -0.5$.

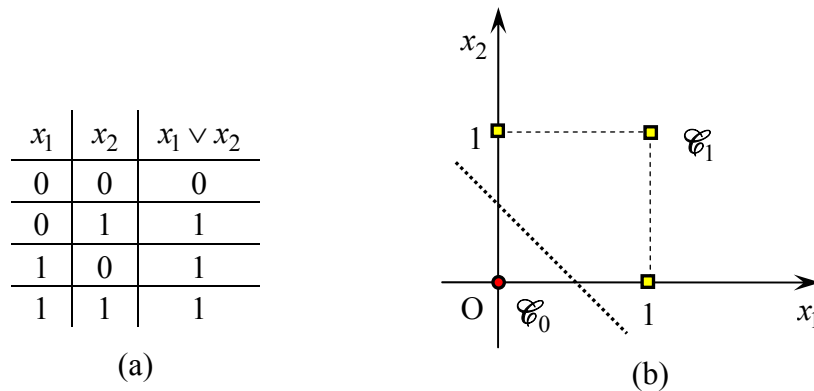


Figura 3.1.3. (a) Tabela de adevăr și (b) reprezentarea grafică a funcției logice OR.

Pentru funcția logică XOR cu tabela de adevăr și reprezentarea grafică din figura 3.1.4, clasele $\mathcal{C}_0$ și $\mathcal{C}_1$ nu sunt liniar separabile, ceea ce arată că această funcție nu poate fi implementată utilizând o rețea neurală de tip perceptron cu un singur strat.

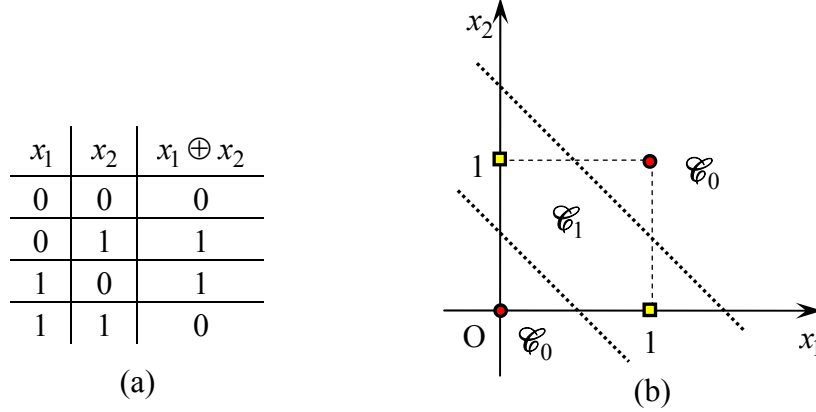


Figura 3.1.4. (a) Tabela de adevăr și (b) reprezentarea grafică a funcției logice XOR.

3.2. Utilizarea rețelelor cu un singur strat

3.2.1. Proprietatea de separare folosită în clasificare

O rețea neuronală cu un singur strat cu n intrări, $\mathbf{x} \in \mathbb{R}^n$, p ieșiri (neuroni) $\mathbf{y} \in \mathbb{R}^p$ și funcții de activare treaptă (unipolară sau bipolară) posedă proprietatea de a partiționa spațiul vectorilor de intrare (adică $\mathbb{R}^n$) în 2^p mulțimi prin p hiperplane. O asemenea rețea va putea fi utilizată în probleme de clasificare *numai dacă* regiunile de decizie sunt mulțimi din $\mathbb{R}^n$ separabile liniar (adică cu ajutorul unor hiperplane).

3.2.2. Obiectivul antrenării rețelei

La intrarea rețelei sunt prezentate N eșantioane, sau *prototipuri*, notate $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$, despre care se presupune că aparțin claselor $\mathcal{C}_j$, $j = 1, \dots, 2^p$, clase ce se pot defini ca regiuni separabile liniar din $\mathbb{R}^n$. Să introducem, pentru simplificarea scrierii ulterioare, notația $M = 2^p$. Evident problema clasificării are sens pentru $N > M$, altminteri situându-ne într-o situație pur banală. În termeni concreți, pentru fixare, vom considera situațiile următoare:

- m_1 vectori din $\mathbb{R}^n$ aparțin clasei $\mathcal{C}_1$, adică:

$$\mathbf{x}_k \in \mathcal{C}_1, k = 1, \dots, m_1; \quad (3.2.1-1)$$

- m_2 vectori din $\mathbb{R}^n$ aparțin clasei $\mathcal{C}_2$, adică:

$$\mathbf{x}_k \in \mathcal{C}_2, k = m_1 + 1, \dots, m_1 + m_2; \quad (3.2.1-2)$$

...

- m_M vectori din $\mathbb{R}^n$ aparțin clasei $\mathcal{C}_M$, adică:

$$\mathbf{x}_k \in \mathcal{C}_M, k = m_1 + m_2 + \dots + m_{M-1} + 1, \dots, m_1 + m_2 + \dots + m_{M-1} + m_M = N. \quad (3.2.1-M)$$

Modalitățile de definire prin expresii analitice a celor $M = 2^p$ clase (adică ecuațiile celor p hiperplane ce realizează separarea) *nu sunt cunoscute aprioric*. Determinarea sub formă analitică a acestor hiperplane reprezintă tocmai *obiectivul antrenării rețelei*, astfel încât, în urma unei antrenări adecvate, parametrii rețelei (ponderi și deplasări) să furnizeze coeficienții ecuațiilor ce definesc hiperplanele de separare.

În acest scop, fiecăreia dintre cele 2^p clase $\mathcal{C}_j$ i se atașează câte un vector $\underline{\mathbf{y}}_j \in \mathbb{R}^p$, $j = 1, \dots, 2^p$, cu elemente binare, astfel:

- valori 0 și 1 – în cazul când funcțiile de activare ale neuronilor sunt de tip treaptă unipolară;
- valori -1 și 1 – în cazul când funcțiile de activare ale neuronilor sunt de tip treaptă bipolară.

Subliniem faptul că prin acest procedeu pot fi construiți exact 2^p vectori distincți, cu elemente binare, ceea ce permite realizarea unei corespondențe biunivoce între vectorii $\underline{\mathbf{y}}_j$ și clasele $\mathcal{C}_j$. Astfel, problema clasificării revine la a antrena rețeaua (adică a determina parametrii rețelei) de așa manieră încât *rețeaua cu parametrii rezultați din antrenare* să asigure satisfacerea următoarelor N egalități ce decurg din (1-1)...(1-M):

$$\mathbf{f}(\mathbf{x}_k) = \begin{cases} \underline{\mathbf{y}}_1, & k = 1, \dots, m_1; \\ \underline{\mathbf{y}}_2, & k = m_1 + 1, \dots, m_1 + m_2; \\ \dots & \\ \underline{\mathbf{y}}_M, & k = m_1 + \dots + m_{M-1} + 1, \dots, N; \end{cases} \quad (3.2.2)$$

Cu alte cuvinte, fiecărui vector prototip $\mathbf{x}_i$, $i = 1, \dots, N$, i se poate atașa un vector țintă (eng. *target*) $\mathbf{z}_i$, cu:

$$\mathbf{z}_i \in \{\underline{\mathbf{y}}_1, \underline{\mathbf{y}}_2, \dots, \underline{\mathbf{y}}_M\}, i = 1, \dots, N, \quad (3.2.3)$$

astfel încât să fie satisfăcute egalitățile (3.2.2). Subliniem faptul că dintre cei N vectori țintă $\mathbf{z}_i$, $i = 1, \dots, N$, numai $M < N$ sunt distincți (coincizând cu unul din vectorii $\underline{\mathbf{y}}_1, \underline{\mathbf{y}}_2, \dots, \underline{\mathbf{y}}_M$).

Astfel, prin antrenare, se va urmări ca:

$$\mathbf{f}(\mathbf{x}_i) = \mathbf{z}_i, i = 1, \dots, N, \quad (3.2.4)$$

unde forma particulară a lui $\mathbf{z}_i$ din mulțimea (3.2.3) este precizată prin egalitatea (3.2.2), corespunzător valorii concrete a indicelui i .

Observații

1° În formularea obiectivului antrenării s-a presupus că se operează cu clase separabile liniar. Totuși, uzual, în practică, această ipoteză de separabilitate liniară *nu poate fi verificată*

înainte de a începe antrenarea. Altfel spus, nu dispunem de un procedeu care să ne asigure, *înainte de efectuarea antrenării*, că vectorii de intrare $\mathbf{x}_i$, $i=1, \dots, N$, pot fi grupați cu satisfacerea condițiilor (3.2.1-1), ..., (3.2.1- M), în $M = 2^p$ clase $\mathcal{C}_j$, $j=1, \dots, M$, separabile liniar. Astfel, *însuși rezultatul antrenării* este cel care *confirmă* sau *infirmă* (prin verificarea satisfacerii celor N egalități (3.2.4)) ipoteza separabilității liniare a claselor, iar în caz de confirmare, sunt furnizate și valorile numerice ale coeficienților ecuațiilor pentru hiperplanele de separare.

2° În cazul când rezultatul antrenării arată că cele N egalități (3.2.4) nu pot fi satisfăcute, nu înseamnă că vectorii $\mathbf{x}_i$, $i=1, \dots, N$, nu *ar putea* fi grupați în $M = 2^p$ clase *oarecare*, separabile liniar, ci semnifică faptul că *modul concret de asignare* a lui $\mathbf{x}_i$ la $\mathcal{C}_j$ definit prin relațiile (1-1), ..., (1- M) *nu permite* separarea liniară a claselor. Deci, *un alt mod de definire* a relațiilor de apartenență (3.2.1-1), ..., (3.2.1- M) *ar putea* (nu în mod sigur) conduce la clase separabile liniar. ■

3.2.3. Algoritmul de antrenare al rețelei

Algoritmul de *antrenare*, *învățare*, sau *instruire* (eng. *training*, *learning*) constă în actualizarea parametrilor rețelei, adică a ponderilor (elementele matricei $\mathbf{W} \in \mathbb{R}^{p \times n}$) și a deplasărilor (elementele vectorului coloană $\mathbf{b} \in \mathbb{R}^p$), cu scopul de a satisface cele N egalități (3.2.4).

Se organizează vectorii prototip $\mathbf{x}_i \in \mathbb{R}^n$, $i=1, \dots, N$, și cei țintă $\mathbf{z}_i \in \mathbb{R}^p$, $i=1, \dots, N$, sub forma matricelor:

$$\mathbf{X} = [\mathbf{x}_1 \dots \mathbf{x}_N] \in \mathbb{R}^{n \times N}, \quad (3.2.5)$$

respectiv:

$$\mathbf{Z} = [\mathbf{z}_1 \dots \mathbf{z}_N] \in \mathbb{R}^{p \times N}. \quad (3.2.6)$$

Pentru algoritmul de antrenare se stabilește de către utilizator un *număr maxim de iterații* sau *epoci*. La *fiecare iterație* a algoritmului de antrenare se parcurg etapele descrise mai jos. Valorile parametrilor rețelei la începerea unei iterații sunt notate prin $\mathbf{W}_{\text{old}} \in \mathbb{R}^{p \times n}$ (pentru ponderi) și $\mathbf{b}_{\text{old}} \in \mathbb{R}^p$ (pentru deplasări).

Etapa 1 (Etapa de prezentare)

Pentru *valorile curente* ale parametrilor rețelei ($\mathbf{W}_{\text{old}} \in \mathbb{R}^{p \times n}$, $\mathbf{b}_{\text{old}} \in \mathbb{R}^p$) se calculează ieșirile rețelei:

$$\mathbf{y}_i = \mathbf{f}(\mathbf{x}_i) = \sigma(\mathbf{W}_{\text{old}} \cdot \mathbf{x}_i + \mathbf{b}_{\text{old}}), \quad i=1, \dots, N, \quad (3.2.7)$$

care se organizează sub forma matricei:

$$\mathbf{Y} = [\mathbf{y}_1 \dots \mathbf{y}_N] \in \mathbb{R}^{p \times N}. \quad (3.2.8)$$

Etapă 2 (Etapă de verificare)

Se calculează vectorii erorilor:

$$\mathbf{e}_i = \mathbf{z}_i - \mathbf{y}_i \in \mathbb{R}^P, \quad i = 1, \dots, N, \quad (3.2.9)$$

ca diferențe dintre vectorii țintă $\mathbf{z}_i$ (aparținând mulțimii (3.2.3)) și vectorii ieșirii la iterația curentă $\mathbf{y}_i$ (definiți conform (3.2.7)).

Algoritmul se oprește dacă este îndeplinită una din următoarele condiții:
($\chi 1$) toți vectorii eroare sunt nuli, adică:

$$\mathbf{e}_i \equiv \mathbf{0}, \quad i = 1, \dots, N; \quad (3.2.10)$$

sau

($\chi 2$) a fost atins numărul maxim de iterații.

La oprire, algoritmul va furniza valorile parametrilor rețelei rezultate din antrenare, care vor fi notate prin $\mathbf{W}_f \in \mathbb{R}^{P \times n}$ (ponderi) și $\mathbf{b}_f \in \mathbb{R}^P$ (deplasări).

Dacă nici una din condițiile ($\chi 1$) sau ($\chi 2$) nu este îndeplinită, se continuă cu etapa următoare a iterației curente.

Etapă 3 (Etapă de actualizare a parametrilor)

Se construiește matricea erorilor:

$$\mathbf{E} = [\mathbf{e}_1 \dots \mathbf{e}_N] = \mathbf{Z} - \mathbf{Y} \in \mathbb{R}^{P \times N}, \quad (3.2.11)$$

cu ajutorul căreia se calculează matricele de actualizare a parametrilor:

- pentru ponderi:

$$\mathbf{D}_W = \mathbf{E} \cdot \mathbf{X}^T \in \mathbb{R}^{P \times n}, \quad (3.2.12-a)$$

- pentru deplasări:

$$\mathbf{D}_b = \mathbf{E} \cdot \underbrace{[1 \dots 1]}_{N \text{ elemente}}^T \in \mathbb{R}^P. \quad (3.2.12-b)$$

Se calculează noile valori ale parametrilor:

- pentru ponderi:

$$\mathbf{W}_{\text{new}} = \mathbf{W}_{\text{old}} + \mathbf{D}_W, \quad (3.2.13-a)$$

- pentru deplasări:

$$\mathbf{b}_{\text{new}} = \mathbf{b}_{\text{old}} + \mathbf{D}_b. \quad (3.2.13-b)$$

Cu aceste calcule se încheie iterația curentă și se trece la o nouă iterație efectuând atribuirile $\mathbf{W}_{\text{new}} \rightarrow \mathbf{W}_{\text{old}}$, $\mathbf{b}_{\text{new}} \rightarrow \mathbf{b}_{\text{old}}$.

Observații

1° Inițializarea parametrilor rețelei (adică $\mathbf{W}_{\text{old}}$ și $\mathbf{b}_{\text{old}}$ la prima iterație a algoritmului) se face cu valori arbitrare.

2° Condiția ($\chi 1$) din Etapa 2 (anularea tuturor vectorilor eroare) poate fi exprimată unitar, prin intermediul erorii pătratice globale (eng. *sum squared error*) $J_{SSE} = \sum_{i=1}^N \mathbf{e}_i^T \mathbf{e}_i$, a

erorii pătratice medii (eng. *mean squared error*) $J_{MSE} = \frac{1}{N} \sum_{i=1}^N \mathbf{e}_i^T \mathbf{e}_i$, sau a erorii medii absolute (eng. *mean absolute error*) $J_{MAE} = \frac{1}{N} \sum_{i=1}^N \frac{1}{p} \sum_{k=1}^p |\mathbf{e}_{i,k}|$, (unde $\mathbf{e}_{i,k}$ notează componenta k a vectorului eroare $\mathbf{e}_i = [\mathbf{e}_{i,1} \dots \mathbf{e}_{i,p}]^T \in \mathbb{R}^p$) sub forma:

$$J_{SSE} = 0, \quad (3.2.14)$$

$$J_{MSE} = 0, \quad (3.2.14')$$

respectiv:

$$J_{MAE} = 0, \quad (3.2.14'')$$

care sunt echivalente cu (3.2.10). Datorită simplității, în programare se preferă utilizarea unui test de forma (3.2.14), (3.2.14') sau (3.2.14''), față de cele N teste formulate în (3.2.10).

3° Este demonstrat faptul că vectorii $\mathbf{x}_i$, $i=1, \dots, N$, pot fi grupați în clasele C_j , separabile liniar, conform (1-1), (1-2), ... , (1- M), *dacă și numai dacă* algoritmul de antrenare va conduce la satisfacerea condiției ($\chi 1$) într-un număr finit de iterații. Din acest motiv, la stabilirea numărului maxim de iterații, utilizatorul trebuie să aibă în vedere o valoare acoperitoare pentru satisfacerea condiției ($\chi 1$), evident, sub rezerva că ($\chi 1$) poate fi efectiv îndeplinită.

4° În cazul când condiția ($\chi 1$) este satisfăcută, valorile $\mathbf{W}_f \in \mathbb{R}^{p \times n}$ și $\mathbf{b}_f \in \mathbb{R}^p$ furnizate în finalul algoritmului de antrenare definesc ecuațiile celor p hiperplane (care separă cele $M = 2^p$ clase) prin intermediul celor p componente scalare ale egalității:

$$\mathbf{W}_f \cdot \mathbf{x} + \mathbf{b}_f = \mathbf{0}. \quad (3.2.15)$$

5° Dacă condiția ($\chi 1$) nu poate fi îndeplinită într-un număr suficient de mare de iterații ale algoritmului de antrenare, rezultă că vectorii $\mathbf{x}_i$, $i=1, \dots, N$, nu pot fi grupați în clasele C_j separabile liniar, conform (1-1), (1-2), ... , (1- M). Este evident că în acest caz eroarea pătratică globală *nu se va anula*. ■

6° Rezultatul antrenării este independent de ordinea în care sunt considerați vectorii prototip. În secțiunea 3.2.2 a fost considerată ipoteza că vectorii prototip sunt ordonați în funcție de clasa din care fac parte numai datorită simplificării expunerii. ■

Exemplul 3.2.1. Ilustrăm aplicarea algoritmului de antrenare prezentat anterior în cazul determinării parametrilor unui perceptron ($p=1$) cu două intrări ($n=2$) ce implementează funcția logică OR (figura 3.1.3). Pentru aceasta considerăm vectorii prototip

$$\mathbf{x}_1 = \begin{bmatrix} 0 \\ 0 \end{bmatrix}, \mathbf{x}_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, \mathbf{x}_3 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \mathbf{x}_4 = \begin{bmatrix} 1 \\ 1 \end{bmatrix},$$

cu care formăm matricea

$$\mathbf{X} = \begin{bmatrix} 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 1 \end{bmatrix}.$$

Matricea țintelor este

$$\mathbf{Z} = \begin{bmatrix} 0 & 1 & 1 & 1 \end{bmatrix}.$$

Inițializăm parametrii perceptronului cu valori nule:

$$\mathbf{W}_{\text{old}} = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{b}_{\text{old}} = 0.$$

Iterația 1: Se calculează:

$$\mathbf{Y} = \begin{bmatrix} 1 & 1 & 1 & 1 \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} -1 & 0 & 0 & 0 \end{bmatrix}, \\ \mathbf{D}_W = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{D}_b = -1 \Rightarrow \mathbf{W}_{\text{new}} = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{b}_{\text{new}} = -1.$$

Iterația 2:

$$\mathbf{W}_{\text{old}} = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{b}_{\text{old}} = -1 \\ \mathbf{Y} = \begin{bmatrix} 0 & 0 & 0 & 0 \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} 0 & 1 & 1 & 1 \end{bmatrix}, \\ \mathbf{D}_W = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{D}_b = 3 \Rightarrow \mathbf{W}_{\text{new}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{new}} = 2.$$

Iterația 3:

$$\mathbf{W}_{\text{old}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{old}} = 2 \\ \mathbf{Y} = \begin{bmatrix} 1 & 1 & 1 & 1 \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} -1 & 0 & 0 & 0 \end{bmatrix} \\ \mathbf{D}_W = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{D}_b = -1 \Rightarrow \mathbf{W}_{\text{new}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{new}} = 1.$$

Iterația 4

$$\mathbf{W}_{\text{old}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{old}} = 1 \\ \mathbf{Y} = \begin{bmatrix} 1 & 1 & 1 & 1 \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} -1 & 0 & 0 & 0 \end{bmatrix}, \\ \mathbf{D}_W = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{D}_b = -1 \Rightarrow \mathbf{W}_{\text{new}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{new}} = 0.$$

Iterația 5:

$$\mathbf{W}_{\text{old}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{old}} = 0 \\ \mathbf{Y} = \begin{bmatrix} 1 & 1 & 1 & 1 \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} -1 & 0 & 0 & 0 \end{bmatrix}, \\ \mathbf{D}_W = \begin{bmatrix} 0 & 0 \end{bmatrix}, \quad \mathbf{D}_b = -1 \Rightarrow \mathbf{W}_{\text{new}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{new}} = -1.$$

Iterația 6:

$$\mathbf{W}_{\text{old}} = \begin{bmatrix} 2 & 2 \end{bmatrix}, \quad \mathbf{b}_{\text{old}} = -1 \\ \mathbf{Y} = \begin{bmatrix} 0 & 1 & 1 & 1 \end{bmatrix}, \quad \mathbf{E} = \begin{bmatrix} 0 & 0 & 0 & 0 \end{bmatrix} \Rightarrow \text{STOP}$$

Condiția χ_1 este satisfăcută. Algoritmul se încheie furnizând valorile parametrilor rețelei rezultate din antrenare, $\mathbf{W}_f = \begin{bmatrix} 2 & 2 \end{bmatrix}$ și $\mathbf{b} = -1$. Ecuația dreptei de separație este $d(\mathbf{x}) = 2x_1 + 2x_2 - 1 = 0$, fiind chiar cea reprezentată grafic în figura 3.1.3(b). ■

3.3. Utilizarea rețelelor cu mai multe straturi

Este demonstrat faptul că rețelele cu mai multe straturi permit realizarea clasificării în condițiile când clasele *nu constituie* regiuni separabile liniar din $\mathbb{R}^n$. În aceste situații, complexitatea rețelei (ca și număr de straturi) este dependentă de proprietățile mulțimilor din $\mathbb{R}^n$ utilizate pentru clasificare (de exemplu, satisfacerea condiției de convexitate). Procedeele de antrenare constituie generalizări ale algoritmului discutat în cazul rețelelor cu un singur strat.

3.4. Facilități software oferite de *Neural Network Toolbox*

În pachetul *Neural Network Toolbox* există funcția `newp` crează o rețea perceptron cu un singur strat. După ce a fost creată, o rețea poate fi antrenată utilizând funcția `train` care implementează algoritmul de antrenare a unei rețele perceptron conform parametrilor de antrenare aleși după creare. Pentru folosirea algoritmului de antrenare pe lot se alege rutina `'trainb'` drept funcție de antrenare a rețelei (`net.trainFcn = 'trainb'`). Valoarea implicită este `net.trainFcn = 'trainc'`.

Pentru simularea comportării unei rețele neurale create anterior se poate utiliza funcția `sim`.

Pentru reprezentarea grafică a vectorilor de intrare bi- și tridimensionali se poate apela funcția `plotpv`. În scopul atașării pe acest grafic a hiperplanelor de separare dintre clase se poate apela funcția `plotpc` (după un apel prealabil al funcției `plotpv`). Modul de apel al funcțiilor descrise mai sus poate fi studiat în urma consultării help-ului aferent.

Pentru aprofundarea noțiunilor prezentate în acest capitol se recomanda studierea problemelor cu soluții analitice 7.3 – 7.5 și a problemelor cu soluții numerice 8.5 – 8.7.

Capitolul 4

APLICAȚII ALE REȚELELOR FEEDFORWARD CU FUNCȚII DE ACTIVARE LINIARE ÎN PROBLEME DE APROXIMARE

Rețelele neuronale feed-forward cu un singur strat cu funcții de activare liniare pot fi utilizate pentru a realiza aproximarea liniară a unor funcții definite cu ajutorul unei mulțimi de vectori prototip și respectiv a unei mulțimi de vectori țintă. Arhitectura rețelei (adică numărul de neuroni, precum și numărul de intrări) este impusă de dimensiunea vectorilor țintă și respectiv dimensiunea vectorilor prototip.

Utilizarea în problemele de aproximare liniară a rețelelor neuronale feed-forward cu funcții de activare liniare se bazează pe antrenarea rețelelor prin algoritmul Widrow-Hoff. Aceste rețele mai poartă denumirea de ADALINE (ADaptive LINear Element), denumire propusă de Widrow și Hoff pentru că algoritmul de antrenare are efectul adaptării (ADA) valorii parametrilor rețelei cu funcții liniare (LIN) în scopul aproximării.

4.1. Transferul liniar intrare-ieșire și proprietatea de aproximare

Pentru o rețea feedforward cu un singur strat cu n intrări, p neuroni cu funcții de activare liniare, transferul intrare-ieșire este descris prin dependența liniară:

$$\mathbf{y} = \mathbf{f}(\mathbf{x}) = \mathbf{W}\mathbf{x} + \mathbf{b}, \quad \mathbf{W} \in \mathbb{R}^{p \times n}, \quad \mathbf{b} \in \mathbb{R}^p, \quad (4.1.1)$$

unde $\mathbf{x} \in \mathbb{R}^n$ notează vectorul de intrare iar $\mathbf{y} \in \mathbb{R}^p$ notează vectorul de ieșire. Acest transfer este parametrizat de ponderile $\mathbf{W} \in \mathbb{R}^{p \times n}$ și deplasările $\mathbf{b} \in \mathbb{R}^p$.

Se consideră o mulțime de N vectori prototip:

$$\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N \in \mathbb{R}^n, \quad (4.1.2)$$

și o mulțime de N vectori țintă:

$$\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_N \in \mathbb{R}^p, \quad (4.1.3)$$

care corespund vectorilor prototip în sensul că există o aplicație:

$$\varphi: \mathbb{R}^n \rightarrow \mathbb{R}^p, \quad (4.1.4)$$

care asigură:

$$\varphi(\mathbf{x}_i) = \mathbf{z}_i, \quad i = 1, \dots, N. \quad (4.1.5)$$

Aplicând funcția (4.1.1) pentru cei N vectori prototip $\mathbf{x}_i$, putem scrie:

$$\begin{aligned} \mathbf{y}_1 &= \mathbf{W}\mathbf{x}_1 + \mathbf{b} = [\mathbf{W} \ \mathbf{b}] \begin{bmatrix} \mathbf{x}_1 \\ 1 \end{bmatrix} = \mathbf{\Theta} \underline{\mathbf{x}}_1, \\ &\vdots \\ \mathbf{y}_N &= \mathbf{W}\mathbf{x}_N + \mathbf{b} = [\mathbf{W} \ \mathbf{b}] \begin{bmatrix} \mathbf{x}_N \\ 1 \end{bmatrix} = \mathbf{\Theta} \underline{\mathbf{x}}_N, \end{aligned} \quad (4.1.6)$$

unde s-au folosit notațiile:

$$\mathbf{\Theta} = [\mathbf{W} \ \mathbf{b}] \in \mathbb{R}^{p \times (n+1)}, \quad (4.1.7)$$

$$\underline{\mathbf{x}}_i = \begin{bmatrix} \mathbf{x}_i \\ 1 \end{bmatrix} \in \mathbb{R}^{n+1}, \quad i = 1, \dots, N. \quad (4.1.8)$$

Partajând pe linii matricea parametrilor $\mathbf{\Theta} \in \mathbb{R}^{p \times (n+1)}$ sub forma:

$$\mathbf{\Theta} = \begin{bmatrix} \boldsymbol{\theta}_1^T \\ \vdots \\ \boldsymbol{\theta}_p^T \end{bmatrix}, \quad \boldsymbol{\theta}_k \in \mathbb{R}^{n+1}, \quad k = 1, \dots, p, \quad (4.1.9)$$

se constată că relația ce corespunde unui indice oarecare i în (4.1.6) poate fi rescrisă astfel:

$$\mathbf{y}_i = \begin{bmatrix} \boldsymbol{\theta}_1^T \underline{\mathbf{x}}_i \\ \vdots \\ \boldsymbol{\theta}_p^T \underline{\mathbf{x}}_i \end{bmatrix} = \begin{bmatrix} \underline{\mathbf{x}}_i^T \boldsymbol{\theta}_1 \\ \vdots \\ \underline{\mathbf{x}}_i^T \boldsymbol{\theta}_p \end{bmatrix} = \mathbf{A}_i \mathbf{v}, \quad i = 1, \dots, N. \quad (4.1.10)$$

În relația (4.1.10) matricea $\mathbf{A}_i$ este definită prin:

$$\mathbf{A}_i = \begin{bmatrix} \underline{\mathbf{x}}_i^T & \mathbf{O}_{1 \times (n+1)} & \cdots & \mathbf{O}_{1 \times (n+1)} \\ \mathbf{O}_{1 \times (n+1)} & \underline{\mathbf{x}}_i^T & \cdots & \mathbf{O}_{1 \times (n+1)} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{O}_{1 \times (n+1)} & \mathbf{O}_{1 \times (n+1)} & \cdots & \underline{\mathbf{x}}_i^T \end{bmatrix} \in \mathbb{R}^{p \times p(n+1)}, \quad i = 1, \dots, N, \quad (4.1.11)$$

iar $\mathbf{v} \in \mathbb{R}^{p(n+1)}$ este vectorul tuturor parametrilor rețelei:

$$\mathbf{v} = \begin{bmatrix} \boldsymbol{\theta}_1 \\ \vdots \\ \boldsymbol{\theta}_p \end{bmatrix} \in \mathbb{R}^{p(n+1)}. \quad (4.1.12)$$

Într-o scriere globală, cele N relații (4.1.10) (toate conținând același vector al parametrilor rețelei, $\mathbf{v} \in \mathbb{R}^{p(n+1)}$) conduc la relația vectorială (compusă din pN egalități scalare):

$$\begin{bmatrix} \mathbf{y}_1 \\ \vdots \\ \mathbf{y}_N \end{bmatrix} = \begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} \mathbf{v}. \quad (4.1.13)$$

Se observă imediat că (4.1.13) descrie comportarea rețelei (pentru valori arbitrare ale parametrilor rețelei) în cazul când la intrare se prezintă vectorii prototip $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$.

Dacă în (4.1.13), vectorii $\mathbf{y}_i \in \mathbb{R}^p$, $i = 1, \dots, N$, se înlocuiesc cu vectorii țintă $\mathbf{z}_i \in \mathbb{R}^p$, $i = 1, \dots, N$, și dacă se schimbă membrul stâng cu cel drept, se obține relația:

$$\begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} \mathbf{v} = \begin{bmatrix} \mathbf{z}_1 \\ \vdots \\ \mathbf{z}_N \end{bmatrix}. \quad (4.1.14)$$

Această relație poate fi privită ca un sistem algebric liniar de pN ecuații cu $p(n+1)$ necunoscute (necunoscutele fiind parametrii rețelei, grupați în vectorul $\mathbf{v} \in \mathbb{R}^{p(n+1)}$). În sistemul respectiv, matricea coeficienților se construiește din elementele vectorilor prototip $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$ (conform celor N egalități (4.1.11)), iar vectorul termenilor liberi se construiește din elementele vectorilor țintă $\mathbf{z}_i \in \mathbb{R}^p$, $i = 1, \dots, N$.

După cum este prezentat în secțiunea A.1.2 din Appendix, sistemul (4.1.14) posedă întotdeauna o soluție în sensul celor mai mici pătrate (unică sau nu), care se va nota prin $\mathbf{v}^* \in \mathbb{R}^{p(n+1)}$.

Pentru un set dat de vectori prototip, $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$, relația (4.1.10) poate fi privită ca exprimând dependența ieșirii rețelei neurale de parametrii acesteia, $\mathbf{y}_i = \mathbf{y}_i(\mathbf{v})$, $i = 1, \dots, N$. Utilizând vectorii de eroare:

$$\mathbf{e}_i(\mathbf{v}) = \mathbf{z}_i - \mathbf{y}_i(\mathbf{v}) \in \mathbb{R}^p, \quad i = 1, \dots, N, \quad (4.1.15)$$

se calculează eroarea pătratică globală (*Sum Squared Error – SSE*):

$$J_{SSE}(\mathbf{v}) = \sum_{i=1}^N [\mathbf{e}_i(\mathbf{v})]^T \mathbf{e}_i(\mathbf{v}), \quad (4.1.16)$$

și eroarea pătratică medie (*Mean Squared Error – MSE*):

$$J_{MSE}(\mathbf{v}) = \frac{1}{N} \sum_{i=1}^N [\mathbf{e}_i(\mathbf{v})]^T \mathbf{e}_i(\mathbf{v}). \quad (4.1.16')$$

Valoarea parametrilor rețelei definiți prin cele $p(n+1)$ elemente ale vectorului $\mathbf{v}^* \in \mathbb{R}^{p(n+1)}$ asigură minimizarea MSE (echivalentă cu minimizarea SSE) pe întreg spațiul $\mathbb{R}^{p(n+1)}$,

$$0 \leq J(\mathbf{v}^*) \leq J(\mathbf{v}), \quad \forall \mathbf{v} \in \mathbb{R}^{p(n+1)}, \quad (4.1.17)$$

(unde J notează una între funcțiile J_{SSE} și J_{MSE}) fiind notat: $\mathbf{v}^* = \arg \min_{\mathbf{v} \in \mathbb{R}^{p(n+1)}} J(\mathbf{v})$.

Vectorului $\mathbf{v}^* \in \mathbb{R}^{p(n+1)}$ îi corespund matricele $\mathbf{W}^* \in \mathbb{R}^{p \times n}$ și $\mathbf{b}^* \in \mathbb{R}^p$ conținând valorile parametrilor rețelei (vezi relațiile (4.1.7), (4.1.9) și (4.1.12)). Astfel, pentru aceste valori $\mathbf{W}^*$ și $\mathbf{b}^*$, rețeaua va furniza prin funcția f din (4.1.1) o aproximare liniară, în sensul celor mai mici pătrate, a funcției ϕ (4.1.5) (care definește corespondența dintre vectorii prototip și țintă).

Observație:

Cu notațiile unde $\tilde{\mathbf{A}} = [\mathbf{A}_1^T \dots \mathbf{A}_N^T]^T$ și $\tilde{\mathbf{z}}^T = [\mathbf{z}_1^T \dots \mathbf{z}_N^T]^T$, sistemul de ecuații algebrice (4.1.14) poate fi adus la forma $\tilde{\mathbf{A}}\mathbf{v} = \tilde{\mathbf{z}}$. Rezolvarea acestui sistem comportă următoarea discuție, care decurge din Teorema Kronecker-Capelli:

1° Dacă:

$$\text{rang} \begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} = \text{rang} \left[\begin{array}{c|c} \mathbf{A}_1 & \mathbf{z}_1 \\ \vdots & \vdots \\ \mathbf{A}_N & \mathbf{z}_N \end{array} \right] < \min(pN, p(n+1)), \quad (4.1.18)$$

atunci sistemul (4.1.14) posedă o *infinitate de soluții exacte* $\mathbf{v}^* \in \mathbb{R}^{p(n+1)}$, pentru care vectorii eroare $\mathbf{e}_i$, $i = 1, \dots, N$, definiți prin (4.1.15) se anulează și implicit, rezultă:

$$J(\mathbf{v}^*) = 0. \quad (4.1.19)$$

2° Dacă:

$$\text{rang} \begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} = \text{rang} \left[\begin{array}{c|c} \mathbf{A}_1 & \mathbf{z}_1 \\ \vdots & \vdots \\ \mathbf{A}_N & \mathbf{z}_N \end{array} \right] = \min(pN, p(n+1)), \quad (4.1.20)$$

atunci sistemul (4.1.14) posedă o *singură soluție exactă* $\mathbf{v}^* = \tilde{\mathbf{A}}^{-1}\tilde{\mathbf{z}}$, pentru care vectorii eroare $\mathbf{e}_i$, $i = 1, \dots, N$, definiți prin (4.1.15) se anulează și, implicit, rezultă satisfacerea egalității (4.1.19).

3° Dacă:

$$\text{rang} \begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} < \text{rang} \left[\begin{array}{c|c} \mathbf{A}_1 & \mathbf{z}_1 \\ \vdots & \vdots \\ \mathbf{A}_N & \mathbf{z}_N \end{array} \right], \quad (4.1.21)$$

atunci sistemul (14) posedă numai *pseudo-soluții* $\mathbf{v}^*$, sau *soluții în sensul celor mai mici pătrate* conform (17), pentru care avem inegalitatea strictă:

$$J(\mathbf{v}^*) > 0. \quad (4.1.22)$$

Se disting următoarele două subcazuri:

a) Dacă

$$\text{rang} \begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} < \min(pN, p(n+1)), \quad (4.1.23)$$

atunci există o *infinitate de pseudo-soluții* satisfăcând (4.1.17). Dintre acestea, soluția

$$\mathbf{v}^* = (\tilde{\mathbf{A}}^T \tilde{\mathbf{A}})^{-1} \tilde{\mathbf{A}}^T \tilde{\mathbf{z}}, \quad (4.1.24)$$

este are norma minimă.

b) Dacă

$$\text{rang} \begin{bmatrix} \mathbf{A}_1 \\ \vdots \\ \mathbf{A}_N \end{bmatrix} = \min(pN, p(n+1)), \quad (4.1.25)$$

atunci există o *singură pseudo-soluție* $\mathbf{v}^*$ satisfăcând (4.1.17), și anume cea dată de (4.1.24).

Această discuție formulată cu privire la vectorul $\mathbf{v}^* \in \mathbb{R}^{p(n+1)}$, se transpune direct în termenii parametrilor rețelei $\mathbf{W}^* \in \mathbb{R}^{p \times n}$ (ponderi) și $\mathbf{b}^* \in \mathbb{R}^p$ (deplasări), parametri care constituie elementele vectorului $\mathbf{v}^* \in \mathbb{R}^{p(n+1)}$. ■

4.2. Utilizarea rețelelor cu un singur strat

4.2.1. Obiectivul antrenării rețelei

Pornind de la *mulțimea vectorilor prototip* (4.1.2) și *mulțimea vectorilor țintă* (4.1.3), prin antrenarea rețelei trebuie să se determine parametrii acesteia $\mathbf{W}^* \in \mathbb{R}^{p \times n}$ și $\mathbf{b}^* \in \mathbb{R}^p$ care să asigure *minimizarea erorii pătratice globale*, adică satisfacerea inegalității (16). Funcția f definită prin (4.1.1) pentru $\mathbf{W} = \mathbf{W}^*$ și $\mathbf{b} = \mathbf{b}^*$ realizează o *aproximare liniară în sensul celor mai mici pătrate* a funcției φ definită conform (4.1.5).

4.2.2. Algoritmul de antrenare Widrow-Hoff

În baza celor discutate în paragraful 4.1, parametrii rețelei $\mathbf{W}^* \in \mathbb{R}^{p \times n}$ și $\mathbf{b}^* \in \mathbb{R}^p$ pot fi determinați prin rezolvarea sistemului algebric (4.1.14) cu pN ecuații și $p(n+1)$ necunoscute. Totuși, din punct de vedere numeric, acest procedeu nu este recomandabil întrucât acuratețea soluției depinde de condiționarea numerică a sistemului (4.1.14). Din acest motiv, se preferă pentru antrenare, *algoritmul iterativ Widrow-Hoff* (denumit și *regula delta*), care se bazează pe *minimizarea erorii pătratice globale prin metoda gradientului*.

Considerând vectorii eroare $\mathbf{e}_i$, $i = 1, \dots, N$, definiți prin (4.1.15), ca funcții de parametri rețelei (care sunt conținuți în vectorul $\mathbf{v} \in \mathbb{R}^{p(n+1)}$) se formează funcția obiectiv MSE $J(\mathbf{v})$ dată de (4.1.16').

Minimizarea prin metoda gradientului a lui $J(\mathbf{v})$ (care înseamnă satisfacerea condiției (4.1.17)) constă în modificarea vectorului $\mathbf{v} \in \mathbb{R}^{p(n+1)}$ la fiecare iterație, în conformitate cu formula:

$$\mathbf{v}_{\text{new}} = \mathbf{v}_{\text{old}} - \eta \nabla J(\mathbf{v}_{\text{old}}), \quad (4.2.1)$$

unde $\eta > 0$ notează *pasul metodei de optimizare*, iar $\nabla J = [\partial J / \partial \mathbf{v}]^T$ notează vectorul gradient, normal la suprafața $J(\mathbf{v})$ definită prin (4.1.16').

Se observă că $J(\mathbf{v})$ este o formă pătratică definită pe spațiul parametrilor, $\mathbb{R}^{p(n+1)}$:

$$\begin{aligned} J(\mathbf{v}) &= \frac{1}{N} \sum_{i=1}^N [\mathbf{e}_i(\mathbf{v})]^T \mathbf{e}_i(\mathbf{v}) = \frac{1}{N} \sum_{i=1}^N [\mathbf{z}_i - \mathbf{A}_i \mathbf{v}]^T [\mathbf{z}_i - \mathbf{A}_i \mathbf{v}] = \\ &= \frac{1}{N} \sum_{i=1}^N [\mathbf{z}_i^T \mathbf{z}_i - 2\mathbf{z}_i^T \mathbf{A}_i \mathbf{v} + \mathbf{v}^T \mathbf{A}_i^T \mathbf{A}_i \mathbf{v}] = c + \mathbf{d}^T \mathbf{v} + \frac{1}{2} \mathbf{v}^T \mathbf{R} \mathbf{v}, \end{aligned} \quad (4.2.2)$$

unde: $\mathbf{R} = 2 \frac{1}{N} \sum_{i=1}^N \mathbf{A}_i^T \mathbf{A}_i = \frac{2}{N} \tilde{\mathbf{A}}^T \tilde{\mathbf{A}} \in \mathbb{R}^{p(n+1) \times p(n+1)}$, $\mathbf{d} = -2 \frac{1}{N} \sum_{i=1}^N \mathbf{A}_i^T \mathbf{z}_i = -\frac{2}{N} \tilde{\mathbf{A}}^T \tilde{\mathbf{z}} \in \mathbb{R}^{p(n+1)}$,

$c = \frac{1}{N} \sum_{i=1}^N \mathbf{z}_i^T \mathbf{z}_i = \frac{1}{N} \tilde{\mathbf{z}}^T \tilde{\mathbf{z}} \in \mathbb{R}$, astfel că $\nabla J(\mathbf{v}) = \mathbf{d} + \mathbf{R} \mathbf{v}$ și $D^2 J(\mathbf{v}) = \mathbf{R}$ (matricea Hessian).

Printr-o exprimare adecvată a vectorului gradient $\nabla J(\mathbf{v})$, relația (4.2.1) poate fi formulată *direct în termenii parametrilor rețelei* $\mathbf{W} \in \mathbb{R}^{p \times n}$ (ponderi) și $\mathbf{b} \in \mathbb{R}^p$ (deplasări), furnizând, astfel, *mecanismul de efectuare a unei iterații* pentru *algoritmul de antrenare Widrow-Hoff*.

Pentru algoritmul de antrenare se stabilesc de către utilizator un *număr maxim de iterații* sau *epoci*, precum și o valoare ε care se dorește a fi atinsă în minimizarea funcției obiectiv (4.1.16') (adică un *prag pentru eroarea pătratică globală*). Atragem atenția asupra faptului că funcția obiectiv $J(\mathbf{v})$ definită prin (4.1.16') satisface condiția (4.1.19), dacă și numai dacă sistemul de ecuații algebrice liniare (4.1.14) posedă o soluție exactă (sau soluții exacte). În cazul când relația (4.1.19) nu este satisfăcută, atingerea unui prag *oarecare* de eroare ε impus de utilizator *nu este întotdeauna posibilă* (întrucât ne situăm în situația caracterizată prin (4.1.22)). Facem observația că utilizând funcția de tip MSE, valoarea de prag poate fi aleasă independent de numărul de eșantioane cu care se lucrează.

De asemenea, se fixează o valoare pentru pasul metodei de optimizare (η), care, în contextul problemei de antrenare, este referită sub denumirea de *rată de învățare* (eng. *learning rate*).

Se organizează vectorii prototip $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$, și cei țintă $\mathbf{z}_i \in \mathbb{R}^p$, $i = 1, \dots, N$, sub forma matricelor:

$$\mathbf{X} = [\mathbf{x}_1 \dots \mathbf{x}_N] \in \mathbb{R}^{n \times N}, \quad (4.2.3)$$

respectiv:

$$\mathbf{Z} = [\mathbf{z}_1 \dots \mathbf{z}_N] \in \mathbb{R}^{p \times N}. \quad (4.2.4)$$

La fiecare iterație a algoritmului Widrow-Hoff se parcurg etapele descrise mai jos.

Valorile parametrilor rețelei la începutul unei iterații sunt notate prin $\mathbf{W}_{\text{old}} \in \mathbb{R}^{p \times n}$ (pentru ponderi) și $\mathbf{b}_{\text{old}} \in \mathbb{R}^p$ (pentru deplasări).

Etapa I (Etapa de prezentare)

Pentru valorile curente ale parametrilor rețelei ($\mathbf{W}_{\text{old}} \in \mathbb{R}^{p \times n}$, $\mathbf{b}_{\text{old}} \in \mathbb{R}^p$) se calculează ieșirile rețelei:

$$\mathbf{y}_i = \mathbf{f}(\mathbf{x}_i) = \mathbf{W}_{\text{old}} \mathbf{x}_i + \mathbf{b}_{\text{old}}, \quad \mathbf{y}_i \in \mathbb{R}^p, \quad i = 1, \dots, N, \quad (4.2.5)$$

care se organizează sub forma matricei:

$$\mathbf{Y} = [\mathbf{y}_1 \dots \mathbf{y}_N] = \mathbf{W}_{\text{old}} \mathbf{X} + \mathbf{b}_{\text{old}} \in \mathbb{R}^{p \times N}. \quad (4.2.6)$$

Etapa II (Etapa de verificare)

Se calculează vectorii eroare $\mathbf{e}_i = \mathbf{e}_i(\mathbf{W}_{\text{old}}, \mathbf{b}_{\text{old}})$, $i = 1, \dots, N$, definiți prin (4.1.15) ca diferența între vectorii țintă $\mathbf{z}_i \in \mathbb{R}^p$, $i = 1, \dots, N$, și vectorii ieșire la iterația curentă $\mathbf{y}_i = \mathbf{y}_i \in \mathbb{R}^p$, $i = 1, \dots, N$.

Algoritmul se oprește dacă este îndeplinită una din următoarele condiții:

($\chi 1$) eroarea pătratică globală este inferioară valorii de prag (eng. *threshold*) ε stabilite la startarea algoritmului, adică:

$$J(\mathbf{W}_{\text{old}}, \mathbf{b}_{\text{old}}) = \frac{1}{N} \sum_{i=1}^N \mathbf{e}_i^T \mathbf{e}_i < \varepsilon; \quad (4.2.7)$$

sau:

($\chi 2$) a fost atins numărul maxim de iterații stabilit la startarea algoritmului.

La oprire, algoritmul va furniza valorile parametrilor rețelei rezultate din antrenare, care vor fi notați prin $\mathbf{W}_f \in \mathbb{R}^{p \times n}$ (ponderi) și $\mathbf{b}_f \in \mathbb{R}^p$ (deplasări).

Dacă nici una din condițiile ($\chi 1$) sau ($\chi 2$) nu este îndeplinită, se continuă cu etapa următoare a iterației curente.

Etapa III (Etapa de actualizare a parametrilor)

Se construiește matricea erorilor:

$$\mathbf{E} = [\mathbf{e}_1 \dots \mathbf{e}_N] = \mathbf{Z} - \mathbf{Y} \in \mathbb{R}^{p \times N}, \quad (4.2.8)$$

cu ajutorul căreia se calculează matricele de actualizare a parametrilor:

- pentru ponderi:

$$\mathbf{D}_W = \mathbf{E} \cdot \mathbf{X}^T \in \mathbb{R}^{p \times n}, \quad (4.2.9-a)$$

- pentru deplasări:

$$\mathbf{D}_b = \mathbf{E} \cdot \underbrace{[1 \dots 1]}_{N \text{ elemente}}^T \in \mathbb{R}^p. \quad (4.2.9-b)$$

Se calculează noile valori ale parametrilor:

- pentru *ponderi*:

$$\mathbf{W}_{\text{new}} = \mathbf{W}_{\text{old}} + \eta \mathbf{D}_W, \quad (4.2.10-a)$$

- pentru *deplasări*:

$$\mathbf{b}_{\text{new}} = \mathbf{b}_{\text{old}} + \eta \mathbf{D}_b. \quad (4.2.10-b)$$

Cu aceste calcule se încheie iterația curentă și se trece la o nouă iterație efectuând atribuirile: $\mathbf{W}_{\text{new}} \rightarrow \mathbf{W}_{\text{old}}$, $\mathbf{b}_{\text{new}} \rightarrow \mathbf{b}_{\text{old}}$.

Observații:

1° Inițializarea parametrilor rețelei (adică $\mathbf{W}_{\text{old}}$ și $\mathbf{b}_{\text{old}}$ la prima iterație a algoritmului) se face cu valori arbitrare. Pentru simplitate, inițializarea se poate face cu valori nule.

2° Este demonstrat faptul că algoritmul Widrow-Hoff converge, după caz, *către soluția unică* sau *către una din soluțiile* sistemului algebric de ecuații liniare (în accepțiunea discuției din paragraful 4.1). Convergența se realizează în condițiile specifice *metodei gradientului* (ca tehnică de optimizare), aplicată formei pătratice (4.2.2), fiind garantată pentru valori ale ratei de învățare mai mici decât valoarea critică:

$$\eta_c = \frac{2}{\lambda_{\max}(\mathbf{R})}, \quad (4.2.11)$$

unde $\lambda_{\max}(\mathbf{R})$ notează valoarea proprie maximă a matricei simetrice pozitiv definite $\mathbf{R}$.

Astfel, se pot formula următoarele observații:

- pentru η inadecvat mic, deplasarea către optim, în spațiul parametrilor, este foarte lentă,
- pentru η inadecvat mare, deplasarea către optim, în spațiul parametrilor, se realizează cu oscilații (eventual cu pierderea stabilității algoritmului).

3° În cazul când algoritmul se încheie prin satisfacerea condiției $(\chi 1)$, acuratețea parametrilor furnizați $\mathbf{W}_f \in \mathbb{R}^{p \times n}$ și $\mathbf{b}_f \in \mathbb{R}^p$ depinde de valoarea de prag ε stabilită pentru eroarea pătratică globală. Uzual, condiția $(\chi 1)$ devine operantă atunci când poate fi satisfăcută relația (4.1.18), adică atunci când valoarea minimă a lui $J(\mathbf{v})$ este 0.

4° Algoritmul Widrow-Hoff, privit ca metodă de căutare a minimului funcției $J(\mathbf{v})$ (4.1.16'), nu se confruntă cu problema eșuării în minime locale, deoarece $J(\mathbf{v})$ este o formă pătratică definită pe spațiul parametrilor $\mathbb{R}^{p(n+1)}$ care posedă numai un punct de minim global egal cu 0, sau o infinitate de astfel de puncte (în accepțiunea discuției din paragraful 4.1). ■

4.3. Facilități software oferite de Neural Network Toolbox

Funcția `newlind` permite determinarea parametrilor unei rețele ADALINE prin rezolvarea sistemului (4.1.14) în sensul celor mai mici pătrate. Forma de apel este

`net = newlind(P, T)`

unde P este matricea $R \times Q$ a vectorilor prototip, iar T este matricea $S \times Q$ a vectorilor țintă (cu notațiile din paragraful 4.2.2, P reprezintă matricea $X \in \mathbb{R}^{n \times N}$ iar T este matricea $Z \in \mathbb{R}^{p \times N}$). Funcția `newlind` returnează rețeaua ADALINE cu un singur strat cu S neuroni. Parametrii rețelei sunt calculați pe baza relației (4.1.24).

Funcția `newlin` poate fi folosită pentru crearea unei rețele ADALINE cu un singur strat. Forma de apel este

`net = newlin(PR, S, ID, LR)`

unde PR este matricea $R \times 2$ corespunzătoare valorilor minime și maxime ale celor R componente ale vectorilor prototip, S reprezintă numărul de neuroni al rețelei create, ID este vectorul corespunzător întârzierilor semnalelor de intrare (valoare implicită nulă) și LR reprezintă rata de învățare ce va fi utilizată în antrenarea rețelei prin aplicarea algoritmului Widrow-Hoff (având valoarea implicită 0.01). Funcția `newlin` returnează rețeaua ADALINE cu un singur strat cu S neuroni având toți parametrii nuli.

După ce a fost creată, o rețea poate fi antrenată utilizând funcția `train` care implementează algoritmul de antrenare a unei rețele perceptron conform parametrilor de antrenare aleși la crearea rețelei. Implicit, pentru actualizarea la fiecare iterație a parametrilor rețelei este apelată funcția `learnwh` care implementează relațiile (4.2.9)-(4.2.10) din algoritmul Widrow-Hoff. Criteriul de performanță utilizat este MSE. Numărul maxim de epoci de antrenare a rețelei și valoarea de prag pentru MSE sunt reprezentate de `net.trainParam.epochs` și, respectiv, `net.trainParam.goal`.

Funcția `maxlinlr`, având forma de apel

`lr = maxlinlr(P)`

poate fi folosită pentru determinarea ratei de învățare maxime care asigură stabilitatea algoritmului de antrenare Widrow-Hoff. Această funcție primește ca argument matricea P a vectorilor prototip (respectiv, cu notațiile din paragraful 4.2.2, matricea $X \in \mathbb{R}^{n \times N}$) și implementează în MATLAB relația (4.1.36) sub forma $\eta_c = 1 / \lambda_{\max}(XX^T)$ în cazul în care se dorește antrenarea unei rețele ADALINE cu deplasare nulă pentru toți neuronii, sau sub forma $\eta_c = 1 / \lambda_{\max}(X_2 X_2^T)$, cu $X_2 = \begin{bmatrix} X \\ \mathbf{1}_{1 \times N} \end{bmatrix}$ (unde $\mathbf{1}_{1 \times N}$ notează vectorul linie cu N elemente egale cu 1) dacă se dorește antrenarea unei rețele ADALINE cu deplasare nenulă.

Pentru simularea rețelei antrenate se poate utiliza funcția `sim`.

În cazul unui neuron cu o singură intrare și funcție de activare liniară, pentru reprezentarea grafică tridimensională a suprafeței erorii pătratice globale (ca funcție de pondere și deplasarea neuronului) se pot utiliza funcțiile `errsurf` și `plotes` (apelate în această ordine).

Pentru aprofundarea noțiunilor prezentate în acest capitol se recomandă studierea problemelor cu soluții analitice 7.6 – 7.8 și a problemelor cu soluții numerice 8.8 – 8.11.

Capitolul 5

APLICAȚII ALE REȚELELOR FEEDFORWARD MULTISTRAT ÎN PROBLEME DE APROXIMARE ȘI CLASIFICARE

Rețelele neuronale feed-forward multistrat cu noduri de intrare sigmoidale și noduri de ieșire liniare pot fi utilizate pentru a realiza aproximarea unor funcții neliniare, definite cu ajutorul unei mulțimi de vectori prototip și respectiv a unei mulțimi de vectori țintă. Numărul de intrări și ieșiri ai rețelei sunt impuse de dimensiunea vectorilor țintă și respectiv - dimensiunea vectorilor prototip. Numărul de neuroni din stratul de intrare (și din stratul ascuns - în cazul rețelelor cu trei straturi) este stabilit de utilizator. Aceste arhitecturi neuronale sunt denumite și rețele MLP (eng. *MultiLayer Perceptron*). Pentru studierea transferului intrare-ieșire pentru acest tip de rețele se recomandă utilizarea funcțiilor disponibile în pachetul **Neural Network Toolbox**. Un accent deosebit se va pune pe vizualizarea grafică a progresului antrenării și anume pe vizualizare grafică tridimensională a suprafeței de eroare și a căutării minimului în cazul unui neuron cu o singură intrare și funcție de activare sigmoidală, respectiv pe vizualizare grafică bidimensională a evoluției procesului de aproximare în cazul unei rețele cu două straturi.

5.1. Proprietatea de aproximare a rețelelor feed-forward cu două straturi

O rețea cu două straturi cu noduri sigmoidale în stratul de intrare și liniare în stratul de ieșire realizează transferul intrare-ieșire prin funcția:

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^p, \quad (5.1.1-a)$$

$$y = f(x) = \sigma^2(W^2 \sigma^1(W^1 x + b^1) + b^2), \quad (5.1.1-b)$$

unde notațiile au următoarea semnificație, în conformitate cu arhitectura reprezentată grafic în figura 5.1.1:

$\mathbf{x} \in \mathbb{R}^n$ – vectorul de intrare;

$\mathbf{y} \in \mathbb{R}^p$ – vectorul de ieșire;

$\mathbf{W}^1 \in \mathbb{R}^{m \times n}$ – matricea ponderilor stratului de intrare;

$\mathbf{b}^1 \in \mathbb{R}^m$ – vectorul deplasărilor stratului de intrare;

$\mathbf{W}^2 \in \mathbb{R}^{p \times m}$ – matricea ponderilor stratului de ieșire;

$\mathbf{b}^2 \in \mathbb{R}^p$ – vectorul deplasărilor stratului de ieșire;

$\sigma^1: \mathbb{R}^m \rightarrow \mathbb{R}^m$, $\sigma^1(\mathbf{u}^1) = [\sigma^1(u_1^1) \dots \sigma^1(u_m^1)]^T$ – funcția vectorială de activare a stratului de intrare, unde $\mathbf{u}^1 = \mathbf{W}^1 \mathbf{x} + \mathbf{b}^1 \stackrel{\text{not}}{=} [u_1^1 \dots u_m^1]^T$ și σ^1 are expresia analitică:

$\sigma^1(u) = 1/(1 + e^{-u})$ în cazul sigmoidului unipolar și, respectiv,

$\sigma^1(u) = (e^u - e^{-u})/(e^u + e^{-u})$ în cazul sigmoidului bipolar;

$\sigma^2: \mathbb{R}^p \rightarrow \mathbb{R}^p$, $\sigma^2(\mathbf{u}^2) = \mathbf{u}^2$ (echivalent cu $\sigma^2(\mathbf{u}^2) = [\sigma^2(u_1^2) \dots \sigma^2(u_p^2)]^T$ și

$\sigma^2(u) = u$) – funcția de activare (liniară) a stratului de ieșire, unde

$\mathbf{u}^2 = \mathbf{W}^2 \mathbf{x}^2 + \mathbf{b}^2 \stackrel{\text{not}}{=} [u_1^2 \dots u_p^2]^T \in \mathbb{R}^p$, $\mathbf{x}^2 \equiv \mathbf{y}^1 = \sigma^1(\mathbf{u}^1)$.

Transferul intrare-ieșire realizat este parametrizat de ponderile $\mathbf{W}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}^2 \in \mathbb{R}^{p \times m}$ și deplasările $\mathbf{b}^1 \in \mathbb{R}^m$, $\mathbf{b}^2 \in \mathbb{R}^p$ rețelei.

Fie funcția:

$$\varphi: \mathbb{R}^n \rightarrow \mathbb{R}^p, \quad (5.1.2)$$

netedă pe o mulțime compactă $S \subseteq \mathbb{R}^n$.

Este demonstrat faptul că, dacă rețeaua neuronală posedă un număr adecvat de noduri m în stratul de intrare, atunci parametrii rețelei pot fi determinați de așa manieră încât $\mathbf{f}$ definită prin (1) să aproximeze pe S funcția φ din (2) cu o precizie impusă (în sensul de „oricât de bună”). Într-o formulare matematică riguroasă, acest rezultat remarcabil poate fi scris sub forma:

$$\forall \varepsilon > 0, \exists m \in \mathbb{N} \text{ și } \mathbf{W}^1 \in \mathbb{R}^{m \times n}, \mathbf{b}^1 \in \mathbb{R}^m, \mathbf{W}^2 \in \mathbb{R}^{p \times m}, \mathbf{b}^2 \in \mathbb{R}^p \text{ a.î.} \quad (5.1.3)$$

$$\|\mathbf{f}(\mathbf{x}) - \varphi(\mathbf{x})\| < \varepsilon, \quad \forall \mathbf{x} \in S.$$

Observație:

Proprietatea de aproximare prezentată mai sus nu furnizează informații concrete privind alegerea numărului de noduri $m \in \mathbb{N}$ și determinarea parametrilor rețelei $\mathbf{W}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{b}^1 \in \mathbb{R}^m$, $\mathbf{W}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}^2 \in \mathbb{R}^p$, ci garantează numai existența acestor valori. ■

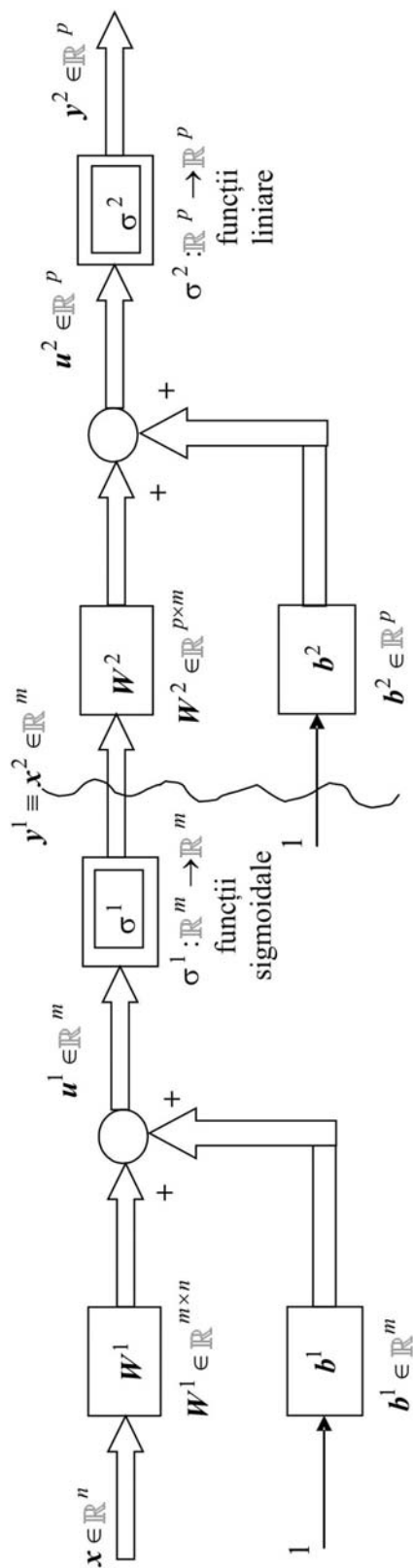


Figura 5.1.1. Schema bloc a unei rețele neuronale cu două straturi, utilizată în probleme de aproximare neliniară

Din punct de vedere practic, exploatarea rezultatului discutat anterior are loc în următorul context. Funcția φ din (5.1.2) este cunoscută într-un *număr finit* N de puncte $\mathbf{x}_i \in S \subseteq \mathbb{R}^n$, pentru care avem:

$$\varphi(\mathbf{x}_i) = \mathbf{z}_i, \quad i = 1, \dots, N. \quad (5.1.4)$$

La intrarea rețelei se prezintă *vectorii prototip*:

$$\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N \in \mathbb{R}^n, \quad (5.1.5)$$

pentru care se consideră *vectorii țintă*:

$$\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_N \in \mathbb{R}^p \quad (5.1.6)$$

definiți conform (5.1.4). Notând prin $\mathbf{y}_i, i = 1, \dots, N$, vectorii de ieșire ai MLP, adică:

$$\mathbf{y}_i = \mathbf{f}(\mathbf{x}_i) \in \mathbb{R}^p, \quad i = 1, \dots, N, \quad (5.1.7)$$

se definesc vectorii eroare prin:

$$\mathbf{e}_i = \mathbf{z}_i - \mathbf{y}_i \in \mathbb{R}^p, \quad i = 1, \dots, N. \quad (5.1.8)$$

Considerând un număr de neuroni m în stratul de intrare al rețelei, vectorii de ieșire și, implicit, vectorii de eroare, depind de parametrii rețelei $\mathbf{W}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{b}^1 \in \mathbb{R}^m$, $\mathbf{W}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}^2 \in \mathbb{R}^p$. Se definește funcția obiectiv MSE (eroarea pătratică medie):

$$J_{MSE}(\mathbf{W}^1, \mathbf{b}^1, \mathbf{W}^2, \mathbf{b}^2) = \frac{1}{N} \sum_{i=1}^N \mathbf{e}_i^T \mathbf{e}_i. \quad (5.1.9)$$

Pentru orice număr de neuroni m în stratul de intrare al rețelei, există valori ale parametrilor rețelei notate prin $\mathbf{W}^{1*} \in \mathbb{R}^{m \times n}$, $\mathbf{b}^{1*} \in \mathbb{R}^m$, $\mathbf{W}^{2*} \in \mathbb{R}^{p \times m}$, $\mathbf{b}^{2*} \in \mathbb{R}^p$, care asigură *minimizarea erorii pătratice medii*:

$$J_{MSE}(\mathbf{W}^{1*}, \mathbf{b}^{1*}, \mathbf{W}^{2*}, \mathbf{b}^{2*}) \leq J_{MSE}(\mathbf{W}^1, \mathbf{b}^1, \mathbf{W}^2, \mathbf{b}^2). \quad (5.1.10)$$

Observație:

Parametrii $\mathbf{W}^{1*} \in \mathbb{R}^{m \times n}$, $\mathbf{b}^{1*} \in \mathbb{R}^m$, $\mathbf{W}^{2*} \in \mathbb{R}^{p \times m}$, $\mathbf{b}^{2*} \in \mathbb{R}^p$ din (5.1.9) nu asigură și satisfacerea condiției (5.1.3) pe întreg compactul $S \subseteq \mathbb{R}^n$. În schimb, inegalitatea (5.1.9) posedă o semnificație numerică concretă, furnizând astfel un obiectiv consistent pentru antrenarea rețelei neuronale. ■

5.2. Utilizarea rețelelor feedforward cu două straturi

Utilizarea rețelelor neuronale feed-forward multistrat în problemele de aproximare neliniare se bazează pe antrenarea rețelelor prin algoritmul de propagare inversă. Denumirea algoritmului provine din faptul că pentru antrenare informațiile sunt procesate în sens invers, de la ieșirea rețelei, către intrare.

5.2.1. Obiectivul antrenării rețelei

Pornind de la mulțimea vectorilor prototip (5.1.5) și mulțimea vectorilor țintă (5.1.6), prin antrenarea unei rețele cu m noduri în stratul de intrare trebuie să se determine parametrii acesteia $\mathbf{W}^{1*} \in \mathbb{R}^{m \times n}$, $\mathbf{b}^{1*} \in \mathbb{R}^m$, $\mathbf{W}^{2*} \in \mathbb{R}^{p \times m}$, $\mathbf{b}^{2*} \in \mathbb{R}^p$ care să asigure minimizarea erorii pătratice globale, adică satisfacerea inegalității (5.1.10). Cu aceste valori ale parametrilor, funcția f definită prin (5.1.1) realizează o aproximare (neliniară) în sensul celor mai mici pătrate a funcției φ , cunoscută prin intermediul eșantioanelor (5.1.4). Valoarea minimă obținută în (5.1.9) depinde (în general) de numărul de noduri m din stratul de intrare al rețelei și deci m va influența calitatea aproximării realizate de rețea pentru funcția φ .

5.2.2. Algoritmul de antrenare prin propagare inversă

Algoritmul de antrenare prin *propagare inversă* (eng. *backpropagation*) se bazează pe minimizarea erorii pătratice medii (MSE) prin metoda gradientului.

Să notăm prin $\mathbf{v} \in \mathbb{R}^{m(n+1)+p(m+1)}$ vectorul parametrilor rețelei construit cu elementele matricelor $\mathbf{W}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}^2 \in \mathbb{R}^{p \times m}$ și ale vectorilor $\mathbf{b}^1 \in \mathbb{R}^m$, $\mathbf{b}^2 \in \mathbb{R}^p$, elemente ce sunt dispuse în vectorul $\mathbf{v}$ conform unei reguli prestabilite (de exemplu, întâi liniile matricei $\begin{bmatrix} \mathbf{W}^1 & \mathbf{b}^1 \end{bmatrix}$, urmând apoi liniile matricei $\begin{bmatrix} \mathbf{W}^2 & \mathbf{b}^2 \end{bmatrix}$).

Considerând vectorii eroare $\mathbf{e}_i$ definiți prin (5.1.8), ca funcții de parametrii rețelei (adică de elementele vectorului $\mathbf{v}$) se definește funcția obiectiv (5.1.9):

$$J(\mathbf{v}) = \frac{1}{N} \sum_{i=1}^N \mathbf{e}_i^T(\mathbf{v}) \mathbf{e}_i(\mathbf{v}). \quad (5.2.1)$$

Minimizarea prin metoda gradientului a funcției $J(\mathbf{v})$ (care înseamnă satisfacerea condiției (5.1.10)) constă în modificarea vectorului parametrilor $\mathbf{v} \in \mathbb{R}^{m(n+1)+p(m+1)}$ la fiecare iterație, în conformitate cu formula:

$$\mathbf{v}_{\text{new}} = \mathbf{v}_{\text{old}} - \eta \nabla J(\mathbf{v}_{\text{old}}), \quad (5.2.2)$$

unde $\eta > 0$ notează *pasul metodei de optimizare*, iar $\nabla J(\mathbf{v}_{\text{old}}) = \left[\partial J / \partial \mathbf{v} \right]^T \Big|_{\mathbf{v}_{\text{old}}}$ notează vectorul gradient (normal la suprafețele de nivel constant ale lui J (5.2.1), care trec prin $\mathbf{v}_{\text{old}}$).

Printr-o exprimare adecvată a vectorului gradient, relația (5.2.2) poate fi formulată direct în termenii parametrilor rețelei $\mathbf{W}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{b}^1 \in \mathbb{R}^m$, $\mathbf{W}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}^2 \in \mathbb{R}^p$, furnizând astfel *mecanismul de efectuare a unei iterații* pentru *algoritmul de antrenare prin propagare inversă*. Termenul de “propagare inversă” se referă la maniera de calcul a valorilor curente ale vectorului gradient care se bazează pe o strategie de propagare a erorilor, definite prin (5.1.8), de la stratul de ieșire, înapoi către stratul de intrare.

Pentru algoritmul de antrenare se stabilesc, de către utilizator, un *număr maxim de iterații* sau *epoci* și o valoare ε care se dorește a fi atinsă în minimizarea funcției obiectiv (5.2.1) (adică un prag pentru eroarea pătratică globală). Atingerea unui prag oarecare de

eroare ε impus de utilizator *nu este întotdeauna posibilă*, deoarece valoarea minimă a lui $J(\mathbf{v})$ definită prin (5.2.1) poate fi strict pozitivă.

De asemenea, se fixează (de către utilizator) o valoare pentru pasul metodei de optimizare $\eta > 0$, care, în contextul problemei de optimizare, este referită sub denumirea de *rată de învățare* (eng. *learning rate*).

Se organizează vectorii prototip $\mathbf{x}_i \in \mathbb{R}^n$, $i = 1, \dots, N$, și cei țintă $\mathbf{z}_i \in \mathbb{R}^p$, $i = 1, \dots, N$, sub forma matricelor:

$$\mathbf{X} = [\mathbf{x}_1 \dots \mathbf{x}_N] \in \mathbb{R}^{n \times N}, \quad (5.2.3)$$

respectiv:

$$\mathbf{Z} = [\mathbf{z}_1 \dots \mathbf{z}_N] \in \mathbb{R}^{p \times N}. \quad (5.2.4)$$

La *fiecare iterație* a algoritmului de propagare inversă se parcurg etapele descrise mai jos. Valorile parametrilor rețelei la începerea unei iterații sunt notate prin $\mathbf{W}_{\text{old}}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}_{\text{old}}^2 \in \mathbb{R}^{p \times m}$ (pentru ponderi) și, respectiv, $\mathbf{b}_{\text{old}}^1 \in \mathbb{R}^m$, $\mathbf{b}_{\text{old}}^2 \in \mathbb{R}^p$ (pentru deplasări).

Etapa I (Etapa de prezentare)

Se prezintă vectorii prototip $\mathbf{x}_i$, $i = 1, \dots, N$, la intrarea rețelei neurale, adică pentru valorile curente ale parametrilor rețelei ($\mathbf{W}_{\text{old}}^1$, $\mathbf{W}_{\text{old}}^2$, $\mathbf{b}_{\text{old}}^1$, $\mathbf{b}_{\text{old}}^2$) se calculează, conform (5.1.1-b), ieșirile rețelei, obținând:

$$\mathbf{y}_i = \mathbf{f}(\mathbf{x}_i) = \sigma^2 \left(\mathbf{W}_{\text{old}}^2 \sigma^1 (\mathbf{W}_{\text{old}}^1 \mathbf{x}_i + \mathbf{b}_{\text{old}}^1) + \mathbf{b}_{\text{old}}^2 \right), \quad i = 1, \dots, N, \quad (5.2.5)$$

care se organizează sub forma matricei:

$$\mathbf{Y} = [\mathbf{y}_1 \dots \mathbf{y}_N] \in \mathbb{R}^{p \times N}. \quad (5.2.6)$$

Etapa II (Etapa de verificare)

Se calculează vectorii erorilor $\mathbf{e}_i$, $i = 1, \dots, N$, definiți prin (5.1.8) ca diferența dintre vectorii țintă $\mathbf{z}_i \in \mathbb{R}^p$ și vectorii ieșirii la iterația curentă $\mathbf{y}_i \in \mathbb{R}^p$, construind matricea

$$\mathbf{E} = [\mathbf{e}_1 \dots \mathbf{e}_N] = \mathbf{Z} - \mathbf{Y} \in \mathbb{R}^{p \times N}. \quad (5.2.7)$$

Algoritmul se oprește dacă este îndeplinită una din următoarele două condiții:
($\chi 1$) eroarea pătratică medie este inferioară pragului ε stabilit la startarea algoritmului, adică:

$$J(\mathbf{W}_{\text{old}}^2, \mathbf{W}_{\text{old}}^1, \mathbf{b}_{\text{old}}^1, \mathbf{b}_{\text{old}}^2) = \frac{1}{N} \sum_{i=1}^N \mathbf{e}_i^T \mathbf{e}_i < \varepsilon; \quad (5.2.8)$$

sau:

($\chi 2$) a fost atins numărul maxim de iterații stabilit la startarea algoritmului.

La oprire, algoritmul va furniza valorile parametrilor rețelei rezultate din antrenare, care vor fi notate prin $\mathbf{W}_f^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}_f^2 \in \mathbb{R}^{p \times m}$ (pentru ponderi) și, respectiv, $\mathbf{b}_f^1 \in \mathbb{R}^m$, $\mathbf{b}_f^2 \in \mathbb{R}^p$ (pentru deplasări).

Dacă nici una din condițiile ($\chi 1$) sau ($\chi 2$) nu este îndeplinită, se continuă cu etapa următoare a iterației curente.

Etapa III (Etapa de calcul a vectorilor delta prin propagare inversă)

Pentru valorile curente ale parametrilor rețelei $\mathbf{W}_{old}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}_{old}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}_{old}^1 \in \mathbb{R}^m$, $\mathbf{b}_{old}^2 \in \mathbb{R}^p$, pentru fiecare vector prototip $\mathbf{x}_i$, $i = 1, \dots, N$, se calculează *matricele Jacobian ale funcțiilor de activare* (vectoriale) ale celor două straturi, notate prin:

$$\mathbf{Q}_i^1 \in \mathbb{R}^{m \times m}, \mathbf{Q}_i^2 \in \mathbb{R}^{p \times p}, i = 1, \dots, N.$$

- pentru stratul de intrare:

$$\mathbf{Q}_i^1 = \text{diag} \{q_i^1(1), \dots, q_i^1(m)\}, \quad (5.2.9-a)$$

unde:

$$q_i^1(k) = (\sigma^1)(u_i^1(k)) = \frac{1}{(\text{Inv } \sigma^1)(y_i^1(k))}, \quad k = 1, \dots, m; \quad (5.2.9-b)$$

- pentru stratul de ieșire:

$$\mathbf{Q}_i^2 = \text{diag} \{q_i^2(1), \dots, q_i^2(m)\}, \quad (5.2.10-a)$$

unde:

$$q_i^2(k) = (\sigma^2)(u_i^2(k)) = \frac{1}{(\text{Inv } \sigma^2)(y_i^2(k))} = 1, \quad k = 1, \dots, m. \quad (5.2.10-b)$$

În relațiile (5.2.9-b) și (5.2.10-b) care definesc elementele matricelor Jacobian, notațiile $(\text{Inv } \sigma^1)$ și $(\text{Inv } \sigma^2)$ au semnificația de inverse ale funcțiilor de activare σ^1 și respectiv σ^2 corespunzătoare celor două straturi. Deoarece funcțiile de activare ale celui de al doilea strat sunt liniare, $\mathbf{Q}_i^2$ este *matricea unitate*, $i = 1, \dots, N$.

Cu ajutorul matricelor Jacobian $\mathbf{Q}_i^1 \in \mathbb{R}^{m \times m}$, $\mathbf{Q}_i^2 \in \mathbb{R}^{p \times p}$ și a vectorilor erorilor $\mathbf{e}_i \in \mathbb{R}^p$, $i = 1, \dots, N$, $i = 1, \dots, N$, se definesc *vectorii delta* (prin procesarea informației de la ieșire, către intrare, procesare numită *propagare inversă*):

- pentru stratul de ieșire:

$$\delta_i^2 = \mathbf{Q}_i^2 \mathbf{e}_i = \mathbf{e}_i \in \mathbb{R}^p, \quad i = 1, \dots, N; \quad (5.2.11)$$

- pentru stratul de intrare:

$$\delta_i^1 = \mathbf{Q}_i^1 \left((\mathbf{W}^2)^T \delta_i^2 \right) \in \mathbb{R}^m, \quad i = 1, \dots, N. \quad (5.2.12)$$

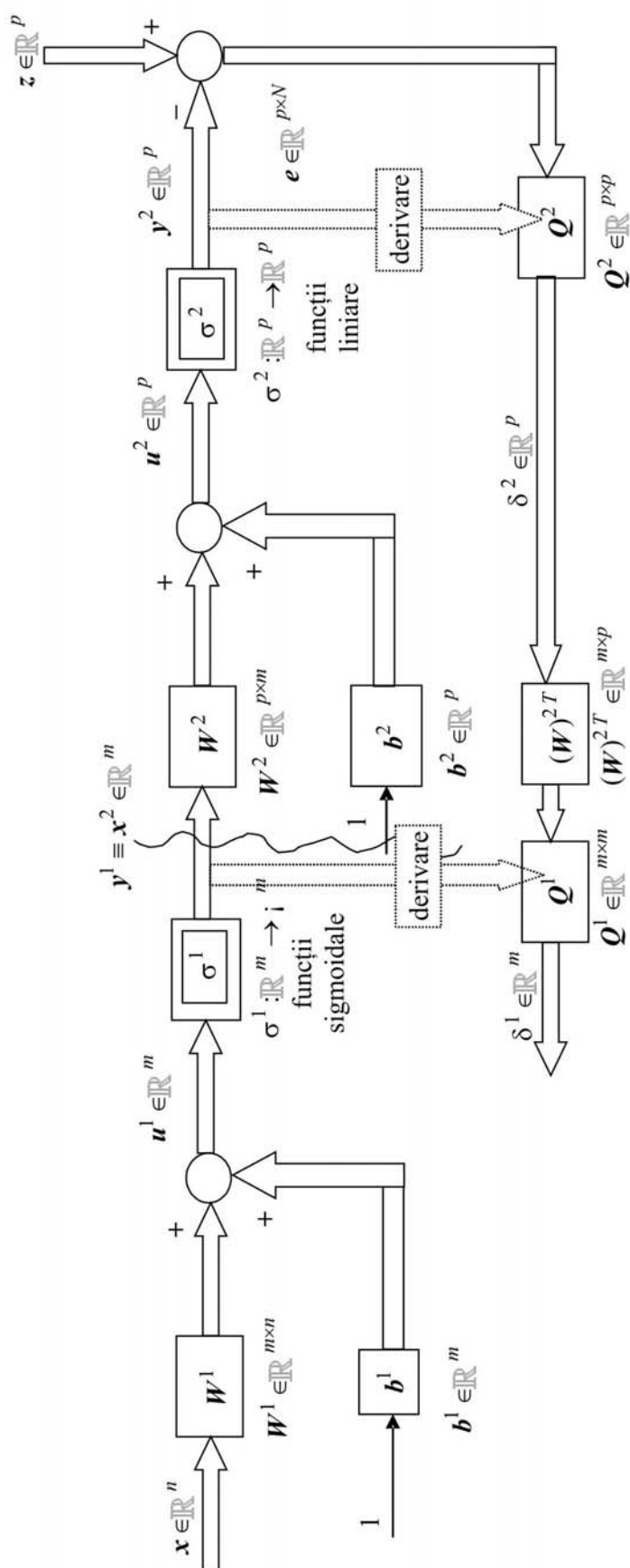


Figura 5.2.1. Schema bloc a procesării în sens invers a informației în antrenarea unei rețele cu două straturi prin algoritmul propagării inverse - calculul vectorilor delta δ^2 și δ^1

Astfel, suntem în măsură să construim *matricele delta* ale celor două straturi:

- pentru stratul de ieșire

$$\Delta^2 = [\delta_i^2 \dots \delta_N^2] \in \mathbb{R}^{p \times N}, \quad (5.2.13)$$

- pentru stratul de intrare

$$\Delta^1 = [\delta_i^1 \dots \delta_N^1] \in \mathbb{R}^{m \times N}. \quad (5.2.14)$$

O reprezentare sub forma de schemă bloc a modului de calcul al *vectorilor delta* prin propagare inversă într-o rețea cu două straturi este dată în figura 5.2.1.

Etapa IV (Etapa de actualizare a parametrilor)

Cu ajutorul matricelor delta se pot calcula *matricele de actualizare* a parametrilor aplicând *regula delta* pe fiecare strat:

- pentru stratul de ieșire

– ponderi

$$D_{W^2} = \Delta^2 (X^2)^T \in \mathbb{R}^{p \times m}, \quad (5.2.15-a)$$

– deplasări

$$D_{b^2} = \Delta^2 \underbrace{[1 \dots 1]^T}_{N \text{ elemente}} \in \mathbb{R}^p; \quad (5.2.15-b)$$

- pentru stratul de intrare

– ponderi

$$D_{W^1} = \Delta^1 X^T \in \mathbb{R}^{m \times n}, \quad (5.2.16-a)$$

– deplasări

$$D_{b^1} = \Delta^1 \underbrace{[1 \dots 1]^T}_{N \text{ elemente}} \in \mathbb{R}^m. \quad (5.2.16-b)$$

În relația (5.2.15-a), matricea $X^2 \in \mathbb{R}^{m \times N}$ colectează vectorii $x_i^2 = y_i^1 \in \mathbb{R}^m$, $i = 1, \dots, N$, care se prezintă la intrarea celui de-al doilea strat, adică:

$$X^2 = [x_1^2 \dots x_N^2] \in \mathbb{R}^{m \times N}. \quad (5.2.17)$$

Pe baza matricelor de actualizare se determină noile valori ale parametrilor rețelei:

- pentru stratul de ieșire

– ponderi

$$W_{\text{new}}^2 = W_{\text{old}}^2 + \eta D_{W^2}, \quad (5.2.18-a)$$

– deplasări

$$b_{\text{new}}^2 = b_{\text{old}}^2 + \eta D_{b^2}; \quad (5.2.18-b)$$

- pentru stratul de intrare

– ponderi

$$W_{\text{new}}^1 = W_{\text{old}}^1 + \eta D_{W^1}, \quad (5.2.19-a)$$

– deplasări

$$\mathbf{b}_{\text{new}}^1 = \mathbf{b}_{\text{old}}^1 + \eta \mathbf{D}_b^1. \quad (5.2.19\text{-b})$$

Cu aceste calcule, iterația curentă se încheie și se trece la o nouă iterație, efectuând atribuirile:

$$\begin{aligned} \mathbf{W}_{\text{new}}^2 &\rightarrow \mathbf{W}_{\text{old}}^2, & \mathbf{b}_{\text{new}}^2 &\rightarrow \mathbf{b}_{\text{old}}^2, \\ \mathbf{W}_{\text{new}}^1 &\rightarrow \mathbf{W}_{\text{old}}^1, & \mathbf{b}_{\text{new}}^1 &\rightarrow \mathbf{b}_{\text{old}}^1. \end{aligned} \quad (5.2.20)$$

5.2.3. Algoritmul de antrenare prin propagare inversă privit ca problemă standard de optimizare

În acest paragraf vom prezenta o demonstrație pentru egalitățile (5.2.15) și (5.2.16), care furnizează matricele de actualizare a parametrilor rețelei la fiecare iterație a algoritmului de antrenare. În acest scop, vom porni de la forma iterației standard a metodei de gradient, exprimată prin (5.2.2).

Aplicarea iterațiilor de tip (5.2.2) pentru funcția obiectiv (5.2.1) permite ajustarea parametrilor rețelei în conformitate cu metoda gradientului care operează în spațiul $\mathbb{R}^{m(n+1)+p(m+1)}$. Vectorii acestui spațiu au forma:

$$\mathbf{v} = [\tilde{\mathbf{w}}_1^1 \dots \tilde{\mathbf{w}}_m^1 \tilde{\mathbf{w}}_1^2 \dots \tilde{\mathbf{w}}_p^2]^T, \quad (5.2.21)$$

unde $\tilde{\mathbf{w}}_k^1$, $k=1, \dots, m$, și $\tilde{\mathbf{w}}_j^2$, $j=1, \dots, p$, notează liniile matricelor parametrilor pentru *primul strat*:

$$\tilde{\mathbf{W}}^1 = [\mathbf{W}^1 \mathbf{b}^1] = \begin{bmatrix} \tilde{\mathbf{w}}_1^1 \\ \dots \\ \tilde{\mathbf{w}}_m^1 \end{bmatrix} \in \mathbb{R}^{m \times (n+1)} \quad (5.2.22)$$

și respectiv pentru al *doilea strat*:

$$\tilde{\mathbf{W}}^2 = [\mathbf{W}^2 \mathbf{b}^2] = \begin{bmatrix} \tilde{\mathbf{w}}_1^2 \\ \dots \\ \tilde{\mathbf{w}}_p^2 \end{bmatrix} \in \mathbb{R}^{p \times (m+1)}. \quad (5.2.23)$$

Scriind iterația (5.2.2) pentru fiecare set de parametri $\tilde{\mathbf{w}}_k^1$, $k=1, \dots, m$, în parte:

$$\tilde{\mathbf{w}}_{k \text{ new}}^1 = \tilde{\mathbf{w}}_{k \text{ old}}^1 - \eta \frac{\partial J}{\partial \tilde{\mathbf{w}}_k^1}, \quad k=1, \dots, m, \quad (5.2.24)$$

și asamblând cele m egalități vectoriale din (5.2.24) sub formă matriceală, obținem pentru *stratul de intrare*:

$$\tilde{\mathbf{W}}_{\text{new}}^1 = \tilde{\mathbf{W}}_{\text{old}}^1 - \eta \begin{bmatrix} \frac{\partial J}{\partial \tilde{\mathbf{w}}_1^1} \\ \dots \\ \frac{\partial J}{\partial \tilde{\mathbf{w}}_m^1} \end{bmatrix}. \quad (5.2.25)$$

Egalitatea (5.2.25) este însă de aceeași formă cu (5.2.19), constatare ce arată că actualizarea parametrilor în stratul de intrare se realizează în *scrierea matriceală* de tip (5.2.19), prin intermediul matricei:

$$\mathbf{D}_{\tilde{\mathbf{W}}^1} = \begin{bmatrix} \mathbf{D}_{\mathbf{W}^1} & \mathbf{D}_{\mathbf{b}^1} \end{bmatrix} = - \begin{bmatrix} \frac{\partial J}{\partial \tilde{\mathbf{w}}_1^1} \\ \dots \\ \frac{\partial J}{\partial \tilde{\mathbf{w}}_m^1} \end{bmatrix} \in \mathbb{R}^{m \times (n+1)}. \quad (5.2.26)$$

În mod analog se arată că actualizarea parametrilor în stratul de ieșire se realizează în *scrierea matriceală* de tip (5.2.18), prin intermediul matricei:

$$\mathbf{D}_{\tilde{\mathbf{W}}^2} = \begin{bmatrix} \mathbf{D}_{\mathbf{W}^2} & \mathbf{D}_{\mathbf{b}^2} \end{bmatrix} = - \begin{bmatrix} \frac{\partial J}{\partial \tilde{\mathbf{w}}_1^2} \\ \dots \\ \frac{\partial J}{\partial \tilde{\mathbf{w}}_p^2} \end{bmatrix} \in \mathbb{R}^{p \times (m+1)}. \quad (5.2.27)$$

Pe de altă parte, scriind funcția obiectiv (5.2.1) sub forma:

$$J(\mathbf{v}) = \sum_{i=1}^N J_i, \quad J_i = \frac{1}{2} \mathbf{e}_i^T \mathbf{e}_i, \quad i = 1, \dots, N, \quad (5.2.28)$$

vectorii linie $\frac{\partial J}{\partial \tilde{\mathbf{w}}_k^1}$, $k = 1, \dots, m$, și $\frac{\partial J}{\partial \tilde{\mathbf{w}}_j^2}$, $j = 1, \dots, p$, ce intră în structura matricelor $\mathbf{D}_{\tilde{\mathbf{W}}^1}$ din (5.2.26) și respectiv $\mathbf{D}_{\tilde{\mathbf{W}}^2}$ din (5.2.27) pot fi exprimați drept:

$$\frac{\partial J}{\partial \tilde{\mathbf{w}}_k^1} = \sum_{i=1}^N \frac{\partial J_i}{\partial \tilde{\mathbf{w}}_k^1}, \quad k = 1, \dots, m, \quad (5.2.29)$$

și respectiv:

$$\frac{\partial J}{\partial \tilde{\mathbf{w}}_j^2} = \sum_{i=1}^N \frac{\partial J_i}{\partial \tilde{\mathbf{w}}_j^2}, \quad j = 1, \dots, p. \quad (5.2.30)$$

Cu această observație, *matricele de actualizare a parametrilor* pot fi puse sub forma:

- pentru stratul de intrare:

$$\mathbf{D}_{\tilde{\mathbf{W}}^1} = \sum_{i=1}^N \mathbf{D}_{\tilde{\mathbf{W}}^1}^i, \quad (5.2.31)$$

cu:

$$\mathbf{D}_{\tilde{\mathbf{w}}^1}^i = - \begin{bmatrix} \frac{\partial J_i}{\partial \tilde{\mathbf{w}}_1^1} \\ \dots \\ \frac{\partial J_i}{\partial \tilde{\mathbf{w}}_m^1} \end{bmatrix} \in \mathbb{R}^{m \times (n+1)}, \quad i = 1, \dots, N; \quad (5.2.32)$$

- pentru stratul de ieșire:

$$\mathbf{D}_{\tilde{\mathbf{w}}^2} = \sum_{i=1}^N \mathbf{D}_{\tilde{\mathbf{w}}^2}^i, \quad (5.2.33)$$

cu:

$$\mathbf{D}_{\tilde{\mathbf{w}}^2}^i = - \begin{bmatrix} \frac{\partial J_i}{\partial \tilde{\mathbf{w}}_1^2} \\ \dots \\ \frac{\partial J_i}{\partial \tilde{\mathbf{w}}_p^2} \end{bmatrix} \in \mathbb{R}^{p \times (m+1)}, \quad i = 1, \dots, N \quad (5.2.34)$$

Aplicând regula derivării funcțiilor compuse pentru *stratul de ieșire* se obține:

$$\frac{\partial J_i}{\partial \tilde{\mathbf{w}}_j^2} = \frac{\partial J_i}{\partial \mathbf{e}_i} \cdot \frac{\partial \mathbf{e}_i}{\partial \mathbf{y}_i} \cdot \frac{\partial \mathbf{y}_i}{\partial \mathbf{u}_i} \cdot \frac{\partial \mathbf{u}_i}{\partial \tilde{\mathbf{w}}_j^2} = -(\boldsymbol{\delta}_i^2)^T \begin{bmatrix} 0 \\ \dots \\ \left(\tilde{\mathbf{x}}_i^2 \right)^T \leftarrow j, \quad j = 1, \dots, p, \quad i = 1, \dots, N, \\ \dots \\ 0 \end{bmatrix} \quad (5.2.35)$$

unde vectorii $\boldsymbol{\delta}_i^2$, $i = 1, \dots, N$, satisfac egalitatea (5.2.11), adică reprezintă *vectorii delta ai stratului de ieșire*. Astfel putem scrie:

$$\mathbf{D}_{\tilde{\mathbf{w}}^2}^i = \begin{bmatrix} \boldsymbol{\delta}_i^2(1) \left(\tilde{\mathbf{x}}_i^2 \right)^T \\ \dots \\ \boldsymbol{\delta}_i^2(p) \left(\tilde{\mathbf{x}}_i^2 \right)^T \end{bmatrix} = \boldsymbol{\delta}_i^2 \left(\tilde{\mathbf{x}}_i^2 \right)^T, \quad i = 1, \dots, N. \quad (5.2.36)$$

Înlocuind în (5.2.33) se obține:

$$\mathbf{D}_{\tilde{\mathbf{w}}^2} = \sum_{i=1}^N \boldsymbol{\delta}_i^2 \left(\tilde{\mathbf{x}}_i^2 \right)^T = \begin{bmatrix} \boldsymbol{\delta}_1^2 & \dots & \boldsymbol{\delta}_N^2 \end{bmatrix} \begin{bmatrix} \left(\tilde{\mathbf{x}}_1^2 \right)^T \\ \dots \\ \left(\tilde{\mathbf{x}}_N^2 \right)^T \end{bmatrix}, \quad (5.2.37)$$

adică s-au demonstrat egalitățile (5.2.15) din *Etapa IV* a algoritmului de antrenare.

Aplicăm acum regula derivării funcțiilor compuse pentru *stratul de intrare* și avem:

$$\frac{\partial J_i}{\partial \tilde{\mathbf{w}}_k^1} = \frac{\partial J_i}{\partial \mathbf{e}_i} \cdot \frac{\partial \mathbf{e}_i}{\partial \mathbf{y}_i} \cdot \frac{\partial \mathbf{y}_i}{\partial \mathbf{u}_i} \cdot \frac{\partial \mathbf{u}_i}{\partial \mathbf{y}_i'} \cdot \frac{\partial \mathbf{y}_i'}{\partial \mathbf{u}_i'} \cdot \frac{\partial \mathbf{u}_i'}{\partial \tilde{\mathbf{w}}_k^1} = -(\delta_i^2)^T \mathbf{W}^2 \mathbf{Q}_i^1 \begin{bmatrix} 0 \\ \dots \\ \tilde{\mathbf{x}}_i^T \\ \dots \\ 0 \end{bmatrix} = -(\delta_i^1)^T \begin{bmatrix} 0 \\ \dots \\ \tilde{\mathbf{x}}_i^T \\ \dots \\ 0 \end{bmatrix} \leftarrow k, \quad (5.2.38)$$

$k = 1, \dots, m$, $i = 1, \dots, N$, unde vectorii δ_i^1 , $i = 1, \dots, N$, satisfac egalitatea (5.2.12), adică reprezintă vectorii delta ai stratului de intrare. Astfel putem scrie:

$$\mathbf{D}_{\tilde{\mathbf{W}}^1}^i = \begin{bmatrix} \delta_i^1(1) \tilde{\mathbf{x}}_i^T \\ \dots \\ \delta_i^1(m) \tilde{\mathbf{x}}_i^T \end{bmatrix} = \delta_i^1 \tilde{\mathbf{x}}_i^T, \quad i = 1, \dots, N. \quad (5.2.39)$$

Înlocuind în (5.2.31), obținem:

$$\mathbf{D}_{\tilde{\mathbf{W}}^1} = \sum_{i=1}^N \delta_i^1 \tilde{\mathbf{x}}_i^T = \begin{bmatrix} \delta_1^1 & \dots & \delta_N^1 \end{bmatrix} \begin{bmatrix} \tilde{\mathbf{x}}_1^T \\ \dots \\ \tilde{\mathbf{x}}_N^T \end{bmatrix}, \quad (5.2.40)$$

adică s-au demonstrat și egalitățile (5.2.16) din *Etapa IV* a algoritmului de antrenare.

Observații:

1° Inițializarea parametrilor rețelei (adică $\mathbf{W}_{\text{old}}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}_{\text{old}}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}_{\text{old}}^1 \in \mathbb{R}^m$, $\mathbf{b}_{\text{old}}^2 \in \mathbb{R}^p$ la prima iterație a algoritmului) se face cu valori *subunitare arbitrare*; în acest mod se asigură sensibilitate maximă pentru funcțiile sigmoideale din primul strat al rețelei.

2° Convergența algoritmului de propagare inversă se realizează în condițiile specifice metodei gradientului (ca tehnică de optimizare), fiind puternic dependentă de valoarea ratei de învățare η și anume:

- pentru η inadecvat mic, deplasarea către optim, în spațiul parametrilor, este foarte lentă,
- pentru η inadecvat mare, deplasarea către optim, în spațiul parametrilor, se realizează cu oscilații (eventual cu pierderea stabilității algoritmului).

3° În cazul când algoritmul se încheie prin satisfacerea condiției ($\chi 1$), acuratețea parametrilor furnizați $\mathbf{W}_f^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}_f^2 \in \mathbb{R}^{p \times m}$ (ponderi) și $\mathbf{b}_f^1 \in \mathbb{R}^m$, $\mathbf{b}_f^2 \in \mathbb{R}^p$ (deplasări) depinde de valoarea pragului ε stabilit pentru eroarea pătratică globală. Valoarea finală a lui J (implicit alegerea valorii de prag ε) trebuie corelată cu magnitudinea valorilor țintă $\mathbf{z}_i$, $i = \overline{1, N}$. De exemplu, dacă $\mathbf{z}_i$ este de ordinul 10^6 în unitățile de măsură, eroarea poate fi aleasă de ordinul $10^4 - 10^3$; nu vom utiliza drept criteriu de antrenare un prag de ordinul 10^{-2} pentru J deoarece acesta nu ar putea fi atins practic niciodată.

4° Suprafața erorii $J(\mathbf{v})$ definită prin (10) uzual posedă *minime locale* în care căutarea minimului poate eșua, dacă rata de învățare are o valoare prea mică. Astfel de minime locale nu pot fi părăsite decât prin schimbarea valorii ratei de învățare, sau prin restartarea algoritmului cu alte valori inițiale ale parametrilor.

5° Numărul de neuroni din stratul de intrare se alege în funcție de “gradul de neliniaritate” care se constată la funcția ϕ din examinarea eșantioanelor (5.1.4). În alegerea arhitecturii rețelei (adică a numărului de neuroni de pe primul strat) pot părea două arhitecturi care trebuie verificate și eliminate:

- a. Subaproximarea (eng. *underfitting*): care are loc pentru un număr prea mic de neuroni. În cazul $n = p = 1$, graficul realizat de rețea nu prezintă un număr suficient de puncte de inflexiune. Pentru a asigura calitatea aproximării este suficient la J să fie foarte mic (nu va fi de aceeași valoare ca atunci când trec prin vecinătatea tuturor punctelor).
- b. Supraaproximarea (eng. *overfitting*): se alege un număr mare de neuroni în cazul funcției scalare $n = p = 1$. Numărul excesiv de neuroni creează prea multe puncte de inflexiune, care nu ar putea exista pentru un fenomen real.

Alegerea adecvată a numărului de neuroni se bazează pe experiență și, eventual, pe teste de factură încercare-eroare (eng. *trial and error*) efectuate pe problema concretă. ■

5.2.4. Strategii de creștere a performanțelor antrenării prin propagare inversă

Pentru îmbunătățirea convergenței algoritmului de antrenare se utilizează modificări ale iterațiilor standard de tipul (5.2.18) și (5.2.19), modificări ce au drept efect creșterea vitezei de căutare și, totodată, fac posibilă evitarea situațiilor patologice de eșuare în minime locale (Sima, Varga, 1986). Astfel de modificări au drept fundament teoretic diferite maniere de a interveni în iterația standard a metodei de gradient (5.2.2).

În scopul prezentării strategiilor de creștere a performanțelor antrenării prin propagare inversă, vom rescrie iterațiile (5.2.18) și (5.2.19) în maniera compactă oferită de notațiile matriceale (5.2.26) și (5.2.27) adică:

- pentru stratul de ieșire:

$$\tilde{\mathbf{W}}_{\text{new}}^2 = \tilde{\mathbf{W}}_{\text{old}}^2 + \eta \mathbf{D}_{\tilde{\mathbf{W}}^2}; \quad (5.2.41)$$

- pentru stratul de intrare:

$$\tilde{\mathbf{W}}_{\text{new}}^1 = \tilde{\mathbf{W}}_{\text{old}}^1 + \eta \mathbf{D}_{\tilde{\mathbf{W}}^1}. \quad (5.2.42)$$

Considerând iterația curentă ca fiind cea de-a k -a în procesul de antrenare, putem reformula (5.2.41) și (5.2.42) într-o manieră generală:

$$\tilde{\mathbf{W}}^2(k+1) = \tilde{\mathbf{W}}^2(k) + \tilde{\mathbf{M}}^2(k), \quad (5.2.43)$$

$$\tilde{\mathbf{W}}^1(k+1) = \tilde{\mathbf{W}}^1(k) + \tilde{\mathbf{M}}^1(k), \quad (5.2.44)$$

unde matricele $\tilde{\mathbf{M}}^2(k)$ și $\tilde{\mathbf{M}}^1(k)$ definesc schimbările parametrilor impuse de procesul de căutare la iterația curentă:

$$\tilde{\mathbf{M}}^2(k) = \eta \mathbf{D}_{\tilde{\mathbf{W}}^2}(k), \quad (5.2.45)$$

$$\tilde{\mathbf{M}}^1(k) = \eta \mathbf{D}_{\tilde{\mathbf{W}}^1}(k). \quad (5.2.46)$$

În relațiile (5.2.45) și (5.2.46) matricele $D_{\tilde{W}^2}(k)$ și $D_{\tilde{W}^1}(k)$ sunt date de (5.2.15) și respectiv (5.2.16) calculate cu vectorii delta de la iterația k , adică:

$$D_{\tilde{W}^2}(k) = \Delta^2(k) \left(\tilde{X}^2 \right)^T, \quad (5.2.47)$$

și, respectiv:

$$D_{\tilde{W}^1}(k) = \Delta^1(k) \tilde{X}^T. \quad (5.2.48)$$

Strategiile de creștere a performanțelor antrenării prin propagare inversă se bazează pe schimbări ale modului de calcul dat de (5.2.45) și (5.2.46) pentru matricele $\tilde{M}^2(k)$ și $\tilde{M}^1(k)$, ce operează în iterația k descrisă prin egalitățile (5.2.43) și (5.2.44).

A. Propagarea inversă cu moment

Introducerea *momentului* revine la a înlocui (5.2.45) și (5.2.46) cu următorul mod de calcul al matricelor $\tilde{M}^2(k)$ și $\tilde{M}^1(k)$:

$$\tilde{M}^2(k) = m\tilde{M}^2(k-1) + (1-m)\eta D_{\tilde{W}^2}(k), \quad (5.2.49)$$

și respectiv:

$$\tilde{M}^1(k) = m\tilde{M}^1(k-1) + (1-m)\eta D_{\tilde{W}^1}(k), \quad (5.2.50)$$

unde momentul m are o valoare subunitară apropiată de 1.

Prin aceasta, direcția de căutare curentă în procesul de optimizare este exprimată drept o *combinație liniară* între direcția gradientului la pasul curent (dată de termenii de pe poziția a doua a sumelor din (5.2.49) și (5.2.50)) și direcția de căutare la pasul precedent (dată de termenii de pe prima poziție a sumelor din (5.2.49) și (5.2.50)). Cu alte cuvinte, căutarea se realizează memorând, de la un pas la altul, vechea direcție, de care se va ține seama (ponderat prin m) în stabilirea noii direcții.

B. Propagarea inversă cu rată de învățare adaptivă

Utilizarea ratei de învățare adaptive revine la a înlocui, în (5.2.45) și (5.2.46), factorul constant η printr-o valoare $\eta(k)$, ajustabilă de la o iterație la alta, adică (5.2.45) și (5.2.46) se transformă în:

$$\tilde{M}^2(k) = \eta(k) D_{\tilde{W}^2}(k), \quad (5.2.51)$$

și, respectiv:

$$\tilde{M}^1(k) = \eta(k) D_{\tilde{W}^1}(k). \quad (5.2.52)$$

Prin aceasta, $\eta(k)$ poate fi crescut atunci când, prin ajustarea parametrilor, valoarea funcției obiectiv J din (5.2.1) scade semnificativ. Astfel, antrenarea se poate realiza cu o viteză mai mare, atâta timp cât direcția de căutare este favorabilă. Când căutarea eșuează în descreșterea lui J din (5.2.1), rata $\eta(k)$ urmează a fi micșorată.

C. Propagarea inversă cu moment și rată de învățare adaptivă

Prin această strategie se combină avantajele celor două procedee prezentate anterior. Astfel, din (5.2.49) și (5.2.51) va rezulta, pentru stratul de ieșire, forma tipică a iterației k sub forma:

$$\tilde{M}^2(k) = m\tilde{M}^2(k-1) + (1-m)\eta(k)D_{\tilde{W}^2}(k). \quad (5.2.53)$$

Analog din (5.2.50) și (5.2.52) va rezulta forma tipică a iterației k pentru ajustarea parametrilor de pe stratul de intrare:

$$\tilde{M}^1(k) = m\tilde{M}^1(k-1) + (1-m)\eta(k)D_{\tilde{W}^1}(k). \quad (5.2.54)$$

5.2.5. Alți algoritmi de antrenare pentru rețele multistrat

După cum am precizat anterior, antrenarea rețelei neurale constă în rezolvarea unei probleme de minimizare corespunzătoare unei funcții obiectiv de tip eroare pătratică medie (având expresia din (5.1.9)). Am mai insistat asupra faptului că ceea ce, în limbajul rețelelor neurale, este cunoscut sub denumirea de “propagare inversă” reprezintă un algoritm de minimizare de tip gradient cu descreștere maximă (steepest descent), dacă utilizăm exprimările mai generale specifice tehnicilor de optimizare. Altfel spus, propagarea inversă este un algoritm de tip gradient cu descreștere maximă, aranjat într-o formulare matematică specială, asociată tipului particular de aplicații care descriu funcționarea rețelelor neurale (și, implicit, care contribuie la definirea funcției obiectiv).

Acest mod de înțelegere a algoritmului de antrenare al rețelei neurale lasă loc unei întrebări firești. Pentru antrenarea unei rețele neurale pot fi utilizate și alte metode de minimizare (împrumutate din “recuzita” tehnicilor de optimizare)?

Răspunsul este afirmativ. În prezent, există algoritmi de antrenare ai rețelelor neurale care au fost obținuți (adaptați) din metodele de optimizare cele mai eficiente (ca de exemplu metoda Levenberg-Marquardt, metodele bazate pe direcții conjugate etc.). În general vorbind, metodele amintite prezintă performanțe numerice mult mai bune decât gradientul cu descreștere maximă, motiv pentru care utilizarea lor în antrenarea rețelelor neurale este mai avantajoasă decât propagarea inversă. Cu toate acestea, atragem atenția că propagarea inversă reprezintă un subiect tratat de orice manual / monografie referitoare la rețele neuronale, întrucât deține o “semnificație istorică” pentru dezvoltarea domeniului, fiind strâns legată de începuturile utilizării rețelelor multistrat.

Pentru a aprofunda problematica metodelor de optimizare implementate drept algoritmi de antrenare supervizată a rețelelor neurale, cititorul este invitat să parcurgă Anexa A a prezentei cărți.

5.3. Utilizarea rețelelor feedforward cu trei straturi

Rețelele feedforward cu trei straturi (primele două cu noduri sigmoidale și cel de ieșire cu noduri liniare) posedă proprietatea de aproximare discutată în paragraful 5.1. Uzual se face apel la astfel de rețele atunci când funcția ce trebuie aproximată φ din (5.1.2) prezintă moduri de variație mai complicate, puse în evidență de eșantioanele (5.1.4).

Antrenarea se realizează prin metoda propagării inverse organizată pe cele trei straturi prin generalizarea algoritmului prezentat în paragraful 5.2.2 aplicând *regula delta* succesiv celor trei straturi. În acest caz, vor trebui calculați *vectorii și matricele delta* pentru toate cele trei straturi, într-o manieră analogă cu (5.2.15) – (5.2.19).

5.4. Utilizarea rețelelor feedforward în probleme de clasificare

După cum a fost prezentat în capitolul 3, rețelele cu funcții de activare treaptă au proprietatea de a separa spațiul ieșirilor prin frontiere de tip hiperplane. Datorită acestei proprietăți, rețelele cu funcție de activare treaptă pot fi utilizate în probleme de clasificare. Un asemenea clasificator poate fi utilizat însă numai pentru clase liniar separabile. În plus, o rețea neurală de tip perceptron cu un strat nu este capabilă să realizeze o clasificare pe problema simplă de tip XOR (care necesită o rețea cu două straturi, după cum se arată în problema 7.5). Detalii privind problematica clasificării sunt accesibile în anexa B.

Aceste deficiențe pot fi depășite dacă se utilizează rețele neurale cu funcții de activare de tip sigmoidal având o arhitectură cu două straturi. Pentru utilizarea unei astfel de arhitecturi, ca și în cazul rețelelor cu funcții de activare de tip treaptă, se parcurg două etape:

- (i) antrenarea rețelei neurale pe baza unui lot reprezentativ, capabil să dea informații relevante despre clase;
- (ii) utilizarea rețelei antrenate drept clasificator.

În continuare vom considera rețele neurale cu funcții de activare de tip sigmoid unipolar, dar în același măsură pot fi utilizate și noduri bipolare. Să presupunem că proiectăm o aplicație pentru a realiza clasificarea punctelor dintr-un domeniu din $\mathbb{R}^n$ în s clase.

O primă variantă este aceea de a considera o rețea cu n intrări și s ieșiri. Clasei $\mathcal{C}_j$ îi asociem vectorul $\underline{y}_j = [0 \dots 1 \dots 0]^T \in \mathbb{R}^s$, pentru $j = 1, 2, \dots, s$. Întrucât ieșirile rețelei sunt date de funcții sigmoide, luând valori în intervalul deschis $(0, 1)$ și doar tinzând asimptotic către capetele acestuia, în vectorii țintă nu pot fi utilizate valorile binare 0 și 1. Pentru ca neuronii din stratul de ieșire să furnizeze valori foarte apropiate fie de 0 fie de 1, ar trebui ca ponderile acestora să ia valori foarte mari. Din acest motiv se preferă relaxarea formei binare cu 0 și 1 a vectorilor țintă, utilizând în locul lor 0 valorile ε și, respectiv, $1 - \varepsilon$, cu $\varepsilon > 0$ suficient de mic (uzual ε se ia de ordinul sutimilor). Astfel, clasei $\mathcal{C}_j$ îi asociem vectorul $\tilde{y}_j = [\varepsilon \dots 1 - \varepsilon \dots \varepsilon]^T \in \mathbb{R}^s$, pentru $j = 1, 2, \dots, s$. Antrenarea rețelei neurale se va realiza pe baza perechilor $(\mathbf{x}^i, \mathbf{z}^i)$, $i = \overline{1, N}$, cu $\mathbf{x}^i \in \mathbb{R}^n$ și $\mathbf{z}^i \in \mathbb{R}^s$, $\mathbf{z}^i = \tilde{y}_j$ dacă $\mathbf{x}^i \in \mathcal{C}_j$. Prin perechile $(\mathbf{x}^i, \mathbf{z}^i)$, $i = \overline{1, N}$, definim o problemă de tip aproximare, contextul teoretic dezvoltat de Cybenko ne asigură succesul antrenării (dacă se consideră un număr suficient de neuroni pe primul strat).

Dacă numărul de clase este o putere a lui 2, de forma $s = 2^p$, se poate utiliza o a doua variantă, similară celei prezentate în capitolul 3. Se consideră o rețea cu n intrări și p ieșiri. Clasei $\mathcal{C}_j$ i se asociază un vector $\underline{y}_j \in \mathbb{R}^p$, $j = 1, \dots, 2^p$, reprezentând codul binar al indicelui

j al clasei, apoi în locul valorilor 0 și 1 se utilizează ε și, respectiv, $1 - \varepsilon$, cu $\varepsilon > 0$ suficient de mic, construind vectorul $\tilde{y}_j \in \mathbb{R}^P$. Antrenarea rețelei neurale se va realiza pe baza perechilor (x^i, z^i) , $i = \overline{1, N}$, cu $x^i \in \mathbb{R}^n$ și $z^i \in \mathbb{R}^P$, $z^i = \tilde{y}_j$ dacă $x^i \in \mathcal{C}_j$.

În cazul în care nodurile rețelei neurale au noduri bipolare, vectorii țintă se formează cu elementele $-1 + \varepsilon$ și $1 - \varepsilon$.

Practica a arătat că stratul de ieșire al rețelei poate avea și noduri de tip liniar. În aceste condiții apartenența la o clasă în etapa de învățare se va realiza cu vectori țintă formați cu valori de 0 și 1 ferme (vectorii notați în paragrafele anterioare cu y_j). Datorită utilizării funcției de activare liniară pe ultimul strat, ieșirile rețelei pot avea și valori mai mici decât 0 sau mai mari decât 1.

Indiferent de tipul funcțiilor de activare utilizate pe stratul de ieșire al rețelei neurale utilizată pentru clasificare, decizia privind apartenența unui vector prototip la o anumită clasă va fi luată pe baza distanței dintre ieșirea rețelei neurale și vectorii țintă corespunzători claselor considerate. Un vector prototip va fi atribuit clasei al cărei vector prototip este situat la distanța cea mai mică de ieșirea rețelei neurale utilizate drept clasificator.

La fel ca și în probleme de aproximare, în faza de antrenare a rețelei în scopul recunoașterii de trăsături sau forme trebuie să fie prezentate rețelei toate elementele semnificative din punct de vedere al aplicației. Rețeaua va trebui “să vadă” eșantioanele reprezentative pentru toate clasele care apar în aplicații (numărul de clase și semnificația claselor este o problemă care precede antrenarea rețelei și trebuie complet soluționată înainte de momentul începerii antrenării). Adică rețeaua nu este implicată în a decide câte clase distincte sunt.

Valorile concrete pe care le furnizează clasificatorul pe ieșirea rețelei neurale, pot fi utilizate pentru a informa asupra „calității clasificării”, (altfel spus, un fel de garanție pe care o oferă rețeaua că separă (atribuie) corect vectorii de intrare pe clase). Distanțe mari față de ε și $1 - \varepsilon$ (în cazul ieșirilor sigmoideale) sau 0, 1 (în cazul ieșirilor liniare) atrag atenția asupra unor elemente clasificate cu o doză mare de risc (de incertitudine) în cazul că în etapa de învățare, rețeaua „nu a văzut” eșantioane apropiate ca trăsături cu aceste elemente. Astfel de elemente necesită o altă procedură de clasificare pentru reducerea incertitudinii sau, dacă apariția lor este frecventă, atrage atenția asupra necesității reantrenării rețelei (de așa manieră încât rețeaua să aibă cunoștință despre trăsăturile aferente elementelor respective).

5.5. Facilități software oferite de Neural Network Toolbox

Funcția `newff` creează o rețea feed-forward. Numărul de straturi și de neuroni din stratul de intrare și din straturile ascunse, precum și funcția de activare a neuronilor fiecărui strat sunt specificate de către utilizator în apelul funcției. La crearea rețelei utilizatorul trebuie să furnizeze și rutina ce va fi utilizată pentru antrenarea acesteia precum și parametrii de antrenare doriți.

Pentru inițializarea cu valori arbitrare a parametrilor unei rețele (ponderi și deplasări) se poate utiliza funcția `rand`. De asemenea, pentru inițializarea parametrilor se poate apela funcția `init`.

În cazul unui neuron cu o singură intrare, pentru reprezentarea grafică tridimensională a suprafeței erorii pătratice globale (ca funcție de ponderea și deplasarea neuronului) se pot utiliza funcțiile `errsurf` și `plotes` (apelate în această ordine).

Funcția `train` antrenează rețeaua conform rutinei de antrenare specificate la creare și a parametrilor de antrenare aleși.

Algoritmul de antrenare prin propagare inversă descris mai sus este implementat de rutina `'traingd'` (se bazează pe metoda *gradient descent* de minimizare a unei funcții – algoritm lent). Dintre parametrii asociați acestui algoritm amintim `net.trainParam.lr` (valoarea ratei de învățare utilizată în antrenare), `net.trainParam.epochs` (numărul de epoci de antrenare), `net.trainParam.goal` (valoarea de prag impusă pentru criteriul de performanță minimizat `net.performFcn`) și `net.trainParam.show` (numărul de epoci de antrenare după care se prezintă grafic progresul antrenării).

Două variante ale acestui algoritm sunt implementate în funcțiile `'traingdm'` (*gradient descent with momentum* – în general, mai rapid decât `traingd`) și `'traingdx'` (*gradient descent with adaptive learning rate* – de asemenea, mai rapid decât `traingd`). Metodele de minimizare bazate pe direcții de căutare conjugate gradientului sunt implementate prin funcțiile de antrenare `'traincgf'` (metoda Fletcher-Reeves – cel mai mic necesare de memorie), `'traincgp'` (metoda Polak-Ribière – convergență mai rapidă în unele cazuri), `'traincgb'` (metoda Powell-Beale – convergență rapidă în general), `'trainscg'` (*scaled conjugate gradient* – algoritm foarte bun pentru probleme generale).

Un algoritm de antrenare mai rapid pentru rețele de dimensiuni moderate este algoritmul Levenberg-Marquardt, implementat prin rutina `'trainlm'` (prezintă facilități de reducere a memoriei utilizate în cazul când setul de antrenare are lungime mare). Funcția `'trainbr'` (*Bayesian regularization*) implementează o variantă modificată a acestui algoritm.

Algoritmul de antrenare de tip quasi-Newton bazat pe metoda Broyden-Fletcher-Goldfarb-Shanno este implementat prin funcția `'trainbfgs'` (necesită stocarea aproximării matricei Hessian și efectuează mai multe calcule la fiecare iterație decât algoritmi de tip gradient descent). Funcția `'trainoss'` (metoda *one step secant*) realizează un compromis între metodele de tip gradient conjugat și cele de tip quasi-Newton.

Pentru simularea rețelei antrenate se poate folosi funcția `sim`.

Pentru aprofundarea noțiunilor prezentate în acest capitol se recomandă studierea problemelor cu soluții analitice 7.9 – 7.16 și a problemelor cu soluții numerice 8.12 – 8.16.

Capitolul 6

APLICAȚII ALE REȚELELOR FEEDFORWARD MULTISTRAT ÎN IDENTIFICAREA ȘI CONTROLUL SISTEMELOR DINAMICE

6.1. Identificarea sistemelor dinamice utilizând rețele neurale

Identificarea sistemelor dinamice bazată pe modele neurale s-a conturat ca direcție de cercetare la începutul anilor '90, drept urmare a investigațiilor matematice asupra proprietăților de aproximare ale rețelelor neuronale. Evoluția noului domeniu a fost impulsionată de contribuțiile remarcabile ale unor nume de prestigiu în Automatică, ca, de exemplu, K. Narendra, K. Parthasarathy, P.J. Gawthrop, T.J. McAvoy, L. Ljung. Pe de altă parte, firmele producătoare de software tehnico-științific au demarat dezvoltarea de facilități specifice rețelelor neuronale, astfel încât apariția în 1992 a primei versiuni a **Neural Network Toolbox** încorporată în mediul MATLAB 4.2 a avut un impact major asupra interesului acordat de către automatizști acestei direcții de cercetare.

Modelele generale de identificare parametrică a sistemelor dinamice (liniare sau neliniare) se pot implementa, atât hard cât și soft, cu ajutorul rețelelor neurale artificiale. Parametrii modelului neural sunt reprezentați de ponderile sinaptice ale neuronilor și deplasările acestora. Teoria generală privind identificarea sistemelor se poate aplica direct pentru obținerea modelului neural ținând cont de echivalențele:

structura modelului = arhitectura rețelei neurale

estimare = antrenare

validare = generalizare

set de date de estimare = set de date de antrenare

set de date de validare = set de date de generalizare

6.1.1. Structura modelelor neurale

Considerăm un sistem dinamic cu o intrare și o ieșire (figura 6.1.1) pentru care se măsoară semnalele de intrare u și de ieșire y la momentele discrete de timp $t_k = kT$, (unde T reprezintă perioada de eșantionare a semnalelor, $k \in \mathbb{N}$), notând $u[k] = u(kT)$, $y[k] = y(kT)$, $k = \overline{1, N}$.



Figura 6.1.1. Sistem dinamic SISO

În problema identificării sistemelor dinamice, unde neliniaritatea care trebuie modelată depinde de timp, problema care se pune este cea a reprezentării explicite a timpului în structura rețelei neurale. Pentru ca o rețea neurală să fie dinamică, ea trebuie să posede o memorie. Ideea de la care s-a pornit la începutul anilor 1990 a fost cea de a se utiliza rețele neurale statice împreună cu un sistem extern de întârziere a semnalelor de intrare în rețea, obținând așa numitele *rețele neurale cu dinamică externă*.

Un interes practic deosebit îl prezintă obținerea modelelor neurale de tip ARX (NNARX) în timp discret, cu perioadă de eșantionare dată:

$$\hat{y}[k] = f(y[k-1], y[k-2], \dots, y[k-n], u[k-d], u[k-d-1], \dots, u[k-d-m]), \quad (6.1.1)$$

unde $d, n \in \mathbb{N}^*$, $m \in \mathbb{N}$, și:

$$f: \mathbb{R}^{n+m+1} \rightarrow \mathbb{R} \quad (6.1.2)$$

este o aplicație (liniară sau neliniară, netedă) ce descrie transferul intrare-ieșire al unei rețele neurale statice cu topologie adecvată. Această funcție poate fi reprezentată cu ajutorul unei rețele neurale statice cu propagare înainte a semnalului cu funcții de activare liniare (ADALINE) sau neliniare (MLP). Rețeaua are ca intrări valorile de la momentele anterioare de timp ale intrării și ieșirii procesului. Valorile semnalelor de intrare și de ieșire la momentele anterioare de timp formează vectorul *regresorilor* care este notat $\phi[k] = [y[k-1], y[k-2], \dots, y[k-n], u[k-d], u[k-d-1], \dots, u[k-d-m]]^T$ și se obține utilizând un sistem de întârziere (*Tapped Delay Line*) plasat în afara rețelei neurale propriu-zise.

Configurația prezentată în figura 6.1.2 este cunoscută și sub denumirea de *modelul serie-paralel* (SPM) deoarece în raport cu intrarea $u[k]$ rețeaua neurală este conectată paralel cu procesul, iar în raport cu ieșirea $y[k]$ conectarea este în serie. Eroarea de predicție $e[k]$, diferența dintre ieșirea reală a procesului $y[k]$ și ieșirea estimată cu ajutorul rețelei neurale $\hat{y}[k]$, este utilizată pentru antrenarea rețelei aplicând algoritmi de antrenare uzuali, independenți de timp. În această figură, z^{-1} reprezintă operatorul de întârziere cu un pas, $z^{-1}u[k] = u[k-1]$, iar ZOH notează extrapolatorul de ordin zero (*Zero Order Hold*) utilizat pentru discretizarea, cu perioada de eșantionare dorită, a unui proces continuu.

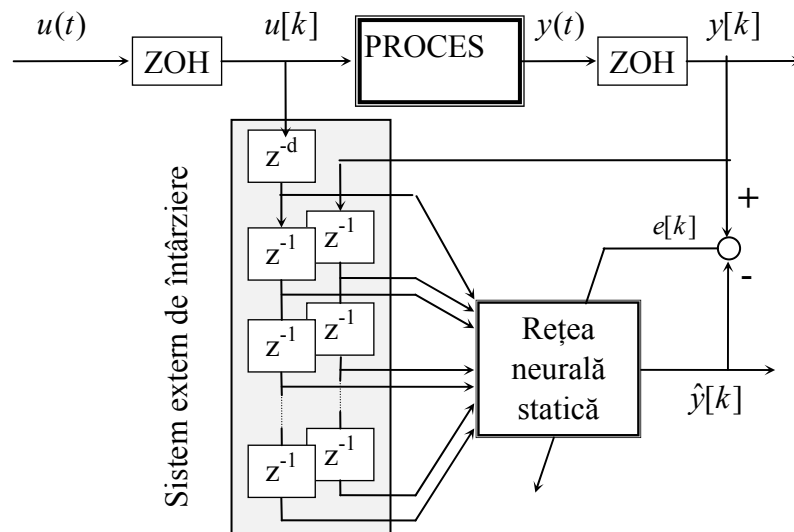


Figura 6.1.2. Modelul neural de tip ARX (NNARX) (modelul serie-paralel)

După antrenarea rețelei neurale, modelul identificat poate fi folosit independent de proces, aducând la intrarea rețelei ieșirea acesteia, printr-o conexiune inversă externă. O astfel de configurație poartă numele de *modelul paralel* și este utilizată pentru simularea proceselor dinamice. Configurația poate fi utilizată și în faza de antrenare a rețelei neurale, modelul obținut în acest fel fiind de fapt modelul neural de tip *output error* (NNOE) (figura 6.1.3)

$$\hat{y}[k] = f(\hat{y}[k-1], \hat{y}[k-2], \dots, \hat{y}[k-n], u[k-d], u[k-d-1], \dots, u[k-d-m]). \quad (6.1.3)$$

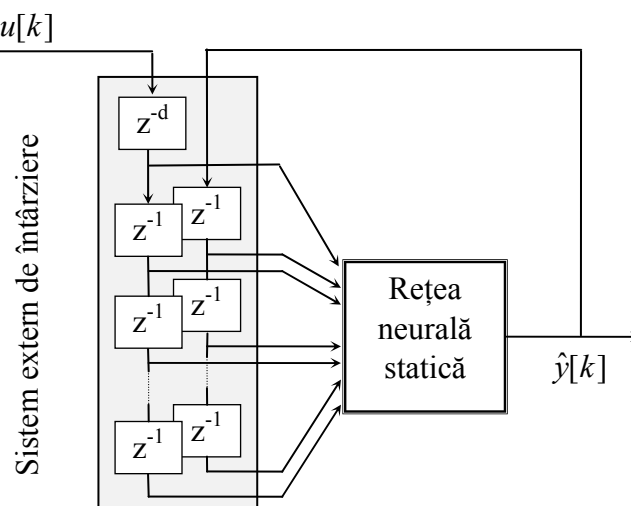


Figura 6.1.3. Modelul neural OE (NNOE) (modelul paralel)

În mod asemănător se pot obține și modelele neurale de tip ARMAX (NNARMAX) sau Box-Jenkins (NNBJ). Dezavantajele modelelor recursive de tip NNOE, NNARMAX sau NNBJ sunt posibila instabilitate a modelului și necesitatea unui efort suplimentar pentru a calcula gradientul ce intervine în procedura iterativă de antrenare a rețelei.

Pentru reprezentarea modelelor neliniare de tip NNARX cu o singură intrare și o singură ieșire Narendra și Parthasarathy (1990) au introdus patru modele. Modelele propuse de ei sunt descrise de ecuațiile:

$$\text{Modelul 1:} \quad \hat{y}[k+1] = \sum_{i=0}^{n-1} \alpha_i y[k-i] + g(u[k], u[k-1], \dots, u[k-m+1]), \quad (6.1.4)$$

$$\text{Modelul 2:} \quad \hat{y}[k+1] = f(y[k], y[k-1], \dots, y[k-n+1]) + \sum_{i=0}^{m-1} \beta_i u[k-i], \quad (6.1.5)$$

$$\text{Modelul 3:} \quad \hat{y}[k+1] = f(y[k], y[k-1], \dots, y[k-n+1]) + g(u[k], u[k-1], \dots, u[k-m+1]), \quad (6.1.6)$$

$$\text{Modelul 4:} \quad \hat{y}[k+1] = f(y[k], y[k-1], \dots, y[k-n+1], u[k], u[k-1], \dots, u[k-m+1]), \quad (6.1.7)$$

Se presupune că funcțiile f și g sunt diferențiabile în raport cu toate argumentele lor. Aceste funcții pot fi approximate cu ajutorul unor rețele neurale statice. Alegerea modelului adecvat depinde de informația despre procesul considerat disponibilă a priori. În modelele 1 (figura 6.1.4) și 2 (figura 6.1.5) se presupune că neliniaritatea procesului implică numai intrarea, respectiv numai ieșirea.

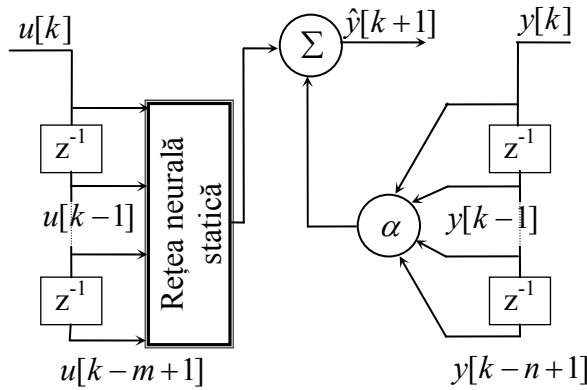


Figura 6.1.4. Reprezentarea modelului 1.

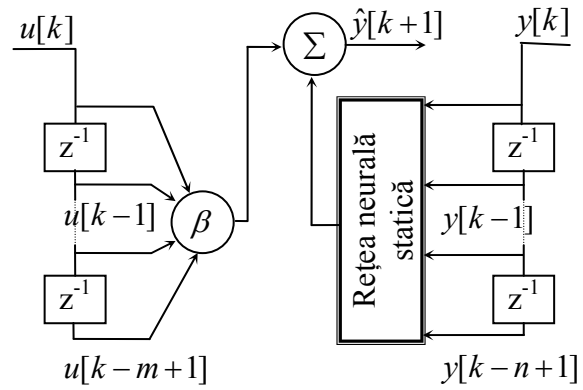


Figura 6.1.5. Reprezentarea modelului 2.

Modelul 3 (figura 6.1.6) rezultă ca suma a două funcții neliniare distincte care depind numai de semnalul de intrare, respectiv ieșire. Modelul 4 (figura 6.1.7), cel mai general, reprezintă proiecția neliniară atât a intrărilor cât și a ieșirilor procesului de identificat.

Modelele de tip NNARX pot fi utilizate pentru probleme simple, dar utilizarea lor nu este avantajoasă pentru probleme ce implică dependența semnalului de ieșire al sistemului de un număr mare de valori trecute ale semnalului de intrare. Pentru modelarea unor asemenea sisteme se recomandă utilizarea modelelor recursive.

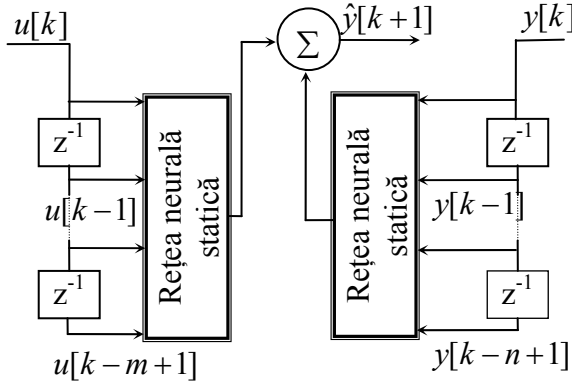


Figura 6.1.6. Reprezentarea modelului 3

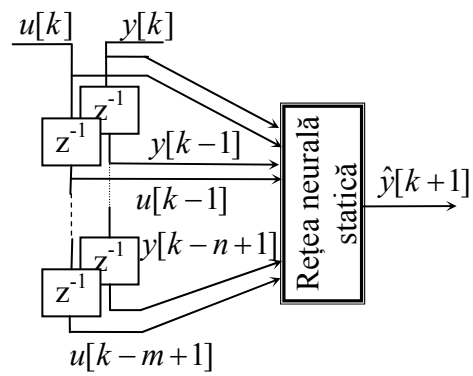


Figura 6.1.7. Reprezentarea modelului 4

O altă posibilitate de modelare a unui sistem dinamic neliniar revine la a construi în prima etapă un model liniar al sistemului. Reziduurile obținute ca diferență între ieșirile actuale ale sistemului și ieșirile estimate pe baza modelului liniar vor conține evident neliniaritățile nemodelate. Pentru a modela aceste neliniarități se construiește apoi un model neural având ca intrări intrările sistemului și reziduurile.

6.1.2. Estimarea modelelor neurale

Aplicația f din relația (6.1.2) este parametrizată de ponderile și deplasările rețelei. Indiferent de topologia rețelei, identificarea pe baza așa-numitei scheme serie-paralel (figura 6.1.2) asigură stabilitatea algoritmului iterativ de minimizare a funcțiilor obiectiv definite cu ajutorul erorii $e[k] = y[k] - \hat{y}[k-1]$.

Formularea concretă a funcțiilor obiectiv este corelată cu modul de obținere și exploatare a datelor experimentale corespunzătoare procesului. În acest sens, literatura evidențiază două categorii de metode de antrenare a rețelei neuronale din schema serie-paralel:

- pe lot (*batch*) sau *off-line*, care utilizează funcții obiectiv de forma (până la o constantă multiplicativă):

$$J = \frac{1}{N} \sum_{k=1}^N e^2[k] \leq \varepsilon, \quad (6.1.8)$$

unde N notează numărul total de elemente disponibile în lot și ε este valoarea de prag dorită;

- pe eşantioane (*pattern*) sau *on-line*, care utilizează funcții obiectiv de forma (până la o constantă multiplicativă):

$$J(i) = \frac{1}{N_E} \sum_{k=i}^{i+N_E-1} e^2[k] \leq \varepsilon, \quad i \in \mathbb{N}, \quad (6.1.9)$$

unde N_E notează numărul de eşantioane conținute în fereastra mobilă de selectare a elementelor, i este indicele iterației curente și ε este valoarea de prag dorită.

Minimizarea operează în spațiul parametrilor rețelei (ponderi și deplasări) și poate fi condusă prin tehnicile clasice de optimizare fără restricții, adaptate la specificul organizării matriceal-vectoriale a parametrilor aplicației f . În limbajul propriu rețelelor neuronale, determinarea parametrilor funcției f ca rezultat al minimizării funcțiilor obiectiv (5.6) sau (5.7) (estimarea modelului parametric) este referită drept proces de antrenare sau învățare.

Observație:

Determinarea modelului neural urmează pașii prezentați în capitolul 2: alegerea regresorilor modelului, a tipului neuronilor din structura rețelei neurale, alegerea numărului de neuroni din stratul intern și apoi determinarea ponderilor rețelei (adică determinarea parametrilor modelului). Principiul care se aplică este de a porni de la simplu la complex (lama lui Occam): se încearcă întâi cu un număr redus de regresori și de neuroni în stratul de intrare și apoi se crește progresiv numărul neuronilor până la un număr rezonabil. Dacă nu se determină un model suficient de bun se crește numărul regresorilor și se încearcă, din nou, valori crescute progresiv ale numărului de neuroni. Dacă nici în acest caz nu se determină un model corespunzător, se alege un alt tip de model pentru sistemul respectiv. ■

6.2. Controlul predictiv al sistemelor dinamice neliniare utilizând modele neurale

Controlul predictiv este de acum un standard larg utilizat în industrie și un număr foarte mare de algoritmi au fost prezentați în literatura de specialitate ca de exemplu: *Generalized Predictive Control* – GPC, *Self - Adaptive Control* – EPSAC și *Unified Predictive Control* – UPC. Marele avantaj al folosirii controlului predictiv provine din utilizarea de către regulator a unui model exact al procesului real, ceea ce duce la îmbunătățirea calității controlului. Scopul algoritmilor de control predictiv bazați pe model (eng. *Model Based Predictive Control* – MBPC) îl constituie găsirea unei secvențe de intrare ce va fi aplicată procesului controlat astfel încât ieșirea sa să urmărească cu eroare cât mai mică o traiectorie dorită (referința) pe durata unui orizont de timp prestabilit (Lazăr, 1999).

Principiul de bază al algoritmilor de control predictiv constă în prezicerea ieșirii procesului, pe baza unui model (obținerea predictorilor), și pe minimizarea unei anumite funcții de cost. Diferențele dintre diverșii algoritmi sunt date de modelul folosit de controller și de forma funcției de cost. În cazul sistemelor liniare, majoritatea modelelor dinamice utilizate în controlul predictiv se obțin prin teste aplicate proceselor sau prin metode de identificare. Pentru controlul predictiv neliniar testarea și identificarea proceselor devin mult mai dificile. Datorită proprietății rețelelor neurale de a fi aproximatori generali, utilizarea unui model neuronal pentru obținerea predictorilor prezintă un foarte mare interes.

O reprezentare grafică a strategiei de control predictiv este prezentată în figura 6.2.1. La fiecare moment de eșantionare k se determină strategia optimă de control, fiind calculate valorile viitoare ale intrării $u[k+1]$, $u[k+2]$, ..., $u[k+N_u]$ (N_u – orizontul de control), astfel

încât ieșirea procesului y să urmărească cu eroare cât mai mică referința r pe intervalul de predicție $[N_1, N_2]$ (N_1 – orizontul minim de predicție, N_2 – orizontul de predicție).

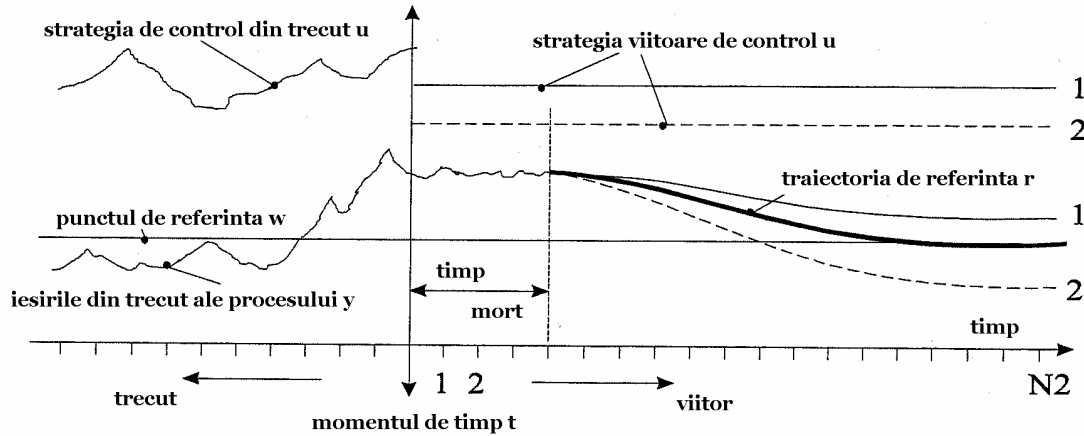


Figura 6.2.1. Strategia controlului predictiv bazat pe model.

Valorile viitoare ale ieșirii procesului sunt determinate folosind un model neural, iar pe baza erorilor între ieșirile prezise și referință este calculată o funcție de cost. Această funcție de cost urmează a fi minimizată cu scopul obținerii unei strategii de control optime ce vor fi aplicate procesului. Una din formele posibile ale funcției de cost, folosită în majoritatea algoritmilor de control predictiv, este dată de relația:

$$J = \sum_{i=N_1}^{N_2} (r[k+i] - y[k+i])^2 + \lambda \sum_{i=1}^{N_u} (u[k+i] - u[k+i-1])^2. \quad (6.2.1)$$

De multe ori, în practică, ponderea λ asupra comenzii are valoarea 0, iar orizontul minim de predicție (N_1) ia valoarea 1.

Funcția de cost J urmează a fi minimizată la orice moment de timp k în funcție de vectorul $\mathbf{u} = [u[k+1], u[k+2], \dots, u[k+N_u]]^T$ ce conține valorile viitoare ale intrării.

Modelul neural folosit în calculul ieșirilor viitoare ale procesului este un model de tip ARX de forma (6.1.1). Pentru vectorii conținând intrările, respectiv ieșirile viitoare ale procesului vom folosi următoarele notații:

$$\begin{aligned} \mathbf{u}[k+i-d] &= [u[k+i-d], u[k+i-d-1], \dots, u[k+i-d-m]]^T, \\ \mathbf{y}[k+i-1] &= [y[k+i-1], y[k+i-2], \dots, y[k+i-n]]^T, \end{aligned} \quad (6.2.2)$$

unde $i = N_1, \dots, N_2$, reprezintă ordinul predictorului, k este momentul curent de timp, iar d , m și n reprezintă parametrii considerați pentru modelul NNARX antrenat: d - timpul mort, m și n ordinele de întârziere ale intrării, respectiv ieșirii. Aplicând vectorii (6.2.2) modelului neural, se obține predictorul de ordin i al ieșirii, $y[k+i]$, după cum rezultă și din figura 6.2.2.

În calculul predictorilor ieșirii, dacă orizontul de predicție este mai mare decât orizontul de control sumat cu timpul mort considerat în modelul procesului ($N_2 > N_u + d$), după iterația $N_u + d$, intrarea va fi considerată constantă, având valoarea $u[k+N_u]$.

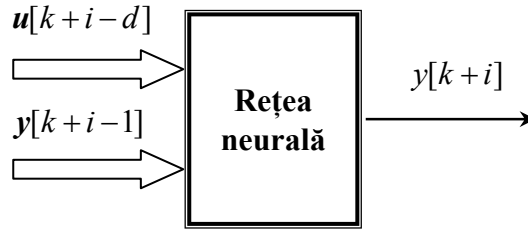


Figura 6.2.2. Calculul predictorilor neurali.

Algoritmul iterativ de control predictiv poate fi ușor înțeles prin descrierea pașilor ce sunt parcurși la orice moment de timp (iterație) k :

- în iterația anterioară (momentul de timp $k-1$), prin minimizarea funcției de cost, a fost obținut vectorul comenzii $\mathbf{u}_{old} = [u[k], u[k+1], \dots, u[k+N_u-1]]^T$;
- primul element al acestui vector, $u[k]$, reprezintă intrarea procesului la momentul de timp curent;
- sunt calculați predictorii ieșirii $y[k+i]$ de ordin $i = N_1, \dots, N_2$, folosind ca intrări ale modelului neural perechile de vectori $\mathbf{u}[k+i-d]$ și $\mathbf{y}[k+i-1]$;
- se calculează funcția de cost J la pasul curent, după formula (6.2.1);
- noua secvență optimă de control se obține prin minimizarea funcției de cost J în raport cu vectorul comenzii $\mathbf{u}_{new} = [u[k+1], u[k+2], \dots, u[k+N_u]]^T$.

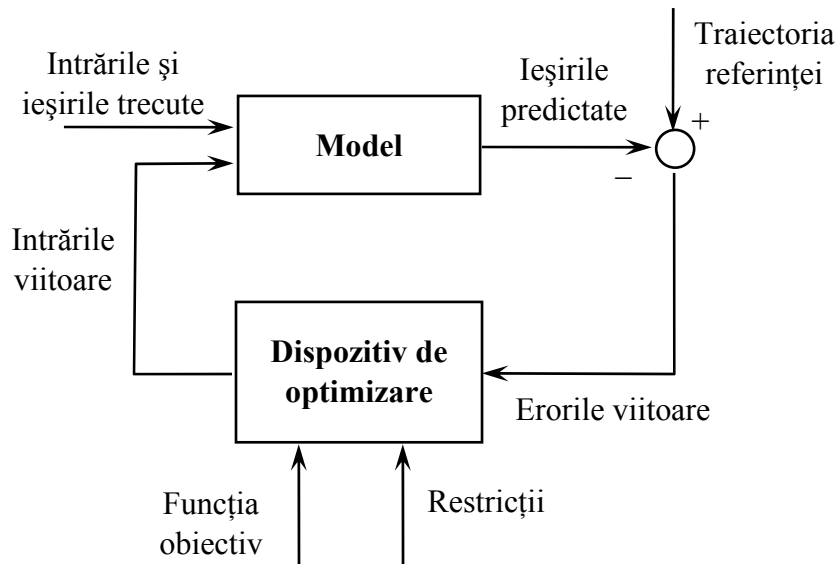


Figura 6.2.3: Structura de bază a unui controller predictiv

La momentul de timp următor ($k+1$), primul element din vectorul $\mathbf{u}_{new}$, $u[k+1]$, este aplicat la intrarea procesului și algoritmul de control predictiv continuă cu iterația $k+1$. La startarea algoritmului, utilizatorul trebuie să aleagă valoarea inițială $u[0]$ a intrării procesului.

Intrările ce urmează a fi aplicate procesului real pot fi supuse unor restricții ce decurg din natura fizică a acestuia. Unii algoritmi de control predictiv permit ca aceste restricții să fie luate în considerare la minimizarea funcției de cost.

Având în vedere principiile MBPC prezentate, structura unui bloc de control ce implementează un astfel de algoritm poate fi prezentată ca în figura 6.2.3:

Viteza de calcul a unui astfel de controller trebuie să fie suficient de ridicată, pentru a permite efectuarea pașilor corespunzători unei iterații a algoritmului de control într-un timp mai mic decât perioada de eșantionare aleasă pentru interfațarea cu procesul real.

6.3. Blocuri Simulink pentru aplicații ale rețelelor neurale în identificarea și conducerea sistemelor

Figura 6.3.1 prezintă o bibliotecă de blocuri Simulink pentru identificarea, simularea și controlul predictiv al sistemelor pe baza modelelor neurale dezvoltată în cadrul Catedrei de Automatică și Informatică Aplicată de la Facultatea de Automatică și Calculatoare din Iași (Kloetzer *et al.*, 2001), (Kloetzer, Păstrăvanu, 2004).

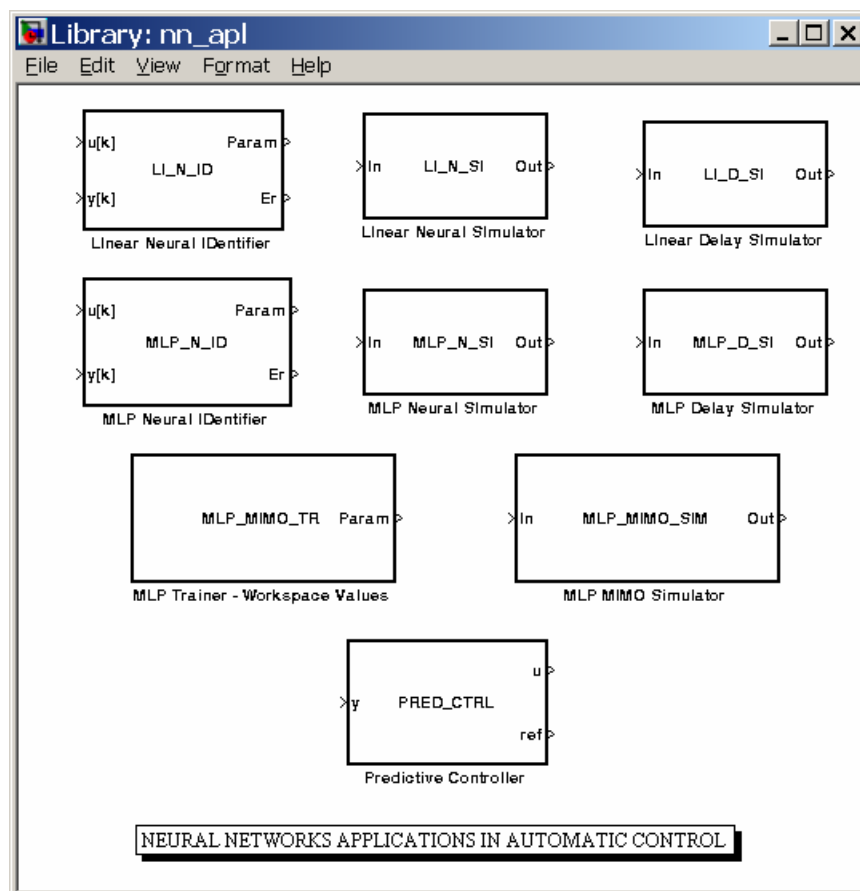


Figura 6.3.1. Bibliotecă de blocuri Simulink pentru identificarea, simularea și controlul predictiv al sistemelor pe baza modelelor neurale.

Dezvoltarea acestei biblioteci de blocuri Simulink pentru identificarea, simularea și controlul predictiv al sistemelor pe baza modelelor neurale a fost motivată de dorința de a ușura descrierea procesului supus identificării (așa cum se va vedea în continuare), lăsând utilizatorului, în același timp, flexibilitate în organizarea simulării. Mai mult chiar, posibilitatea executării simulărilor în timp real, prin utilizarea Real-Time Workshops, extinde aria de aplicabilitate a procedurilor de identificare la procese reale, interfațate adecvat prin cartele de achiziție.

Biblioteca Simulink din figura 6.3.1 constă din blocuri ce permit identificarea și simularea modelelor neurale cu o singură intrare și o singură ieșire (SISO) de forma (6.1.1) și anume:

- (i) pentru funcția f liniară:
 - blocul **LI_N_ID** (*Linear Neural IDentifier*) folosit pentru antrenarea unei rețele ADALINE,
 - blocul **LI_N_SI** (*Linear Neural Simulator*) folosit pentru simularea transferului intrare-ieșire a unui model neural cu topologie ADALINE, identificat anterior,
 - blocul **LI_D_SI** (*Linear Delay Simulator*) folosit pentru simularea transferului intrare-ieșire a unui model neural de tip ADALINE, cu timp mort, identificat anterior;
- (ii) pentru funcția f neliniară:
 - blocul **MLP_N_ID** (*MLP Neural IDentifier*) servește la antrenarea rețelelor de tip MLP,
 - blocul **MLP_N_SI** (*MLP Neural Simulator*) servește pentru simularea transferului intrare-ieșire a unui model neural cu topologie MLP, identificat anterior,
 - blocul **MLP_D_SI** (*MLP Delay Simulator*) utilizat pentru simularea transferului intrare-ieșire a unui model neural cu topologie MLP, cu timp mort, identificat anterior.

Pentru identificarea și simularea modelelor neurale neliniare cu mai multe intrări și ieșiri (MIMO), această bibliotecă Simulink mai include:

- blocul **MLP_MIMO_TR** (*MLP MIMO Trainer*) servește la antrenarea rețelelor de tip MLP utilizând matrice din spațiul de lucru Matlab,
- blocul **MLP_MIMO_SI** (*MLP MIMO Simulator*) servește pentru simularea transferului intrare-ieșire a unui model neural cu topologie MLP, identificat anterior, cu ajutorul blocului **MLP_MIMO_TR**.

În plus, blocul **PRED_CTRL** permite implementarea unui controller predictiv pentru un sistem SISO utilizând un model neural, identificat anterior cu folosind blocul **MLP_N_ID**.

Casetele de dialog ale blocurilor **LI_N_ID** și **MLP_N_ID** pentru identificarea modelelor SISO sunt prezentate în figurile 6.3.2 și, respectiv, 6.3.4. Aceste blocuri permit atât antrenarea pe lot (*off-line*), în concordanță cu (6.1.8), cât și antrenare pe eșantioane (*on-line*), în concordanță cu (6.1.9), cu un număr de eșantioane N_E specificat de utilizator în caseta de dialog. Utilizatorul poate seta valoarea de prag pentru eroarea pătratică medie utilizată în antrenare și modul în care se formează vectorii regresorilor.

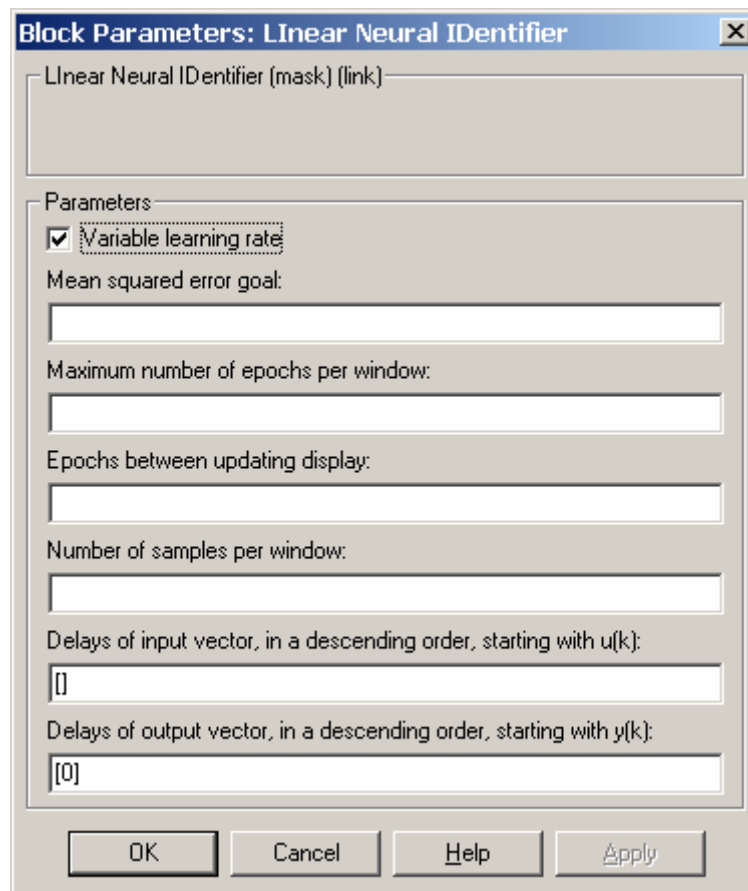


Figura 6.3.2. Caseta de dialog a blocului **LIN_N_ID**.

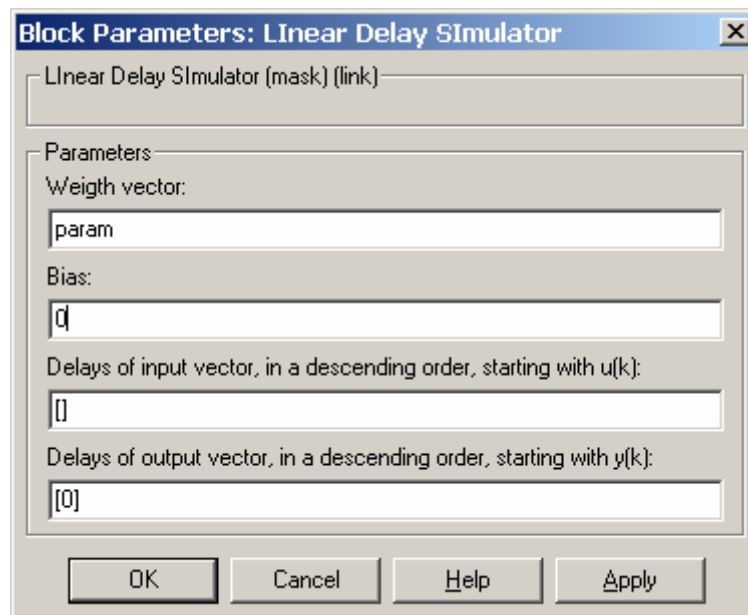


Figura 6.3.3. Caseta de dialog a blocului **LIN_D_SI**.

Block Parameters: MLP Neural Identifier [X]

MLP Neural Identifier (mask) (link)

Parameters

Number of input neurons:

Minimum input value:

Maximum input value:

Learning method: Levenberg-Marquardt ▾

Mean squared error goal:

Maximum number of epochs per window:

Epochs between updating display:

Number of samples per window:

Delays of input vector, in a descending order, starting with $u(k)$:

Delays of output vector, in a descending order, starting with $y(k)$:

☐ Initialize using INIT function

OK Cancel Help Apply

Figura 6.3.4. Caseta de dialog a blocului *MLP_N_ID*.

Block Parameters: MLP Neural Simulator [X]

MLP Neural Simulator (mask) (link)

Parameters

Parameter vector (resulting from training):

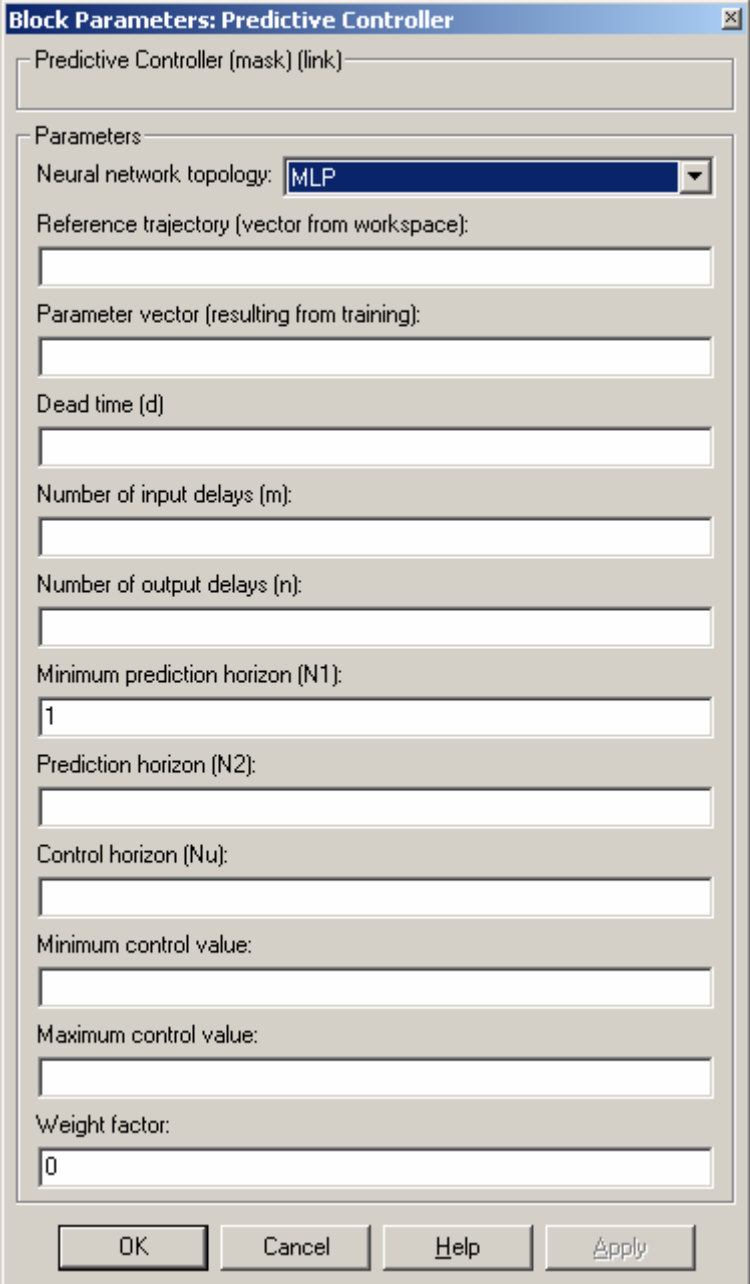
Number of input neurons:

OK Cancel Help Apply

Figura 6.3.5. Caseta de dialog a blocului *MLP_N_SI*.

Casetele de dialog ale blocurilor **LIN_D_SI** și **MLP_N_SI**, utilizate în simulare, sunt prezentate în figurile 6.3.3 și, respectiv, 6.3.5. Parametrii rețelelor pot fi setați manual sau pot fi încărcăți automat din spațiul de lucru al Matlab-ului.

În găsirea regresorilor optimi pentru vectorii de intrare și ieșire este recomandată o strategie *bottom-up*, în sensul că arhitectura rețelei va fi inițializată cu un număr redus de regresori. Pentru blocul de tip MLP se va porni cu un număr redus de neuroni pe stratul de intrare. Aceste valori vor incrementate succesiv până la obținerea rezultatului dorit.



The image shows a MATLAB dialog box titled "Block Parameters: Predictive Controller". It contains several input fields and a dropdown menu for configuring a predictive controller block. The "Neural network topology" dropdown is set to "MLP". The "Minimum prediction horizon (N1)" is set to 1, and the "Weight factor" is set to 0. Other fields like "Reference trajectory", "Parameter vector", "Dead time", "Number of input delays", "Number of output delays", "Prediction horizon (N2)", "Control horizon (Nu)", "Minimum control value", and "Maximum control value" are empty.

Parameter	Value
Predictive Controller (mask) (link)	
Neural network topology	MLP
Reference trajectory (vector from workspace)	
Parameter vector (resulting from training)	
Dead time (d)	
Number of input delays (m)	
Number of output delays (n)	
Minimum prediction horizon (N1)	1
Prediction horizon (N2)	
Control horizon (Nu)	
Minimum control value	
Maximum control value	
Weight factor	0

Figura 6.3.6. Caseta de dialog a blocului **PRED_CTRL**.

Blocul **PRED_CTRL** poate fi folosit pentru controlul proceselor neliniare ce au drept model neuronal o rețea de tip MLP. Caseta de dialog a acestui bloc este prezentată în figura 6.3.6, parametrii ce pot fi aleși de utilizator privind atât rețeaua neurală (presupusă deja antrenată), cât și algoritmul de control predictiv. Vectorul ce conține parametrii modelului neuronal are aceeași organizare ca și vectorul obținut în urma identificării folosind blocurile **MLP_N_ID**. Valorile întârzierilor intrării și ieșirii, precum și valoarea timpului mort, trebuie să fie aceleași cu cele considerate la crearea și antrenarea modelului NNARX.

Parametrilor corespunzători algoritmului de control predictiv ce pot fi setați de către utilizator sunt în concordanță cu notațiile utilizate în paragraful 6.2. Traectoria referinței este presupusă a fi cunoscută apriori, fiind specificată sub forma unui vector.

Ieșirea 'u' a blocului de control furnizează valoarea intrării de control la fiecare pas de simulare, iar ieșirea 'ref' poate fi utilă pentru a aprecia modul în care ieșirea procesului urmărește referința dorită pe durata simulării.

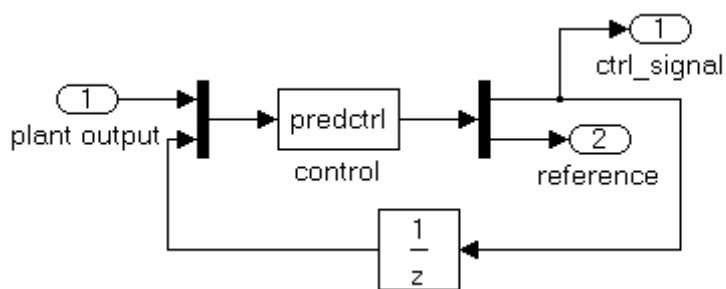


Figura 6.3.7. Structura internă a blocului **PRED_CTRL**.

În structura internă a blocului **PRED_CTRL** (figura 6.3.7) intră funcția-S denumită **predctrl** care este responsabilă atât de organizarea datelor interne pe baza parametrilor specifici modelului neuronal și algoritmului de control, cât și de minimizarea funcției de cost J (6.2.1). La fiecare pas de simulare (iterație a algoritmului de control predictiv) această funcție apelează o rutină **cost.m**, care realizează calculul predictorilor neuronali și al funcției de cost. Pentru crearea modelului neuronal s-au folosit funcțiile specifice din *Neural Network Toolbox*, iar pentru minimizarea funcției de cost s-a făcut apel la funcția **fmincon** din *Optimal Toolbox*, care permite impunerea unor limite (restricții) pentru variabila în raport cu care se efectuează minimizarea. La fiecare iterație a algoritmului este memorat vectorul obținut prin minimizarea lui J , iar acest vector va constitui punctul de start al minimizării la iterația următoare.

Conectarea blocului într-o schemă Simulink în vederea simulării controlului unui oarecare proces dinamic neliniar (pentru care se dispune de un model neuronal) este ilustrată în figura 6.3.8. În timpul simulării, poate fi urmărită evoluția intrării de control și urmărirea referinței de către ieșirea procesului.

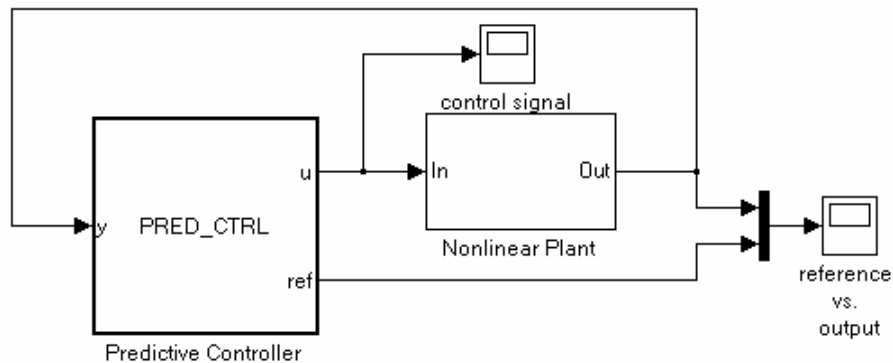


Figura 6.3.8. Modul de conectare a blocului **PRED_CTRL** într-o schemă Simulink.

6.4. Studii de caz

6.4.1. Exemple ilustrative pentru construcția modelelor

Deși utilizarea modelelor neuronale este atractivă îndeosebi în cazul neliniar, luăm în discuție, pentru început, construcția unui model liniar, deoarece acesta permite comparații relevante (de natură parametrică) cu reprezentările standard de tip funcție de transfer, discretă în timp. Al doilea exemplu ilustrează construcția unui model neliniar pentru care discuția privind validitatea se limitează doar la analiza comparativă a răspunsurilor obținute de la proces și respectiv model, în condițiile acelorași semnale de test aplicate la intrare. În nici unul dintre cazuri nu abordăm problematica validării modelelor cu ajutorul testelor statistice.

Pentru ambele exemple, funcționarea procesului supus identificării este simulată în mediul Simulink, antrenarea rețelelor nefăcând însă uz de cunoașterea ordinului modelelor de simulare (ipoteză absolut firească în desfășurarea aplicațiilor reale). Totodată, pentru fiecare din exemple au fost considerate, separat, două situații ce urmează a fi detaliate mai jos:

- (a) ieșirea procesului nu este afectată de perturbații;
- (b) ieșirea procesului este afectată aditiv de perturbații de tip Gaussian care pot atinge 1% din amplitudinea maximă a semnalului util (corespunzând, de exemplu, zgomotului de conversie al unui convertor A/D cu rezoluție scăzută).

A. Construcția unui model liniar

În figura 6.4.1 se prezintă schema Simulink (aferentă situației (b) menționate la începutul secțiunii curente) utilizată pentru construirea unui model liniar de forma (6.1.1) cu perioada de eșantionare $T=1$ secundă, cu ajutorul blocului **LI_N_ID**. Procesul ce face obiectul identificării este descris în Simulink prin funcția de transfer continuă în timp

$$G(s) = \frac{400e^{-15s}}{300s^2 + 40s + 1}, \quad (6.4.1)$$

procesul fiind excitat cu o sursă de zgomot alb de bandă limitată, adecvat aleasă. Singurele informații presupuse cunoscute apriori sunt comportarea liniară a procesului și valoarea timpul mort (identificabil, de exemplu, de pe un răspuns la semnal treaptă), adică $d = 15$ în (6.4.1) pentru $T = 1$ secundă.

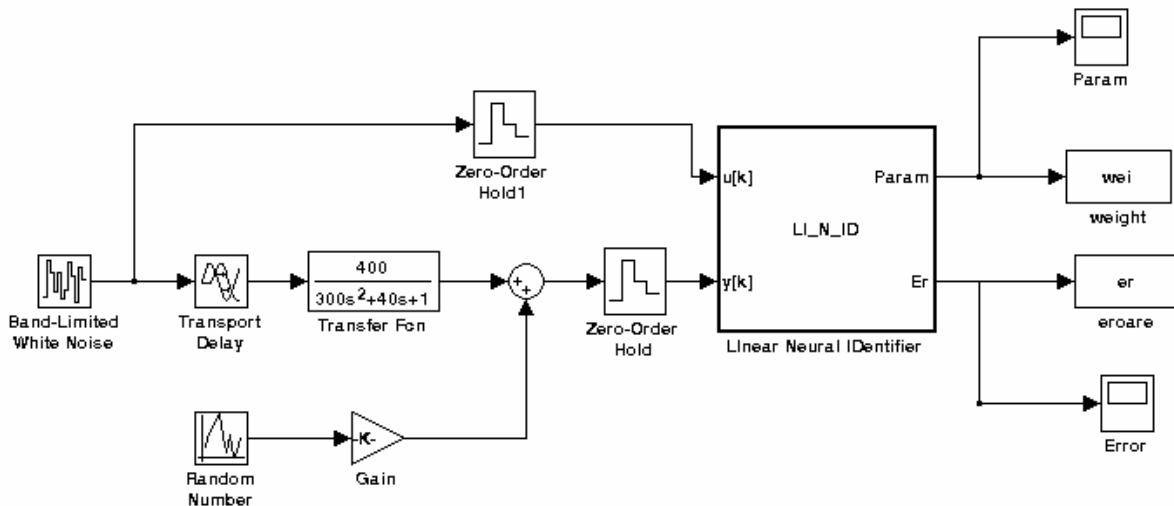


Figura 6.4.1. Schema Simulink de identificare a modelului liniar corespunzător exemplului din secțiunea A.

Selectăm drept complet edificatoare doar o parte din rezultate, prezentate în tabelul 6.4.1 pe care le exprimăm (dată fiind liniaritatea modelului) sub formă de funcții de transfer discrete (ale căror coeficienți sunt, până la un semn cunoscut, chiar parametrii furnizați de blocul **LI_N_ID**). Această exprimare facilitează analiza calității modelului provenit din identificare, prin comparație directă cu discretizată cu $T=1$ secundă a lui $G(s)$, obținută prin metoda Euler predictor

$$\tilde{G}(z^{-1}) = z^{-15} \frac{0.6378z^{-1} + 0.6101z^{-2}}{1 - 1.872z^{-1} + 0.8752z^{-2}}. \quad (6.4.2)$$

Toate cazurile comentate se referă la aceleași condiții de antrenare on-line cu $N_E=50$ eșantioane în funcția obiectiv J din (6.1.9), maxim 100 epoci/fereastră și rată de învățare variabilă. Rezultatele din tabelul 6.4.1 evidențiază rolul jucat de valoarea obținută în antrenare pentru funcția obiectiv J în aprecierea convergenței parametrilor rețelei și, totodată, în evaluarea calității modelului identificat.

Figurile 6.4.2.a și 6.4.2.b furnizează reprezentările grafice ale evoluțiilor parametrilor rețelei pentru o durată de 150 de secunde corespunzătoare cazului (b) pentru $n=2$ $m=2$ și, respectiv, pentru $n=2$ și $m=1$. Aceste reprezentări grafice ilustrează totodată și capacitatea blocurilor Simulink utilizate de a furniza on-line informații cantitative ce caracterizează progresul antrenării. Mai subliniem ca elemente importante în construcția modelului neuronal

faptul că în toate situațiile de convergență a parametrilor coeficientul termenului în z^0 al numărătorului (în exprimarea sub formă de funcție de transfer) rezultă numeric nul (raportat la precizia de calcul).

Tabelul 6.4.1. *Prezentare sintetică a rezultatelor antrenării pentru opt teste de identificare selectate ca relevante pentru secțiunea 6.4.1.*

				Cazul (a) - Ieșire neperturbată
d	n	m	$J(150)$	Exprimare sub formă de funcție de transfer
15	2	1	10^{-2}	Parametrii nu converg la valori staționare
15	2	2	10^{-9}	$z^{-15} \frac{0.6377z^{-1} + 0.6102z^{-2}}{1 - 1.872z^{-1} + 0.8751z^{-2}}$
15	3	2	10^{-2}	Parametrii nu converg la valori staționare
15	3	3	10^{-9}	$z^{-15} \frac{0.6378z^{-1} + 0.9392z^{-2} + 0.3148z^{-3}}{1 - 1.356z^{-1} - 0.0907z^{-2} + 0.4515z^{-3}} =$ $= z^{-15} \frac{0.6378(z + \mathbf{0.5160})(z + 0.9566)}{(z + \mathbf{0.5159})(z - 0.9665)(z - 0.9054)}$
				Cazul (b) - Ieșire perturbată
d	n	m	$J(150)$	Exprimare sub formă de funcție de transfer
15	2	1	10^{-2}	Parametrii nu converg la valori staționare
15	2	2	10^{-5}	$z^{-15} \frac{0.637z^{-1} + 0.6158z^{-2}}{1 - 1.8711z^{-1} + 0.8742z^{-2}}$
15	3	2	10^{-2}	Parametrii nu converg la valori staționare
15	3	3	10^{-5}	$z^{-15} \frac{0.6378z^{-1} + 0.9448z^{-2} + 0.3144z^{-3}}{1 - 1.3552z^{-1} - 0.0922z^{-2} + 0.4521z^{-3}} =$ $= z^{-15} \frac{0.6378(z + \mathbf{0.5048})(z + 0.9766)}{(z + \mathbf{0.5167})(z - 0.9043)(z - 0.9676)}$

Redundanța în stabilirea numărului de regresori conduce la apariția perechilor pol-zero cu valori numerice identice (raportate la precizia de calcul), fapt ilustrat în tabelul 6.4.1 pentru $n = m = 3$ în ambele cazuri (a) și (b).

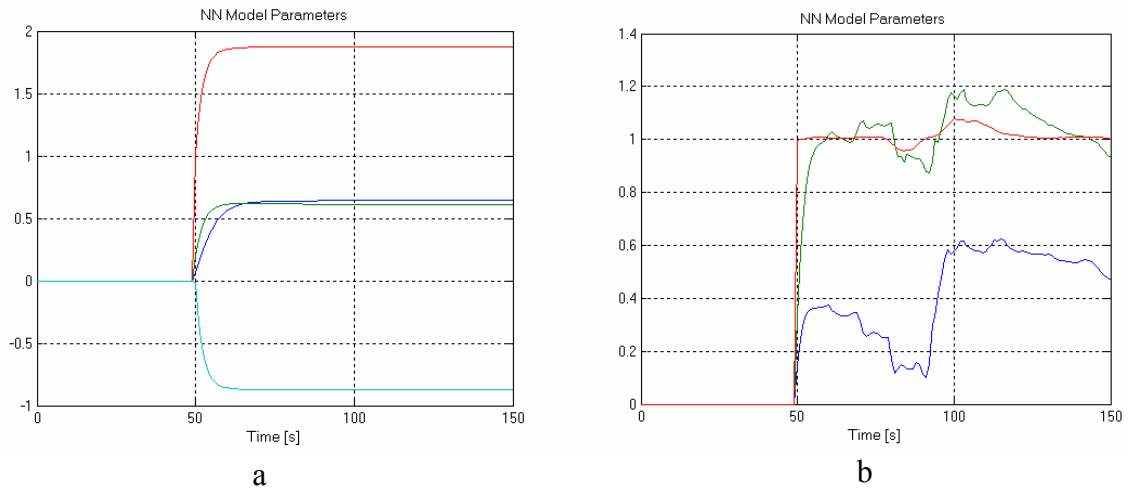


Figura 6.4.2. Evoluția parametrilor rețelei neuronale pentru exemplul din secțiunea 6.4.1.A corespunzând unui interval de 150 secunde: a) Cazul $n=2$ și $m=2$ (parametrii converg), b) Cazul $n=2$ și $m=1$ (parametrii nu converg).

B. Construcția unui model neliniar

Exemplul prezentat în continuare ilustrează utilizarea blocului **MLP_N_ID** în identificarea unui model neliniar de tipul (6.1.1) cu perioada de eșantionare $T = 1$ secundă. Procesul ce urmează a fi identificat este simulat în Simulink pe baza următoarei ecuații cu diferențe (recomandată în literatură ca exemplu relevant pentru identificare neurală (Liu, *et al.*, 1998)):

$$y(k) = 2.5 \frac{y(k-2)y(k-1)}{1 + y^2(k-1) + y^2(k-2)} + 0.3 \cos(0.5[y(k-1) + y(k-2)]) + 1.2u(k-1). \quad (6.4.3)$$

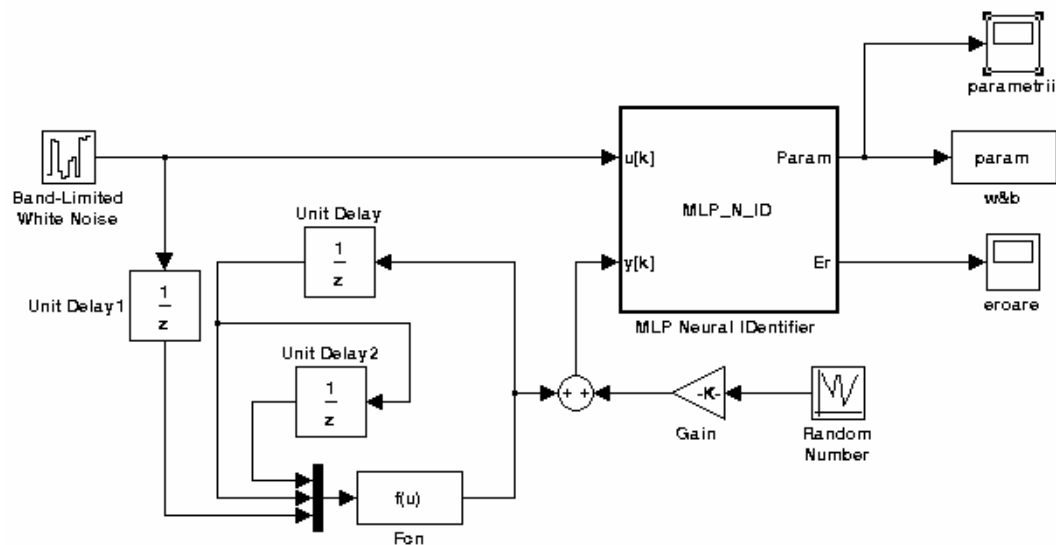


Figura 6.4.3. Schema Simulink de identificare a modelului neliniar corespunzător exemplului din secțiunea B.

Figura 6.4.3 prezintă diagrama Simulink utilizată pentru identificare. Semnalul de intrare este de tip „zgomot alb” adecvat ales. Toate cazurile ce vor fi comentate (extrase dintr-un set amplu de teste) se referă la aceleași condiții de antrenare on-line cu $N_E = 200$ eșantioane în funcția obiectiv J din (6.1.9), utilizând algoritmul Levenberg-Marquardt. Criteriile de oprire a antrenării au fost fie atingerea unui număr maxim de 100 epoci/fereastră sau o acuratețe de aproximare $\varepsilon = 10^{-9}$ pentru funcția obiectiv (6.1.9). Stratul de intrare conține 10 neuroni sigmoidali. Rezultatele obținute pentru $d = 0$, $m \geq 1$, $n \geq 2$, în (6.1.1) sunt comparabile, din punct de vedere al acurateței, pentru cazurile (a) și (b) având valori de ordinul 10^{-5} pentru $J(250)$ în cazul (a) și respectiv 10^{-3} în cazul (b). Cazurile cu d, m, n ce nu îndeplinesc condițiile menționate anterior pot fi ușor separate ele furnizând valori pentru $J(250)$ de ordinul $10^{-1} - 10^{-2}$.

Pentru a ilustra calitatea modelului neural rezultat din identificare, figura 6.4.4 prezintă grafice comparative pentru răspunsul procesului și a modelului pentru o intrare sinusoidală cu amplitudinea 0.7 (figura 6.4.4.a) și un semnal de intrare aleator, uniform distribuit în intervalul $[-1,1]$ (figura 6.4.4.b). Parametrii modelului neural sunt cei obținuți în situația (b) cu $d=0$, $m=1$, $n=2$ (adică numărul minim de regresori în (6.1.1) capabili să asigure valori favorabile pentru $J(250)$). Semnalele de test utilizate nu au fost „văzute” în timpul identificării. Folosirea numărului exact de regresori prezenți în ecuația (6.4.3) ($d=1$, $m=0$, $n=2$) nu a adus îmbunătățiri vizibile în acuratețea de aproximare a modelului obținut în comparație cu o serie de teste nedetaliată în această lucrare.

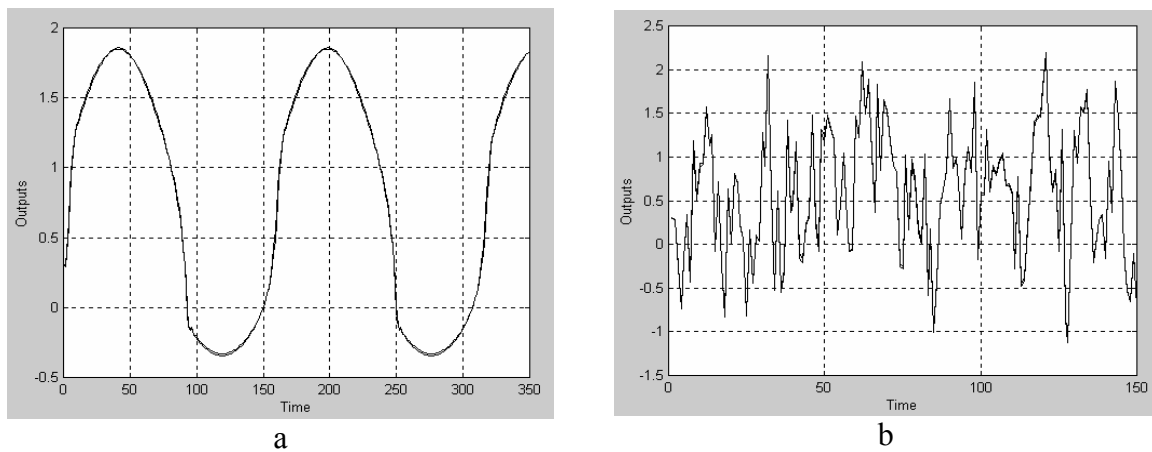


Figura 6.4.4: Grafice comparative cu răspunsurile procesului (linie continuă) și modelului neural MLP (linie întreruptă): a) intrare sinusoidală de amplitudine 0.7; b) semnal aleator uniform distribuit în $[-1,1]$.

În general vorbind, testele efectuate au arătat că modelele neurale nu sunt afectate semnificativ de folosirea regresorilor suplimentari în comparație cu folosirea numărului exact de regresori specificat anterior. Totuși, redundanța crește complexitatea modelului neural fără nici un beneficiu real, de aceea e preferabil de căutat un model apropiat de cel minimal.

6.4.2. Exemplu de utilizare a controlului predictiv bazat pe model neural

Modelul neliniar identificat în secțiunea precedentă, corespunzător sistemului descris de ecuația cu diferențe (6.4.2) este utilizat în continuare pentru a proiecta un controller predictiv pentru acest sistem utilizând blocul **PRED_CTRL**. Schema Simulink utilizată este prezentată în figura 6.4.5. Parametrii utilizați de algoritmul de control predictiv au următoarele valori: orizontul minim de predicție $N_1 = 1$, orizontul de predicție $N_2 = 3$, orizontul de control $N_u = 2$, factorul de ponderare $\lambda = 0$ și valoarea inițială a intrării $u_{\text{init}} = 0$.

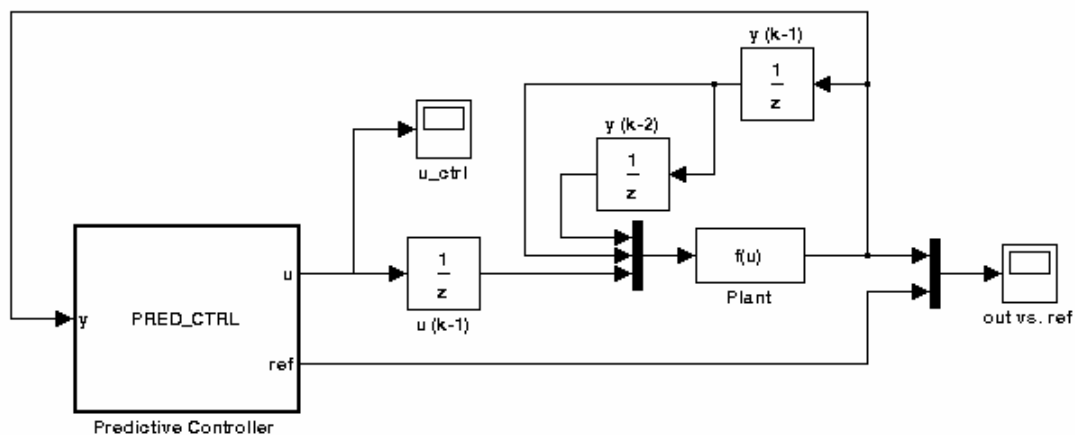


Figura 6.4.5. Modelul Simulink utilizat pentru controlul predictiv al sistemului neliniar (6.4.2).

Figura 6.4.6 ilustrează urmărirea unei referințe sinusoidale, cu amplitudine 1, fază 0, offset 0,2 și frecvență 0,005 Hz, iar figura 6.4.6 urmărirea unei referințe de formă unghiulară, cu valori cuprinse în intervalul $[-1,1]$.

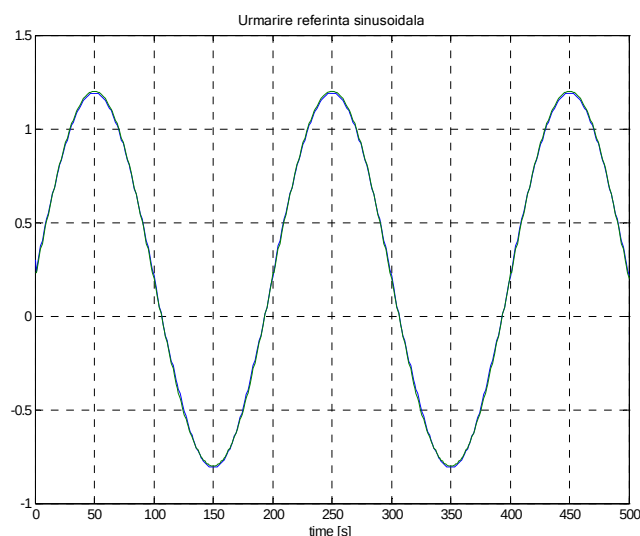


Figura 6.4.6. Grafice comparative ale unei referințe sinusoidale și a ieșirii sistemului neliniar controlat de un regulator predictiv.

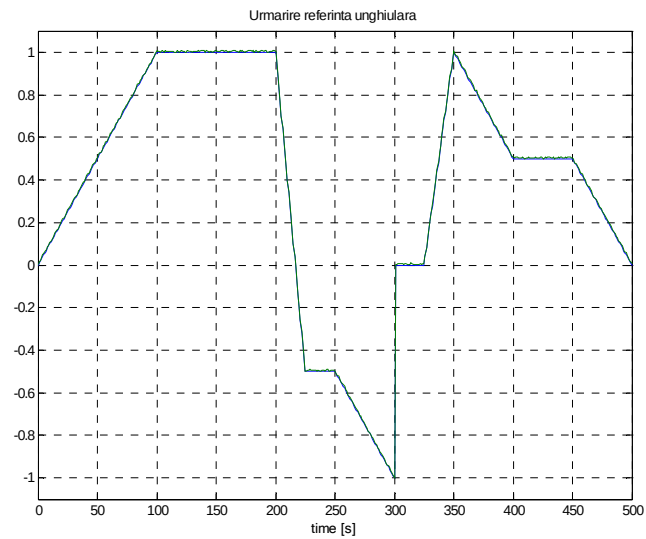


Figura 6.4.7. Grafice comparative ale unei referințe unghiulare și a ieșirii sistemului neliniar controlat de un regulator predictiv.

Pentru aprofundarea noțiunilor prezentate în acest capitol se recomandă studierea problemelor cu soluții analitice 7.17 – 7.20 și a problemelor cu soluții numerice 8.17, 8.18.

Capitolul 7

PROBLEME CU SOLUȚII ANALITICE

Problema 7.1.

Se consideră un neuron cu o singură intrare, cu funcția de activare liniară. La intrarea neuronului se aplică semnalul $x(t) = \cos \omega t$. Să se determine parametrii neuronului în cazul când la ieșire se obțin semnalele:

a) $y(t) = \cos^2 \frac{\omega t}{2}$;

b) $y(t) = -\cos^2 \frac{\omega t}{2}$;

c) $y(t) = \sin^2 \frac{\omega t}{2}$;

d) $y(t) = -\sin^2 \frac{\omega t}{2}$.

Indicații și răspunsuri:

Cu notațiile din figura 2.1.1, pentru un neuron cu ponderea w , deplasarea b și funcție de activare liniară, $\sigma(u) = u$, dependența ieșirii y a neuronului de intrarea x a acestuia este descrisă de:

$$y(x) = \sigma(u(x)) = wx + b, \quad \forall x \in \mathbb{R}.$$

a) Relația precedentă devine:

$$y(t) = wx(t) + b, \quad \forall t \in \mathbb{R} \Leftrightarrow \cos^2 \frac{\omega t}{2} = w \cos \omega t + b, \quad \forall t \in \mathbb{R}.$$

Ținând cont că $\cos \omega t = 2 \cos^2 (\omega t/2) - 1$, rezultă:

$$\cos^2 \frac{\omega t}{2} = w \left[2 \cos^2 \frac{\omega t}{2} - 1 \right] + b, \quad \forall t \in \mathbb{R} \Rightarrow \begin{cases} 2w = 1 \\ -w + b = 0 \end{cases} \Rightarrow \begin{cases} w = \frac{1}{2} \\ b = \frac{1}{2} \end{cases}$$

b) Similar, se ajunge la:

$$-\cos^2 \frac{\omega t}{2} = w \left[2 \cos^2 \frac{\omega t}{2} - 1 \right] + b, \forall t \in \mathbb{R} \Rightarrow \begin{cases} 2w = -1 \\ -w + b = 0 \end{cases} \Rightarrow \begin{cases} w = -\frac{1}{2} \\ b = -\frac{1}{2} \end{cases}$$

c) Utilizând relația $y(t) = \sin^2(\omega t/2) = 1 - \cos^2(\omega t/2)$, rezultă:

$$1 - \cos^2 \frac{\omega t}{2} = w \left[2 \cos^2 \frac{\omega t}{2} - 1 \right] + b, \forall t \in \mathbb{R} \Rightarrow \begin{cases} 2w = -1 \\ -w + b = 1 \end{cases} \Rightarrow \begin{cases} w = -\frac{1}{2} \\ b = \frac{1}{2} \end{cases}$$

d) Analog cazului c) se obține:

$$-1 + \cos^2 \frac{\omega t}{2} = w \left[2 \cos^2 \frac{\omega t}{2} - 1 \right] + b, \forall t \in \mathbb{R} \Rightarrow \begin{cases} 2w = 1 \\ -w + b = -1 \end{cases} \Rightarrow \begin{cases} w = \frac{1}{2} \\ b = -\frac{1}{2} \end{cases}$$

Reprezentările grafice cerute în enunțul problemei se obțin imediat.

Problema 7.2.

Se consideră un neuron cu o singură intrare și funcția de activare treaptă unipolară, la intrarea la intrarea căruia se aplică semnalul $x(t) = \sin 2\pi t$, $t > 0$. Să se exprime analitic și să se reprezinte grafic semnalul de ieșire $y(t)$ în următoarele cazuri:

- | | |
|------------------------|-------------------------|
| (i) $w = 1, b = 0$; | (ii) $w = -1, b = 0$; |
| (iii) $w = 1, b = 1$; | (iv) $w = -1, b = 1$; |
| (v) $w = 1, b = -1$; | (vi) $w = -1, b = -1$. |

Indicații și răspunsuri:

Ieșirea neuronului este dată de $y(t) = \sigma(wx(t) + b)$, unde σ este funcția de activare de tip treaptă unipolară:

$$\sigma(u) = \begin{cases} 0; & u < 0, \\ 1; & u \geq 0. \end{cases}$$

Se obține astfel

$$y(t) = \begin{cases} 0, & \text{dacă } w \sin 2\pi t + b < 0, \\ 1, & \text{dacă } w \sin 2\pi t + b \geq 0. \end{cases}$$

$$(i) \ w = 1 \text{ și } b = 0 \Rightarrow y(t) = \sigma(\sin 2\pi t) = \begin{cases} 1, & t \in [0, 1/2] + k, k \in \mathbb{Z}, \\ 0, & \text{în rest;} \end{cases}$$

$$(ii) \ w = -1 \text{ și } b = 0 \Rightarrow y(t) = \sigma(-\sin 2\pi t) = \begin{cases} 1, & t \in [1/2, 1] + k, k \in \mathbb{Z}, \\ 0, & \text{în rest;} \end{cases}$$

$$(iii) \ w = 1 \text{ și } b = 1 \Rightarrow y(t) = \sigma[\sin 2\pi t + 1] = 1, \forall t \in \mathbb{R} ;$$

$$(iv) \ w = -1 \text{ și } b = 1 \Rightarrow y(t) = \sigma[-\sin 2\pi t + 1] = 1, \forall t \in \mathbb{R} ;$$

- (v) $w = 1$ și $b = -1 \Rightarrow y(t) = \sigma[\sin 2\pi t - 1] = \begin{cases} 1, & t = k + 1/4, k \in \mathbb{Z}, \\ 0, & \text{în rest}; \end{cases}$
- (vi) $w = -1$ și $b = -1 \Rightarrow y(t) = \sigma[-\sin 2\pi t - 1] = \begin{cases} 1, & t = k + 3/4, k \in \mathbb{Z}, \\ 0, & \text{în rest}. \end{cases}$

Reprezentările grafice cerute în enunțul problemei se obțin imediat.

Problema 7.3.

Se consideră un neuron cu două intrări și funcția de activare de tip treaptă bipolară. Vectorii de la intrarea neuronului sunt de forma: $\mathbf{x} = [x_1 \ x_2]^T \in \mathbb{R}^2$.

Planul (x_1, x_2) poate fi separat în două semiplane în funcție de valorile ieșirii y a rețelei astfel: într-unul din aceste semiplane $y = -1$, iar în celălalt $y = 1$.

Să se determine parametrii neuronului astfel încât să aibă loc următoarele separări:

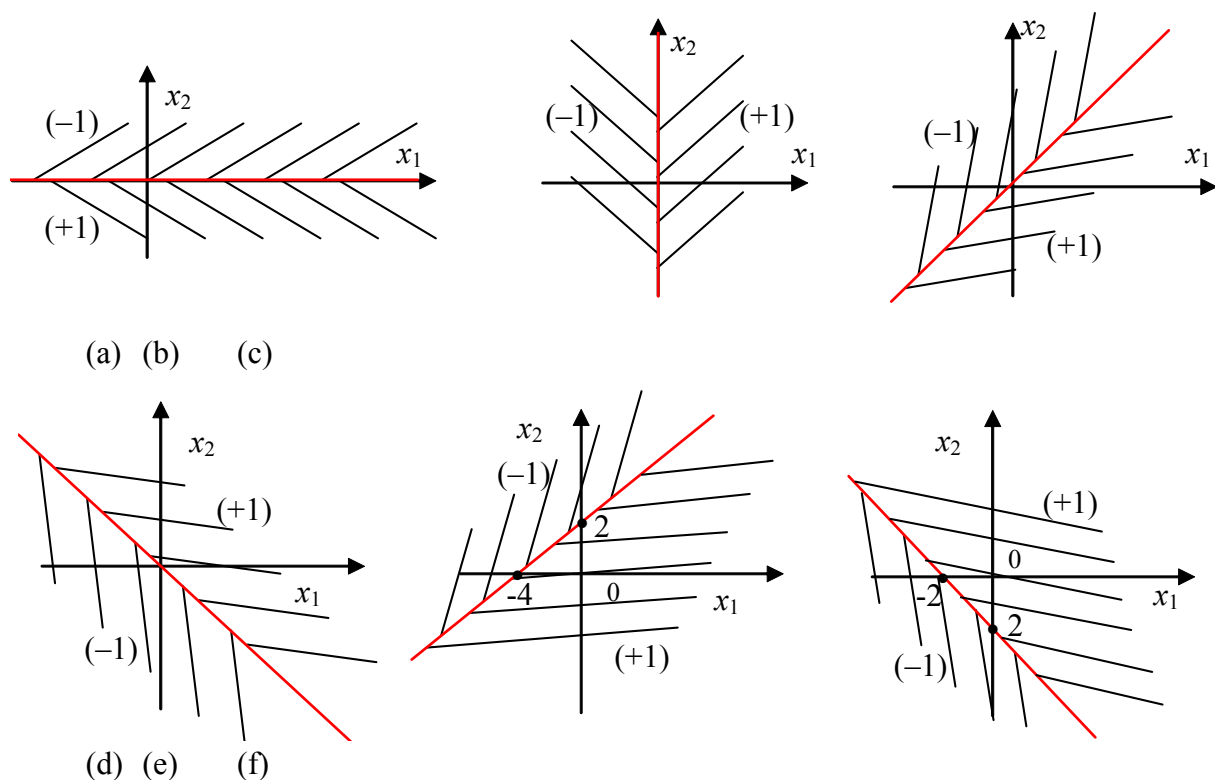


Figura 7.3.1. Separarea planului (x_1, x_2) în regiuni.

Indicații și răspunsuri:

Expresia funcției de activare de tip treaptă bipolară este:

$$\sigma(u) = \begin{cases} -1, & u < 0, \\ 1, & u \geq 0. \end{cases}$$

Cu notațiile din figura 3.1.1, vectorul ponderilor $\mathbf{w} = [w_1 \ w_2] \in \mathbb{R}^{1 \times 2}$ și deplasarea $b \in \mathbb{R}$, pentru neuronul cu două intrări se obține ieșirea:

$$y(x_1, x_2) = -1 \Leftrightarrow w_1 x_1 + w_2 x_2 + b < 0,$$

$$y(x_1, x_2) = +1 \Leftrightarrow w_1 x_1 + w_2 x_2 + b \geq 0.$$

Ecuția dreptei care separă planul (x_1, x_2) în cele două semiplane în funcție de valorile ieșirii y a rețelei este

$$w_1 x_1 + w_2 x_2 + b = 0. \quad (\text{S})$$

(a) Ecuția dreptei de separație este $x_2 = 0$, $\forall x_1 \in \mathbb{R}$. Înlocuind în (S) se obține $w_1 x_1 + b = 0$, $\forall x_1 \in \mathbb{R}$, de unde rezultă $w_1 = 0$, $b = 0$, astfel că $u = w_2 x_2$. Din figura (a) se observă că pentru $x_2 > 0$ ieșirea neuronului are valoarea $y = -1$, care corespunde lui $u < 0$, echivalent cu $w_2 x_2 < 0$. Rezultă că separarea planului (x_1, x_2) în modul prezentat în figura (a) poate fi realizată de un neuron pentru care ponderea w_2 are o valoare negativă arbitrară $w_2 < 0$.

(b) Analog rezultă că $w_2 = 0$, $b = 0$. Se observă că $x_1 < 0$ pentru $u < 0$, de unde se obține $w_1 > 0$.

(c) Ecuția dreptei de separație este $x_2 = x_1$, $x_1 \in \mathbb{R}$. Înlocuim pe x_1, x_2 în expresia lui y și obținem: $w_1 x_1 + w_2 x_2 + b = 0, (\forall) x_1 = x_2 = x \in R \Rightarrow (w_1 + w_2)x + b = 0, (\forall) x \in R$. De aici rezultă că $w_1 = -w_2$; $b = 0$. Pentru $x_1 = 0$ și $x_2 > 0$ avem $y = -1$, de unde rezultă că $w_2 x_2 < 0$, deci $w_2 < 0$.

(d) Ecuția dreptei de separație este $x_2 = -x_1$, $x_1 \in \mathbb{R}$. Similar cazului (c) se obține $w_1 = w_2 > 0$, $b = 0$.

(e) $x_2 = \frac{1}{2} x_1 + 2$, $x_1 \in \mathbb{R}$, ceea ce conduce la $w_1 = -\frac{1}{2} w_2$, $b = -2w_2$. Pentru $x_1 = x_2 = 0$ ieșirea neuronului trebuie să fie $+1$, astfel că $w_2 < 0$.

(f) Ecuția dreptei de separație este $x_2 = -x_1 - 2$, $x_1 \in \mathbb{R}$, de unde rezultă $w_1 = w_2$ și $b = 2w_2$. Pentru $x_1 = x_2 = 0$ ieșirea neuronului trebuie să fie $+1$, astfel că $w_2 > 0$.

Problema 7.4.

Se consideră un perceptron cu o singură intrare și funcția de activare unipolară. La intrarea acestuia se aplică vectorii prototip: $x^1 = 1$, $x^2 = 3$, $x^3 = 0$.

Se definesc vectorii eroare: $e^i = z^i - y^i, i = 1, 2, 3$, pentru vectorii țintă z^i , $i = 1, 2, 3$, asociați celor 3 vectori prototip.

- a) Să se discute în planul parametrilor rețelei (b – abscisă, w – ordonată) valorile luate de erorile e^i , $i = 1, 2, 3$ și să se precizeze regiunea pentru care se realizează $e^1 = 0$, $e^2 = 0$, $e^3 = 0$.

Se vor avea în vedere următoarele două seturi de valori ale vectorilor țintă:

(i) $z^1 = 0, z^2 = 1, z^3 = 0$;

(ii) $z^1 = 0, z^2 = 1, z^3 = 1$.

- b) Să se explice diferențele dintre rezultatele obținute în cazurile a) și b) de la punctul I.

Indicații și răspunsuri:

I) Erorile sunt date de $e^i = z^i - \sigma(wx^i + b)$, $i = \overline{1, 3}$, cu $\sigma(u) = \begin{cases} 0, & u < 0, \\ 1, & u \geq 0. \end{cases}$

a) Pentru vectori țintă: $z^1 = 0, z^2 = 1, z^3 = 0$, se obține

$$e^1 = \begin{cases} 0, & w + b < 0, \\ -1, & w + b \geq 0; \end{cases} \quad e^2 = \begin{cases} 0, & 3w + b \geq 0, \\ 1, & 3w + b < 0; \end{cases} \quad e^3 = \begin{cases} 0, & b < 0, \quad w \in \mathbb{R}, \\ -1, & b \geq 0, \quad w \in \mathbb{R}. \end{cases}$$

Regiunea pentru care se realizează $e^1 = 0$, $e^2 = 0$, $e^3 = 0$ este prezentată în figura 7.4.1.

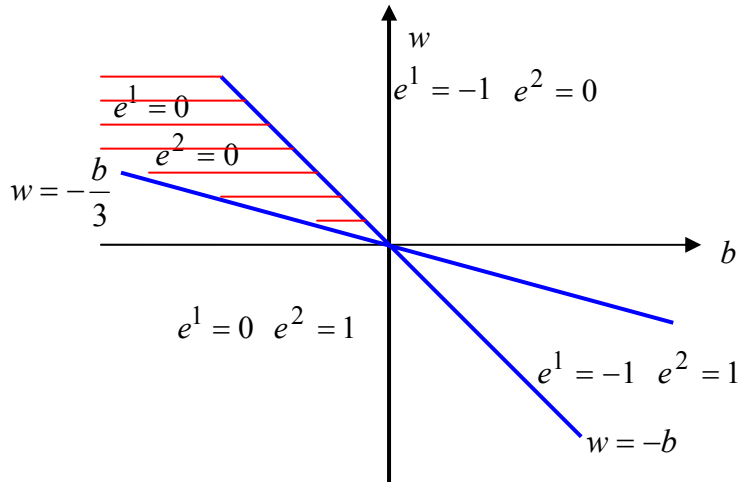


Figura 7.4.1. Reprezentarea grafică a regiunii din planul parametrilor pentru care $e^1 = 0$, $e^2 = 0$, $e^3 = 0$ pentru vectorii țintă din cazul (a).

b) Pentru vectori țintă: $z^1 = 0, z^2 = 1, z^3 = 1$, rezultă

$$e^1 = \begin{cases} 0, & w + b < 0, \\ -1, & w + b \geq 0; \end{cases} \quad e^2 = \begin{cases} 0, & 3w + b \geq 0, \\ 1, & 3w + b < 0; \end{cases} \quad e^3 = \begin{cases} 0, & b \geq 0, \quad w \in \mathbb{R}, \\ 1, & b < 0, \quad w \in \mathbb{R}. \end{cases}$$

II. Diferențele dintre cele două cazuri sunt datorate modificării vectorului țintă z^3 .

Problema 7.5.

Se consideră funcția logică XOR având tabela de adevăr și reprezentarea grafică din figura 7.5.1. Funcția stabilește apartenența vectorilor $\mathbf{x} = [x_1 \ x_2]^T$, $x_1, x_2 \in \{0, 1\}$, la două clase definite prin $\mathcal{C}_0 = \{(x_1, x_2) \in P \mid x_1 \oplus x_2 = 0\}$ și $\mathcal{C}_1 = \{(x_1, x_2) \in P \mid x_1 \oplus x_2 = 1\}$.

Să se arate că cele două clase pot fi separate cu ajutorul unei rețele perceptron cu 2 straturi, având 2 neuroni în stratul de intrare.

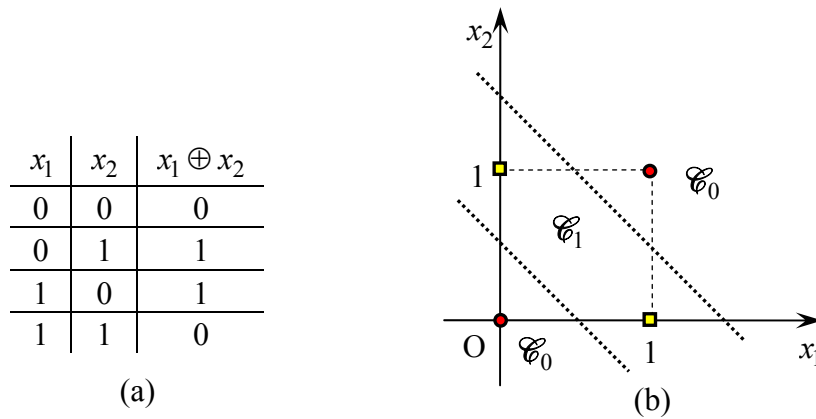


Figura 7.5.1. (a) Tabela de adevăr și (b) reprezentarea grafică a funcției logice XOR.

Indicații și răspunsuri:

Se utilizează exprimarea funcției XOR sub forma echivalentă $x_1 \oplus x_2 = (x_1 \wedge \overline{x_2}) \vee (\overline{x_1} \wedge x_2)$.

Pe stratul de intrare plasăm 2 neuroni care implementează funcțiile logice $x_1 \wedge \overline{x_2}$ respectiv $\overline{x_1} \wedge x_2$. Stratul de ieșire va conține un singur neuron cu 2 intrări ce implementează funcția OR. Ținând cont de chestiunile prezentate în exemplul 3.1.1 se obțin următoarele valori:

- pentru neuronul 1 de pe stratul de intrare: $w_{1,1}^1 = 1$, $w_{1,2}^1 = -1$, $b_1^1 = -0.5$;
- pentru neuronul 2 de pe stratul de intrare: $w_{2,1}^1 = -1$, $w_{2,2}^1 = 1$, $b_2^1 = -0.5$;
- pentru neuronul de pe stratul de ieșire: $w_1^2 = w_2^2 = 1$, $b^2 = -0.5$.

Problema 7.6.

În planul $\mathbb{R}^2$ se consideră următoarele două perechi de puncte (x_i, z_i) , $i = 1, 2$: $(0, 0)$, $(1, 2)$.

- a) Să se determine ecuația unei drepte de forma $y = wx + b$ care trece prin punctele (x_i, z_i) , $i = \overline{1, 2}$.
- b) Se consideră un neuron liniar cu o singură intrare unde se prezintă valorile x_1 , x_2 . Neuronul este conectat cu un comparator (ca în figură) care furnizează valorile: $e_1 = z_1 - f(x_1)$, $e_2 = z_2 - f(x_2)$, atunci când pe intrarea cu (+) a comparatorului se prezintă z_1 , respectiv z_2 .

Se consideră eroarea pătratică globală $J = e_1^2 + e_2^2$, ca funcție de care depinde de parametrii neuronului w și b .

- Să se reprezinte grafic suprafața $J(w, b)$.
- Să se arate că această suprafață posedă un singur punct de minim global și să se precizeze valorile w^* , b^* corespunzătoare acestui punct.

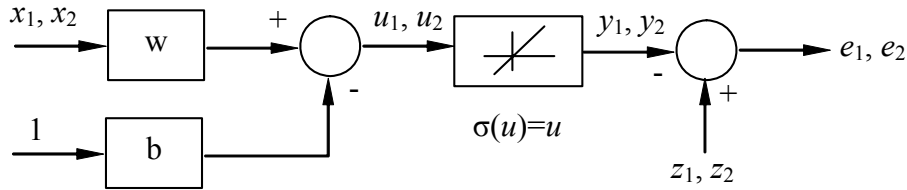


Figura 7.6.1. Schema bloc pentru generarea semnalelor de eroare din problema 7.5.

Indicații și răspunsuri:

a) Ecuația dreptei de forma $y = wx + b$ ce trece prin punctele date se determină din relațiile:

$$\begin{cases} 0 = 0 \cdot w + b \\ 2 = w + b \end{cases} \Rightarrow \begin{cases} w = 2 \\ b = 0 \end{cases}$$

b) Pentru neuronul liniar cu o intrare, $y(x_i) = wx_i + b$, $e_i = z_i - (x_i w - b)$, $i = 1, 2$. Eroarea pătratică globală este $J(w, b) = e_1^2 + e_2^2 = b^2 + (2 - w + b)^2 = 2b^2 + 4 + w^2 - 4w + 4b - 2wb$.

(ii) Punctele de extrem ale funcției J se determină egalând cu 0 derivatele parțiale ale acesteia în raport cu cei doi parametri, $\frac{\partial J}{\partial w} = 2w - 4 - 2b$ și $\frac{\partial J}{\partial b} = 4b + 4 - 2w$, ceea ce conduce la sistemul

$$\begin{cases} 2w - 4 - 2b = 0 \\ 4b + 4 - 2w = 0 \end{cases} \Leftrightarrow \begin{cases} 2w - 2b = 4 \\ -2w + 4b = -4 \end{cases} \Rightarrow \begin{cases} b^* = 0 \\ w^* = 2 \end{cases}$$

Matricea Hessian a funcției J în punctul $b^* = 0$, $w^* = 2$, este

$$H = \begin{bmatrix} \frac{\partial^2 J}{\partial w^2} & \frac{\partial^2 J}{\partial w \partial b} \\ \frac{\partial^2 J}{\partial b \partial w} & \frac{\partial^2 J}{\partial b^2} \end{bmatrix} = \begin{bmatrix} 2 & -2 \\ -2 & 4 \end{bmatrix}.$$

Deoarece valorile proprii ale lui H sunt $\lambda_{1,2} = 3 \pm \sqrt{5} > 0$, matricea H este pozitiv definită. Rezultă că punctul $b^* = 0$, $w^* = 2$ este punct de minim global al funcției J .

Problema 7.7.

Se consideră amplificatorul operațional din figură, unde $x_1, x_2, \dots, x_n$ sunt semnale de intrare, iar y este semnal de ieșire.

Să se modeleze funcționarea amplificatorului utilizând o rețea neuronală cu funcții de activare liniare (se va dimensiona rețeaua și se vor preciza valorile parametrilor).

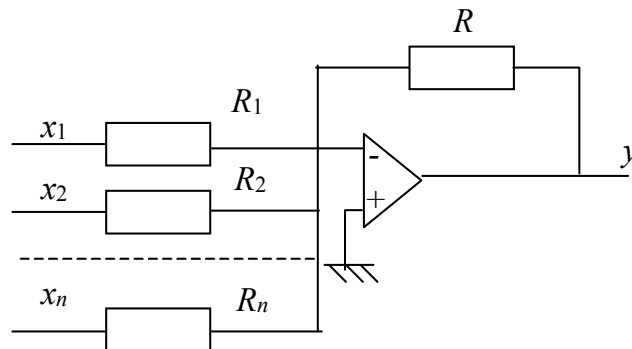


Figura 7.7.1. Schema bloc a unui amplificator operațional cu n intrări.

Indicații și răspunsuri:

Funcționarea amplificatorului operațional este descrisă de ecuația:

$$\sum_{i=1}^n \frac{x_i}{R_i} = -\frac{y}{R},$$

unde x_i , $i=1, n$, și y reprezintă tensiunea de la intrarea i , ($i=1, n$) și, respectiv, ieșirea amplificatorului operațional.

Legătura dintre ieșirea y a unui neuron liniar și cele n intrări ale sale este:

$$y = \sum_{i=1}^n x_i w_i + b.$$

O rețea neuronală formată dintr-un neuron liniar cu n intrări cu parametrii $w_i = -\frac{R}{R_i}$ și

$b = 0$ va modela funcționarea amplificatorului operațional.

Problema 7.8.

Se consideră un neuron cu o singură intrare și funcția de activare liniară. Pentru ce valori ale parametrilor neuronului ieșirea acestuia aproximează, în sensul celor mai mici pătrate, funcția φ cunoscută în următoarele 3 puncte: $\varphi(-1) = -1$, $\varphi(0) = 0$, $\varphi(1) = 2$

Indicații și răspunsuri:

Notăm $x_1 = -1$, $x_2 = 0$, $x_3 = 1$, și $z_i = \varphi(x_i)$, $i = 1, 2, 3$.

Pentru neuronul liniar cu parametrii w și b , notăm $y_i = \sigma(wx_i + b)$, $e_i = z_i - y_i$, $i = 1, 2, 3$.

Parametrii neuronului vor fi determinați astfel încât să fie minimizată eroarea pătratică:

$$J(w, b) = \sum_{i=1}^3 e_i^2 = (-1 + w - b)^2 + (-b)^2 + (2 - w - b)^2 = 2w^2 + 3b^2 - 6w - 2b + 5.$$

Punctul de minim al funcției se obține ca soluție a sistemului de ecuații:

$$\begin{cases} \frac{\partial J}{\partial w} = 0 \\ \frac{\partial J}{\partial b} = 0 \end{cases} \Rightarrow \begin{cases} 4w - 6 = 0 \\ 6b - 2 = 0 \end{cases} \Rightarrow \begin{cases} w = 3/2 \\ b = 1/3 \end{cases}$$

Se calculează matricea hessian și se observă că este pozitiv definită, deci J are un unic punct de minim (global).

Problema 7.9.

Se consideră un neuron cu o singură intrare și funcția de activare σ derivabilă. La intrarea neuronului se aplică $x \in \mathbb{R}$, iar $y \in \mathbb{R}$ este comparat cu $z \in \mathbb{R}$ dat (fixat), furnizând eroarea

$e = z - y$ conform figurii. Se consideră funcția: $J = \frac{1}{2}e^2$.

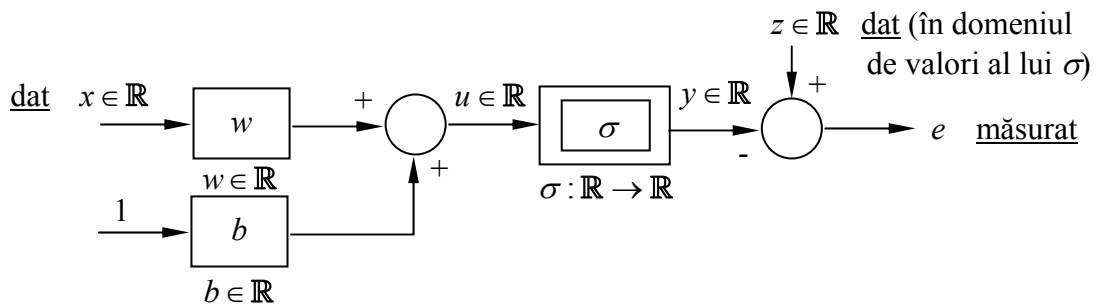


Figura 7.9.1. Schema bloc pentru generarea semnalelor de eroare din problema 7.9.

- Să se exprime $\frac{\partial J}{\partial w}$ și $\frac{\partial J}{\partial b}$ în funcție de e , x și $\sigma'(u)$.
- Să se exprime $\sigma'(u)$ ca o funcție de argument y .
- Folosind rezultatele obținute la subpunctele a) și b) să se exprime $\frac{\partial J}{\partial w}$ și $\frac{\partial J}{\partial b}$ ca depinzând de e , x și valoarea unei funcții de argument y .

Indicații și răspunsuri:

$$\begin{aligned} \text{a) } J(w, b) &= \frac{1}{2}e^2 = \frac{1}{2}[z - \sigma(wx + b)]^2 \Rightarrow \frac{\partial J}{\partial w} = \frac{\partial J}{\partial e} \frac{\partial e}{\partial \sigma} \frac{\partial \sigma}{\partial u} \frac{\partial u}{\partial w} = -e\sigma'(u)x \\ \frac{\partial J}{\partial b} &= \frac{\partial J}{\partial e} \frac{\partial e}{\partial \sigma} \frac{\partial \sigma}{\partial u} \frac{\partial u}{\partial b} = -e\sigma'(u) \end{aligned}$$

$$\text{b) } \left. \begin{aligned} \frac{\partial \sigma}{\partial y} &= \frac{\partial \sigma}{\partial u} \cdot \frac{\partial u}{\partial y} \\ \frac{\partial \sigma}{\partial y} &= 1 \end{aligned} \right\} \Rightarrow \frac{\partial \sigma}{\partial u} = \frac{1}{\frac{\partial u}{\partial y}} \Leftrightarrow \sigma'(u) = \frac{1}{(\text{Inv}\sigma)'(y)}$$

$$\text{c) Notând } f(y) = \frac{1}{(\text{Inv}\sigma)'(y)}, \text{ rezultă } \begin{aligned} \frac{\partial J}{\partial w} &= -ef(y)x \\ \frac{\partial J}{\partial b} &= -ef(y) \end{aligned}$$

Problema 7.10.

Se consideră funcția $\sigma: \mathbb{R} \rightarrow \mathbb{R}$, $y = \sigma(u)$. Să se exprime $\sigma'(u)$ ca o funcție de y pentru următoarele situații:

- a) $\sigma(u) = \text{logsig}(u)$,
- b) $\sigma(u) = \text{tansig}(u)$.

Indicații și răspunsuri:

$$\begin{aligned} \text{a) } \sigma(u) &= \frac{1}{1+e^{-u}} = y \Rightarrow \sigma'(u) = \frac{1}{(\text{Inv}\sigma)'(y)}, \text{ cu } \text{Inv}\sigma(y) = \ln\left(\frac{y}{1-y}\right) \Rightarrow \sigma'(u) = y(1-y) \\ \text{b) } \sigma(u) &= \frac{e^u - e^{-u}}{e^u + e^{-u}} = y \Rightarrow \sigma'(u) = \frac{1}{(\text{Inv}\sigma)'(y)}, \text{ cu } \text{Inv}\sigma(y) = \frac{1}{2} \ln\left(\frac{1+y}{1-y}\right) \Rightarrow \sigma'(u) = 1-y^2 \end{aligned}$$

Problema 7.11.

a) Se consideră $F: \mathbb{R}^n \rightarrow \mathbb{R}$, $F(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{x}$. Să se demonstreze că: $\frac{\partial F}{\partial \mathbf{x}} = \mathbf{x}^T$.

b) Se consideră funcția $F: \mathbb{R}^n \rightarrow \mathbb{R}^n$, $F(\mathbf{x}) = \mathbf{a} - \mathbf{x}$, cu $\mathbf{a} \in \mathbb{R}^n$ un vector constant. Să se demonstreze că: $\frac{\partial F}{\partial \mathbf{x}} = -\mathbf{I}_n$, unde $\mathbf{I}_n$ notează matricea unitate de ordinul n .

c) Se consideră funcția $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$, $f(\mathbf{x}) = \mathbf{A}\mathbf{x}$, cu $\mathbf{A} = \begin{bmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_n \end{bmatrix}$, $\mathbf{a}_k \in \mathbb{R}^{1 \times n}$, $k = \overline{1, n}$. Să se

calculeze matricele Jacobian $\left[\frac{\partial f}{\partial \mathbf{a}_k} \right]$, $k = \overline{1, n}$.

Indicații și răspunsuri:

$$\text{a) } F(\mathbf{x}) = \frac{1}{2} \sum_{i=1}^n x_i^2 \text{ cu } \mathbf{x} = [x_1 \ x_2 \ \dots \ x_n]^T.$$

$$\frac{\partial F}{\partial x_i} = x_i, i = \overline{1, n}, \Rightarrow \frac{\partial F}{\partial \mathbf{x}} = \left[\frac{\partial F}{\partial x_1} \ \frac{\partial F}{\partial x_2} \ \dots \ \frac{\partial F}{\partial x_n} \right] = [x_1 \ x_2 \ \dots \ x_n] = \mathbf{x}^T$$

$$\text{b) } F(\mathbf{x}) = \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix} - \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} a_1 - x_1 \\ a_2 - x_2 \\ \vdots \\ a_n - x_n \end{bmatrix} = \begin{bmatrix} f_1(\mathbf{x}) \\ f_2(\mathbf{x}) \\ \vdots \\ f_n(\mathbf{x}) \end{bmatrix}; \text{ cu } f_i(\mathbf{x}) = a_i - x_i, i = \overline{1, n}.$$

$$\frac{\partial F}{\partial \mathbf{x}} = \begin{bmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} & \dots & \frac{\partial f_1}{\partial x_n} \\ \frac{\partial f_2}{\partial x_1} & \frac{\partial f_2}{\partial x_2} & \dots & \frac{\partial f_2}{\partial x_n} \\ \vdots & \vdots & & \vdots \\ \frac{\partial f_n}{\partial x_1} & \frac{\partial f_n}{\partial x_2} & \dots & \frac{\partial f_n}{\partial x_n} \end{bmatrix} = \begin{bmatrix} -1 & 0 & \dots & 0 \\ 0 & -1 & \dots & 0 \\ \vdots & \vdots & & \vdots \\ 0 & 0 & \dots & -1 \end{bmatrix} = -\mathbf{I}_n$$

$$\text{c) } f(\mathbf{x}) = \mathbf{Ax}, f(\mathbf{x}) = \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix} \mathbf{x}, \quad a_k \in R^{1 \times n}, \quad k = \overline{1, n} \Rightarrow f(\mathbf{x}) = \begin{bmatrix} a_1 \mathbf{x} \\ a_2 \mathbf{x} \\ \vdots \\ a_n \mathbf{x} \end{bmatrix} = \begin{bmatrix} f_1(\mathbf{x}) \\ f_2(\mathbf{x}) \\ \vdots \\ f_n(\mathbf{x}) \end{bmatrix}$$

Cu notațiile $\mathbf{a}_k = [a_k^1 \ a_k^2 \ \dots \ a_k^n]$, $k = \overline{1, n}$, $\mathbf{x} = [x_1 \ x_2 \ \dots \ x_n]^T$; se obține:

$$\frac{\partial f}{\partial \mathbf{a}_k} = \begin{bmatrix} \frac{\partial f_1}{\partial a_k^1} & \frac{\partial f_1}{\partial a_k^2} & \dots & \frac{\partial f_1}{\partial a_k^n} \\ \frac{\partial f_2}{\partial a_k^1} & \frac{\partial f_2}{\partial a_k^2} & \dots & \frac{\partial f_2}{\partial a_k^n} \\ \vdots & \vdots & & \vdots \\ \frac{\partial f_k}{\partial a_k^1} & \frac{\partial f_k}{\partial a_k^2} & \dots & \frac{\partial f_k}{\partial a_k^n} \\ \vdots & \vdots & & \vdots \\ \frac{\partial f_n}{\partial a_k^1} & \frac{\partial f_n}{\partial a_k^2} & \dots & \frac{\partial f_n}{\partial a_k^n} \end{bmatrix} = \begin{bmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & & \vdots \\ x_1 & x_2 & \dots & x_n \\ \vdots & \vdots & & \vdots \\ 0 & 0 & \dots & 0 \end{bmatrix} = \begin{bmatrix} \mathbf{O}_{(k-1) \times n} \\ \mathbf{x}^T \\ \mathbf{O}_{(n-k) \times n} \end{bmatrix}, \quad k = \overline{1, n}.$$

Problema 7.12.

Se consideră un neuron liniar cu o singură intrare și un neuron tansigmoïdal cu o singură intrare, ambii având deplasări nule. Intrarea neuronilor este comună, notată x , iar ieșirile lor sunt conectate la un comparator, furnizând $e(x)$, ca în figură:

Să se discute soluțiile ecuației $e(x) = 0$ pentru $w_1 \in \mathbb{R}$, $w_2 \in \mathbb{R}$ (soluțiile nu se vor calcula efectiv!). În planul parametrilor (w_1, w_2) – cu w_1 abscisă și w_2 ordonată – se vor delimita regiunile care conțin numere diferite de soluții distincte.

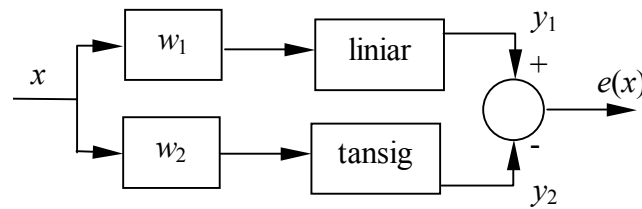


Figura 7.12.1. Schema bloc a generării semnalelor de eroare din problema 7.12.

Indicații și răspunsuri:

$$e(x) = y_1(x) - y_2(x) = w_1 \cdot x - w_2 \cdot \frac{e^x - e^{-x}}{e^x + e^{-x}} \Rightarrow e'(x) = w_1 - w_2 \cdot \frac{4}{(e^x + e^{-x})^2}$$

Pentru a discuta soluțiile ecuației studiem semnul lui $e'(x)$.

Cazul 1. Dacă $\frac{w_1}{w_2} < 0$ atunci $e'(x)$ are semn constant pe $\mathbb{R}$.

a) Dacă $w_1 > 0$ și $w_2 < 0$ atunci $e'(x) > 0$ și $e(x)$ este strict crescătoare pe $\mathbb{R}$, de la $-\infty$ la $+\infty$. Ecuația $e(x) = 0$ are o singură soluție.

b) Dacă $w_1 < 0$ și $w_2 > 0$ atunci $e'(x) < 0$ și $e(x)$ este strict descrescătoare pe $\mathbb{R}$, de la $+\infty$ la $-\infty$. Ecuația $e(x) = 0$ are o singură soluție..

Cazul 2. Dacă $\frac{w_2}{w_1} > 0$ atunci ecuația $e'(x) = 0$ are soluții, ceea ce arată că funcția $e(x)$

are puncte de extrem (de minim sau de maxim).

$$\text{Din } e'(x) = 0 \text{ se obține } w_1 = \frac{4 \cdot w_2}{(e^x + e^{-x})^2} \Leftrightarrow (e^x + e^{-x})^2 = \frac{4 \cdot w_2}{w_1}.$$

$$\text{Rezultă } e^x + e^{-x} = 2\sqrt{\frac{w_2}{w_1}} \text{ sau } e^x + e^{-x} = -2\sqrt{\frac{w_2}{w_1}}.$$

$$\text{a) } e^x + e^{-x} = 2\sqrt{\frac{w_2}{w_1}}$$

Notăm $e^x = t$ și obținem relația $t + \frac{1}{t} = 2\sqrt{\frac{w_2}{w_1}} \Rightarrow t^2 - 2t\sqrt{\frac{w_2}{w_1}} + 1 = 0$. Se rezolvă ecuația de gradul 2 în t . După ce se determină cele două rădăcini ale acesteia se revine la $e^x = t$.

$$\text{b) } e^x + e^{-x} = -2\sqrt{\frac{w_2}{w_1}}$$

Notăm $e^x = t$ și obținem relația $t + \frac{1}{t} = -2\sqrt{\frac{w_2}{w_1}} \Rightarrow t^2 + 2t\sqrt{\frac{w_2}{w_1}} + 1 = 0$. Se rezolvă ecuația de gradul 2 în t . După ce se determină cele două rădăcini ale acesteia se revine la $e^x = t$.

Problema 7.13.

La intrarea unui neuron tansigmoïdal se aplică vectorul prototip $x \in \mathbb{R}$ pentru care vectorul țintă este $z = 1$. Se calculează vectorul delta δ aferent perechii $(x, 1)$ și se obține $\delta = 1$. Să se determine parametrii neuronului care conduc la această soluție.

Indicații și răspunsuri:

Fie funcție de activare: $\sigma(u) = \frac{e^u - e^{-u}}{e^u + e^{-u}}$. Vectorul delta aferent perechii $(x, 1)$ este $\delta = 1$.

Eroarea este $e = z - y = z - \sigma(u) = \frac{2e^{-u}}{e^u + e^{-u}}$ și $\delta = qe = \sigma'(u)e$, cu $\sigma'(u) = \frac{4}{(e^u + e^{-u})^2} \Rightarrow$

$\delta = \frac{8e^{-u}}{(e^u + e^{-u})^3}$. Din $\delta = 1$, rezultă $8e^{-u} = (e^u + e^{-u})^3$, care se rezolvă în raport cu e^u .

Problema 7.14.

Se consideră următoarele trei puncte de forma $(x^i, z^i) \in \mathbb{R}^2$, $i = \overline{1, 3}$: $(-1, -1/2)$, $(0, 0)$, $(1, 1/2)$.

- Să se arate că există $\omega, \phi \in \mathbb{R}$ astfel încât funcția $g(x) = \sin(\omega x + \phi)$ trece prin cele trei puncte date.
- Să se arate că există ponderea w și deplasarea b pentru un neuron liniar, $w, b \in \mathbb{R}$, astfel încât funcția sa intrare-ieșire trece prin cele trei puncte date.
- Să se arate că există ponderea w și deplasarea b pentru un neuron tansigmoïdal, $w, b \in \mathbb{R}$, astfel încât funcția sa intrare-ieșire trece prin cele trei puncte date.

Indicații și răspunsuri:

- a) Se ajunge la sistemul de ecuații
$$\begin{cases} \sin(-\omega + \phi) = -1/2, \\ \sin(\phi) = 0 \Rightarrow \phi = k\pi, k \in \mathbb{Z} \\ \sin(\omega + \phi) = 1/2 \end{cases}$$

Ținând cont de simetria și periodicitatea funcției sinus se obțin două soluții:

$$\omega = \arcsin\left(\frac{1}{2}\right) \Rightarrow \omega = \frac{\pi}{6}, \text{ și } \omega = \arcsin\left(\frac{1}{2}\right) \Rightarrow \omega = \frac{\pi}{6}.$$

- b) În cazul unui neuron liniar se ajunge la sistemul de ecuații
$$\begin{cases} -1/2 = -w + b \\ 0 = b \\ 1/2 = w + b \end{cases} \Rightarrow \begin{cases} b = 0 \\ w = 1/2 \end{cases}$$

- c) Pentru un neuron tansigmoïdal avem: $u = wx + b$, $\sigma(u) = \frac{e^u - e^{-u}}{e^u + e^{-u}}$. Cu notația $e^u = a > 0$, rezultă:

$$-1/2 = \frac{a-1/a}{a+1/a} \Rightarrow -1/2 = \frac{a^2-1}{a^2+1} \Rightarrow -a^2-1 = 2a^2-2 \Rightarrow 3a^2-1=0 \Rightarrow a = \frac{1}{\sqrt{3}}$$

$$1/2 = \frac{a-1/a}{a+1/a} \Rightarrow 1/2 = \frac{a^2-1}{a^2+1} \Rightarrow a^2+1 = 2a^2-2 \Rightarrow a^2-3=0 \Rightarrow a = \sqrt{3}$$

$$0 = \frac{a-1/a}{a+1/a} \Rightarrow 0 = a^2-1 \Rightarrow a = 1$$

$$\text{În final se obține } \begin{cases} -\ln(\sqrt{3}) = -w + b \\ \ln(\sqrt{3}) = w + b \\ \ln(1) = 0 = b \end{cases} \Rightarrow \begin{cases} w = \ln(\sqrt{3}) \\ b = 0 \end{cases}$$

Problema 7.15.

Se consideră trei puncte de forma $(x^i, z^i) \in \mathbb{R}^2$, $i = \overline{1, 3}$: $(-1, -1/2)$, $(0, 0)$, $(1, 3/4)$.

a.(i). Să se arate că nu există $\omega, \varphi \in \mathbb{R}$ astfel încât funcția $g(x) = \sin(\omega x + \varphi)$ să treacă prin cele trei puncte date.

(ii). Se caută funcția $g(x)$ care minimizează $J(\omega, \varphi) = \frac{1}{2} \sum_{i=1}^3 [g(x^i) - z^i]^2$. Să se scrie

iterația generală a algoritmului de gradient care minimizează criteriul $J(\omega, \varphi)$, sub

$$\text{forma: } \begin{bmatrix} \omega_{\text{new}} \\ \varphi_{\text{new}} \end{bmatrix} = F \left(\begin{bmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{bmatrix} \right).$$

b.(i). Să se arate că nu există ponderea w și deplasarea b pentru un neuron liniar, $w, b \in \mathbb{R}$, astfel încât funcția sa intrare-ieșire $f(x)$ să treacă prin cele trei puncte date.

(ii). Se caută $f(x)$ care minimizează $J(w, b) = \frac{1}{2} \sum_{i=1}^3 [f(x^i) - z^i]^2$. Să se scrie iterația

generală a algoritmului de antrenare a neuronului.

c. Se reia punctul b pentru un neuron tansigmoïdal.

Indicații și răspunsuri:

a.(i). Demonstrăm prin reducere la absurd, ținând cont de paritatea funcției cosinus (pară).

Se ajunge la

$$\begin{cases} \sin(-\omega + \varphi) = -1/2 \Rightarrow \sin(\omega - \varphi) = \sin(\omega) \cos(-\varphi) + \cos(\omega) \sin(-\varphi) = \sin(\omega) \cos(\varphi) - \cos(\omega) \sin(\varphi) = 1/2 \\ \sin(\varphi) = 0 \Rightarrow \varphi = k\pi, k \in \mathbb{Z} \\ \sin(\omega + \varphi) = \sin(\omega) \cos(\varphi) + \cos(\omega) \sin(\varphi) = \sin(\omega) \cos(\varphi) = 3/4 \end{cases}$$

$\Rightarrow 3/4 = 1/2 \quad (F) \Rightarrow$ nu există $\omega, \varphi \in \mathbb{R}$ astfel încât funcția $g(x) = \sin(\omega x + \varphi)$ să treacă prin cele trei puncte.

(ii). Relațiile corespunzătoare unei iterații sunt: $\omega_{\text{new}} = \omega_{\text{old}} + \eta \frac{\partial J}{\partial \omega}$, $\varphi_{\text{new}} = \varphi_{\text{old}} + \eta \frac{\partial J}{\partial \varphi}$,

unde $J = \sum_{i=1}^3 J^i$, cu $J^i = \frac{1}{2}(e^i)^2$, $e^i = y^i - z^i = g(x^i) - z^i$, $i = 1, 2, 3$.

Se obține:

$$\begin{aligned} \frac{\partial J}{\partial \omega} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} &= \left[\sin \left(\begin{bmatrix} 1 & -1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) + \frac{1}{2} \right] \cos \left(\begin{bmatrix} 1 & -1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) + \\ &= \left[\sin \left(\begin{bmatrix} 1 & 1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) - \frac{3}{4} \right] \cos \left(\begin{bmatrix} 1 & 1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) \\ \frac{\partial J}{\partial \varphi} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} &= \left[\sin \left(\begin{bmatrix} -1 & 1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) + 1/2 \right] \cos \left(\begin{bmatrix} -1 & 1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) + \frac{1}{2} \sin \left(\begin{bmatrix} 0 & 2 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) \\ &+ \left[\sin \left(\begin{bmatrix} 1 & 1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) - 3/4 \right] \cos \left(\begin{bmatrix} 1 & 1 \end{bmatrix} \begin{pmatrix} \omega_{\text{old}} \\ \varphi_{\text{old}} \end{pmatrix} \right) \end{aligned}$$

Similar se rezolvă punctele b) și c) ale problemei.

Problema 7.16.

Se consideră sistemul din figura de mai jos:

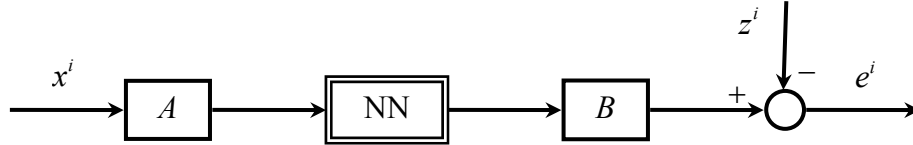


Figura 7.16.1. Schema bloc a generării semnalelor de eroare din problema 7.16.

în care $(x^i, z^i) \in \mathbb{R}^2$, $i = \overline{1, N}$, sunt perechi cunoscute de numere reale, A și B sunt două valori reale cunoscute, iar NN este un neuron cu funcția de activare sigmoidală, ponderea w și deplasarea b .

Fie funcția de cost $J = \frac{1}{2} \sum_{i=1}^N (e^i)^2$ care trebuie minimizată în raport cu w și b . Să se scrie

o iterație a algoritmului de gradient de forma $\begin{bmatrix} w_{\text{new}} \\ b_{\text{new}} \end{bmatrix} = F \left(\begin{bmatrix} w_{\text{old}} \\ b_{\text{old}} \end{bmatrix} \right)$ pentru cazurile:

- funcția de activare a neuronului este logsigmoidală;
- funcția de activare a neuronului este tansigmoidală.

Indicații și răspunsuri:

Se aplică relațiile din paragraful 5.2.2.

Problema 7.17.

Să se reprezinte următoarele sisteme discrete sub forma unei rețele neuronale dinamice constituită dintr-o ADALINE și întârzierile adecvate:

$$\text{a) } G_1(z) = 10 \frac{(z+1)(z-1)}{\left(z - \frac{1}{2}\right)\left(z - \frac{1}{3}\right)}; \quad \text{b) } G_2(z) = \frac{z^{-5} - 0.3z^{-6}}{1 - 0.5z^{-1} + 0.06z^{-2}};$$

$$\text{c) } G_3(z) = \frac{b_2 z^2 + b_1 z + b_0}{a_3 z^3 + a_2 z^2 + a_1 z + a_0}, \quad a_i, b_j \neq 0, \quad i = \overline{0,3}, \quad j = \overline{0,2};$$

$$\text{d) } G_4(z) = z^{-3} G_3(z).$$

Indicații și răspunsuri:

$$\text{a) } G_1(z) = 10 \frac{z^2 - 1}{z^2 - \frac{5}{6}z + \frac{1}{6}} \Rightarrow G_1(z^{-1}) = 10 \frac{1 - z^{-2}}{1 - \frac{5}{6}z^{-1} + \frac{1}{6}z^{-2}}$$

$y[k]$ depinde de: $y[k-1], y[k-2], u[k], u[k-2]$.

$$\text{b) } G_2(z) = \frac{z^{-5} - 0.3z^{-6}}{1 - 0.5z^{-1} + 0.06z^{-2}}$$

$y[k]$ depinde de: $y[k-1], y[k-2], u[k-5], u[k-6]$

$$\text{c) } G_3(z^{-1}) = \frac{z^3(b_2 z^{-1} + b_1 z^{-2} + b_0 z^{-3})}{z^3(a_3 + a_2 z^{-1} + a_1 z^{-2} + a_0 z^{-3})}, \quad a_i, b_j \neq 0, \quad i = \overline{0,3}, \quad j = \overline{0,2}$$

$y[k]$ depinde de: $y[k-1], y[k-2], y[k-3], u[k-1], u[k-2], u[k-3]$

$$\text{d) } G_4(z) = z^{-3} G_3(z)$$

$y[k]$ depinde de: $y[k-1], y[k-2], y[k-3], u[k-4], u[k-5], u[k-6]$

Problema 7.18.

Se consideră rețeaua neuronală cu dinamică externă din figura 7.18.1, unde rețeaua ADALINE are parametrii:

$$\mathbf{w} = [2 \quad 0 \quad -18 \quad \eta \quad -1], \quad b = 0,$$

iar q^{-1} notează operatorul de întârziere în domeniul timp.

- Să se determine funcția de transfer a rețelei.
- Să se discute stabilitatea IMEM în funcție de valoarea parametrului $\eta \in \mathbb{R}$.
- Pentru valorile lui η pentru care sistemul este stabil IMEM să se arate că există rețele dinamice cu o configurație mai simplă care realizează același transfer intrare-ieșire și să se reprezinte grafic aceste configurații.

Indicații și răspunsuri:

- Pentru rețeaua ADALINE cu parametrii $\mathbf{w} = [2 \quad 0 \quad -18 \quad \eta \quad -1]$ și $b = 0$ se obține

$$y[k] = w_1 u[k] + w_2 u[k-1] + w_3 u[k-2] + w_4 y[k-1] + w_5 y[k-2] + b \Rightarrow$$

$$\begin{aligned}
&\Rightarrow y[k] = 2u[k] - 18u[k-2] + \eta y[k-1] - y[k-2] \Leftrightarrow \\
&\Leftrightarrow (1 - \eta z^{-1} + z^{-2})y[k] = (2 - 18z^{-2})u[k]. \\
\text{Rezultă } G(z^{-1}) &= \frac{2 - 18z^{-2}}{1 - \eta z^{-1} + z^{-2}} \Rightarrow G(z) = \frac{2z^2 - 18}{z^2 - \eta z + 1}
\end{aligned}$$

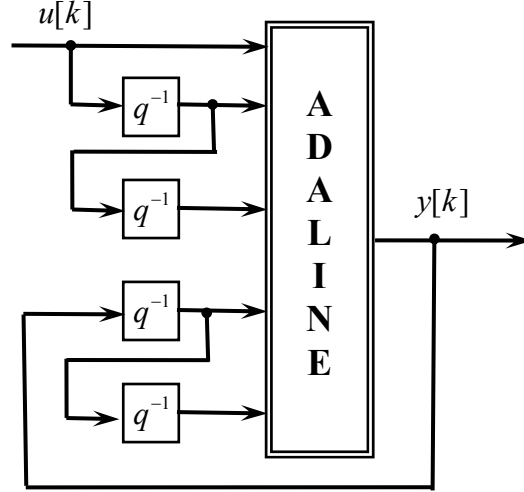


Figura 7.18.1. Rețeaua neuronală dinamică considerată în problema 7.18.

b) Un sistem dinamic în timp discret este stabil IMEM dacă și numai dacă toți polii sunt plasați în interiorul cercului unitate. Polii rețelei neurale cu dinamică externă sunt soluțiile

ecuației: $z^2 - \eta z + 1 = 0$. Cu notația $\Delta = \eta^2 - 4$, polii lui $G(z)$ sunt $z_{1/2} = \frac{\eta \pm \sqrt{\Delta}}{2}$.

1. Dacă $\eta \in (-\infty, -2] \cup [2, \infty) \Rightarrow \Delta \geq 0$ ecuația va avea soluții reale și una dintre rădăcini va fi supraunitară în modul.

$$\left. \begin{aligned} z_1 &= \frac{\eta + \sqrt{\Delta}}{2} = \frac{\eta}{2} + \frac{\sqrt{\Delta}}{2} \\ \eta \in [2, \infty) &\Rightarrow \frac{\eta}{2} \in [1, \infty) \end{aligned} \right\} \Rightarrow z_1 \geq 1, \quad \left. \begin{aligned} z_1 &= \frac{\eta - \sqrt{\Delta}}{2} = \frac{\eta}{2} - \frac{\sqrt{\Delta}}{2} \\ \eta \in (-\infty, -2] &\Rightarrow \frac{\eta}{2} \in (-\infty, -1] \end{aligned} \right\} \Rightarrow z_1 \leq -1$$

În aceste situații sistemul nu este stabil IMEM; cu excepția cazului în care polul care nu este poziționat în interiorul cercului unitate este egal cu unul dintre zerourile sistemului.

Zerourile sistemului sunt soluțiile ecuației: $2(z^2 - 9) = 0 \Rightarrow 2(z - 3)(z + 3) = 0$

Se observă că:

$$\begin{aligned}
&z_1 = \frac{\eta + \sqrt{\Delta}}{2} = \frac{\eta}{2} + \frac{\sqrt{\Delta}}{2} \\
\text{(i)} \quad &\left. \begin{aligned} \eta \in [2, \infty) &\Rightarrow \frac{\eta}{2} \in [1, \infty) \\ z_1 &= 3 \end{aligned} \right\} \Rightarrow \frac{\eta}{2} + \frac{\sqrt{\Delta}}{2} = 3 \Rightarrow \eta + \sqrt{\eta^2 - 4} = 6
\end{aligned}$$

$$\eta \in [2, 6) \Rightarrow \eta^2 - 4 = \eta^2 - 12\eta + 36 \Rightarrow 12\eta = 40 \Rightarrow \eta = \frac{10}{3}$$

$$\text{Polul sistemului este: } p = \frac{\eta - \sqrt{\eta^2 - 4}}{2} = \frac{\eta + \eta - 6}{2} = \eta - 3 = \frac{1}{3} < 1$$

$$\text{/ii) } \left. \begin{array}{l} z_1 = \frac{\eta - \sqrt{\Delta}}{2} = \frac{\eta}{2} - \frac{\sqrt{\Delta}}{2} \\ \eta \in (-\infty, -2] \Rightarrow \frac{\eta}{2} \in (-\infty, -1] \\ z_1 = -3 \end{array} \right\} \Rightarrow \frac{\eta}{2} - \frac{\sqrt{\Delta}}{2} = -3 \Rightarrow \eta - \sqrt{\eta^2 - 4} = -6$$

$$\eta \in (-6, -2] \Rightarrow \eta^2 - 4 = \eta^2 + 12\eta + 36 \Rightarrow 12\eta = -40 \Rightarrow \eta = -\frac{10}{3}$$

$$\text{Polul sistemului este: } p = \frac{\eta + \sqrt{\eta^2 - 4}}{2} = \frac{\eta + \eta + 6}{2} = \eta + 3 = -\frac{1}{3} > -1$$

2. Dacă $\eta \in (-2, 2) \Rightarrow \Delta < 0$ sistemul va avea poli complex conjugați cu parte reală subunitară în modul. Pentru ca sistemul să fie stabil IMEM polii $z_{1,2} = \frac{\eta \pm j\sqrt{-\Delta}}{2}$ trebuie să fie satisfacă inecuația $(\text{Re}(z_i))^2 + (\text{Im}(z_i))^2 < 1$, adică

$$\left(\frac{\eta}{2}\right)^2 + \left(\frac{\sqrt{4 - \eta^2}}{2}\right)^2 < 1 \Leftrightarrow \frac{\eta^2 + 4 - \eta^2}{4} < 1 \Leftrightarrow 1 < 1 (F)$$

În acest caz sistemul nu este stabil IMEM ci se află la limita de stabilitate.

Dacă $\eta \in (-\infty, -2] \cup [2, \infty) - \{-10/3, 10/3\} \Rightarrow$ sistemul nu este stabil IMEM.

Dacă $\eta \in \{-10/3, 10/3\} \Rightarrow$ sistemul este stabil IMEM.

Dacă $\eta \in (-2, 2) \Rightarrow$ sistemul se afla la limita de stabilitate (nu este stabil IMEM).

c) Dacă $\eta = 10/3$ rezultă

$$G(z^{-1}) = \frac{2 - 18z^{-2}}{1 - \eta z^{-1} + z^{-2}} = \frac{2(z-3)(z+3)}{(z-3)(z-1/3)} \Rightarrow G(z) = \frac{2(z+3)}{z-1/3} = \frac{2+6z^{-1}}{1-1/3z^{-1}}.$$

Deci $y(k)$ depinde de: $y(k-1), u(k), u(k-1)$.

Dacă $\eta = -10/3$

$$G(z^{-1}) = \frac{2 - 18z^{-2}}{1 - \eta z^{-1} + z^{-2}} = \frac{2(z-3)(z+3)}{(z+3)(z+1/3)} \Rightarrow G(z) = \frac{2(z-3)}{z+1/3} = \frac{2-6z^{-1}}{1+1/3z^{-1}}.$$

Deci $y(k)$ depinde de: $y(k-1), u(k), u(k-1)$

Problema 7.19.

Se consideră o instalație având funcția de transfer

$$G(s) = \frac{5e^{-20s}}{10s+1},$$

pentru care se determină un model discret folosind o rețea neuronală dinamică constituită dintr-o ADALINE și întârzierile necesare.

Să se reprezinte grafic structura rețelei neuronale dinamice și să se precizeze parametrii ADALINE considerând că discretizarea lui $G(s)$ se realizează exact, adică

$$G_d(z) = (1 - z^{-1}) \cdot Z \left\{ \frac{G(s)}{s} \right\},$$

pentru:

- perioada de eșantionare de 1 sec;
- perioada de eșantionare de 2 sec.

Indicații și răspunsuri:

$$\frac{G(s)}{s} = e^{-20s} F(s), \text{ cu } F(s) = \frac{5}{s(10s+1)} \Rightarrow Z \left\{ \frac{G(s)}{s} \right\} = Z \{ e^{-20s} \} Z \{ F(s) \} = z^{-\frac{20}{T}} Z \{ F(s) \}.$$

$$F(z) = Z \{ F(s) \} = \sum \operatorname{Rez} F(s) \frac{z}{z - e^{sT}} \Big|_{\text{poles} F(s)}$$

$$\left. \begin{aligned} R_1 &= \lim_{s \rightarrow 0} \frac{z}{z - e^{sT}} \frac{5}{s(10s+1)} s = \frac{5z}{z-1} \\ R_2 &= \lim_{s \rightarrow -\frac{1}{10}} \frac{z}{z - e^{sT}} \frac{5}{s(10s+1)} (s+1/10) = \frac{z}{z - e^{-\frac{T}{10}}} \frac{5}{-\frac{1}{10}} = \frac{-5z}{z - e^{-\frac{T}{10}}} \end{aligned} \right\} \Rightarrow F(z) = \frac{5z}{z-1} - \frac{5z}{z - e^{-\frac{T}{10}}}$$

$$G_d(z) = (1 - z^{-1}) z^{-\frac{20}{T}} \left(\frac{5z}{z-1} - \frac{5z}{z - e^{-\frac{T}{10}}} \right)$$

a) Pentru $T = 1$ sec :

$$G_d(z) = (1 - z^{-1}) z^{-20} \left(\frac{5z}{z-1} - \frac{5z}{z - \frac{1}{\sqrt[10]{e}}} \right) = 5(1 - z^{-1}) z^{-19} \left(\frac{1}{z-1} - \frac{1}{z-0.9} \right) = \frac{0.5z^{-20}}{(z-0.9)} \Leftrightarrow$$

$$\Leftrightarrow G_d(z^{-1}) = \frac{0.5z^{-21}}{(1-0.9z^{-1})} \Rightarrow y[k] = 0.9y[k-1] + 0.5u[k-21]$$

Modelul neuronal liniar constă dintr-o rețea ADALINE cu parametrii $\mathbf{w} = [0.5 \ 0.9]$, $b = 0$, la intrările căreia se aplică $y[k-1]$ și $u[k-21]$.

b) Pentru $T = 2$ sec

$$G_d(z) = (1 - z^{-1})z^{-10} \left(\frac{5z}{z-1} - \frac{5z}{z - \frac{1}{\sqrt[3]{e}}} \right) = \frac{(z-1)}{z} z^{-10} 5z \left(\frac{0.2}{(z-1)(z-0.8)} \right) = \frac{z^{-10}}{(z-0.8)} \Leftrightarrow$$

$$\Leftrightarrow G_d(z^{-1}) = \frac{z^{-11}}{(1-0.8z^{-1})} \Rightarrow y[k] = 0.9y[k-1] + u[k-11].$$

Modelul neuronal liniar constă dintr-o rețea ADALINE cu parametrii $\mathbf{w} = [1 \ 0.8]$, $b = 0$, la intrările căreia se aplică $y[k-1]$ și $u[k-11]$.

Parametrii rețelei ADALINE sunt: $\mathbf{w} = [1 \ 0.8]$; $b = 0$.

Problema 7.20.

Se consideră un proces având funcția de transfer

$$G(s) = \frac{30s + 5}{35s^2 + 12s + 1},$$

pentru care se determină un model discret folosind o rețea neuronală dinamică alcătuită dintr-o ADALINE și întârzierile adecvate.

Să se reprezinte grafic structura rețelei neuronale dinamice și parametrii ADALINE presupunând că discretizarea lui $G(s)$ se realizează pentru o perioadă de eșantionare $T = 1$ sec, folosind:

- metoda dreptunghiului în avans, care folosește $s \approx \frac{z-1}{T}$;
- metoda dreptunghiului în întârziere, care folosește $s \approx \frac{z-1}{Tz}$;
- metoda trapezului, care folosește $s \approx \frac{2}{T} \frac{z-1}{z+1}$.

Indicații și răspunsuri:

$$a) G(z) = \frac{30(z-1)+5}{35(z-1)^2 + 12(z-1)+1} = \frac{30z-25}{35z^2-58z-10} \Leftrightarrow G(z^{-1}) = \frac{30/35z^{-1}-25/35z^{-2}}{1-58/35z^{-1}-10/35z^{-2}}.$$

Modelul neuronal liniar constă dintr-o rețea ADALINE cu parametrii $b = 0$ și $\mathbf{w} = 1/35[30 \ -25 \ 58 \ 10]$, la intrările căreia se aplică $y[k-1]$, $y[k-2]$, $u[k-1]$, $u[k-2]$.

$$b) G(z) = \frac{30 \frac{z-1}{z} + 5}{35 \left(\frac{z-1}{z} \right)^2 + 12 \left(\frac{z-1}{z} \right) + 1} = \frac{30z-30+5z}{35(z-1)^2 + 12z(z-1) + z^2} \frac{z^2}{z} = \frac{35z^2-30z}{48z^2-82z+35} \Leftrightarrow$$

$$\Leftrightarrow G(z^{-1}) = \frac{35/48-30/48z^{-1}}{1-82/48z^{-1}+35/48z^{-2}}$$

Modelul neuronal liniar constă dintr-o rețea ADALINE cu parametrii $b = 0$ și $\mathbf{w} = 1/48[35 \ -30 \ 82 \ -35]$, la intrările căreia se aplică $y[k-1]$, $y[k-2]$, $u[k]$, $u[k-1]$.

$$\begin{aligned}
 \text{c) } G(z) &= \frac{60 \frac{z-1}{z+1} + 5}{140 \left(\frac{z-1}{z+1} \right)^2 + 24 \left(\frac{z-1}{z+1} \right) + 1} = \frac{60z - 60 + 5z + 5}{140(z-1)^2 + 24(z+1)(z-1) + (z+1)^2} \frac{(z+1)^2}{z+1} \Leftrightarrow \\
 \Leftrightarrow G(z) &= \frac{65z^2 + 10z - 55}{165z^2 - 278z + 117} \Leftrightarrow G(z^{-1}) = \frac{65/165 + 10/165 z^{-1} - 55/165 z^{-2}}{1 - 278/165 z^{-1} + 117/165 z^{-2}}
 \end{aligned}$$

Modelul neuronal liniar constă dintr-o rețea ADALINE cu parametrii $b=0$ și $\mathbf{w} = 1/165[65 \ 10 \ -55 \ 278 \ 117]$, la intrările căreia se aplică $y[k-1]$, $y[k-2]$, $u[k]$, $u[k-1]$, $u[k-2]$.

Capitolul 8

PROBLEME CU SOLUȚII NUMERICE

Problema 8.1.

Se consideră un neuron cu o intrare (vezi figura 2.1.1). Să se elaboreze un program MATLAB pentru a studia modul de compunere a funcțiilor, considerând la intrarea neuronului valori $x \in [-v, v]$, cu $v = 10$.

Se vor prezenta grafic următoarele trei funcții:

1. $u(x)$ (relația (4)) pentru $x \in [-v, v]$. Pasul de discretizare al intervalului $[-v, v]$ se va alege astfel încât să asigure calitatea plotării;
2. $\sigma(x)$ pentru $u \in [u_{\min}, u_{\max}]$, acest interval reprezentând mulțimea valorilor lui $u(x)$, în condițiile de la punctul 1. Pasul de discretizare al intervalului $[u_{\min}, u_{\max}]$, se va alege astfel încât să asigure calitatea plotării;
3. $f(x) = \sigma(u(x))$ (relația (3)) pentru $x \in [-v, v]$.

Se vor considera cazurile funcției de activare de tip:

- a) sigmoid bipolar;
- b) treaptă unipolară.

Pentru fiecare din cazurile a) și b), se vor studia următoarele situații:

- | | |
|--------------------------|---------------------------|
| (i) $w = 1, b = 0$; | (ii) $w = 1, b = b^*$; |
| (iii) $w = w^*, b = 0$; | (iv) $w = w^*, b = b^*$; |

cu $w^* = \pm 0.5$ și $b^* = \pm 1$.

Se vor comenta legăturile existente între graficele rezultate la (i) - (iv), iar atunci când este posibil, aceste legături vor fi utilizate pentru *trasarea directă* a graficelor, fără a mai face apel la programul elaborat în MATLAB.

Indicații și răspunsuri:

De exemplu, pentru $w = 0.5$ și $b = 1$ se obțin graficele din figura 8.1.1.

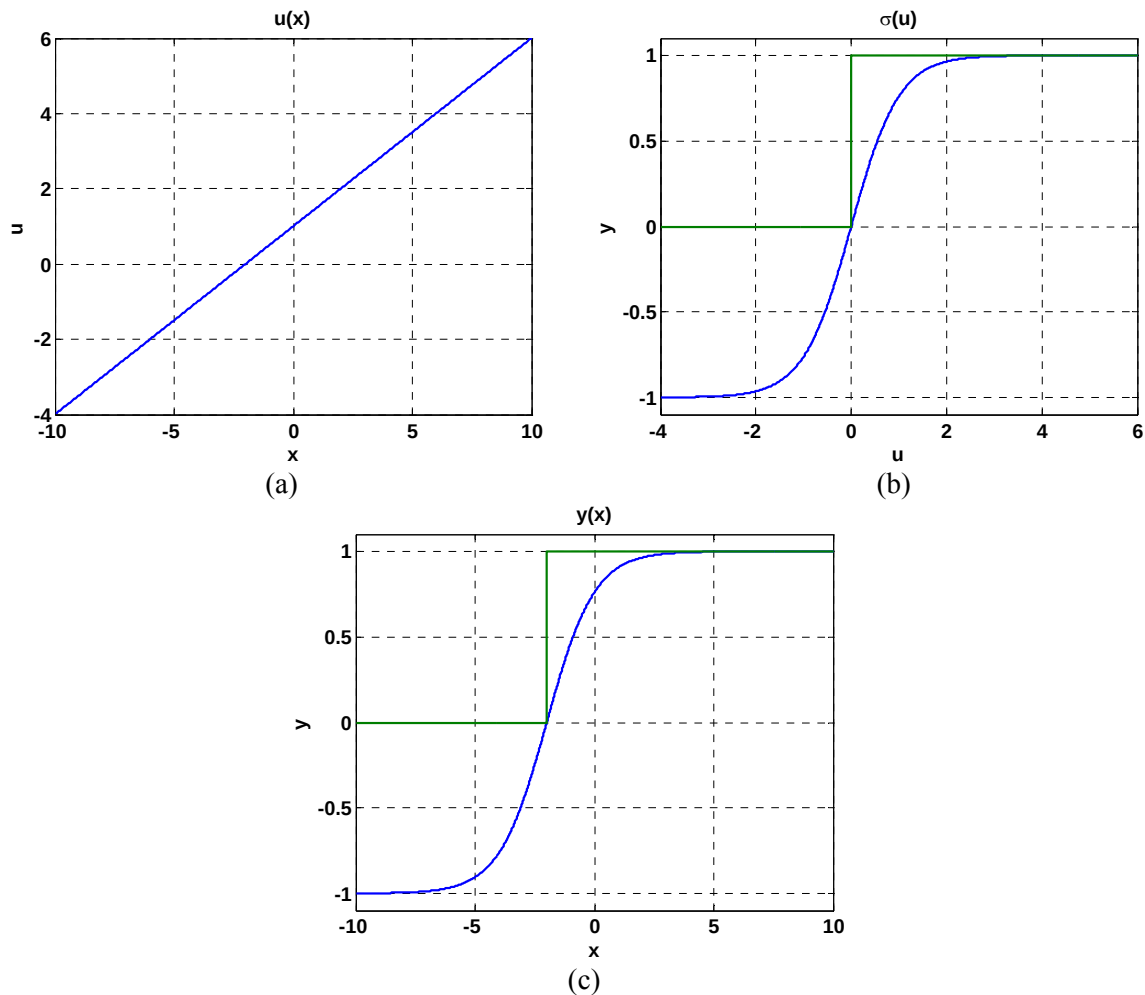


Figura 8.1.1. Reprezentări grafice ale dependenței: (a) $u(x)$, (b) $\sigma(u)$ și (c) $y(x)$, pentru neuroni cu parametrii $w = 0.5$ și $b = 1$.

Aceste figuri au fost obținute utilizând următorul program MATLAB:

```
% Problema 8.1
clear all; close all; clc

v= 5; vs= 0.01;
x= -v:vs:v;
% introducerea parametrilor neuronilor
w= input('pondere w = ');
b= input('deplasare b= ');
% calculul iesirilor
u= w*x+b;
y1= tansig(u); % sigmoid bipolar
y2= hardlim(u); % treapta unipolara
```

```
% reprezentari grafice
figure(1); plot(x, u);
title('u(x)'); xlabel('x'); ylabel('u'); grid

uv= min(u):vs:max(u);
y1v= tansig(uv);
y2v= hardlim(uv);

figure(2); plot(uv, y1v, uv, y2v)
title('\sigma(u)'); xlabel('u'); ylabel('y')
axis([min(u) max(u) -1.1 1.1]); grid

figure(3); plot(x, y1, x, y2)
title('y(x)'); xlabel('x'); ylabel('y');
axis([-v v -1.1 1.1]); grid
```

Problema 8.2.

Se consideră un neuron cu două intrări (vezi figura 2.1.2 pentru $n = 2$). Să se elaboreze un program MATLAB pentru a studia modul de compunere a funcțiilor, considerând la intrarea neuronului, vectorii: $\mathbf{x} = [x_1 \ x_2]^T$, pentru $(x_1, x_2) \in [-v_1, v_1] \times [-v_2, v_2]$, cu $v_1 = v_2 = 10$.

Se vor prezenta grafic următoarele trei funcții:

1. $u(\mathbf{x})$ pentru $(x_1, x_2) \in [-v_1, v_1] \times [-v_2, v_2]$;
2. $\sigma(u)$ pentru $u \in [u_{\min}, u_{\max}]$, acest interval reprezentând mulțimea în care $u(\mathbf{x})$ ia valori, în condițiile de la punctul 1. Pasul de discretizare us al intervalului $[u_{\min}, u_{\max}]$, se va alege astfel încât să asigure calitatea plotării.
3. $f(\mathbf{x}) = \sigma(u(\mathbf{x}))$ pentru $(x_1, x_2) \in [-v_1, v_1] \times [-v_2, v_2]$.

Se vor considera cazurile funcției de activare:

- a) sigmoid bipolar;
- b) treaptă unipolară.

Pentru fiecare din cazurile a) și b), se vor studia următoarele situații:

- | | |
|--|---|
| (i) $\mathbf{w} = [1 \ 0]$, $b = 0$; | (ii) $\mathbf{w} = [1 \ 0]$, $b = b^*$; |
| (iii) $\mathbf{w} = [0 \ 1]$, $b = 0$; | (iv) $\mathbf{w} = [0 \ 1]$, $b = b^*$; |
| (v) $\mathbf{w} = [1 \ 1]$, $b = 0$; | (vi) $\mathbf{w} = [1 \ 1]$, $b = b^*$; |
| (vii) $\mathbf{w} = [w_1^* \ w_2^*]$, $b = 0$; | (viii) $\mathbf{w} = [w_1^* \ w_2^*]$, $b = b^*$; |

cu $w_1^*, w_2^* \in \{\pm 0.5, \pm 1, \pm 2\}$ și $b^* \in \{0, \pm 2\}$.

Se vor comenta legăturile existente între graficele existente la (i) - (viii), iar atunci când este posibil, aceste legături vor fi utilizate pentru *trasarea directă* a graficelor, fără a mai face apel la programul elaborat în MATLAB.

Indicații și răspunsuri:

Pentru $w_1 = 0.5$, $w_2 = -0.5$ și $b = 0$ se obțin graficele din figurile 8.2.1 și 8.2.2.

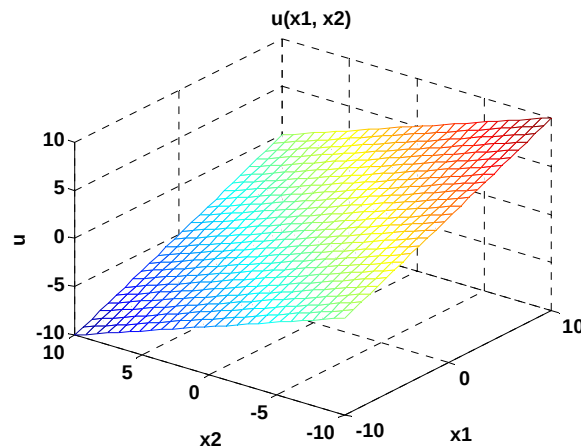


Figura 8.2.1. Dependența $u(x_1, x_2)$ pentru un neuron cu 2 intrări.

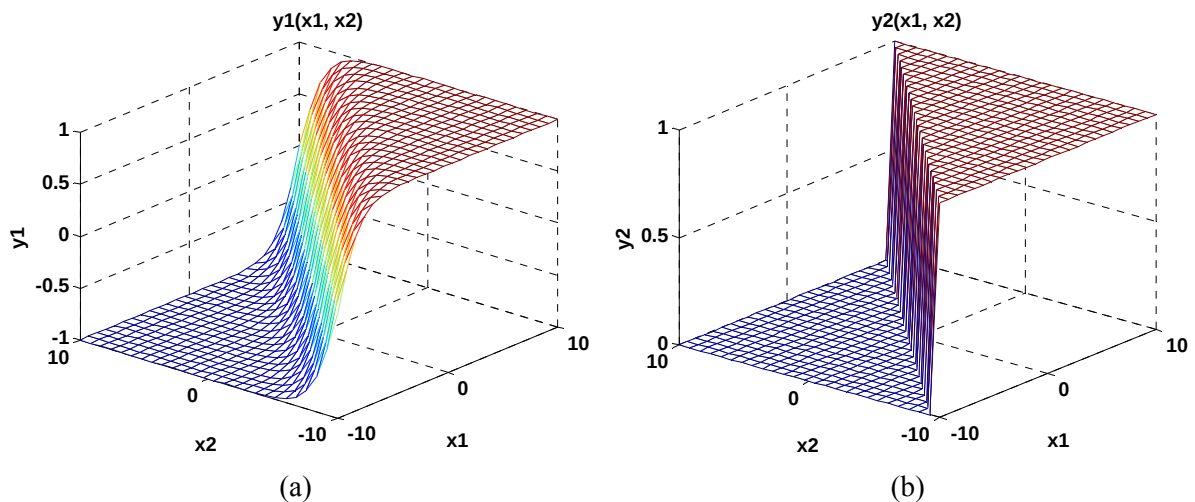


Figura 8.2.2. Dependența $y(x_1, x_2)$ pentru un neuron cu funcție de activare de tip
(a) sigmoid bipolar și (b) treaptă unipolară.

Aceste figuri au fost obținute utilizând următorul program MATLAB:

```
% Problema 8.2
clear all; close all; clc

v1= 10; vs1= 0.1; x1= -v1:vs1:v1;
v2= 10; vs2= 0.1; x2= -v2:vs2:v2;
% introducerea parametrilor neuronilor
w1= input('pondere w1 = ');
w2= input('pondere w2 = ');
b= input('deplasare b= ');
% calculul iesirilor
```

```

[X1, X2]= meshgrid(x1, x2);
U= w1*X1 + w2*X2+b;
Y1= tansig(U);
Y2= double(hardlim(U)); % hardlim returneaza valoare logica!!
% reprezentari grafice
figure(1); mesh(X1, X2, U)
title('u(x1, x2)'); xlabel('x1'); ylabel('x2'); zlabel('u')

uv= min(min(U)):vs1:max(max(U));
y1v= tansig(uv); y2v= hardlim(uv);

figure(2); plot(uv, y1v, uv, y2v)
title('\sigma(u)'); xlabel('u'); ylabel('y')
axis([min(min(U)) max(max(U)) -1.1 1.1]); grid

figure(3); mesh(X1, X2, Y1)
title('y1(x1, x2)'); xlabel('x1'); ylabel('x2'); zlabel('y1')

figure(4); mesh(X1, X2, Y2)
title('y2(x1, x2)'); xlabel('x1'); ylabel('x2'); zlabel('y2')

```

Problema 8.3.

Se consideră o rețea neuronală cu un singur strat cu doi neuroni (vezi figura 2.1.3 pentru $n=1$, $p=2$). Să se elaboreze un program MATLAB pentru a studia modul de compunere a funcțiilor, considerând la intrarea rețelei, valorile $x \in [-v, v]$, cu $v=10$.

Se vor prezenta grafic următoarele șase funcții:

1. $u_1(x)$ și $u_2(x)$ pentru $x \in [-v, v]$;
2. $\sigma_1(u_1)$ pentru $u_1 \in [u_{1\min}, u_{1\max}]$, acest interval reprezentând mulțimea valorilor lui $u(x)$, în condițiile de la punctul 1. Pasul de discretizare al intervalului $[u_{1\min}, u_{1\max}]$, se va alege astfel încât să asigure calitatea plotării.
3. $\sigma_2(u_2)$ pentru $u_2 \in [u_{2\min}, u_{2\max}]$, acest interval reprezentând mulțimea valorilor lui $u_2(x)$, în condițiile de la punctul 1. Pasul de discretizare al intervalului $[u_{2\min}, u_{2\max}]$, se va alege astfel încât să asigure calitatea plotării.
4. $f_1(x) = \sigma_1(u_1(x))$ și $f_2(x) = \sigma_2(u_2(x))$, pentru $x \in [-v, v]$.

Se vor considera cazurile funcției de activare:

- a) sigmoid bipolar;
- b) treaptă unipolară.

Pentru fiecare din cazurile a) și b), se vor studia următoarele situații:

- (i) $W = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$, $b = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$; (ii) $W = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$, $b = \begin{bmatrix} b_1^* \\ b_2^* \end{bmatrix}$; (iii) $W = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$, $b = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$;
- (iv) $W = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$, $b = \begin{bmatrix} b_1^* \\ b_2^* \end{bmatrix}$; (v) $W = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$, $b = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$; (vi) $W = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$, $b = \begin{bmatrix} b_1^* \\ b_2^* \end{bmatrix}$;

$$(vii) \quad \mathbf{W} = \begin{bmatrix} w_1^* \\ w_2^* \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}; \quad (viii) \quad \mathbf{W} = \begin{bmatrix} w_1^* \\ w_2^* \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} b_1^* \\ b_2^* \end{bmatrix};$$

cu $w_1^*, w_2^* \in \{0, \pm 0.5, \pm 1, \pm 2\}$ și $b^* \in \{0, \pm 2\}$.

Se vor comenta legăturile existente între graficele existente la (i) – (viii), iar atunci când este posibil, aceste legături vor fi utilizate pentru *trasarea directă* a graficelor, fără a mai face apel la programul elaborat în MATLAB. De asemenea, se vor comenta legăturile cu rezultatele obținute la Problema 8.1.

Indicații și răspunsuri:

Programul MATLAB este asemănător celui prezentat la problema 8.1.1.

Problema 8.4.

Se consideră o rețea neurală cu două straturi, cu câte un neuron pe fiecare strat (vezi figura 2.2.2 pentru $n_1 = 1$, $p_1 = n_2 = 1$, $p_2 = 1$). Să se elaboreze un program MATLAB pentru a studia modul de compunere a funcțiilor, considerând la intrarea rețelei, valorile $x \in [-v, v]$, cu $v = 10$. Se vor prezenta grafic următoarele trei funcții:

1. $y^1(x^1)$ pentru $x \in [-v, v]$
2. $y^2(x^2)$ pentru $x^2 \in [y_{\min}, y_{\max}]$, acest interval reprezentând mulțimea valorilor lui $y^1(x^1)$, în condițiile de la punctul 1. Pasul de discretizare Δy al intervalului $[y_{\min}, y_{\max}]$ se va alege astfel încât să asigure calitatea plotării.
3. $y^2(x^1)$ pentru $x \in [-v, v]$.

Se vor considera cazurile funcției de activare:

- a) sigmoid bipolar pe ambele straturi;
- b) treaptă unipolară pe ambele straturi.

Pentru fiecare din cazurile a) și b), se vor considera valorile $w^1, w^2 \in \{0, \pm 0.5, \pm 1, \pm 2\}$ și $b^1, b^2 \in \{0, \pm 2\}$. Se vor comenta legăturile cu rezultatele obținute la Problema 8.1.

Indicații și răspunsuri:

Programul MATLAB este asemănător celui prezentat la problema 8.1.1.

Problema 8.5.

Se consideră un perceptron cu două intrări care se utilizează pentru a clasifica următorii 5 vectori bidimensionali:

$$\mathbf{x}_1 = \begin{bmatrix} 0.3 \\ -0.4 \end{bmatrix}, \quad \mathbf{x}_2 = \begin{bmatrix} -0.4 \\ -0.5 \end{bmatrix}, \quad \mathbf{x}_3 = \begin{bmatrix} 0.5 \\ 1.5 \end{bmatrix}, \quad \mathbf{x}_4 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, \quad \mathbf{x}_5 = \begin{bmatrix} -0.3 \\ 0.5 \end{bmatrix},$$

în două clase, notate $\mathcal{C}_1$ și $\mathcal{C}_2$, astfel:

$$\mathbf{x}_1 \in \mathcal{C}_2, \quad \mathbf{x}_2 \in \mathcal{C}_1, \quad \mathbf{x}_3 \in \mathcal{C}_2, \quad \mathbf{x}_4 \in \mathcal{C}_1, \quad \mathbf{x}_5 \in \mathcal{C}_1.$$

Să se elaboreze un program MATLAB care să creeze și să antreneze perceptronul pentru un număr de iterații stabilit de către utilizator. Se vor reprezenta grafic vectorii de intrare și se va trasa dreapta definită prin parametrii perceptronului antrenat. În finalul antrenării se va verifica, prin simulare, dacă rețeaua antrenată permite realizarea clasificării.

Se vor raporta următoarele rezultate ale experimentelor:

- (i) valorile parametrilor inițiali ai rețelei;
- (ii) valorile parametrilor rețelei la sfârșitul antrenării;
- (iii) reprezentarea grafică a vectorilor de intrare și dreapta de separare a claselor definită de rețea, în urma antrenării;
- (iv) numărul de iterații cât a durat antrenarea;
- (v) reprezentarea grafică a erorii pătratice globale ca funcție de numărul de iterații parcurse.

Indicații și răspunsuri:

Se rulează următorul program MATLAB.

```
% Problema 8.5
clear all; close all; clc

% vectori prototip
x1= [0.3; -0.4];    x2= [-0.4; -0.5];
x3= [0.5; 1.5];    x4= [0; 1];
x5= [-0.3; 0.5];
P= [x1 x2 x3 x4 x5];

% vectori tinta
z1= 0; z2= 1;
T= [z2 z1 z2 z1 z1];

% reprezentare grafica a vectorilor de clasificat
figure(1); plotpv(P, T); pause(1);

% creare retea neurala
PR= [min(P')' max(P')'];
net = newp(PR, 1);
w= net.IW{1,1};    % pondere
b= net.b{1};       % deplasare
fprintf('\n Valori initiale parametri retea neurala:');
fprintf('\n\t ponderi: %6.2f, %6.2f \n\t deplasare: %6.2f\n', w(1), w(2), b);

figure(1); plotpc(w, b); pause(1);

net.trainParam.epochs = 1; % numar de epoci pentru antrenare

for i=1:5,
    fprintf('\n Iteratia %d : ', i);
    [net, tr, Y, E] = train(net,P,T); % antrenare retea 1 epoca
    w= net.IW{1,1};    % pondere
```

```

b= net.b{1};      % deplasare
fprintf('\n Valori parametri retea neurala:');
fprintf('\n\t ponderi: %6.2f, %6.2f \n\t deplasare: %6.2f', w(1), w(2), b);
perf = mse(E); % indicator de performanta
fprintf('\n Valoare indicator performanta MSE: %6.2f\n', perf);
figure(1); clf;
plotpv(P, T); plotpc(w, b);
str= ['Iteratia ', num2str(i)]; title(str);
pause(1)
end

```

Parametrii rețelei neurale sunt inițializați cu valori nule. După o iterație (1 epocă de antrenare a rețelei neurale) se obțin parametrii $w = [0.9 \ 1]$, $b = -1$, rețeaua realizând separarea vectorilor prototip prezentată în figura 8.5.1.(a). După 5 iterații (5 epoci de antrenare a rețelei neurale) se ajunge la valorile $w = [2 \ -0.3]$, $b = 0$, separarea vectorilor prototip fiind prezentată în figura 8.5.1.(b).

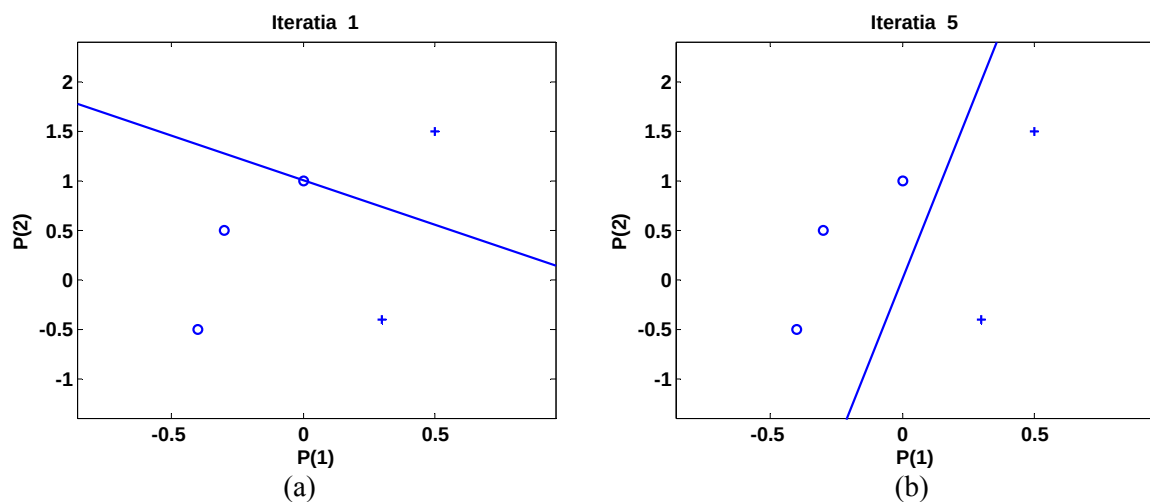


Figura 8.5.1. Separarea vectorilor prototip realizată de rețeaua neurală antrenată:
(a) 1 epocă, (b) 5 epoci.

Problema 8.6.

Să se reia problema 8.5 pentru cazul când se urmărește clasificarea celor 5 vectori astfel:

$$x_1 \in \mathcal{C}_2, x_2 \in \mathcal{C}_1, x_3 \in \mathcal{C}_2, x_4 \in \mathcal{C}_1, x_5 \in \mathcal{C}_2.$$

Indicații și răspunsuri:

În programul prezentat la problema 8.5 se modifică matricea vectorilor țintă în conformitate cu noua separare dorită a vectorilor prototip:


```
% vectori tinta
z1= 0; z2= 1;
T= [z2 z1 z2 z1 z2];
```

Se observă că impunând această clasificare, vectorii prototip nu sunt liniar separabili, astfel că separare dorită nu poate fi realizată cu un perceptron.

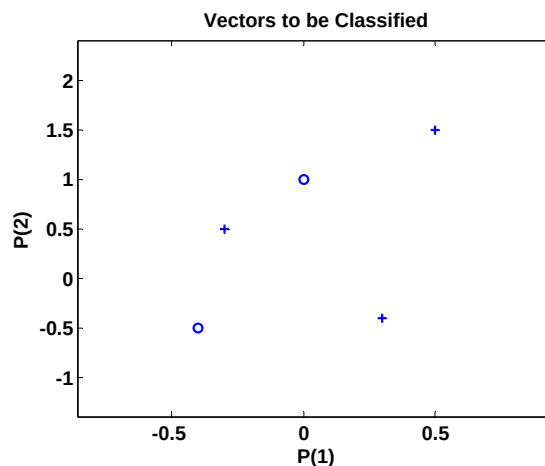


Figura 8.6.1. Clasificare vectorilor prototip impusă în problema 8.6.

Problema 8.7.

Se consideră o rețea perceptron cu un singur strat cu doi neuroni, având două intrări, care se utilizează pentru a clasifica următorii 7 vectori bidimensionali:

$$\mathbf{x}_1 = \begin{bmatrix} 0.8 \\ 1.6 \end{bmatrix}, \mathbf{x}_2 = \begin{bmatrix} 1.25 \\ 1 \end{bmatrix}, \mathbf{x}_3 = \begin{bmatrix} 0.3 \\ 0.5 \end{bmatrix}, \mathbf{x}_4 = \begin{bmatrix} -1 \\ -1.25 \end{bmatrix}, \mathbf{x}_5 = \begin{bmatrix} 0 \\ 0.2 \end{bmatrix}, \mathbf{x}_6 = \begin{bmatrix} -0.3 \\ 0.8 \end{bmatrix}, \mathbf{x}_7 = \begin{bmatrix} -0.5 \\ -1.5 \end{bmatrix},$$

în patru clase, notate $\mathcal{C}_1, \mathcal{C}_2, \mathcal{C}_3, \mathcal{C}_4$, astfel:

$$\mathbf{x}_1 \in \mathcal{C}_1, \mathbf{x}_2 \in \mathcal{C}_2, \mathbf{x}_3 \in \mathcal{C}_3, \mathbf{x}_4 \in \mathcal{C}_4, \mathbf{x}_5 \in \mathcal{C}_3, \mathbf{x}_6 \in \mathcal{C}_3, \mathbf{x}_7 \in \mathcal{C}_4.$$

Să se elaboreze un program MATLAB care să creeze și să antreneze rețeaua.

Se vor reprezenta grafic vectorii de intrare și se vor trasa cele două drepte definite prin parametrii rețelei antrenate. În finalul antrenării se va verifica, prin simulare, dacă rețeaua antrenată permite realizarea clasificării.

Se vor raporta următoarele rezultate ale experimentelor:

- (i) valorile parametrilor inițiali ai rețelei;
- (ii) valorile parametrilor rețelei la sfârșitul antrenării;
- (iii) reprezentarea grafică a vectorilor de intrare și dreptele de separare a claselor definite de rețea, în urma antrenării;
- (iv) numărul de iterații cât a durat antrenarea;

- (v) reprezentarea grafică a erorii pătratice globale ca funcție de numărul de iterații parcurse.

Pentru antrenarea rețelei neurale se vor considera următorii vectori țintă:

$$(A) \ z_1 = \begin{bmatrix} 0 \\ 0 \end{bmatrix}, z_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, z_3 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, z_4 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}; \quad (B) \ z_1 = \begin{bmatrix} 0 \\ 0 \end{bmatrix}, z_2 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}, z_3 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, z_4 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}.$$

Ce observați?

Indicații și răspunsuri:

Pentru vectorii țintă considerați în cazul (A) se utilizează programul MATLAB prezentat mai jos. Graficele din figura 8.7.1 prezintă evoluția indicatorului de performanță MAE pe parcursul antrenării rețelei și clasificarea realizată de rețeaua antrenată.

Utilizând vectorii țintă precizați la punctul (B) nu se poate realiza separarea dorită.

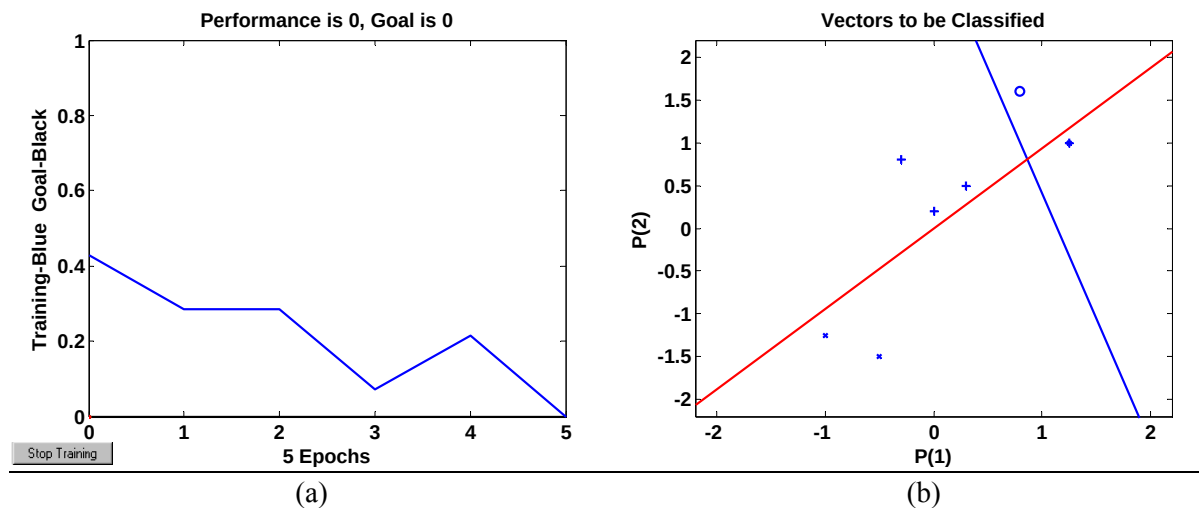


Figura 8.7.1. (a) Evoluția indicatorului de performanță MAE pe parcursul antrenării rețelei.
(b) Clasificarea realizată de rețeaua antrenată.

```
% Problema 8.7.a
clear all; close all; clc

% vectori prototip
x1= [0.8; 1.6];   x2= [1.25; 1];
x3= [0.3; 0.5];   x4= [-1; -1.25];
x5= [0; 0.2];     x6= [-0.3; 0.8];
x7= [-0.5; -1.5];
P= [x1 x2 x3 x4 x5 x6 x7];

% vectori tinta
z1= [0; 0];       z2= [0; 1];
z3= [1; 0];       z4= [1; 1];
T= [z1 z2 z3 z4 z3 z3 z4];
```

```

% reprezentare grafica a vectorilor de clasificat
figure(1); plotpv(P, T)
% legend('class C_1', 'class C_2', 'class C_3', 'class C_4', 4)
axis([-2.2 2.2 -2.2 2.2])

% creare retea neurala
PR= [min(P)' max(P)'];
net = newp(PR, 2); % retea cu 2 neuroni

% parametri initiali retea
disp(' Valori initiale parametri retea:');
w= net.IW{1,1} % ponderile celor 2 neuroni
b= net.b{1,1} % deplasările celor 2 neuroni

% antrenare retea
net.trainParam.epochs = 10; % numar de epoci de antrenare
net.trainFcn= 'trainb';
[net, tr, Y, E] = train(net,P,T);
% parametri finali retea
disp(' Valori finale parametri retea:');
w= net.IW{1,1}
b= net.b{1,1}

perf = mse(E); % indicator de performanta
fprintf('\n Valoare indicator performanta MSE: %6.2f\n', perf);

% reprezentare grafica a clasificarii realizate de retea antrenata
figure(1); hold on
plotpc(w, b)

```

Problema 8.8.

Se va studia și lansa în execuție programul MATLAB *arna_wh*, prezentat mai jos, care antrenează un neuron ADALINE cu o singură intrare, utilizând algoritmul Widrow-Hoff, pentru vectorii prototip:

$$x_1 = 0, \quad x_2 = 1,$$

și, respectiv, vectorii țintă:

$$z_1 = 0, \quad z_2 = 2.$$

1. Să se determine valorile parametrilor neuronului utilizând funcția *newlind*.
2. Să se determine valorile parametrilor neuronului făcând modificările necesare în programul *arna_wh* în următoarele situații:

- a) - număr de iterații maxime 30,
- pragul erorii 0.01,
- rata de învățare dată de funcția *maxlinlr*;
- b) - număr de iterații maxime 30,
- pragul erorii 0,

- rata de învățare dată de funcția `maxlinlr`;
- c) - număr de iterații maxime 40,
 - pragul erorii 0.0001,
 - rata de învățare 0.2;
- d) - număr de iterații maxime 40,
 - pragul erorii 0.01,
 - rata de învățare 0.6;
- e) - număr de iterații maxime 20,
 - pragul erorii 0.01,
 - rata de învățare 0.8.

Pentru fiecare situație se vor preciza valorile selectate inițial pentru parametrii rețelei, valorile rezultate în final pentru parametrii rețelei, numărul de iterații cât a durat optimizarea și valoarea finală a erorii pătratice globale. Se va preciza condiția care determină încheierea antrenării. De asemenea se va schița suprafața erorii pătratice globale și traseul parcurs în planul parametrilor.

În cazul când se consideră relevant, se pot testa mai multe condiții inițiale pentru fiecare din situațiile a) - e).

Indicații și răspunsuri:

```
% Problema 8.8, program ARNA_WH

clear all; close all; clc

% vectori prototip
x1= 0;  x2= 1;
P= [x1 x2];

% vectori tinta
z1= 0;      z2= 2;
T= [z1 z2];

%%%%%%%%
% a. proiectare retea utilizand functia newlind
net= newlind(P, T);
% parametri retea
fprintf('\n\nValorile ale parametrilor rețelei proiectate cu NEWLIND:');
fprintf('\n\t pondere: %g\n\t deplasare: %g', net.IW{1,1}, net.b{1});
% simulare
Y= sim(net, P);
% reprezentare grafica
figure(1);
plot(P, T, '*r'), hold on; plot(P, Y, 'd-b')
xlabel('vectori prototip P');
ylabel('vectori tinta T');
fprintf('\n\n<ENTER...>\n'); pause
```

```

%%%%%%%%%
% b. crearea unei rețele ADALINE ce va fi antrenată cu rata de învățare
% lr= maxlinlr(P)      % valoarea maximă ce asigură convergența
% algoritmului de antrenare sau cu
lr= input('Rata de învățare: '); % valoarea introdusă de utilizator

lin= newlin(minmax(P), 1, [0], lr);
    % neuronul liniar creat este pregătit pt a fi antrenat cu subrutina "learnwh"
    % (se verifică valorile cerute prin apelarea "help learnwh")
    % eroarea este dată de funcția "mse" (eroarea medie pătratică)

% parametrii specifici antrenării
iteratii=input('Nr de iteratii= ');
lin.trainParam.epochs=1; % la fiecare iteratie rețeaua va fi antrenată 1 epocă
    % pentru a putea reprezenta grafic evoluția parametrilor
lin.trainParam.goal=input('Eroarea medie pătratică dorită= ');
lin.trainParam.show=inf; % eroarea nu este afișată în timpul antrenării
    % eroarea va fi memorată la fiecare pas și afișarea va fi făcută la sfârșit

% suprafața erorii
w_range = -2:0.4:6;    b_range = -4:0.4:4;
figure;
ES = errsurf(P,T,w_range,b_range,'purelin')/length(T); % valorile erorii medii pătratice
plotes(w_range,b_range,ES);

% alegerea cu mouse-ul a parametrilor inițiali ai rețelei
subplot(1,2,2);
h = text(mean(get(gca,'xlim')),mean(get(gca,'ylim')),'* Alege cu mouse-ul *');
set(h,'horizontal','center', 'fontweight','bold');
[lin.IW{1,1},lin.b{1}] = ginput(1);
    % valorile alese cu mouse-ul devin parametrii inițiali ai rețelei
fprintf('\n\nValorile inițiale ale parametrilor rețelei:');
fprintf('\n\t pondere: %g\n\t deplasare: %g', lin.IW{1,1}, lin.b{1});
delete(h);

% antrenarea rețelei și reprezentarea la fiecare pas (epocă)
% a parametrilor pe suprafața de eroare
Err(1)= mse(T-sim(lin,P)); % eroarea medie pătratică înainte de începerea antrenării;
    % vectorul Err va conține evoluția erorii
Wi(1)= lin.IW{1,1};      % vectorul care va conține evoluția ponderilor
bi(1)= lin.b{1};        % vectorul care va conține evoluția deplasărilor
h= plotep(lin.IW{1,1}, lin.b{1}, Err(1));
    % parametrii inițiali reprezentați pe suprafața de eroare
    % (cu eroarea corespunzătoare)
fprintf('\nEroarea medie pătratică inițială = %g\n', Err(1)) ;

for i= 1:iteratii,
    [lin, tr, Y, E] = train(lin,P,T);
    h= plotep(lin.IW{1,1}, lin.b{1}, tr.perf(1),h);
        % evoluția parametrilor pe suprafața de eroare
    Err(i)= tr.perf(1); % eroarea pt noile valori ale parametrilor

```

```

        Wi(i+1)= lin.IW{1,1};
        bi(i+1)= lin.b{1};
        if Err(i) <= lin.trainParam.goal
            break; % iesire pe conditia de atingere a pragului de eroare impus
        end
    end

    iteratii=i; % numarul de epoci cat a durat antrenarea
    fprintf('\n\nAntrenarea a durat %g epoci\n', iteratii);
    fprintf('\n\nValorile finale ale parametrilor retelei:');
    fprintf('\n\t pondere: %g\n\t deplasare: %g', lin.IW{1,1}, lin.b{1});
    fprintf('\nEroarea medie patratica finala = %g\n', Err(end)) ;

    % evolutia erorii
    disp('Apasati o tasta pentru a vedea evolutia erorii functie de nr de epoci...'); pause;
    figure;
    semilogy((1:iteratii), Err);
    xlabel('iteratii');
    ylabel('MSE');
    title('Eroarea medie patratica functie de numarul de iteratii');

    % dinamica aproximarii
    disp('Apasati o tasta pentru a vedea aproximarea initiala...'); pause;
    figure;
    plot(P,T,'r+'); % punctele prin care se doreste aproximarea
    xlabel('vectori prototip P');
    ylabel('vectori tinta T');
    hold on;
    h= plot(P, Wi(1)*P+bi(1), 'd-b');

    disp('Apasati o tasta pentru a vedea evolutia dinamica a aproximarii...'); pause;
    for i=1:iteratii
        pause(0.5);
        delete(h);
        h= plot(P, Wi(i+1)*P+bi(i+1), 'd-b');
        title(['Iteratia ' num2str(i)])
    end

    hold off;

```

Prin proiectare directă se obține rețeaua ADALINE cu parametrii $w = 2$ și $b = 0$. Pentru cei doi vectori prototip (denumire generică, având de fapt valori scalare) utilizați se observă că ieșirile rețelei coincid cu valorile țintă impuse (figura 8.8.1), funcția aproximată de rețeaua neurală fiind efectiv liniară.

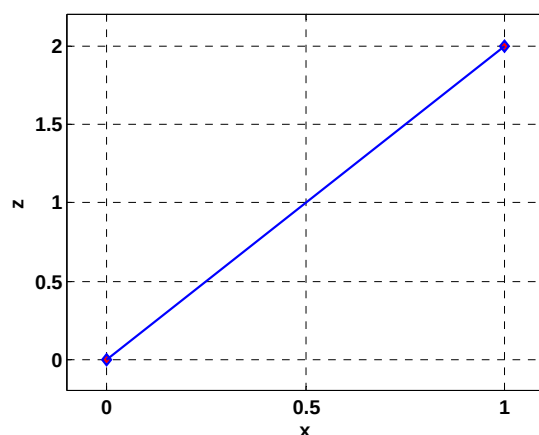


Figura 8.8.1. Comparație între valorile țintă impuse și ieșirile rețelei ADALINE proiectate corespunzătoare valorilor prototip considerate în problema 8.8.

Figura 8.8.2 prezintă evoluția parametrilor rețelei ADALINE pe suprafața de eroare pe parcursul a 20 de epoci de antrenare cu rata de învățare $\eta = 0.6$. Evoluția erorii pătratice medii funcție de iterație este ilustrată în figura 8.8.3. Valorile parametrilor rețelei ADALINE la care se ajunge în urma antrenării sunt $w = 2.01191$ și $b = -0.00732$, care asigură o eroare medie pătratică de 2.0995×10^{-5} .

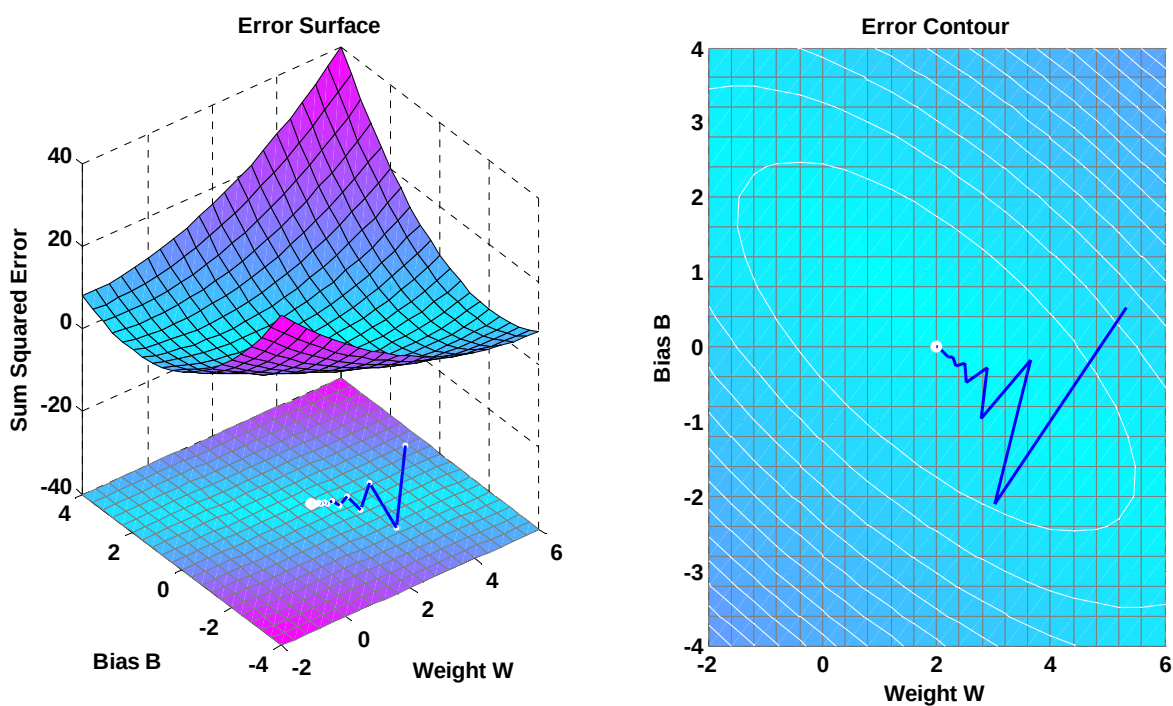


Figura 8.8.2. Evoluția parametrilor rețelei ADALINE pe suprafața de eroare pe parcursul a 20 de epoci de antrenare cu rata de învățare $\eta = 0.6$.

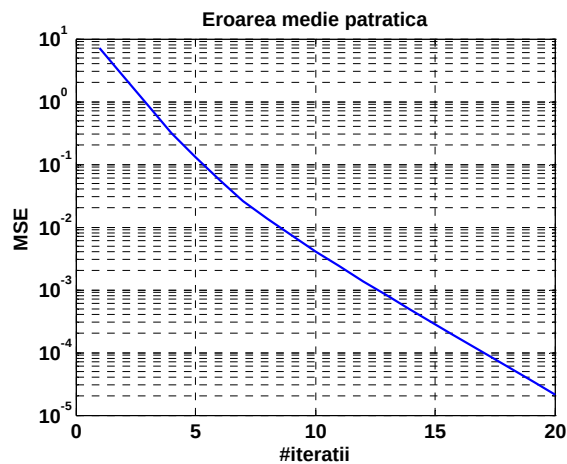


Figura 8.8.3. Evoluția erorii pătratice medii pe parcursul a 20 de epoci de antrenare cu rata de învățare $\eta = 0.6$ pentru rețeaua ADALINE din problema 8.8.

Problema 8.9.

Se va relua Problema 8.8 pentru un singur vector prototip $x_1 = 1$ și, respectiv, un singur vector țintă $z_1 = 2$, pentru situațiile 2.a) și 2.e).

Indicații și răspunsuri:

În programul prezentat la problema 8.8 se modifică în mod corespunzător matricea vectorilor prototip și cea a vectorilor țintă. Se observă că există o infinitate de rețele ADALINE care satisfac cerințele impuse în problemă.

Problema 8.10.

Se va relua Problema 8.8 pentru vectorii prototip:

$$x_1 = 0, \quad x_2 = 1, \quad x_3 = -1,$$

și, respectiv, vectorii țintă:

$$z_1 = 0, \quad z_2 = 2, \quad z_3 = -1,$$

pentru situațiile 2.a), 2.c) și 2.d).

Indicații și răspunsuri:

În programul prezentat la problema 8.8 se modifică în mod corespunzător matricea vectorilor prototip și cea a vectorilor țintă.

```
% vectori prototip
x1= 0;      x2= 1; x3= -1;
P= [x1 x2 x3];
% vectori tinta
z1= 0;      z2= 2; z3= -1;
T= [z1 z2 z3];
```

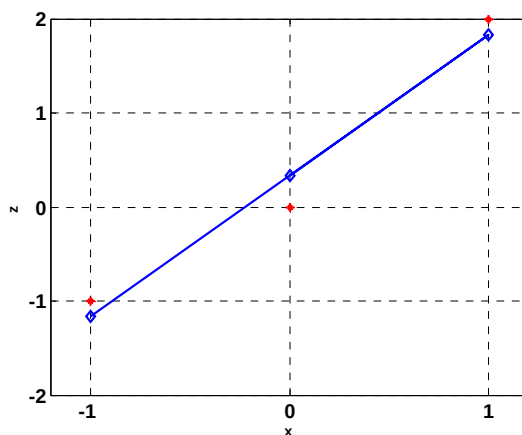



Figura 8.10.1. Comparație între valorile țintă impuse și ieșirile rețelei ADALINE proiectate corespunzătoare valorilor prototip considerate în problema 8.10.

Prin proiectare directă se obține rețeaua ADALINE cu parametrii $w = 1.5$ și $b = 0.33$. Această rețea aproximează cel mai bine, în sensul celor mai mici pătrate, funcția neliniară precizată prin cele trei perechi prototip-țintă precizată în enunțul problemei.

Problema 8.11.

Se consideră amplificatorul operațional din fig. 1 care posedă în reacție rezistența $R = 100\text{ k}\Omega$ și la intrare rezistențele R_1 , R_2 , R_3 .

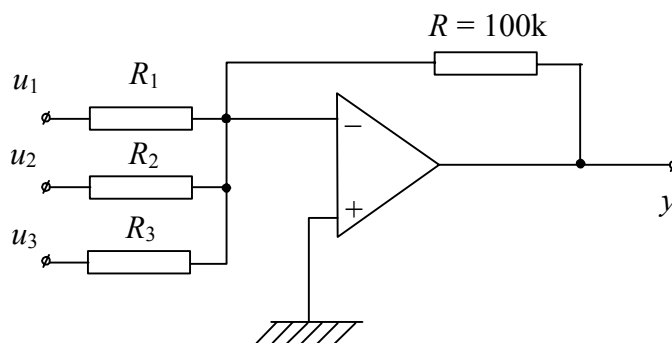


Figura 8.11.1. Schema amplificatorului operațional utilizat în Problema 8.11.

La intrarea u_1 se aplică o tensiune constantă de 2V. Pe intervalul de timp cuprins între 0 s și 5 s, pe intrările u_2 și u_3 se aplică două semnale de frecvență joasă, care, eșantionate cu perioada $T = 0.01$ s, sunt memorate în fișierul de date **opamp_wh.mat** în vectorii **uuu2** și respectiv **uuu3**. Semnalul obținut la ieșirea y , eșantionat cu aceeași perioadă, este memorat în vectorul **yyy** din fișierul de date **opamp_wh.mat**.

Utilizând datele din fișierul **opamp_wh.mat** să se determine valorile rezistențelor necunoscute R_1 , R_2 , R_3 prin următoarele metode:

- a) obținând rețeaua ADALINE prin proiectare directă, apelând funcția **newlind**.
 b) antrenând o rețea ADALINE cu arhitectură adecvată, utilizând funcțiile **newlin** și **train**;

Indicații și răspunsuri:

Pentru estimarea valorilor solicitate în enunțul problemei se ține cont de dependența ieșirii amplificatorului operațional de cele trei intrări ale sale:

$$y(u_1, u_2, u_3) = -\frac{R}{R_1}u_1 - \frac{R}{R_2}u_2 - \frac{R}{R_3}u_3.$$

Astfel, rețeaua ADALINE ce modelează funcționarea acestui amplificator operațional va avea 1 neuron cu 3 intrări, ale cărei ponderi reprezintă $w_1 = -R/R_1$, $w_2 = -R/R_2$, respectiv $w_3 = -R/R_3$. Deplasarea neuronului ar trebui să fie nulă.

```
% Problema 8.11
clear all; close all; clc
load opamp_wh

t= 0:0.01:5;
% vectori prototip
uuu1= 2*ones(size(uuu2));
P= [ uuu1; uuu2; uuu3];
% vectori tinta
T= yyy;
% reprezentare grafica a intrarilor si iesirii functie de timp
figure(1); plot(t, P); legend ('u1', 'u2', 'u3')
figure(2); plot(t, T);

% a. proiectare retea ADALINE utilizand functia newlind
net= newlind(P, T);
% parametri retea proiectata
wd= net.IW{1}; bd= net.b{1};
% valori estimate ale rezistentelor R1, R2 si R3
R= 100;
R1d= -R/wd(1)
R2d= -R/wd(2)
R3d= -R/wd(3)

% b. creare retea ADALINE si antrenare
lr= maxlinlr(P) % rata de invatare

lin= newlin(minmax(P),1,[0],lr);

% parametrii specifici antrenarii
lin.trainParam.epochs= input('Nr de epoci de antrenare = ');
lin.trainParam.goal= input('Eroarea medie patratica dorita= ');
lin.trainParam.show= 5;

lin = train(lin,P,T);
```

```
% parametri retea antrenata
wt= lin.lw{1};
bt= lin.b{1};
% valori estimate ale rezistentelor R1, R2 si R3
R1t= -R/wt(1)
R2t= -R/wt(2)
R3t= -R/wt(3)
```

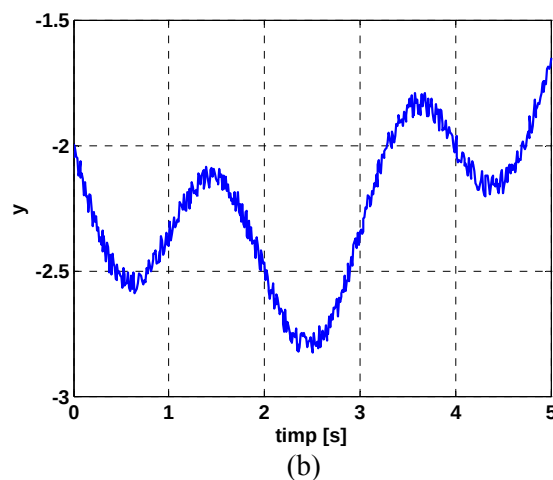
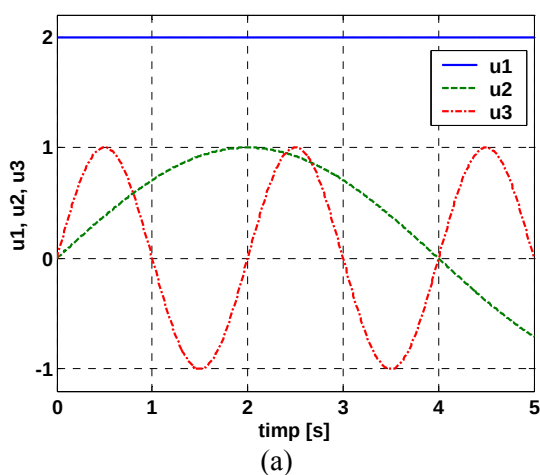


Figura 8.11.2. Reprezentarea grafică a semnalelor (a) de intrare și (b) de ieșire corespunzătoare amplificatorului operațional studiat în problema 8.11.

Reprezentarea grafică a semnalelor de intrare și de ieșire pentru amplificatorul operațional este prezentată în figura 8.11.2.

Prin proiectarea directă a rețelei ADALINE se obțin valorile $R_1^d = 99.85$, $R_2^d = 201.47$ și $R_3^d = 301.21$. Pentru antrenarea rețelei neurale, rata maximă de învățare are valoarea 4.7356×10^{-4} . După 100 de epoci de antrenare se ajunge la valoarea $MSE = 0.000816$ și se obțin valorile $R_1^t = 124.73$, $R_2^t = 202.56$ și $R_3^t = 302.125$.

Mărirea numărului de epoci de antrenare nu aduce îmbunătățiri substanțiale în calitate aproximării valorii rezistenței R_1 . Aceasta se datorează faptului că pe una dintre liniile matricei prototipurilor sunt elemente constante.

Problema 8.12.

Se va studia și lansa în execuție programul MATLAB *arna_bp1*, prezentat mai jos, care antrenează (prin aplicarea regulii delta) un neuron cu funcție sigmoidală bipolară, pentru a aproxima o funcție neliniară φ definită prin vectorii prototip:

$$x_i = \frac{(i-1)\pi}{12}, \quad i = 1, \dots, 7,$$

și vectorii țintă:

$$z_i = \sin(x_i), \quad i = 1, \dots, 7.$$

Să se opereze modificările necesare în program asupra ratei de învățare și a numărului de iterații pentru a realiza antrenarea, pornind de la condițiile inițiale situate în următoarele domenii:

1. $(w_0, b_0) \in [-4.5, -1.5] \times [-4.5, -1.5]$
2. $(w_0, b_0) \in [-1.5, 1.5] \times [-4.5, -1.5]$
3. $(w_0, b_0) \in [1.5, 4.5] \times [-4.5, -1.5]$
4. $(w_0, b_0) \in [-4.5, -1.5] \times [-1.5, 1.5]$
5. $(w_0, b_0) \in [-1.5, 1.5] \times [-1.5, 1.5]$
6. $(w_0, b_0) \in [1.5, 4.5] \times [-1.5, 1.5]$
7. $(w_0, b_0) \in [-4.5, -1.5] \times [1.5, 4.5]$
8. $(w_0, b_0) \in [-1.5, 1.5] \times [1.5, 4.5]$
9. $(w_0, b_0) \in [1.5, 4.5] \times [1.5, 4.5]$

În fiecare din situațiile (1) - (9) se vor preciza:

- a) valorile inițiale ale parametrilor, reprezentarea grafică a aproximării realizate de rețea pe baza acestor valori și eroarea corespunzătoare;
- b) valorile finale ale parametrilor (rezultate din antrenare), reprezentarea grafică a aproximării realizate de rețea pe baza acestor valori și eroarea corespunzătoare;
- c) numărul total de iterații cât a durat antrenarea și reprezentarea grafică a evoluției erorii ca funcție de numărul de iterații;
- d) traiectoria în planul parametrilor (w, b) – reprezentare grafică;
- e) valorile utilizate pentru parametrii de antrenare comentate prin prisma convergenței algoritmului și a formei suprafeței erorii. Se vor raporta și explica eventualele eșecuri în antrenare, datorate alegerii neadecvate a parametrilor de antrenare.

Indicații și răspunsuri:

```
% Problema 8.12, program ARNA_BP1
clear all; close all; clc

% definirea prototipurilor
P = 0:pi/12:pi/2;
% definirea tintelor
T = sin(P);

% reprezentarea grafica a suprafatei erorii
w_range = -4.5:0.5:4.5;
b_range = -4.5:0.5:4.5;
figure
ES = errsurf(P, T, w_range, b_range, 'tansig')/length(T); %valorile erorii medii patratice
plotes(w_range, b_range, ES, [50 25]);

% crearea neuronului tansigmoidal, ce urmeaza a fi antrenat cu
% subrutina 'learnkd' (metoda gradientului)
sigm=newff(minmax(P), [1], {'tansig'}, 'traingd');
```

```

% parametrii specifici antrenarii
sigm.trainParam.epochs= 1; % antrenarea este efectuata epoca cu epoca
iteratii= input('Numar de epoci de antrenare: ');
sigm.trainParam.goal= 1.5e-4;
sigm.trainParam.show= inf % eroarea nu este afisata in timpul antrenarii
% eroarea va fi memorata la fiecare pas si afisarea va fi facuta la sfirsit
sigm.trainParam.lr= input('Valoarea ratei de invatare: ');

% selectarea valorilor initiale ale parametrilor neuronului cu mouse-ul
subplot(1,2,2);
h = text(mean(get(gca,'xlim')), mean(get(gca,'ylim')), '* Alege cu mouse-ul *');
set(h,'horizontal','center', 'fontweight','bold');
[sigm.IW{1,1}, sigm.b{1}] = ginput(1);
% valorile alese cu mouse-ul devin parametrii initiali ai rețelei

fprintf('\n Valori initiale parametri retea neurala:');
fprintf('\n\t pondere: %6.2f\n\t deplasare: %6.2f\n', sigm.IW{1,1}, sigm.b{1});
delete(h);

pause % Apasa o tasta pentru a antrena neuronul...

% Antrenarea rețelei si reprezentarea la fiecare pas (epoca)
% a parametrilor pe suprafata de eroare
Err(1)= mse(T-sim(sigm, P));
% eroarea medie patratica inainte de inceperea antrenarii;
% vectorul Err va contine evolutia erorii
Wi(1)= sigm.IW{1,1}; % vectorul care va contine evolutia ponderilor
bi(1)= sigm.b{1}; % vectorul care va contine evolutia deplasarilor
h= plotep(sigm.IW{1,1}, sigm.b{1}, Err(1));
% parametrii initiali reprezentati pe suprafata de eroare

for i=1:iteratii,
    [sigm,tr]= train(sigm,P,T);
    h= plotep(sigm.IW{1,1}, sigm.b{1}, tr.perf(1), h);
    %evolutia parametrilor pe suprafata de eroare
    title(['Iteratia: ' num2str(i)]);
    Err(i+1)= tr.perf(1); %eroarea pt noile valori ale parametrilor
    Wi(i+1)= sigm.IW{1,1};
    bi(i+1)= sigm.b{1};
    if (Err(i)<=sigm.trainParam.goal),
        break; %iesire pe conditia de atingere a pragului de eroare impus
    end
end

iteratii= i; %numarul de epoci cit a durat antrenarea
fprintf('\n\nAntrenarea a durat %g iteratii\n', iteratii);
fprintf('\n\nValorile finale ale parametrilor rețelei:');
fprintf('\n\t pondere: %g\n\t deplasare: %g', sigm.IW{1,1}, sigm.b{1});
fprintf('\nEroarea medie patratica finala = %g\n', Err(end));
%evolutia erorii

```

```

disp('Apasati o tasta pentru a vedea evolutia erorii functie de nr de epoci...');
pause;
figure;
semilogy((0:iteratii), Err);
xlabel('#iteratii');
ylabel('MSE');
title('Eroarea medie patratica functie de numarul de iteratii');

%dinamica aproximarii
disp('Apasati o tasta pentru a vedea aproximarea initiala...');
pause;

figure;
plot(P,T,'r+'); %punctele prin care se doreste aproximarea
xlabel('vectori prototip P');
ylabel('vectori tinta T');
hold on;
h=plot(P,tansig(Wi(1)*P+bi(1)));
title('Aproximarea initiala');

disp('Apasati o tasta pentru a vedea evolutia dinamica a aproximarii...');
pause;

for i=1:iteratii,
    pause(0.2);
    delete(h);
    h=plot(P, tansig(Wi(i+1)*P+bi(i+1)),'b*-');
    title(['Iteratia:',num2str(i)]);
end

hold off;

```

Figura 8.12.1 prezintă evoluția parametrilor neuronului tansigmoidal pe suprafața de eroare pe parcursul a 100 de epoci de antrenare cu rata de învățare $\eta = 0.9$. Figurile 8.12.2.(a) și 8.12.3.(a) prezintă evoluția erorii medii pătratice MSE pe parcursul a 100 de epoci de antrenare cu rata de învățare $\eta = 0.9$ și, respectiv, $\eta = 0.2$, pornind de la aceleași valori inițiale ale parametrilor neuronului. Aceste grafice permit vizualizarea efectului pe care îl are rata de învățare asupra progresului antrenării.

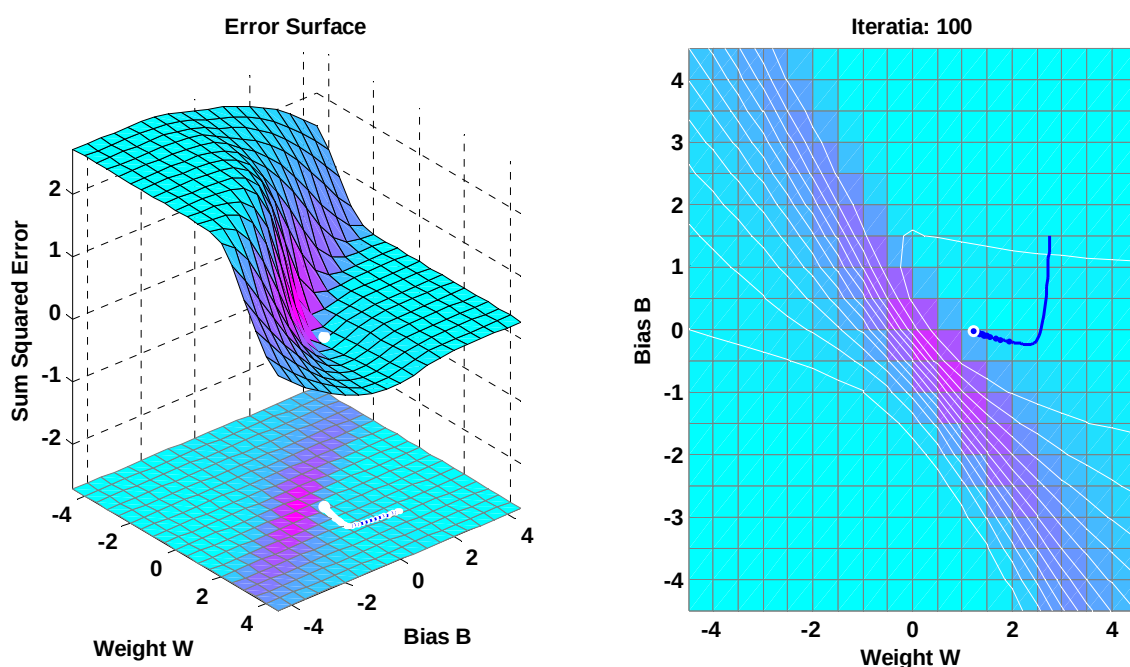


Figura 8.12.1. Evoluția parametrilor neuronului tansigmoidal pe suprafața de eroare pe parcursul a 100 de epoci de antrenare cu rata de învățare $\eta = 0.9$.

Figurile 8.12.2.(b) și 8.12.3.(b) prezintă aproximarea realizată de neuronul rezultat după 100 de epoci de antrenare cu rata de învățare $\eta = 0.9$ și, respectiv, $\eta = 0.2$. Aceste grafice permit vizualizarea semnificației MSE în aprecierea calității aproximării realizate.

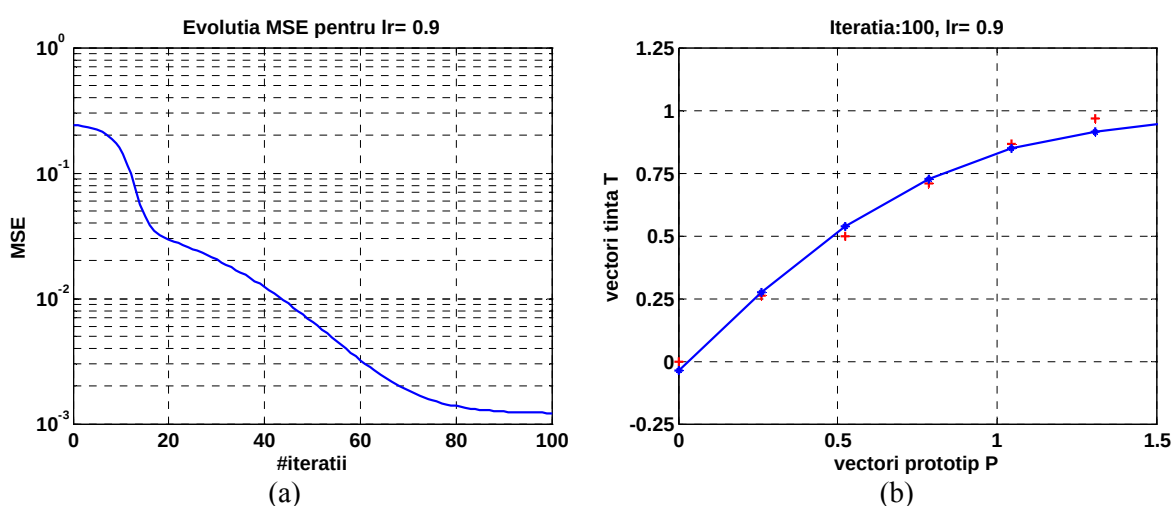


Figura 8.12.2. (a) Evoluția erorii medii pătratice MSE pe parcursul a 100 de epoci de antrenare cu rata de învățare $\eta = 0.9$. (b) Aproximarea realizată de neuronul antrenat.

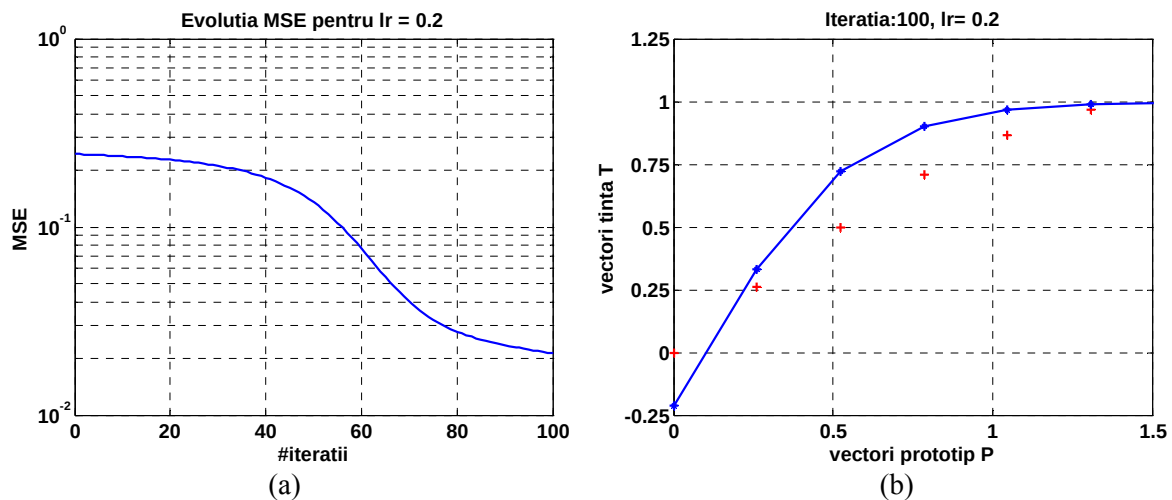


Figura 8.12.3. (a) Evoluția erorii medii pătratice MSE pe parcursul a 100 de epoci de antrenare cu rata de învățare $\eta = 0.2$. (b) Aproximarea realizată de neuronul antrenat.

Problema 8.13.

Să se utilizeze programul *arna_bp1* prezentat în problema 8.12 (cu modificările necesare considerate) pentru a selecta o regiune din planul parametrilor (w, b) și o gamă de valori a ratei de învățare care să asigure antrenarea eficientă pentru perechile de vectori (prototip, țintă):

$$x_i = \frac{\pi}{2} + \frac{(i-1)\pi}{12}; \quad z_i = \sin(x_i); \quad i = 1, \dots, 7.$$

Să se comenteze criteriile care au condus la selectarea respectivă, prin referire la convergența algoritmului și forma suprafeței erorii.

Indicații și răspunsuri:

În programul *arna_bp1* se modifică în mod corespunzător matricea vectorilor prototip și cea a vectorilor țintă.

```
% definirea prototipurilor P:
P = pi/2 : pi/12 : pi;
% definirea tintelor T:
T = sin(P);
```

Problema 8.14.

Se va studia și lansa în execuție programul MATLAB *arna_bp2* care antrenează o rețea cu două straturi (cu noduri tansigmoidale pe primul strat și liniare pe al doilea) pentru a aproxima o funcție neliniară $\varphi: \mathbb{R} \rightarrow \mathbb{R}$ cunoscută prin intermediul perechilor de vectori prototip-țintă:

$$(x_i, z_i), \quad z_i = \varphi(x_i), \quad i = 1, \dots, N.$$

Se consideră vectorii țintă:

$$x_i = (i-1)/50, \quad i = 1, \dots, 51,$$

și vectorii prototip:

$$z_i = \varphi(x_i) = 1 + \exp(-x_i) \sin(2\pi x_i - \frac{\pi}{2}), \quad i = 1, \dots, 51.$$

Antrenarea rețelei se va face folosind mai întâi metoda gradientului (rutina '**traingd**'), iar apoi algoritmul Levenberg-Marquardt (rutina '**trainlm**'). În fiecare din cele două cazuri se vor studia următoarele aspecte:

a) Operând modificările necesare în programul **arna_bp2** să se studieze rolul următorilor factori asupra eficienței antrenării rețelei:

- numărul de neuroni în stratul de intrare;
- rata de învățare (doar în cazul folosirii rutinei '**traingd**').

Aprecierea eficienței se va face în termenii erorii pătratice globale obținute în urma antrenării și a numărului de iterații necesare pentru atingerea unui prag impus pentru eroare.

Se va comenta modul de alegere a valorilor inițiale pentru parametrii rețelei.

b) Pentru testele de antrenare eficiente realizate, se vor preciza valorile parametrilor rețelei rezultați în urma antrenării. De asemenea se va testa capacitatea de generalizare a rețelei antrenate considerând câteva perechi prototip-țintă (x, z) corespunzătoare funcției φ „nevăzute” de rețea pe parcursul antrenării (corespunzătoare unor valori x situate în afara intervalului $[0, 1]$ folosit în antrenare).

Indicații și răspunsuri:

```
% Problema 8.14, program ARNA_BP2
clear all; close all; clc;

% definirea prototipurilor P
P = 0:0.02:1;

% definirea tintelor T
T = 1 + exp(-P).*sin(2*pi*P - pi/2);

% reprezentarea grafica a perechilor (prototip, tinta)
figure; plot(P, T, 'r+');
xlabel('Vectori prototip P');
ylabel('Vectori tinta T');
title('Functia neliniara de aproximat')

pause % Tasteaza pentru a dimensiona reteaua
neur= input('\n\nNumarul de neuroni de intrare: ');

disp('Selectare metoda de antrenare: ');
met = 0;
```

```

while ~((met == 1) || (met == 2)),
    met= input(' 1 - gradient, 2 - Levenberg-Marquardt...');
end
if met == 1, % pentru antrenare prin metoda gradientului
    sigm= newff(minmax(P),[neur 1],{'tansig' 'purelin'},'traingd');
    sigm.trainParam.lr= input('\nValoarea ratei de invatare: ');
else % pentru antrenare prin metoda Levenberg-Marquardt
    sigm= newff(minmax(P),[neur 1],{'tansig' 'purelin'},'trainlm');
end

disp('Valorile initiale ale parametrilor:')
disp('  Stratul de intrare:')
sigm.IW{1,1}
sigm.b{1}
disp('  Stratul de iesire:')
sigm.LW{2,1}
sigm.b{2}

iteratii= input('\nNumar de iteratii: ');
sigm.trainParam.epochs= 1;
sigm.trainParam.goal= 2e-4;
sigm.trainParam.show= inf;

figure(1);
hold on;
h= plot(P,sim(sigm,P));
title('Aproximarea initiala');
Err(1)= mse(T-sim(sigm,P));
    %eroarea medie patratica inainte de inceperea antrenarii;
    % vectorul "Err" va contine evolutia erorii

% antrenarea rețelei
echo off
for i= 1:iteratii,
    [sigm, tr]= train(sigm,P,T);
    Err(i+1)= tr.perf(1);
    pause(0.1);
    delete(h);
    title(strcat('Iteratia:', num2str(i)));
    h= plot(P, sim(sigm,P));
    if Err(i)<= sigm.trainParam.goal
        break;
    end
end

%evolutia erorii
disp('Apasati o tasta pentru a vedea evolutia erorii functie de numarul de epoci...');
pause;

iteratii= i; %numarul de epoci cat a durat antrenarea
fprintf('Eroarea medie patratica obtinuta este %g\n\n', Err(i));

```

```

fprintf("\n\nAntrenarea a durat %g iteratii\n", iteratii);

disp('Valorile finale ale parametrilor:')
disp('  Stratul de intrare:')
disp('  Ponderi:')
sigm.IW{1,1}
disp('  Deplasari:')
sigm.b{1}

disp('  Stratul de iesire:')
disp('  Ponderi:')
sigm.LW{2,1}
disp('  Deplasari:')
sigm.b{2}

figure;
semilogy((0:iteratii), Err);
xlabel('iteratii');
ylabel('eroarea');
title('Eroarea medie patratica functie de numarul de iteratii');

disp('Apasati o tasta pentru a vizualiza capacitatea de generalizare a rețelei
antrenate...');
pause;

% verificare capacitate de generalizare a rețelei antrenate
Pg1 = -0.1:0.02:0;
Tg1 = 1+exp(-Pg1).*sin(2*pi*Pg1-pi/2);
Pg2 = 1:0.02:1.1;
Tg2 = 1+exp(-Pg2).*sin(2*pi*Pg2-pi/2);

figure(1); hold on
plot(Pg1, Tg1, 'r*'); plot(Pg2, Tg2, 'r*');
Yg1= sim(sigm,Pg1); Yg2= sim(sigm,Pg2);
plot(Pg1, Yg1, 'g-'); plot(Pg2, Yg2, 'g-');
hold off;

```

Programul prezentat mai sus poate fi folosit pentru a efectua analiza cerută.

De exemplu, pentru o rețea cu 7 neuroni pe stratul de intrare, impunând $2e-4$ ca valoare de prag pentru eroarea pătratică medie, antrenarea prin metoda Levenberg-Marquardt durează 17 epoci, ajungându-se la eroarea pătratică medie 0.000117845.

Figura 8.14.1.(a) ilustrează evoluția erorii medii pătratice MSE pe parcursul a 17 de epoci de antrenare utilizând algoritmul Levenberg-Marquardt pentru rețeaua cu 7 neuroni pe stratul de intrare. Aproximarea realizată de rețeaua antrenată poate fi vizualizată în figura 8.14.1.(b).

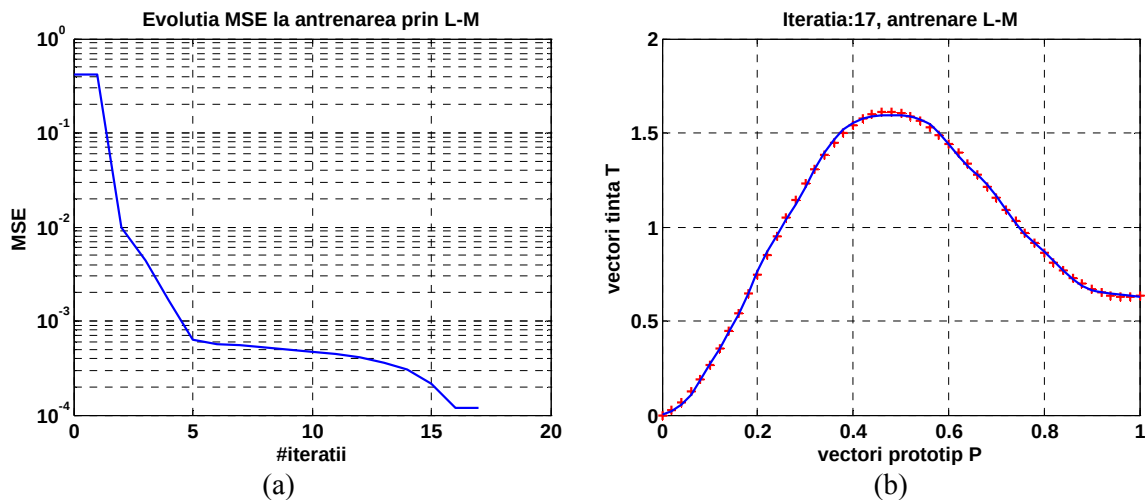


Figura 8.14.1. (a) *Evoluția erorii medii pătratice MSE pe parcursul a 17 de epoci de antrenare utilizând algoritmul Levenberg-Marquardt pentru o rețea cu 7 neuroni de intrare.*
 (b) *Aproximarea realizată de rețeaua antrenată.*

Problema 8.15.

Să se reia Problema 3 pentru aceiași vectori țintă și vectorii prototip:

$$z_i = \varphi(x_i) = 1 + \exp(-x_i) \sin(8\pi x_i - \frac{\pi}{2}), \quad i = 1, \dots, 51.$$

Indicații și răspunsuri:

În programul *arna_bp2* se modifică în mod corespunzător matricea vectorilor țintă.

```
% definirea tintelor T
T = 1 + exp(-P).*sin(8*pi*P - pi/2);
```

Problema 8.16.

Se consideră imaginea RGB din figura 8.16.1 reprezentând un preparat microscopic dintr-un aliaj de cupru care conține trei elemente diferite, și anume: cupru (cromaticitate maron închis), crom (cromaticitate maron deschis) și plumb (cromaticitate albastră).

Să se determine procente din aria totală a imaginii pe care le ocupă fiecare din cele trei elemente.

Indicații și răspunsuri:

Datorită complexității problemei, soluția acesteia este prezentată în continuare foarte sumar. Pentru elaborarea programelor MATLAB, în afară de funcțiile din pachetul Neural Networks Toolbox sunt necesare funcții specifice prelucrării de imagini, din Image Processing Toolbox.

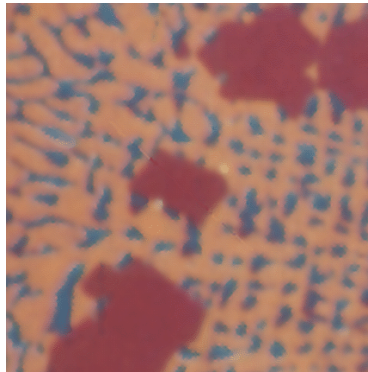


Figura 8.16.1. *Imaginea color a preparatului microscopic din aliaj de cupru.*

Pentru clasificarea pixelilor imaginii considerate se ține cont de faptul că fiecare pixel p al imaginii color este caracterizat prin:

- (i) un vector cu două componente, notat $g(p)$, care exprimă coordonatele geometrice ale pixelului;
- (ii) un vector cu două componente, $c(p)$, care exprimă coordonatele de culoare din cubul RGB.

Folosind coordonatele geometrice, să definim mulțimea pixelilor G_k care corespund regiunilor ocupate de elementul e_k , $k = 1, 2, 3$, unde e_1 semnifică cuprul, e_2 semnifică cromul și e_3 semnifică plumbul. Evident, mulțimile G_k sunt disjuncte.

În cubul RGB (scalat $[0,1] \times [0,1] \times [0,1]$) definim mulțimea de culori asociată fiecărui element e_k , $k = 1, 2, 3$.

$$C_k = \left\{ \bigcup_p c(p) \mid g(p) \in G_k \right\}. \quad (8.16.1)$$

Ipoteză de lucru: Cel puțin din punct de vedere teoretic, se poate considera că mulțimile C_1 , C_2 , C_3 sunt distincte. Satisfacerea practică a acestei ipoteze depinde de calitatea imaginii color.

Algoritm

Pasul 1. Se selectează eșantioane de culoare reprezentative pentru cele trei tipuri regiuni și se definesc mulțimile de eșantioane $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$. (Aceasta constituie cea mai bună posibilitate de a obține informații parțiale despre mulțimile de culori C_1, C_2, C_3 care sunt necunoscute ca valori concrete ocupate în cubul RGB. Se elimină vectorii identici din fiecare mulțime de culoare. În baza ipotezei de lucru considerate, $\tilde{C}_k \subseteq C_k$, $k = 1, 2, 3$ și, drept urmare mulțimile de eșantioane $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$ sunt disjuncte.)

Figura 8.16.2 ilustrează o modalitate de alegere a mulțimilor de eșantioane $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$. Figura 8.16.3 prezintă plasarea culorilor din $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$ în cubul RGB. Pentru reprezentarea grafică se utilizează convenția: $\tilde{C}_1$ - galben, $\tilde{C}_2$ - roșu, $\tilde{C}_3$ - bleu.

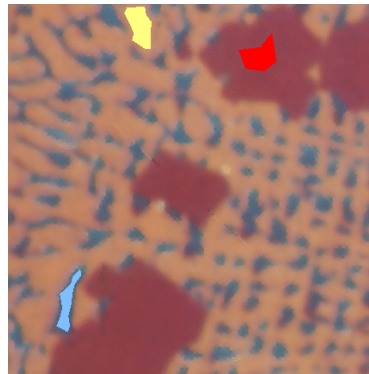


Figura 8.16.2. Alegerea eșantioanelor de culoare din regiuni reprezentative ale imaginii.

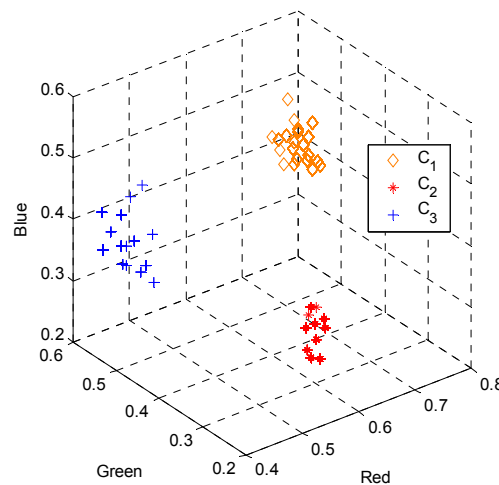


Figura 8.16.3. Plasarea culorilor din mulțimile de eșantioane $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$ în cubul RGB scalat $[0,1] \times [0,1] \times [0,1]$.

Pasul 2. Se antrenează o rețea neurală cu arhitectură adecvată astfel încât să clasifice în 3 clase culorile din mulțimile de eșantioane $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$.

Pasul 3. Se utilizează rețeaua antrenată pentru a clasifica toți pixelii din imagine în cele 4 clase, corespunzătoare lui $\tilde{C}_1, \tilde{C}_2, \tilde{C}_3$.

Pasul 4. Se construiește o nouă imagine digitală în care se definesc următoarele 3 regiuni

$$\hat{G}_k = \left\{ \bigcup_p g(p) \mid c(p) \in \text{class } k \right\}, \quad k = 1, 2, 3, \quad (2)$$

care estimează regiunile ideale de interes G_k . Se utilizează 3 culori distincte pentru a reprezenta regiunile estimate $\tilde{G}_1, \tilde{G}_2, \tilde{G}_3$ (strategie care asigură o bună vizualizare). În limbajul specific procesării imaginilor, operația prin care se obțin regiunile $\tilde{G}_1, \tilde{G}_2, \tilde{G}_3$ ca estimări ale regiunilor ideale G_1, G_2, G_3 poartă denumirea de “segmentare”.

Figura 8.16.4 prezintă clasificarea realizată de o rețea cu 8 neuroni în primul strat, antrenată în Pasul 3 prin algoritmul Levenberg-Marquardt până la o eroare pătratică medie $4.87e-010$ (obținută după 16 iterații).

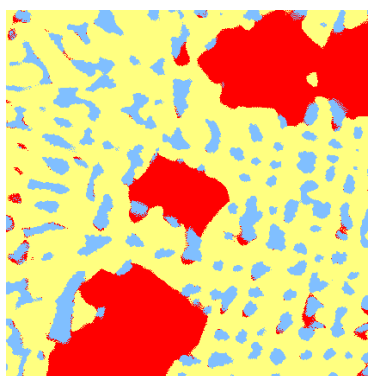


Figura 8.16.4. Separarea în clase (regiuni de interes realizată de o rețea cu 8 neuroni în primul strat și antrenată prin algoritmul Levenberg-Marquardt).

Pasul 5. Se măsoară ariile regiunilor estimate $\tilde{G}_1, \tilde{G}_2, \tilde{G}_3$ și se calculează procente ocupate de fiecare.

Pentru segmentarea prezentată în figura 8.16.4 se obțin următoarele procente: 59.27 % pentru $\tilde{G}_1$ (galben), 23.55 % pentru $\tilde{G}_2$ (roșu) și 17.18 % pentru $\tilde{G}_3$ (bleu). Ajungem astfel la concluzia că aliajul considerat conține 59.27 % crom, 23.55 % cupru și 17.18 % plumb.

Problema 8.17.

Pe baza exemplului prezentat în secțiunea 6.4.1.A să se construiască modelul Simulink pentru identificarea sistemului liniar descris de funcția de transfer

$$G(s) = \frac{125e^{-3s}}{100s^2 + 40s + 3},$$

considerând perioada de eșantionare a semnalelor de 1 secundă.

Să se ruleze experimentele de identificare off-line și on-line în următoarele condiții:

- ieșirea procesului nu este afectată de perturbații;
- ieșirea procesului este afectată aditiv de perturbații de tip Gaussian care pot atinge 1% din amplitudinea maximă a semnalului util (corespunzând, de exemplu, zgomotului de conversie al unui convertor A/D cu rezoluție scăzută).

Corespunzător ecuației (6.1.1), se vor lua în discuție următoarele cazuri:

- (i) $d = 4, n = 2, m = 1$;
- (ii) $d = 4, n = 2, m = 2$;
- (iii) $d = 4, n = 3, m = 2$;
- (iv) $d = 4, n = 3, m = 3$;

Pentru fiecare din cazurile (i)-(iv) se va fixa valoarea de prag a erorii pătratice medii utilizată în antrenarea rețelei (*mean squared error goal*) la valorile 10^{-2} și apoi 0; în fiecare caz în parte se va comenta legătura dintre parametrii rețelei neurale și coeficienții funcției de transfer discrete corespunzătoare sistemului, obținută prin apelarea funcției Matlab `c2d`.

După rularea unui experiment de identificare, se va rula, pentru validarea modelului, un experimentul de simulare, considerând ca semnal de intrare în sistem și în modelul neural obținut atât un semnal de tip zgomot alb, cât și un semnal determinist, de tip treaptă unitară sau sinusoidal.

Indicații și răspunsuri:

Se utilizează modelele Simulink prezentate în figurile 8.17.1 și 8.17.2.

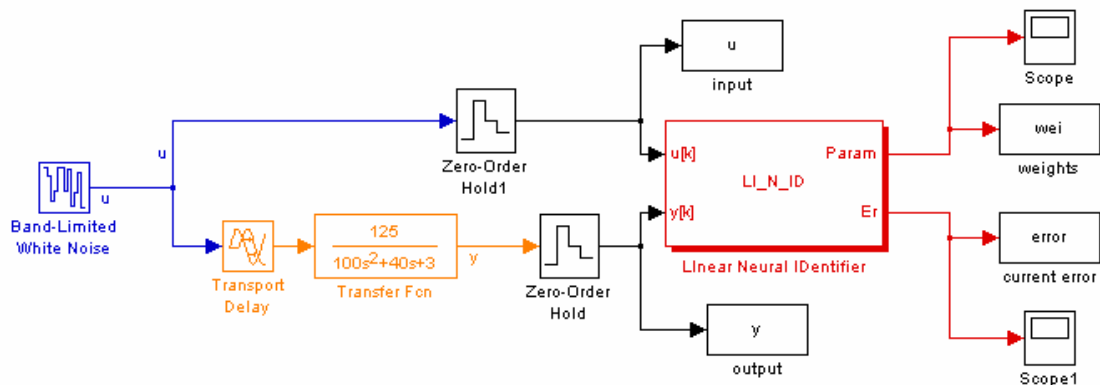


Figura 8.17.1. Exemplu de construire a modelului Simulink pentru identificarea sistemului dinamic din problema 8.17.

Pentru discretizarea funcției de transfer date se utilizează secvența de comenzi MATLAB de mai jos.

```
% funcția de transfer în timp continuu
s= tf('s');
Gc= 125 / (100*s*s + 40*s + 3);
Gc.ioDelay= 3

% discretizare cu perioada de esantionare T=1
Gd= c2d(Gc, 1, 'zoh')
```

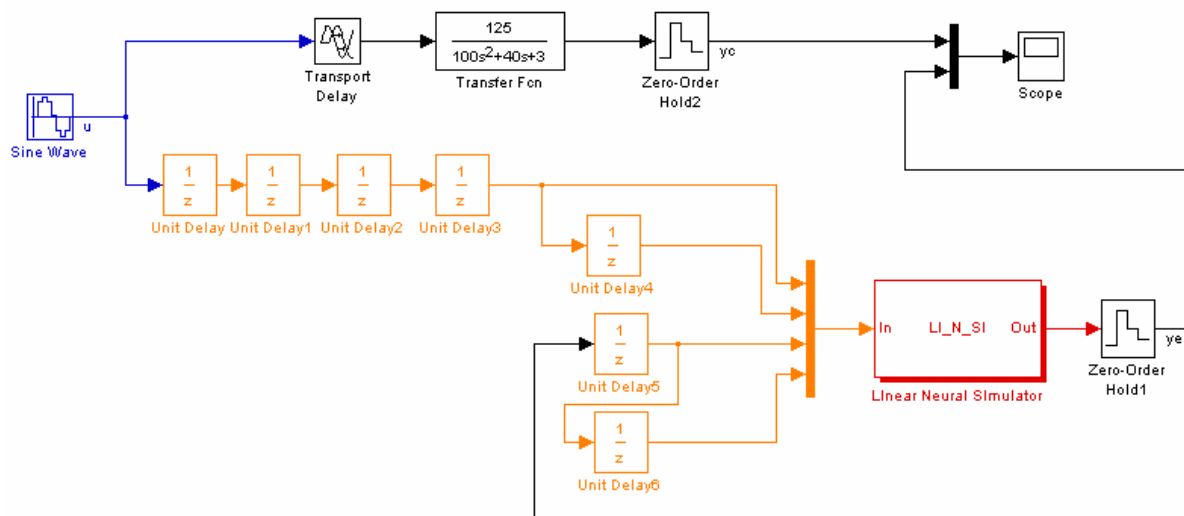



Figura 8.17.2. Exemplu de construire a modelului Simulink pentru simularea sistem dinamic utilizând modelul neural identificat.

Se obține astfel

$$G_d(z) = z^{-3} \frac{0.54z + 0.47}{z^2 - 1.64z + 0.67} = z^{-4} \frac{0.54 + 0.47z^{-1}}{1 - 1.64z^{-1} + 0.67z^{-2}},$$

ceea ce arată că în timp discret sistemul linear considerat este descris de ecuația cu diferențe:

$$(1 - 1.64z^{-1} + 0.67z^{-2})y[k] = z^{-4}(0.54 + 0.47z^{-1})u[k]$$

adică:

$$y[k] = 1.64y[k-1] - 0.67y[k-2] + 0.54u[k-4] + 0.47u[k-5].$$

Parametrii modelului neural obținut prin identificare pot fi comparați cu parametrii modelului exact obținut prin discretizare.

Problema 8.18.

Se consideră sistemul nelinier în timp discret descris de ecuația cu diferențe (pentru o perioadă de eșantionare $T=1$ secundă):

$$y(k) = \frac{y[k-1] + y[k-2]}{1 + 2y[k-1]y[k-2]} - 1.5u(k-1) + u[k-2].$$

1. Folosind blocurile Simulink prezentate în secțiunea 6.3, să se obțină un model neuronal de tip MLP pentru procesul dat, efectuând experimente de identificare off-line și on-line în următoarele condiții:

- ieșirea procesului nu este afectată de perturbații;
- ieșirea procesului este afectată aditiv de perturbații de tip Gaussian care pot atinge 1% din amplitudinea maximă a semnalului util (corespunzând, de exemplu, zgomotului de conversie al unui convertor A/D cu rezoluție scăzută).

Corespunzător modelului NNARX, se va lua în discuție atât cazul când întârzierile rețelei neuronale sunt egale cu întârzierile reale ale procesului, cât și cazul introducerii unor regresori suplimentari în modelul neuronal.

2. Calitatea modelelor neuronale obținute va fi apreciată din punct de vedere al analizei comparative a răspunsurilor obținute de la proces și respectiv model, pentru următoarele cazuri:

- (a) semnal de intrare aleator cu distribuție uniformă în intervalul $[-1, 1]$;
- (b) semnal de intrare aleator binar de nivele 1 și -1.

3. Modelele obținute vor fi folosite pentru a realiza controlul predictiv al procesului considerat în următoarele cazuri:

- (a) referință sinusoidală, cu amplitudine 1, fază inițială 0 și pulsație $2\pi/250$ rad/sec.;
- (b) referință în formă de dinți de ferăstrău de amplitudine 1, offset 0.5 și perioadă 100 secunde.

Pentru algoritmul de control predictiv, se vor considera următorii parametrii: orizont de control $N_u=2$, orizont minim de predicție $N_l=1$, orizont de predicție $N_2=3$ și ponderea asupra comenzii $\lambda=0,2$.

Indicații și răspunsuri:

Se utilizează blocurile descrise în secțiunea 6.3 pentru a construi modele Simulink similare celor prezentate în secțiunea 6.4.

Anexe

Anexa A

METODE DE OPTIMIZARE UTILIZATE CA ALGORITMI DE ANTRENARE SUPERVIZATĂ A REȚELELOR NEURALE

A.1. Introducere

Se consideră o funcție $f: \mathbb{R}^n \rightarrow \mathbb{R}$ a cărei valoare minimă pe o mulțime $\Omega \subseteq \mathbb{R}^n$ dorim să o determinăm. Se obține astfel problema de optimizare:

$$\begin{aligned} &\text{minimizarea } f(\mathbf{x}) \\ &\text{cu restricția } \mathbf{x} \in \Omega. \end{aligned} \tag{A.1.1}$$

Funcția f este denumită *funcție de cost* sau *funcție obiectiv*, iar Ω reprezintă mulțimea valorilor fezabile ale vectorului de decizie $\mathbf{x} = [x_1, x_2, \dots, x_n]^T \in \mathbb{R}^n$.

În cazul în care $\Omega \subset \mathbb{R}^n$, (A.1.1) reprezintă forma generală a unei probleme de optimizare cu restricții. Dacă $\Omega = \mathbb{R}^n$ se obține o problemă de optimizare fără restricții.

În multe situații concrete, mulțimea restricțiilor este precizată sub forma:

$$\Omega = \left\{ \mathbf{x} \in \mathbb{R}^n \mid \mathbf{h}(\mathbf{x}) = 0, \mathbf{g}(\mathbf{x}) \leq 0 \right\},$$

unde funcțiile $\mathbf{h}: \mathbb{R}^n \rightarrow \mathbb{R}^m$ și $\mathbf{g}: \mathbb{R}^n \rightarrow \mathbb{R}^p$ sunt cunoscute.

Definiția A.1.1. Fie $f: \mathbb{R}^n \rightarrow \mathbb{R}$ și $\Omega \subset \mathbb{R}^n$.

Un punct $\mathbf{x}^* \in \Omega$ este *punct de minim local* pentru f peste Ω dacă $\exists \varepsilon > 0$ astfel încât $f(\mathbf{x}) \geq f(\mathbf{x}^*)$, pentru orice $\mathbf{x} \in \Omega \setminus \{\mathbf{x}^*\}$ satisfăcând $\|\mathbf{x} - \mathbf{x}^*\| < \varepsilon$.

Un punct $\mathbf{x}^* \in \Omega$ este *punct de minim global* pentru f peste Ω dacă $f(\mathbf{x}) \geq f(\mathbf{x}^*)$ pentru orice $\mathbf{x} \in \Omega \setminus \{\mathbf{x}^*\}$. ■

Punctul de minim $\mathbf{x}^*$ reprezintă argumentul care minimizează funcția f (pe Ω sau pe întreg spațiu $\mathbb{R}^n$) și se notează $\mathbf{x}^* = \arg \min_{\mathbf{x} \in \Omega} f(\mathbf{x})$ sau $\mathbf{x}^* = \arg \min f(\mathbf{x})$.

Dacă în definiția A.1 se înlocuiește “ $\geq$ ” cu “ $>$ ” atunci punctul $\mathbf{x}^*$ reprezintă un *punct de minim (local sau global) strict*.

Figura A.1.1 prezintă graficul unei funcții $f: \mathbb{R} \rightarrow \mathbb{R}$. Punctele x_1^* și x_2^* sunt puncte de minim global, respectiv local, strict; iar x_3^* este punct minim local nestrict.

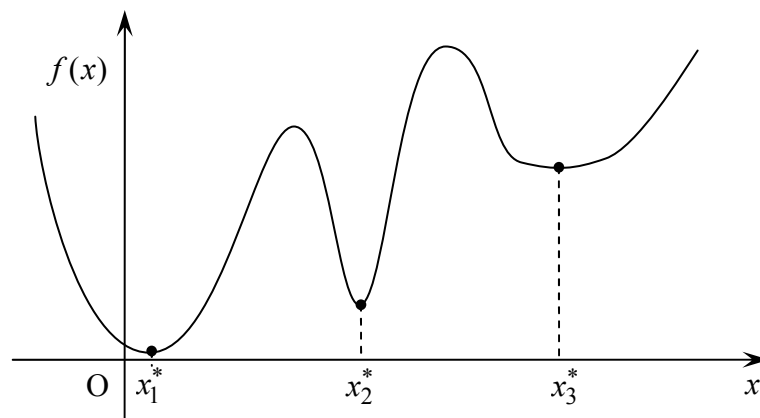


Figura A.1.1. Puncte de minim ale unei funcții reale.

A.1.1. Condiții de minim local

Pentru prezentarea unor condiții necesare (și suficiente) de minim, este nevoie de câteva cunoștințe preliminare.

Definiția A.1.2. Se consideră funcția $f: \mathbb{R}^n \rightarrow \mathbb{R}$, de două ori derivabilă pe $\mathbb{R}^n$.

Gradientul funcției scalare f într-un punct $\mathbf{x}$ este vectorul coloană format cu derivatele parțiale de ordinul 1 ale lui f ,

$$\nabla f(\mathbf{x}) = \left[\frac{\partial f}{\partial x_1}(\mathbf{x}) \quad \frac{\partial f}{\partial x_2}(\mathbf{x}) \quad \cdots \quad \frac{\partial f}{\partial x_n}(\mathbf{x}) \right]^T \quad (\text{A.1.2})$$

iar matricea Hessian a lui f este formată cu derivatele parțiale de ordinul 2 ale lui f .

$$\mathbf{H}_f(\mathbf{x}) = \begin{bmatrix} \frac{\partial^2 f}{\partial x_1^2}(\mathbf{x}) & \frac{\partial^2 f}{\partial x_1 \partial x_2}(\mathbf{x}) & \cdots & \frac{\partial^2 f}{\partial x_1 \partial x_n}(\mathbf{x}) \\ \frac{\partial^2 f}{\partial x_2 \partial x_1}(\mathbf{x}) & \frac{\partial^2 f}{\partial x_2^2}(\mathbf{x}) & \cdots & \frac{\partial^2 f}{\partial x_2 \partial x_n}(\mathbf{x}) \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1}(\mathbf{x}) & \frac{\partial^2 f}{\partial x_n \partial x_2}(\mathbf{x}) & \cdots & \frac{\partial^2 f}{\partial x_n^2}(\mathbf{x}) \end{bmatrix} \left(\triangleq \nabla^2 f(\mathbf{x}) \right) \quad \blacksquare \quad (\text{A.1.3})$$

Dacă funcția f este de clasă $C_{\mathbb{R}^n}^2$ (adică f este de două ori derivabilă și derivata a doua a sa este continuă pe $\mathbb{R}^n$) atunci, conform *teoremei Clairaut-Schwarz* are loc egalitatea derivatelor parțiale de ordinul 2

$$\frac{\partial^2 f}{\partial x_i \partial x_j}(\mathbf{x}) = \frac{\partial^2 f}{\partial x_j \partial x_i}(\mathbf{x}), \quad \forall \mathbf{x} \in \mathbb{R}^n, \quad i, j = \overline{1, n},$$

astfel că matricea Hessian este simetrică.

Exemplul A.1.1.

- a) O funcție liniară pe $\mathbb{R}^n$ este definită printr-o relație de forma $f(\mathbf{x}) = \mathbf{d}^T \mathbf{x}$, unde $\mathbf{d} = [d_1 \ d_2 \ \dots \ d_n]^T \in \mathbb{R}^n$. Gradientul acestei funcții este $\nabla f(\mathbf{x}) = \mathbf{d}$.
- b) Considerăm funcția definită prin $f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c$, unde $\mathbf{A} \in \mathbb{R}^{n \times n}$, $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$. Dacă matricea $\mathbf{A}$ este simetrică, funcția f se numește *formă pătratică* și $\nabla f(\mathbf{x}) = \mathbf{A} \mathbf{x} + \mathbf{b}$ iar $\mathbf{H}_f(\mathbf{x}) = \mathbf{A}$. Dacă matricea $\mathbf{A}$ nu este simetrică, atunci $\nabla f(\mathbf{x}) = \frac{1}{2}(\mathbf{A} + \mathbf{A}^T) \mathbf{x} + \mathbf{b}$ și $\mathbf{H}_f(\mathbf{x}) = \frac{1}{2}(\mathbf{A} + \mathbf{A}^T)$. ■

Definiția A.1.3. Un vector $\mathbf{d} \in \mathbb{R}^n$, $\mathbf{d} \neq 0$ reprezintă o *direcție fezabilă* în punctul $\mathbf{x} \in \Omega$ dacă există $\alpha_0 > 0$ astfel încât $\mathbf{x} + \alpha \mathbf{d} \in \Omega$, pentru orice $\alpha \in [0, \alpha_0]$. ■

În figura A.1.2 direcția $\mathbf{d}_1$ este fezabilă în punctul $\mathbf{x}$, în timp ce $\mathbf{d}_2$ nu este fezabilă. Se observă că dacă $\mathbf{x}$ este punct interior mulțimii Ω atunci orice direcție este fezabilă în $\mathbf{x}$.

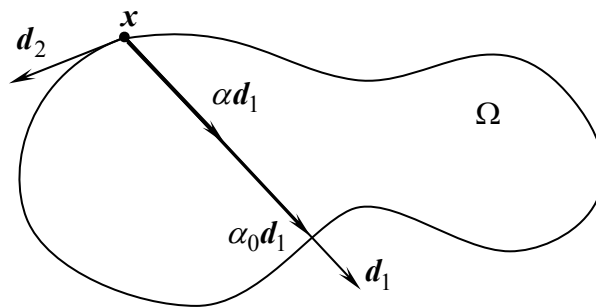


Figura A.1.2. Ilustrarea grafică a conceptului de direcție fezabilă într-un punct $\mathbf{x}$.

Definiția A.1.4. Fie $f: \mathbb{R}^n \rightarrow \mathbb{R}$ și $\mathbf{d} \in \mathbb{R}^n$ o direcție fezabilă în punctul $\mathbf{x} \in \Omega$. Derivata direcțională a lui f în direcția $\mathbf{d}$, notată $\frac{\partial f}{\partial \mathbf{d}}$, este definită prin:

$$\frac{\partial f}{\partial \mathbf{d}}(\mathbf{x}) = \lim_{\alpha \rightarrow 0} \frac{f(\mathbf{x} + \alpha \mathbf{d}) - f(\mathbf{x})}{\alpha} \quad \blacksquare \text{ (A.1.4)}$$

Dacă vectorul $\mathbf{d}$ este unitar, adică $\|\mathbf{d}\| = 1$, atunci $\partial f / \partial \mathbf{d}$ reprezintă rata de creștere a lui f în direcția $\mathbf{d}$ și poate fi calculată considerând $f(\mathbf{x} + \alpha \mathbf{d})$ ca funcție de α :

$$\frac{\partial f}{\partial \mathbf{d}}(\mathbf{x}) = \frac{d}{d\alpha} f(\mathbf{x} + \alpha \mathbf{d})|_{\alpha=0} = [\nabla f(\mathbf{x})]^T \cdot \mathbf{d} = \mathbf{d}^T \cdot \nabla f(\mathbf{x}). \quad \text{(A.1.4')}$$

Următoarele propoziții precizează condiții necesare (și suficiente) de minim.

Teorema A.1.1. Fie $\Omega \subset \mathbb{R}^n$ și $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C_{\mathbb{R}^n}^1$. Dacă $\mathbf{x}^*$ este punct de minim local al lui f peste Ω , atunci pentru orice direcție $\mathbf{d}$ fezabilă în $\mathbf{x}^*$ este satisfăcută condiția

$$\mathbf{d}^T \cdot \nabla f(\mathbf{x}^*) \geq 0. \quad \blacksquare \text{ (A.1.5)}$$

Corolar A.1.1. Fie $\Omega \subset \mathbb{R}^n$ și $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C_{\mathbb{R}^n}^1$. Dacă $\mathbf{x}^*$ este punct de minim local al lui f și $\mathbf{x}^*$ este punct interior mulțimii Ω atunci:

$$\nabla f(\mathbf{x}^*) = 0. \quad \blacksquare \text{ (A.1.6)}$$

Teorema A.1.2. Fie $\Omega \subset \mathbb{R}^n$ și $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C_{\mathbb{R}^n}^2$, $\mathbf{x}^*$ un punct de minim local al lui f peste Ω și $\mathbf{d}$ o direcție fezabilă în $\mathbf{x}^*$. Dacă $\mathbf{d}^T \cdot \nabla f(\mathbf{x}^*) = 0$, atunci

$$\mathbf{d}^T \mathbf{H}_f(\mathbf{x}^*) \mathbf{d} \geq 0, \quad \text{(A.1.7)}$$

unde $\mathbf{H}_f(\mathbf{x}^*)$ reprezintă matricea Hessian a lui f calculată în $\mathbf{x}^*$. $\blacksquare$

Corolar A.1.2. Fie $\mathbf{x}^*$ punct interior mulțimii $\Omega \subset \mathbb{R}^n$ și $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C_{\mathbb{R}^n}^2$. Dacă $\mathbf{x}^*$ este punct de minim local al lui f atunci:

$$\nabla f(\mathbf{x}^*) = 0 \quad \text{(A.1.8)}$$

și matricea Hessian $\mathbf{H}_f(\mathbf{x}^*)$ este pozitiv semidefinită ($\mathbf{H}_f(\mathbf{x}^*) \geq 0$), adică

$$\mathbf{d}^T \mathbf{H}_f(\mathbf{x}^*) \mathbf{d} \geq 0, \quad \forall \mathbf{d} \in \mathbb{R}^n \setminus \{0\}. \quad \blacksquare \text{ (A.1.9)}$$

Condițiile necesare prezentate în Corolarul A.1.2 nu sunt însă și suficiente, așa cum se poate observa în cazul funcției $f_1: \mathbb{R} \rightarrow \mathbb{R}$, $f_1(x) = x^3$, pentru punctul $x^* = 0$ (figura A.1.3.a) care satisface $f_1'(0) = f_1''(0) = 0$. O situație similară este cea a funcției $f_2: \mathbb{R}^2 \rightarrow \mathbb{R}$, $f_2(\mathbf{x}) = x_1^2 - x_2^2$, în punctul $\mathbf{x}^* = [0 \ 0]^T$ (figura A.1.3.b) care reprezintă un *punct în șa*.

Următoarea teoremă prezintă condiții suficiente de minim local.

Teorema A.1.3. Fie $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $f \in C_{\mathbb{R}^n}^2$, și $\mathbf{x}^*$ un punct interior mulțimii $\Omega \subset \mathbb{R}^n$.

Dacă

$$\nabla f(\mathbf{x}^*) = 0, \quad (\text{A.1.10})$$

și matricea Hessian în punctul $\mathbf{x}^*$ este pozitiv definită ($\mathbf{H}_f(\mathbf{x}^*) > 0$), adică

$$\mathbf{d}^T \mathbf{H}_f(\mathbf{x}^*) \mathbf{d} > 0, \quad \forall \mathbf{d} \in \mathbb{R}^n \setminus \{0\}, \quad (\text{A.1.11})$$

atunci $\mathbf{x}^*$ este punct de minim local strict al funcției f . ■

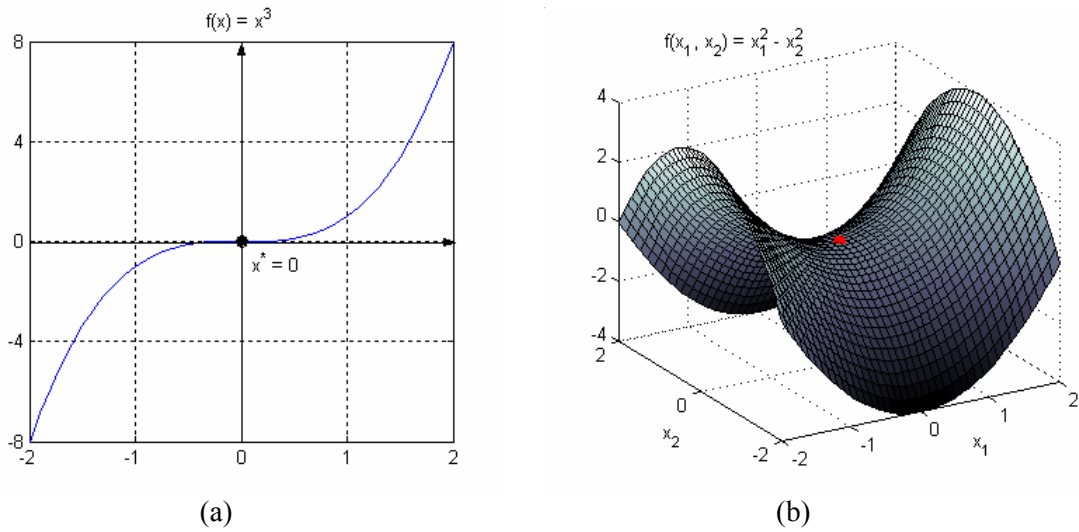


Figura A.1.3. Ilustrarea grafică a conceptului de punct în așa: (a) pentru funcția $f(x) = x^3$ în punctul $\mathbf{x}^* = 0$, (b) pentru funcția $f(x_1, x_2) = x_1^2 - x_2^2$ în punctul $\mathbf{x}^* = [0 \ 0]^T$.

Exemplul A.1.2. Considerăm funcția $f: \mathbb{R}^2 \rightarrow \mathbb{R}$, $f(\mathbf{x}) = x_1^2 + x_2^2$, a cărei reprezentare grafică este prezentată în figura A.1.4.

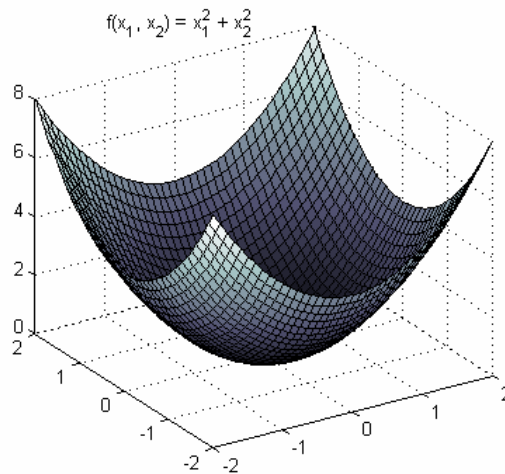


Figura A.1.4. Reprezentarea grafică a funcției $f(x_1, x_2) = x_1^2 + x_2^2$.

Gradientul lui f într-un punct oarecare este $\nabla f(\mathbf{x}) = [2x_1 \ 2x_2]^T$. Condiția necesară de minim (A.1.10) conduce la $\nabla f(\mathbf{x}^*) = 0$ dacă și numai dacă $\mathbf{x}^* = [0 \ 0]^T$. Matricea Hessian este constantă pe $\mathbb{R}^2$, $\mathbf{H}_f = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix}$, simetrică și cu valori proprii strict pozitive, $\lambda_1 = \lambda_2 = 2 > 0$, deci este pozitiv definită. Aplicând Teorema A.3 rezultă că $\mathbf{x}^* = [0 \ 0]^T$ este punct de minim local al funcției f pe orice mulțime Ω ce conține $\mathbf{x}^*$ în interior. În plus $\mathbf{x}^*$ este punct de minim global al lui f pe $\mathbb{R}^2$. ■

Exemplul A.1.3. O generalizare a cazului din exemplul precedent este reprezentată de forma pătratică

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, \quad f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c, \quad (\text{A.1.12})$$

unde $\mathbf{A} \in \mathbb{R}^{n \times n}$ este matrice simetrică, $\mathbf{A}^T = \mathbf{A}$, $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$. Pentru această funcție, $\nabla f(\mathbf{x}) = \mathbf{A} \mathbf{x} + \mathbf{b}$ iar $\mathbf{H}_f(\mathbf{x}) = \mathbf{A}$.

Dacă matricea $\mathbf{A}$ este pozitiv definită, atunci funcția f admite un unic punct de minim global pe $\mathbb{R}^n$, care se determină prin rezolvarea ecuației $\nabla f(\mathbf{x}) = 0$, ceea ce conduce la $\mathbf{x}^* = -\mathbf{A}^{-1} \mathbf{b}$. Dacă matricea $\mathbf{A}$ este negativ definită (adică $-\mathbf{A}$ este pozitiv definită), atunci funcția f admite un unic punct de maxim global pe $\mathbb{R}^n$, care satisface $\nabla f(\mathbf{x}) = 0$, adică $\mathbf{x}^* = -\mathbf{A}^{-1} \mathbf{b}$. În cazul în care nici $\mathbf{A}$ nici $-\mathbf{A}$ nu sunt pozitiv definite nu se poate spune nimic despre punctele de extrem ale funcției f . ■

A.1.2. Rezolvarea ecuației $\mathbf{A} \mathbf{x} = \mathbf{b}$

În numeroase cazuri problema determinării parametrilor necunoscuți ai unui model revine la rezolvarea unei ecuații liniare în vectorul $\mathbf{x} \in \mathbb{R}^n$, de forma

$$\mathbf{A} \mathbf{x} = \mathbf{b} \quad (\text{A.1.13})$$

unde $\mathbf{A} \in \mathbb{R}^{m \times n}$ și $\mathbf{b} \in \mathbb{R}^m$. În funcție de dimensiunile și rangul matricei sistemului sunt posibile următoarele situații:

a) Dacă $m = n$ și matricea $\mathbf{A}$ este nesară, adică $\det \mathbf{A} \neq 0$, atunci ecuația (A.1.13) admite soluția unică $\mathbf{x}^* = \mathbf{A}^{-1} \mathbf{b}$.

b) Considerăm cazul $m \geq n$ și $\text{rang } \mathbf{A} = n$. Dacă $\text{rang } \mathbf{A} < \text{rang} [\mathbf{A} \ \mathbf{b}]$ (se spune că matricea $\mathbf{A}$ este de rang complet pe coloane), atunci ecuația (A.1.13) nu admite soluții exacte.

În această situație, se caută o soluție $\mathbf{x}^*$ care să minimizeze norma euclidiană a erorii $\mathbf{e} = \mathbf{A} \mathbf{x} - \mathbf{b}$. Vectorul $\mathbf{x}^*$ care satisface condiția $\|\mathbf{A} \mathbf{x} - \mathbf{b}\|_2^2 \geq \|\mathbf{A} \mathbf{x}^* - \mathbf{b}\|_2^2$, pentru orice $\mathbf{x} \in \mathbb{R}^n$, notat $\mathbf{x}^* = \arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{A} \mathbf{x} - \mathbf{b}\|_2^2$, poartă denumirea de *soluție în sensul celor mai mici pătrate* a ecuației (A.1.13).

Pentru a determina această soluție considerăm funcția

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, f(\mathbf{x}) = \|\mathbf{Ax} - \mathbf{b}\|_2^2 = (\mathbf{Ax} - \mathbf{b})^T \cdot (\mathbf{Ax} - \mathbf{b}) = \mathbf{x}^T \mathbf{A}^T \mathbf{A} \mathbf{x} - 2\mathbf{b}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{b}, \quad (\text{A.1.14})$$

care este pătratică, de forma (A.1.12), și pentru care calculăm

$$\nabla f(\mathbf{x}) = 2\mathbf{A}^T \mathbf{A} \mathbf{x} - 2\mathbf{A}^T \mathbf{b}$$

$$\mathbf{H}_f(\mathbf{x}) = 2\mathbf{A}^T \mathbf{A}$$

Din condiția $\text{rang } \mathbf{A} = n$ rezultă că $\text{rang}(\mathbf{A}^T \mathbf{A}) = n$, astfel că ecuația $\nabla f(\mathbf{x}) = 0$ are soluție unică

$$\mathbf{x}^* = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{b}. \quad (\text{A.1.15})$$

Se observă că matricea simetrică $\mathbf{A}^T \mathbf{A}$ este pozitiv definită (deoarece $\mathbf{d}^T \mathbf{A}^T \mathbf{A} \mathbf{d} = \|\mathbf{Ad}\|_2^2 > 0$ pentru orice $\mathbf{d} \neq 0$), deci, conform teoremei A.3, punctul $\mathbf{x}^*$ (A.1.15) este unicul punct de minim al funcției f .

În particular, dacă $m = n$ soluția (A.1.15) devine $\mathbf{x}^* = \mathbf{A}^{-1} \mathbf{b}$.

Valoarea minimă a funcției (A.1.14) este

$$f(\mathbf{x}^*) = (\mathbf{Ax}^* - \mathbf{b})^T \cdot (\mathbf{Ax}^* - \mathbf{b}) = (\mathbf{Ax}^*)^T \mathbf{e}^* - \mathbf{b}^T \mathbf{e}^* = -\mathbf{b}^T \mathbf{e}^*,$$

unde $\mathbf{e}^* = \mathbf{Ax}^* - \mathbf{b}$, deoarece $(\mathbf{Ax}^*)^T \mathbf{e}^* = \frac{1}{2} \nabla f(\mathbf{x}^*) = 0$.

c) Considerăm cazul $m \leq n$ și $\text{rang } \mathbf{A} = m$ (se spune că matricea $\mathbf{A}$ este de rang complet pe linii), ecuația $\mathbf{Ax} = \mathbf{b}$ admite o infinitate de soluții. În aceasta situație ne interesează determinarea acestei soluții care este mai apropiată de origine, adică minimizează $\|\mathbf{x}\|_2$, sau, altfel spus, soluția problemei de optimizare

$$\text{minimizare } \|\mathbf{x}\|_2 \quad (\text{A.1.16})$$

cu restricția $\mathbf{Ax} = \mathbf{b}$

Soluția unică a acestei probleme este dată de

$$\mathbf{x}^* = \mathbf{A}^T (\mathbf{A} \mathbf{A}^T)^{-1} \mathbf{b}. \quad (\text{A.1.17})$$

Pentru a demonstra acest fapt, considerăm o soluție $\mathbf{x} \neq \mathbf{x}^*$ a ecuației $\mathbf{Ax} = \mathbf{b}$. Se observă că

$$\begin{aligned} \|\mathbf{x}\|_2^2 &= \|(\mathbf{x} - \mathbf{x}^*) + \mathbf{x}^*\|_2^2 = [(\mathbf{x} - \mathbf{x}^*) + \mathbf{x}^*]^T \cdot [(\mathbf{x} - \mathbf{x}^*) + \mathbf{x}^*] = \\ &= \|\mathbf{x} - \mathbf{x}^*\|_2^2 + \|\mathbf{x}^*\|_2^2 + 2\mathbf{x}^{*T} \cdot (\mathbf{x} - \mathbf{x}^*). \end{aligned}$$

Arătăm în continuare că $\mathbf{x}^{*T} \cdot (\mathbf{x} - \mathbf{x}^*) = 0$ deoarece:

$$\begin{aligned} \mathbf{x}^{*T} \cdot (\mathbf{x} - \mathbf{x}^*) &= [\mathbf{A}^T (\mathbf{A} \mathbf{A}^T)^{-1} \mathbf{b}]^T [\mathbf{x} - \mathbf{A}^T (\mathbf{A} \mathbf{A}^T)^{-1} \mathbf{b}] = \mathbf{b}^T (\mathbf{A} \mathbf{A}^T)^{-1} \mathbf{A} [\mathbf{x} - \mathbf{A}^T (\mathbf{A} \mathbf{A}^T)^{-1} \mathbf{b}] = \\ &= \mathbf{b}^T (\mathbf{A} \mathbf{A}^T)^{-1} [\mathbf{Ax} - (\mathbf{A} \mathbf{A}^T) (\mathbf{A} \mathbf{A}^T)^{-1} \mathbf{b}] = \mathbf{b}^T (\mathbf{A} \mathbf{A}^T)^{-1} [\mathbf{b} - \mathbf{b}] = 0 \end{aligned}$$

În consecință, pentru orice soluție $\mathbf{x} \neq \mathbf{x}^*$ a ecuației $\mathbf{Ax} = \mathbf{b}$, cum $\|\mathbf{x} - \mathbf{x}^*\|_2 \neq 0$, rezultă că $\|\mathbf{x}\|_2^2 = \|\mathbf{x} - \mathbf{x}^*\|_2^2 + \|\mathbf{x}^*\|_2^2 > \|\mathbf{x}^*\|_2^2$, ceea ce trebuia demonstrat.

În particular, dacă $m = n$ soluția (A.1.17) devine $\mathbf{x}^* = \mathbf{A}^{-1}\mathbf{b}$.

Observație. Fie $\mathbf{A} \in \mathbb{R}^{m \times n}$ o matrice dată. Matricea $\mathbf{A}^\dagger \in \mathbb{R}^{n \times m}$ se numește *pseudo-inversa Morre-Penrose* (inversa generalizată) a matricei $\mathbf{A}$ dacă

$$\mathbf{A} \mathbf{A}^\dagger \mathbf{A} = \mathbf{A}$$

și există matricele $\mathbf{U} \in \mathbb{R}^{n \times n}$, $\mathbf{V} \in \mathbb{R}^{m \times m}$, astfel încât

$$\mathbf{A}^\dagger = \mathbf{U} \mathbf{A}^T, \quad \mathbf{A}^\dagger = \mathbf{A}^T \mathbf{V}.$$

Relația $\mathbf{A}^\dagger = \mathbf{U} \mathbf{A}^T$ arată că fiecare linie din matricea $\mathbf{A}^\dagger$ este o combinație liniară a liniilor din $\mathbf{A}^T$ iar $\mathbf{A}^\dagger = \mathbf{A}^T \mathbf{V}$ arată că fiecare coloană din $\mathbf{A}^\dagger$ este o combinație liniară a coloanelor din $\mathbf{A}^T$.

Se poate demonstra că, dacă există o pseudo-inversă Morre-Penrose $\mathbf{A}^\dagger$ a matricei $\mathbf{A}$, atunci este unică. Pseudo-inversa unei matrice are următoarele proprietăți:

$$\begin{aligned} (\mathbf{A}^T)^\dagger &= (\mathbf{A}^\dagger)^T, \quad (\mathbf{A}^\dagger)^\dagger = \mathbf{A}, \\ \mathbf{A} \mathbf{A}^\dagger \mathbf{A} &= \mathbf{A}, \quad \mathbf{A}^\dagger \mathbf{A} \mathbf{A}^\dagger = \mathbf{A}^\dagger, \\ (\mathbf{A} \mathbf{A}^\dagger)^T &= \mathbf{A} \mathbf{A}^\dagger, \quad (\mathbf{A}^\dagger \mathbf{A})^T = \mathbf{A}^\dagger \mathbf{A}. \end{aligned}$$

În cazul în care $\mathbf{A} \in \mathbb{R}^{m \times n}$, cu $m \geq n$ și $\text{rang } \mathbf{A} = n$, atunci matricea $\mathbf{A}^\dagger = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T$ este pseudo-inversa Morre-Penrose a matricei $\mathbf{A}$.

Dacă $\mathbf{A} \in \mathbb{R}^{m \times n}$, cu $m \leq n$ și $\text{rang } \mathbf{A} = m$, matricea $\mathbf{A}^\dagger = \mathbf{A}^T (\mathbf{A} \mathbf{A}^T)^{-1}$ este pseudo-inversa Morre-Penrose a matricei $\mathbf{A}$.

Cu aceste observații, soluția (A.1.15) în sensul celor mai mici pătrate obținută în cazul $m \geq n$ cu $\text{rang } \mathbf{A} = n$, și soluția (A.1.17) care minimizează $\|\mathbf{x}\|_2$ obținută în cazul $m \leq n$ cu $\text{rang } \mathbf{A} = m$, pot fi puse sub forma

$$\mathbf{x}^* = \mathbf{A}^\dagger \mathbf{b}. \quad \blacksquare \quad (\text{A.1.18})$$

A.1.3. Estimatorul celor mai mici pătrate

În numeroase situații, ieșirea unui model, $\mathbf{y} \in \mathbb{R}^p$, depinde de intrarea $\mathbf{u} \in \mathbb{R}^m$ printr-o relație liniară în parametrii $\mathbf{x} = [x_1, x_2, \dots, x_n]^T \in \mathbb{R}^n$:

$$\mathbf{y} = \mathbf{f}_1(\mathbf{u})x_1 + \mathbf{f}_2(\mathbf{u})x_2 + \dots + \mathbf{f}_n(\mathbf{u})x_n, \quad (\text{A.1.19})$$

în care funcțiile $\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_n : \mathbb{R}^m \rightarrow \mathbb{R}^p$ sunt cunoscute. Ecuația (A.1.19) poate fi scrisă sub forma matriceală:

$$\mathbf{y} = \mathbf{M}(\mathbf{u})\mathbf{x} \quad (\text{A.1.20})$$

în care matricea $\mathbf{M}(\mathbf{u}) \in \mathbb{R}^{p \times n}$ este dată de $\mathbf{M}(\mathbf{u}) = [\mathbf{f}_1(\mathbf{u}) \ \mathbf{f}_2(\mathbf{u}) \ \dots \ \mathbf{f}_n(\mathbf{u})]$.

Pentru determinarea parametrilor modelului, de regulă se efectuează experimente care să asigure obținerea unui set de perechi de valori intrare-ieșire $(\mathbf{u}_i, \mathbf{y}_i)$, $i = \overline{1, N}$. Înlocuind

fiecare asemenea pereche în ecuația (A.1.20) se obține:

$$\begin{cases} \mathbf{y}_1 = \mathbf{M}_1 \mathbf{x} \\ \mathbf{y}_2 = \mathbf{M}_2 \mathbf{x} \\ \dots \\ \mathbf{y}_N = \mathbf{M}_N \mathbf{x} \end{cases}$$

unde $\mathbf{M}_i = \mathbf{M}(\mathbf{u}_i)$, $i = \overline{1, N}$. Cu notațiile

$$\mathbf{A} = \begin{bmatrix} \mathbf{M}_1 \\ \mathbf{M}_2 \\ \vdots \\ \mathbf{M}_N \end{bmatrix} \in \mathbb{R}^{Np \times n}, \quad \mathbf{b} = \begin{bmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \\ \vdots \\ \mathbf{y}_N \end{bmatrix} \in \mathbb{R}^{Np},$$

sistemul de ecuații precedent poate fi scris în forma echivalentă

$$\mathbf{Ax} = \mathbf{b}. \quad (\text{A.1.21})$$

Astfel, problema determinării parametrilor modelului (A.1.19) revine la rezolvarea unei ecuații matriceale liniare de forma (A.1.13). În general, numărul N de perechi de date pe baza cărora se dorește determinarea parametrilor $\mathbf{x}$ este mare, astfel încât matricea $\mathbf{A}$ satisface condiția $\text{rang } \mathbf{A} = n$ și se ajunge la cazul prezentat în secțiunea A.1.2.b.

Vectorul

$$\mathbf{x}^* = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{b},$$

de forma (A.1.15), care asigură minimizarea erorii pătratice

$$f(\mathbf{x}) = \frac{1}{2} \|\mathbf{Ax} - \mathbf{b}\|_2^2 = \frac{1}{2} \sum_{i=1}^N \|\mathbf{M}_i \mathbf{x} - \mathbf{y}_i\|_2^2 \quad (\text{A.1.22})$$

poarta denumirea de *estimator al celor mai mici pătrate* (eng. *Least-Squares Estimator – LSE*) corespunzător modelului (A.1.19).

A.1.4. Forma generală a metodelor iterative de optimizare

Considerăm o funcție $f: \mathbb{R}^n \rightarrow \mathbb{R}$ de clasă $C_{\mathbb{R}^n}^2$ al cărei punct de minim notat $\mathbf{x}^*$ presupunem că există și ne propunem să-l determinăm. În general, funcția f poate fi neliniară în argument, astfel că pentru determinarea minimului acesteia se utilizează o procedură iterativă care să permită explorarea domeniului de definiție al funcției într-un mod cât mai eficient.

Ideea de bază a metodelor iterative este următoarea: fiind dat un punct inițial $\mathbf{x}_0 \in \mathbb{R}^n$ se generează o secvență $\{\mathbf{x}_k\}_{k \in \mathbb{N}}$ pe baza unei reguli de forma $\mathbf{x}_{k+1} = \mathbf{x}_{k+1}(\mathbf{x}_k)$, $k = 1, 2, \dots$. Un algoritm iterativ este acceptabil atunci când este garantată convergența secvenței generate către soluția optimală a problemei. Dacă condiția de convergență este satisfăcută, procedura iterativă se oprește atunci când o anumită condiție de stop este îndeplinită.

Fie $\mathbf{x}_k$ punctul determinat la sfârșitul iterației $k-1$. La iterația k , pornind din $\mathbf{x}_k$ se determină direcția $\mathbf{d}_k$ de căutare a minimului și factorul $\eta_k > 0$ care precizează lungimea pasului ce urmează a fi făcut pornind din $\mathbf{x}_k$ în direcția $\mathbf{d}_k$. Noul punct $\mathbf{x}_{k+1}$ se determină utilizând relația

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k. \quad (\text{A.1.23})$$

Dacă algoritmul vizează găsirea punctului de minim al funcției f , acest punct ar trebui să satisfacă inegalitatea

$$f(\mathbf{x}_{k+1}) = f(\mathbf{x}_k + \eta_k \mathbf{d}_k) < f(\mathbf{x}_k). \quad (\text{A.1.24})$$

În literatura neuro-fuzzy, factorul $\eta > 0$ este denumit *rată de învățare* (eng. *learning rate*).

Criteriul de stop este testat la începutul fiecărei iterații. Unul dintre criteriile de stop utilizate în metodele iterative de căutare a minimului unei funcții este

$$\|\nabla f(\mathbf{x}_k)\| < \varepsilon \quad (\text{A.1.25})$$

unde $\|\cdot\|$ notează o normă vectorială convenabil aleasă (de regulă norma euclidiană) și $\varepsilon > 0$ reprezintă toleranța impusă la startarea algoritmului. Dacă condiția (A.1.25) este satisfăcută, vectorul gradient în punctul curent $\nabla f(\mathbf{x}_k)$ tinde la zero și secvența iterativă $\{\mathbf{x}_k\}_{k \in \mathbb{N}}$ tinde către un punct staționar al funcției f .

Alte criterii de stop frecvent utilizate sunt:

$$\|\mathbf{x}_k - \mathbf{x}_{k-1}\| < \varepsilon, \quad (\text{A.1.26})$$

$$|f(\mathbf{x}_k) - f(\mathbf{x}_{k-1})| < \varepsilon. \quad (\text{A.1.27})$$

Criteriile (A.1.26), (A.1.27) se utilizează atunci când este de așteptat ca algoritmul să convergă rapid, în timp ce criteriul (A.1.25) este utilizat atunci când algoritmul converge mai lent. Criteriile de stop de forma

$$\frac{\|\mathbf{x}_k - \mathbf{x}_{k-1}\|}{\|\mathbf{x}_{k-1}\|} < \varepsilon \quad (\text{A.1.28})$$

$$\frac{|f(\mathbf{x}_k) - f(\mathbf{x}_{k-1})|}{|f(\mathbf{x}_{k-1})|} < \varepsilon \quad (\text{A.1.29})$$

testează valori relative, independente de unitățile de măsură.

Structura de bază a unui algoritm de optimizare iterativ este următoarea:

Algoritmul O (algoritm de bază).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$.

Pas 1. (Testare criterii de stop)

1.1 Dacă $\|\nabla f(\mathbf{x}_k)\| < \varepsilon$, treci la Pasul 6.

1.2. Dacă $k = N$, treci la Pasul 6.

Pas 2. (Determinarea direcției de căutare) Utilizând o anumită schemă iterativă se determină direcția $\mathbf{d}_k$ de descreștere a funcției f .

- Pas 3. (Determinarea lungimii pasului) Se determină factorul η_k care precizează lungimea pasului, astfel încât să fie satisfăcută condiția de descreștere
- $$f(\mathbf{x}_k + \eta_k \mathbf{d}_k) < f(\mathbf{x}_k).$$
- Pas 4. (Actualizare) Se calculează următorul punct
- $$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k.$$
- Pas 5. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.
- Pas 6. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

Evident că la pasul 1.1 în locul criteriului (A.1.25) poate fi utilizat oricare dintre criteriile de stop (A.1.26)–(A.1.29) prezentate anterior.

Diferențele dintre diferite metode de minimizare constau în modul de determinare a direcțiilor succesive de căutare. Odată determinată această direcție, unii algoritmi consideră valori constante ale lui η , alți algoritmi „adaptează” lungimea pasului la fiecare iterație.

A.1.5. Metode de căutare liniară

Metodele de optimizare unidimensională, numite și *metode de căutare liniară*, stau la baza metodelor de optimizare a funcțiilor multivariabile. Acest fapt poate fi observat la pasul 3 al algoritmului general de optimizare unde factorul $\eta > 0$ este frecvent determinat astfel încât să fie minimizată funcția $\varphi(\eta) = f(\mathbf{x}_k + \eta \mathbf{d}_k)$ cu $\eta \geq 0$.

Pentru fixarea cadrului de discuție, considerăm o funcție $\varphi: \mathbb{R} \rightarrow \mathbb{R}$, care admite un punct de minim $\eta^* \in [0, +\infty)$, adică

$$\varphi(\eta^*) = \min_{\eta \geq 0} \varphi(\eta). \quad (\text{A.1.30})$$

Dacă există un interval închis $[a, b] \subset [0, +\infty)$ astfel încât $\eta^* \in [a, b]$ (fără ca poziția efectivă a punctului de minim să fie cunoscută), intervalul $[a, b]$ este numit interval de căutare (sau interval de incertitudine) pentru problema de minimizare unidimensională $\min_{\eta \geq 0} \varphi(\eta)$.

Una dintre cele mai frecvent utilizate metode de căutare liniară este *metoda secțiunii de aur* (eng. *golden search method*). Această metodă poate fi aplicată în ipoteza că funcția φ este unimodală pe intervalul de căutare $[a, b]$, adică este strict descrescătoare pe intervalul $[a, \eta^*]$ și strict crescătoare pe intervalul $[\eta^*, b]$. În această ipoteză, oricare ar fi două puncte $\lambda, \mu \in [a, b]$, cu $\lambda < \mu$, au loc următoarele implicații:

- dacă $\mu < \eta^*$, atunci $\varphi(\lambda) > \varphi(\mu)$;
- dacă $\lambda > \eta^*$, atunci $\varphi(\lambda) < \varphi(\mu)$;

Ipoteza de unimodalitate nu presupune continuitatea sau derivabilitatea funcției de minimizat. În plus, utilizând această ipoteză, intervalul de căutare poate fi restrâns eliminând acele porțiuni care, cu certitudine, nu pot conține punctul de minim. Aceasta se realizează observând că, dacă funcția φ este unimodală pe intervalul $[a, b]$ și considerăm $\lambda, \mu \in [a, b]$, cu $\lambda < \mu$, au loc următoarele implicații:

- dacă $\varphi(\lambda) \leq \varphi(\mu)$, atunci $[a, \mu]$ este un interval unimodal pentru φ ;
- dacă $\varphi(\lambda) \geq \varphi(\mu)$, atunci $[\lambda, b]$ este un interval unimodal pentru φ .

Pe baza acestor observații, principiul de bază al metodei secțiunii de aur este următorul. Considerăm că la iterația k intervalul de căutare este $[a_k, b_k]$, interval pe care funcția φ este unimodală. Considerăm două puncte $\lambda_k, \mu_k \in [a_k, b_k]$, cu $\lambda_k < \mu_k$, și aplicăm observația precedentă:

- dacă $\varphi(\lambda_k) \leq \varphi(\mu_k)$, atunci $[a_k, \mu_k]$ este un interval unimodal pentru φ și setăm $a_{k+1} = a_k$, $b_{k+1} = \mu_k$;
- dacă $\varphi(\lambda_k) \geq \varphi(\mu_k)$, atunci $[\lambda_k, b_k]$ este un interval unimodal pentru φ și setăm $a_{k+1} = \lambda_k$, $b_{k+1} = b_k$.

Pentru a reduce numărul de evaluări ale funcției φ la fiecare iterație, punctele λ_k și μ_k se aleg astfel încât să fie satisfăcute următoarele condiții:

- $\mu_k - a_k = b_k - \lambda_k$,
- rata de reducere a intervalului de incertitudine la fiecare iterație să fie aceeași, $b_{k+1} - a_{k+1} = \rho(b_k - a_k)$, $\rho \in (0, 1)$,
- la fiecare iterație să fie introdus numai un singur punct nou.

Ținând cont de primele două cerințe se ajunge la

$$\begin{aligned}\lambda_k &= a_k + (1 - \rho)(b_k - a_k), \\ \mu_k &= a_k + \rho(b_k - a_k).\end{aligned}\tag{A.1.31}$$

Aplicând aceste relații pentru intervalul $[a_{k+1}, b_{k+1}]$ și impunând condiția ca la o iterație să fie introdus numai un singur punct nou, se ajunge la concluzia că factorul $\rho \in (0, 1)$ trebuie să fie ales astfel încât $\rho^2 = 1 - \rho$. Singura soluție pozitivă a acestei ecuații fiind

$$\rho = \frac{\sqrt{5} - 1}{2} \approx 0.618,\tag{A.1.32}$$

relațiile (A.1.31) pot fi rescrise în forma

$$\begin{aligned}\lambda_k &= a_k + 0.382(b_k - a_k), \\ \mu_k &= a_k + 0.618(b_k - a_k).\end{aligned}\tag{A.1.33}$$

Din acest motiv metoda secțiunii de aur mai este numită și *metoda 0.618*.

Structura metodei secțiunii de aur pentru minimizarea unei funcții φ unimodale pe un interval $[a, b] \subset [0, +\infty)$ este următoarea:

Algoritmul GSM (Golden Section Method).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează intervalul inițial de căutare $[a_0, b_0]$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$. Se calculează

$$\begin{aligned}\lambda_0 &= a_0 + 0.382(b_0 - a_0), \\ \mu_0 &= a_0 + 0.618(b_0 - a_0).\end{aligned}$$

- Pas 1. (Testare criterii de stop)
1.1 Dacă $b_k - a_k < \varepsilon$, treci la Pasul 6.
1.2. Dacă $k = N$, treci la Pasul 6.
- Pas 2. (Compararea valorilor funcției) Dacă $\varphi(\lambda_k) \leq \varphi(\mu_k)$ treci la pasul 3, altfel, dacă $\varphi(\lambda_k) > \varphi(\mu_k)$ treci la pasul 4.
- Pas 3. (Cazul 1) Se setează
 $a_{k+1} = \lambda_k$, $b_{k+1} = b_k$, $\lambda_{k+1} = \mu_k$, $\varphi(\lambda_{k+1}) = \varphi(\mu_k)$.
Se calculează noul punct $\mu_{k+1} = a_{k+1} + 0.618(b_{k+1} - a_{k+1})$ și valoarea funcției în acest punct, $\varphi(\mu_{k+1})$. Se trece la pasul 5.
- Pas 4. (Cazul 2) Se setează
 $a_{k+1} = a_k$, $b_{k+1} = \mu_k$, $\mu_{k+1} = \lambda_k$, $\varphi(\mu_{k+1}) = \varphi(\lambda_k)$.
Se calculează noul punct $\lambda_{k+1} = a_{k+1} + 0.382(b_{k+1} - a_{k+1})$ și valoarea funcției în acest punct, $\varphi(\lambda_{k+1})$. Se trece la pasul 5.
- Pas 5. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.
- Pas 6. (Finalizare) Se returnează punctul final $\eta_f = \frac{1}{2}(\lambda_k + \mu_k)$.

În afara metodei secțiunii de aur, pentru căutare liniară mai există numeroase metode bazate pe interpolare pătratică sau cubică, sau metode de căutare inexactă. Aceste metode presupun însă că funcția criteriu satisface unele condiții suplimentare, de exemplu să fie de două ori derivabilă.

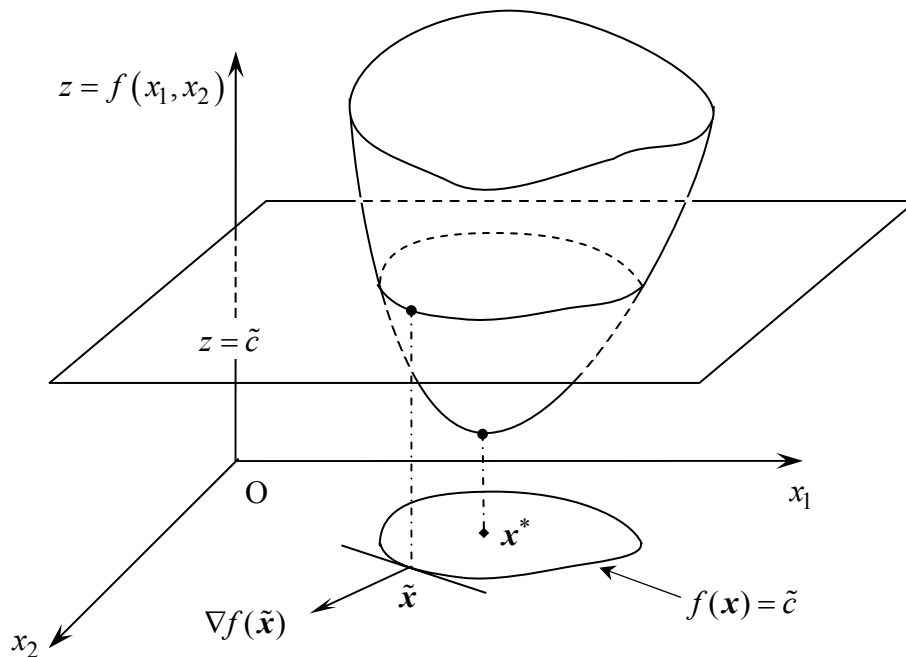


Figura A.2.1. Interpretarea geometrică a gradientului într-un punct pentru o funcție reală de două variabile.

A.2. Metode de optimizare bazate pe gradient

O categorie importantă de metode de minimizare se bazează pe calcularea gradientului funcției obiectiv în punctul curent pentru determinarea direcției de căutare, fiind denumite *metode de gradient*. Din această categorie fac parte: *metoda de descrescere maximă* (eng. *steepest – descent method*) și *metodele de tip Gauss – Newton*. Aceste metode sunt frecvent utilizate în probleme de aproximare a modelelor neliniare, fiind astfel în strânsă legătură cu tehnicile neuro-fuzzy.

Metodele bazate pe gradient utilizează proprietățile geometrice ale gradientului unei funcții într-un punct $\tilde{x}$ pentru care $f(\tilde{x}) = \tilde{c}$. Dacă gradientul funcției f în punctul $\tilde{x}$ este nenul, $\nabla f(\tilde{x}) \neq 0$, atunci $\nabla f(\tilde{x})$ este un vector ortogonal pe vectorul tangent la o curbă netedă arbitrară ce trece prin $\tilde{x}$ pe mulțimea de nivel constant $f(x) = \tilde{c}$ (figura A.2.1). Reamintim faptul că, fiind dată o constantă $c \in \mathbb{R}$, mulțimea de nivel c a funcției f este mulțimea punctelor $x \in \mathbb{R}^n$ ce satisfac $f(x) = c$.

Ținând cont de această interpretare, pentru ca o direcție de căutare a minimului $d \neq 0$ să fie fezabilă în punctul $\tilde{x}$ trebuie ca $d^T \cdot \nabla f(\tilde{x}) < 0$ (figura A.2.2). Se observă că direcția $d = -\nabla f(\tilde{x})$ este fezabilă în $\tilde{x}$.

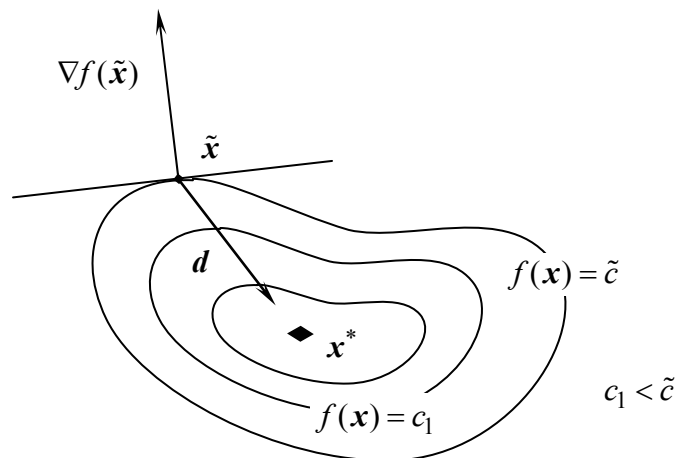


Figura A.2.1. Alegerea unei direcții fezabile de căutare a minimului.

A.2.1. Metoda steepest-descent (de descrescere maximă)

În această metodă, la fiecare iterație este ales pasul ηd astfel încât să fie asigurată descrescerea maximă a funcției obiectiv în direcția d considerată. Concret, la iterația k , pornind din punctul x_k este efectuată o căutare a minimului în direcția opusă gradientului, $d_k = -\nabla f(x_k)$. Mărimea pasului este determinată astfel încât să fie minimizată funcția: $\varphi_k(\eta) = f(x_k - \eta \nabla f(x_k))$, alegând

$$\eta_k = \arg \min_{\eta \geq 0} f(\mathbf{x}^{(k)} - \eta \nabla f(\mathbf{x}^k)). \quad (\text{A.2.1})$$

Se obține astfel punctul

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \eta_k \nabla f(\mathbf{x}_k), \quad (\text{A.2.2})$$

din care va fi efectuată o nouă căutare în același mod la următoarea iterație.

Se poate demonstra că dacă secvența $\mathbf{x}_0, \mathbf{x}_1, \dots$ a fost determinată prin algoritmul de descreștere maximă pentru minimizarea unei funcții f pornind din $\mathbf{x}_0$, atunci pentru orice $k \in \mathbb{N}$ vectorul $\mathbf{x}_{k+1} - \mathbf{x}_k$ este ortogonal pe $\mathbf{x}_{k+2} - \mathbf{x}_{k+1}$ (vezi figura A.2.2).

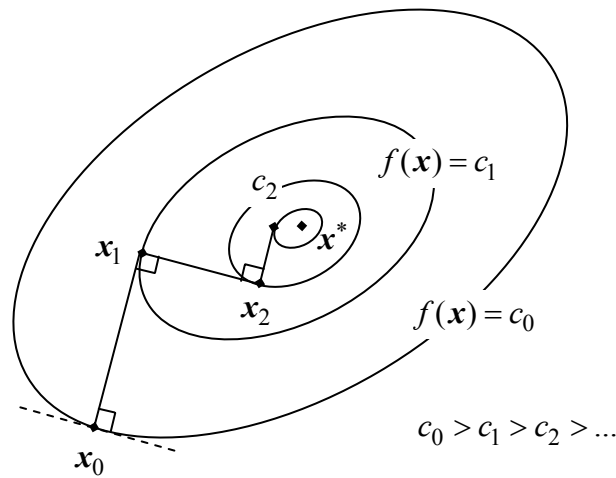


Figura A.2.2. Ilustrarea primelor iterații ale algoritmului steepest-descent.

Așa cum s-a văzut mai sus, dacă $\nabla f(\mathbf{x}_k) \neq 0$, utilizând acest algoritm, din $\mathbf{x}_k$ se determină noul punct $\mathbf{x}_{k+1}$ care satisface condiția de descreștere $f(\mathbf{x}_{k+1}) < f(\mathbf{x}_k)$.

Dacă se ajunge la un punct $\mathbf{x}_k$ pentru care $\nabla f(\mathbf{x}_k) = 0$, atunci $\mathbf{x}_{k+1} = \mathbf{x}_k$. Cum punctul $\mathbf{x}_k$ satisface condiția necesară de minim, nu mai are sens să aplicăm, algoritmul în continuare. În practică, în locul condiției $\nabla f(\mathbf{x}_k) = 0$, pentru oprirea algoritmului iterativ, se utilizează criterii de stop de forma (A.1.25)–(A.1.29).

Algoritmului steepest-descent de minimizare a unei funcții obiectiv $f: \mathbb{R}^n \rightarrow \mathbb{R}$, de clasa $C_{\mathbb{R}^n}^1$ este următoarea:

Algoritmul SD (metoda steepest descent).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$.

Pas 1. (Testare criterii de stop) Fie $\mathbf{g}_k = \nabla f(\mathbf{x}_k)$.

1.1 Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 5.

1.2. Dacă $k = N$, treci la Pasul 5.

Pas 2. (Determinarea lungimii pasului) Se determină factorul η_k astfel încât

$$f(\mathbf{x}_k - \eta_k \mathbf{g}_k) = \min_{\eta > 0} f(\mathbf{x}_k - \eta \mathbf{g}_k).$$

Pas 3. (Actualizare) Se calculează următorul punct

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \eta_k \mathbf{g}_k.$$

Pas 4. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.

Pas 5. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

Exemplul A.2.1. Pentru a ilustra aplicarea metodei steepest-descent considerăm forma pătratică

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c, \quad (\text{A.2.3})$$

pentru care matricea $\mathbf{A} \in \mathbb{R}^{n \times n}$ este simetrică și pozitiv definită ($\mathbf{A}^T = \mathbf{A} > 0$), $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$. Așa cum a fost prezentat în exemplul A.1.3, funcția f admite unic punct de minim, $\mathbf{x}^* = -\mathbf{A}^{-1} \mathbf{b}$, care reprezintă soluția ecuației $\nabla f(\mathbf{x}) = 0$.

Aplicând algoritmul steepest-descent, la iterația k în punctul curent $\mathbf{x}_k$, se calculează gradientul $\mathbf{g}_k = \nabla f(\mathbf{x}_k) = \mathbf{A} \mathbf{x}_k + \mathbf{b}$ și presupunem că satisface condiția $\mathbf{g}_k \neq 0$ (adică $\mathbf{x}_k$ nu este punct de minim pentru f).

Factorul $\eta_k > 0$ la iterația curentă se determină astfel încât $\eta_k = \arg \min_{\eta \geq 0} \varphi_k(\eta)$, unde

$$\begin{aligned} \varphi_k(\eta) &= f(\mathbf{x}_k - \eta \mathbf{g}_k) = \\ &= \frac{1}{2} \mathbf{g}_k^T \mathbf{A} \mathbf{g}_k \eta^2 - [\mathbf{g}_k^T \mathbf{A} \mathbf{x}_k + \mathbf{g}_k^T \mathbf{b}] \eta + \left[\frac{1}{2} \mathbf{x}_k^T \mathbf{A} \mathbf{x}_k + \mathbf{b}^T \mathbf{x}_k + c \right] = \\ &= \frac{1}{2} \mathbf{g}_k^T \mathbf{A} \mathbf{g}_k \eta^2 - \mathbf{g}_k^T \mathbf{g}_k \eta + f(\mathbf{x}_k). \end{aligned}$$

Funcția $\varphi_k(\eta)$ are gradul 2 în η și coeficientul termenului dominant pozitiv ($\frac{1}{2} \mathbf{g}_k^T \mathbf{A} \mathbf{g}_k > 0$ deoarece matricea $\mathbf{A}$ este pozitiv definită), astfel că admite minim în

$$\eta_k = \frac{\mathbf{g}_k^T \mathbf{g}_k}{\mathbf{g}_k^T \mathbf{A} \mathbf{g}_k}.$$

Astfel, la iterația k a algoritmului steepest descent, dacă $\mathbf{g}_k \neq 0$, se calculează următorul punct utilizând

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \frac{\mathbf{g}_k^T \mathbf{g}_k}{\mathbf{g}_k^T \mathbf{A} \mathbf{g}_k} \mathbf{g}_k. \quad (\text{A.2.4})$$

Dacă pentru minimizarea funcției pătratice (A.2.3) se utilizează un factor $\eta > 0$ constant, convergența algoritmului este garantată numai dacă

$$0 < \eta < \frac{2}{\lambda_{\max}(\mathbf{A})} \quad (\text{A.2.5})$$

unde $\lambda_{\max}(\mathbf{A})$ notează valoarea proprie maximă a matricei $\mathbf{A}$ (care este strict pozitivă datorită faptului că $\mathbf{A}$ este pozitiv definită). ■

A.2.2. Metoda lui Newton

Metoda lui Newton de minimizarea a unei funcții obiectiv f se bazează pe aproximarea lui f printr-o formă pătratică și minimizarea acestei forme. Astfel, la fiecare iterație direcția de căutare este determinată utilizând a doua derivată a funcției obiectiv f .

Presupunem că funcția $f: \mathbb{R}^n \rightarrow \mathbb{R}$ al cărei punct de minim dorim să-l determinăm printr-o procedură iterativă este de clasă $C_{\mathbb{R}^n}^2$. La iterația k , punctul curent $\mathbf{x}_k$ este astfel încât matricea Hessian în $\mathbf{x}_k$, notată $\mathbf{H}_k = \nabla^2 f(\mathbf{x}_k)$, este pozitiv definită. Dacă punctul curent $\mathbf{x}_k$ este suficient de apropiat de punctul minim local $\mathbf{x}^*$, funcția f poate fi aproximată în vecinătatea lui $\mathbf{x}_k$ prin

$$f(\mathbf{x}_k + \mathbf{s}) \approx q_k(\mathbf{s}) = f(\mathbf{x}_k) + \mathbf{g}_k^T \mathbf{s} + \frac{1}{2} \mathbf{s}^T \mathbf{H}_k \mathbf{s}$$

unde $\mathbf{g}_k = \nabla f(\mathbf{x}_k)$ și $\mathbf{s} = \mathbf{x} - \mathbf{x}_k$. Această expresie reprezintă dezvoltarea funcției f în serie Taylor din care termenii de ordin mai mare ca 2 au fost neglijați.

Punctul de minim al acestei forme pătratice se determină egalând cu 0 gradientul său (vezi exemplul A.1.3). Astfel se obține $\mathbf{H}_k(\mathbf{x} - \mathbf{x}_k) + \mathbf{g}_k = 0$, care, în ipoteza că matricea $\mathbf{H}_k$ este nesingulară, admite soluția

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \mathbf{H}_k^{-1} \mathbf{g}_k \quad (\text{A.2.6})$$

Pasul acestei metode, având forma $\mathbf{d}_k = -\mathbf{H}_k^{-1} \mathbf{g}_k$, este denumit *pas Newton*.

Metoda Newton de minimizare a unei funcții obiectiv $f: \mathbb{R}^n \rightarrow \mathbb{R}$, de clasă $C_{\mathbb{R}^n}^2$, este următoarea:

Algoritmul Nw (metoda lui Newton).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$.

Pas 1. (Testare criterii de stop) Fie $\mathbf{g}_k = \nabla f(\mathbf{x}_k)$.

1.1 Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 5.

1.2. Dacă $k = N$, treci la Pasul 5.

Pas 2. (Determinarea pasului) Se calculează $\mathbf{H}_k = \nabla^2 f(\mathbf{x}_k)$ și se determină vectorul $\mathbf{s}_k$ soluție a ecuației

$$\mathbf{H}_k \mathbf{s} = -\mathbf{g}_k. \quad (\text{A.2.7})$$

Pas 3. (Actualizare) Se calculează următorul punct

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k.$$

Pas 4. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.

Pas 5. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

Dacă funcția obiectiv f este pătratică și $\mathbf{H}_0 = \nabla^2 f(\mathbf{x}_0)$ este pozitiv definită, metoda lui Newton conduce la punctul de minim într-un singur pas. Dacă funcția f nu este pătratică, atunci ar putea fi necesară efectuarea unui număr mai mare de iterații. Criteriul de stop poate fi oricare dintre cele prezentate anterior.

Dezavantajul major al metodei Newton îl reprezintă necesitatea calculării inversei matricei Hessian operație care, pe lângă faptul că presupune un număr mare de operații, poate introduce probleme numerice datorită erorilor de rotunjire.

Exemplul A.2.2. Pentru a ilustra aplicarea metodei lui Newton considerăm forma pătratică dată prin ecuația (A.2.3)

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, \quad f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c,$$

pentru care matricea $\mathbf{A} \in \mathbb{R}^{n \times n}$ este simetrică și pozitiv definită ($\mathbf{A}^T = \mathbf{A} > 0$), $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$, utilizată și în exemplul A.2.1. Așa cum a fost prezentat în exemplul A.1.3, funcția f admite unic punct de minim, $\mathbf{x}^* = -\mathbf{A}^{-1} \mathbf{b}$.

Aplicând metoda lui Newton, la iterația $k = 0$ se pornește din punctul inițial $\mathbf{x}_0$ și se calculează gradientul $\mathbf{g}_0 = \nabla f(\mathbf{x}_0) = \mathbf{A} \mathbf{x}_0 + \mathbf{b}$. Presupunem că $\mathbf{x}_0 \neq \mathbf{x}^*$, astfel că $\mathbf{g}_0 \neq 0$. Matricea Hessian $\mathbf{H}_0 = \nabla^2 f(\mathbf{x}_0) = \mathbf{A}$ este pozitiv definită. La pasul 2 al algoritmului se rezolvă ecuația (A.2.7), care devine

$$\mathbf{A} \mathbf{s}_0 = -(\mathbf{A} \mathbf{x}_0 + \mathbf{b}),$$

astfel că se obține $\mathbf{s}_0 = -\mathbf{x}_0 - \mathbf{A}^{-1} \mathbf{b}$. Utilizând această expresie la pasul 3 pentru actualizare rezultă

$$\mathbf{x}_1 = \mathbf{x}_0 + \mathbf{s}_0 = -\mathbf{A}^{-1} \mathbf{b} = \mathbf{x}^*.$$

Aceasta arată că, aplicând metoda lui Newton pentru minimizarea unei funcții pătratice se ajunge în punctul de minim într-o singură iterație. ■

A.2.3. Metoda Gauss-Newton

În numeroase aplicații se ajunge la rezolvarea problemei celor mai mici pătrate în formă neliniară, și anume minimizarea unei funcții de forma

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, \quad f(\mathbf{x}) = \frac{1}{2} \mathbf{r}^T(\mathbf{x}) \cdot \mathbf{r}(\mathbf{x}) = \frac{1}{2} \sum_{k=1}^m [r_k(\mathbf{x})]^2, \quad (\text{A.2.8})$$

unde $\mathbf{r}: \mathbb{R}^n \rightarrow \mathbb{R}^m$, $\mathbf{r}(\mathbf{x}) = [r_1(\mathbf{x}) \ r_2(\mathbf{x}) \ \dots \ r_m(\mathbf{x})]^T$, este o funcție neliniară. Dacă funcția $\mathbf{r}(\mathbf{x})$ este liniară, $\mathbf{r}(\mathbf{x}) = \mathbf{A} \mathbf{x} - \mathbf{b}$, se ajunge la estimatorul celor mai mici pătrate prezentat în secțiunea A.1.3. La o asemenea problemă se poate ajunge, de exemplu, dacă dorim să rezolvăm un sistem de m ecuații neliniare

$$r_i(\mathbf{x}) = 0, \quad i = 1, 2, \dots, m.$$

Pentru a aplica metoda lui Newton de minimizare a funcției (A.2.8) este nevoie să calculăm mai întâi gradientul funcției f într-un punct oarecare. Pentru aceasta calculăm

$$\frac{\partial f}{\partial x_i}(\mathbf{x}) = \sum_{k=1}^m r_k(\mathbf{x}) \cdot \frac{\partial r_k}{\partial x_i}(\mathbf{x}) = \left[\frac{\partial \mathbf{r}}{\partial x_i}(\mathbf{x}) \right]^T \cdot \mathbf{r}(\mathbf{x})$$

unde

$$\frac{\partial \mathbf{r}}{\partial x_i}(\mathbf{x}) = \left[\frac{\partial r_1}{\partial x_i}(\mathbf{x}) \quad \frac{\partial r_2}{\partial x_i}(\mathbf{x}) \quad \dots \quad \frac{\partial r_m}{\partial x_i}(\mathbf{x}) \right]^T, \quad i = \overline{1, n}.$$

Se introduce matricea

$$\mathbf{J}(\mathbf{x}) = \left[\frac{\partial \mathbf{r}}{\partial x_1}(\mathbf{x}) \quad \dots \quad \frac{\partial \mathbf{r}}{\partial x_n}(\mathbf{x}) \right] = \begin{bmatrix} \frac{\partial r_1}{\partial x_1}(\mathbf{x}) & \dots & \frac{\partial r_1}{\partial x_n}(\mathbf{x}) \\ \vdots & \ddots & \vdots \\ \frac{\partial r_m}{\partial x_1}(\mathbf{x}) & \dots & \frac{\partial r_m}{\partial x_n}(\mathbf{x}) \end{bmatrix} \in \mathbb{R}^{m \times n}, \quad (\text{A.2.9})$$

care este numită *matricea Jacobian* a funcției vectoriale $\mathbf{r}$. Cu această notație, gradientul funcției f (A.2.8) devine:

$$\nabla f(\mathbf{x}) = \mathbf{J}(\mathbf{x})^T \cdot \mathbf{r}(\mathbf{x}). \quad (\text{A.2.10})$$

Pentru calcularea matricei Hessian a funcției f este nevoie de derivatele parțiale de ordinul 2 ale lui f

$$\frac{\partial^2 f}{\partial x_i \partial x_j}(\mathbf{x}) = \frac{\partial}{\partial x_j} \left[\sum_{k=1}^m r_k(\mathbf{x}) \cdot \frac{\partial r_k}{\partial x_i}(\mathbf{x}) \right] = \sum_{k=1}^m \frac{\partial r_k}{\partial x_j}(\mathbf{x}) \cdot \frac{\partial r_k}{\partial x_i}(\mathbf{x}) + \sum_{k=1}^m r_k(\mathbf{x}) \cdot \frac{\partial^2 r_k}{\partial x_i \partial x_j}(\mathbf{x}),$$

ceea ce conduce la

$$\mathbf{H}(\mathbf{x}) = \mathbf{J}(\mathbf{x})^T \cdot \mathbf{J}(\mathbf{x}) + \mathbf{S}(\mathbf{x}). \quad (\text{A.2.11})$$

unde matricea $\mathbf{S}(\mathbf{x}) = [s_{ij}(\mathbf{x})] \in \mathbb{R}^{n \times n}$ are elementele

$$s_{ij}(\mathbf{x}) = \sum_{k=1}^m r_k(\mathbf{x}) \cdot \frac{\partial^2 r_k}{\partial x_i \partial x_j}(\mathbf{x}), \quad i, j = \overline{1, n}. \quad (\text{A.2.12})$$

Cu relațiile (A.2.10) și (A.2.11), formula de actualizare (A.2.6) care intervine în aplicarea metodei lui Newton devine:

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \left[\mathbf{J}(\mathbf{x}_k)^T \mathbf{J}(\mathbf{x}_k) + \mathbf{S}(\mathbf{x}_k) \right]^{-1} \mathbf{J}(\mathbf{x}_k)^T \mathbf{r}(\mathbf{x}_k), \quad (\text{A.2.13})$$

În multe aplicații matricea $\mathbf{S}(\mathbf{x})$ poate fi neglijată în formula (A.2.13) deoarece elementele sale au valori foarte mici. Se obține astfel *metoda Gauss-Newton* de rezolvare a problemei celor mai mici pătrate în formă neliniară care utilizează formula de actualizare:

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \left[\mathbf{J}(\mathbf{x}_k)^T \mathbf{J}(\mathbf{x}_k) \right]^{-1} \mathbf{J}(\mathbf{x}_k)^T \mathbf{r}(\mathbf{x}_k). \quad (\text{A.2.14})$$

Facem precizarea că pentru a putea aplica această relație este necesar ca matricea Jacobian $\mathbf{J}(\mathbf{x}_k)$ să fie de rang complet pe coloane, adică $\text{rang } \mathbf{J}(\mathbf{x}_k) = n$ (având implicit $n \leq m$).

În concluzie, putem formula în continuare algoritmul Gauss-Newton de minimizare a funcției (A.2.8).

Algoritmul GN (algoritmul Gauss-Newton).

- Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$.
- Pas 1. (Testare criterii de stop) Se calculează $\mathbf{r}_k = \mathbf{r}(\mathbf{x}_k)$, $\mathbf{J}_k = \mathbf{J}(\mathbf{x}_k)$ și $\mathbf{g}_k = \mathbf{J}_k^T \mathbf{r}_k$.
 1.1 Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 5.
 1.2. Dacă $k = N$, treci la Pasul 5.
- Pas 2. (Determinarea pasului) Se determină vectorul $\mathbf{s}_k$ soluție a ecuației
- $$\mathbf{J}_k^T \mathbf{J}_k \mathbf{s} = -\mathbf{g}_k.$$
- Pas 3. (Actualizare) Se calculează următorul punct
- $$\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k.$$
- Pas 4. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.
- Pas 5. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

A.2.4. Metode de tip Newton modificate

Metoda clasică a lui Newton, definită prin ecuația (A.2.6), este o metodă locală. Dacă punctul curent este departe de punctul de minim nu este garantat faptul că matricea Hessian este pozitiv definită, astfel că ecuația (A.2.7) poate să nu furnizeze direcția de descreștere a funcției f . Aceasta se datorează neglijării termenilor de grad mare din dezvoltarea în serie Taylor a funcției obiectiv.

Metoda lui Newton cu pas adaptiv

O modificare simplă a metodei Newton constă în adaptarea pasului de căutare la fiecare iterație, modificând ecuația (A.2.6) în

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \eta_k \mathbf{H}_k^{-1} \mathbf{g}_k \quad (\text{A.2.15})$$

unde factorul $\eta_k > 0$ este selectat în așa fel încât să fie minimizată funcția f .

Concret, metoda Newton cu pas adaptiv pentru minimizarea unei funcții obiectiv $f: \mathbb{R}^n \rightarrow \mathbb{R}$, de clasă $C_{\mathbb{R}^n}^2$, este următoarea:

Algoritmul NwLS (metoda lui Newton & Line Search).

- Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$.
- Pas 1. (Testare criterii de stop) Fie $\mathbf{g}_k = \nabla f(\mathbf{x}_k)$.
 1.1 Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 6.
 1.2. Dacă $k = N$, treci la Pasul 6.
- Pas 2. (Determinarea direcției de căutare) Se calculează $\mathbf{H}_k = \nabla^2 f(\mathbf{x}_k)$ și se determină vectorul $\mathbf{s}_k$ soluție a ecuației

$$\mathbf{H}_k \mathbf{s} = -\mathbf{g}_k.$$

- Pas 3. (Determinarea lungimii pasului) Se determină factorul η_k astfel încât
- $$f(\mathbf{x}_k + \eta_k \mathbf{s}_k) = \min_{\eta > 0} f(\mathbf{x}_k + \eta \mathbf{s}_k). \quad (\text{A.2.16})$$
- Pas 4. (Actualizare) Se calculează următorul punct
- $$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{s}_k.$$
- Pas 5. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.
- Pas 6. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

Pentru rezolvarea problemei (A.2.16) se poate utiliza o metodă de căutare liniară, de exemplu metoda secțiunii de aur.

În practică, și metoda Gauss-Newton este frecvent combinată cu o metodă de căutare liniară, utilizând în locul relației de actualizare (A.2.14) o formulă de tipul (A.2.15), și anume

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \eta_k \left[\mathbf{J}(\mathbf{x}_k)^T \mathbf{J}(\mathbf{x}_k) \right]^{-1} \mathbf{J}(\mathbf{x}_k)^T \mathbf{r}(\mathbf{x}_k),$$

unde factorul $\eta_k > 0$ este selectat în așa fel încât să fie minimizată funcția f (A.2.8).

O altă posibilitate de adaptare a pasului la fiecare iterație este aceea de a înjumătăți valoarea factorului η la fiecare iterație,

$$\eta_{k+1} = \frac{1}{2} \eta_k, \quad k = 0, 1, 2, \dots,$$

cu η_0 ales 1 sau puțin mai mic. Este demonstrat faptul că metoda de înjumătățire a pasului converge către punctul de minim local.

Metoda Levenberg-Marquardt

Dacă matricea Hessian nu este pozitiv definită direcția Newton poate să indice către un punct de maxim sau către un punct în șa. Pentru a evita asemenea situații, se aduna la matricea Hessian $\mathbf{H}$ o matrice pozitiv definită $\mathbf{P}$ pentru a obține o matrice pozitiv definită. Aceasta metodă a fost introdusă pentru prima dată de Levenberg și Marquardt în problema celor mai mici pătrate. Dacă se alege $\mathbf{P} = \lambda \mathbf{I}$, cu $\mathbf{I}$ matricea unitate, pentru valori suficient de mari ale factorului $\lambda > 0$ matricea $\mathbf{H} + \lambda \mathbf{I}$ este pozitiv definită.

Astfel, ecuația utilizată la pasul de actualizare este

$$\mathbf{x}_{k+1} = \mathbf{x}_k - (\mathbf{H}_k + \lambda_k \mathbf{I})^{-1} \mathbf{g}_k. \quad (\text{A.2.17})$$

În practică, la fiecare iterație se inițializează λ_k cu o valoare mică ce este apoi crescută până când este satisfăcută condiția de descreștere $f(\mathbf{x}_{k+1}) < f(\mathbf{x}_k)$.

Pentru $\lambda \rightarrow 0$ se obține metoda clasică a lui Newton în timp ce pentru $\lambda \rightarrow \infty$ rezultă metoda steepest-descent.

Cele două metode pot fi combinate, considerând relația de actualizare de forma

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \eta_k (\mathbf{H}_k + \lambda \mathbf{I})^{-1} \mathbf{g}_k. \quad (\text{A.2.18})$$

A.2.5. Metode quasi-Newton

În cadrul metodelor de tip quasi-Newton se construiește iterativ o aproximare $\mathbf{M}$ a matricei Hessian $\mathbf{H}$ sau a inversei sale $\mathbf{H}^{-1}$ pe măsură ce algoritmul progresa. În continuare ilustrăm pe scurt cel de-al doilea tip de metode.

Cu notațiile uzuale, în vecinătatea punctului de minim se consideră aproximarea

$$\mathbf{g}_{k+1} - \mathbf{g}_k \approx \mathbf{H}_{k+1} (\mathbf{x}_{k+1} - \mathbf{x}_k).$$

Pornind de la relația

$$\Delta \mathbf{x}_k = \mathbf{M}_{k+1} \Delta \mathbf{g}_k, \quad (\text{A.2.19})$$

unde $\Delta \mathbf{x}_k = \mathbf{x}_{k+1} - \mathbf{x}_k$, $\Delta \mathbf{g}_k = \mathbf{g}_{k+1} - \mathbf{g}_k$, se construiește aproximarea $\mathbf{M}_{k+1}$ a matricei $\mathbf{H}_{k+1}^{-1}$. Inițial se consideră $\mathbf{M}_0 = \mathbf{I}$. Metodele bazate pe acest principiu diferă între ele prin relațiile de recurență utilizate pentru actualizarea matricei $\mathbf{M}_k$ la fiecare iterație. Două dintre metodele frecvent utilizate sunt Davidson-Fletcher-Powell (DFP) și Broyden-Fletcher-Golfarb-Shanno (BFGS) care utilizează relațiile:

$$(\text{DFP}) \quad \mathbf{M}_{k+1} = \mathbf{M}_k + \frac{\Delta \mathbf{x}_k \cdot \Delta \mathbf{x}_k^T}{\Delta \mathbf{x}_k^T \cdot \Delta \mathbf{g}_k} - \frac{\mathbf{M}_k \cdot \Delta \mathbf{g}_k \cdot \Delta \mathbf{g}_k^T \cdot \mathbf{M}_k}{\Delta \mathbf{g}_k^T \cdot \mathbf{M}_k \cdot \Delta \mathbf{g}_k}, \quad (\text{A.2.20})$$

$$(\text{BFGS}) \quad \mathbf{M}_{k+1} = \left(\mathbf{I} - \frac{\Delta \mathbf{x}_k \cdot \Delta \mathbf{g}_k^T}{\Delta \mathbf{x}_k^T \cdot \Delta \mathbf{g}_k} \right) \mathbf{M}_k \left(\mathbf{I} - \frac{\Delta \mathbf{g}_k \cdot \Delta \mathbf{x}_k^T}{\Delta \mathbf{x}_k^T \cdot \Delta \mathbf{g}_k} \right) + \frac{\Delta \mathbf{x}_k \cdot \Delta \mathbf{x}_k^T}{\Delta \mathbf{x}_k^T \cdot \Delta \mathbf{g}_k}. \quad (\text{A.2.21})$$

Structura generală a unui algoritm de tip quasi-Newton este următoarea:

Algoritm qNw (algoritm general quasi-Newton).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, aproximarea inițială $\mathbf{M}_0 \in \mathbb{R}^{n \times n}$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații $N \in \mathbb{N}$.

Pas 1. (Testare criterii de stop) Fie $\mathbf{g}_k = \nabla f(\mathbf{x}_k)$.

1.1 Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 7.

1.2. Dacă $k = N$, treci la Pasul 7.

Pas 2. (Determinarea direcției de căutare) Se calculează vectorul $\mathbf{s}_k$

$$\mathbf{s}_k = -\mathbf{M}_k \mathbf{g}_k.$$

Pas 3. (Determinarea lungimii pasului) Se determină factorul η_k astfel încât

$$f(\mathbf{x}_k + \eta_k \mathbf{s}_k) = \min_{\eta > 0} f(\mathbf{x}_k + \eta \mathbf{s}_k).$$

Pas 4. (Actualizare $\mathbf{x}_{k+1}$) Se calculează următorul punct

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{s}_k.$$

Pas 5. (Actualizare $\mathbf{M}_{k+1}$) Se calculează matricea $\mathbf{M}_{k+1}$ astfel încât să fie satisfăcută ecuația quasi-Newton (A.2.19).

Pas 6. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.

Pas 7. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

A.3. Metode bazate pe direcții conjugate

A.3.1. Metoda direcțiilor conjugate

Metodele de gradient conjugat reprezintă un compromis între metodele de tip Newton și metoda de descreștere maximă. Metodele de gradient conjugat necesită numai derivatele de ordinul 1 nefiind necesară calcularea derivatelor de ordinul 2, așa cum se întâmplă în cazul metodelor de tip Newton. Metodele bazate pe direcții conjugate sunt mai răspândite decât metoda steepest descent, dar mai lente decât metodele de tip Newton. Factorul decisiv asupra eficienței unei metode iterative de căutare este direcția de căutare actualizată la fiecare iterație.

Introducem în continuare câteva noțiuni necesare în dezvoltarea teoretică.

Fie $\mathbf{Q} \in \mathbb{R}^{n \times n}$ o matrice simetrică și pozitiv definită. Se spune că direcțiile $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_m$, $m \leq n-1$, sunt $\mathbf{Q}$ -conjugate dacă $\mathbf{d}_i^T \mathbf{Q} \mathbf{d}_j = 0$, pentru orice $i \neq j$. În particular, pentru $\mathbf{Q} = \mathbf{I}_n$, se obține definiția unei mulțimi de vectori ortogonali.

Se poate demonstra că dacă direcțiile $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_m \in \mathbb{R}^n \setminus \{0\}$, $m \leq n-1$, sunt $\mathbf{Q}$ -conjugate, atunci sunt liniar independente.

Algoritmul general de minimizare a unei funcții $f: \mathbb{R}^n \rightarrow \mathbb{R}$ prin metoda direcțiilor conjugate este prezentat în continuare.

Algoritmul GCDM (metoda generală a direcțiilor conjugate).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, matricea pozitiv definită $\mathbf{Q} > 0$ și valoarea de prag $\varepsilon > 0$.

Se calculează $\mathbf{g}_0 = \nabla f(\mathbf{x}_0)$ și se alege $\mathbf{d}_0$ astfel încât $\mathbf{d}_0^T \mathbf{g}_0 < 0$

Pas 1. (Testare criteriu de stop) Fie $\mathbf{g}_k = \nabla f(\mathbf{x}_k)$. Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 6.

Pas 2. (Determinarea lungimii pasului) Se determină factorul η_k astfel încât

$$f(\mathbf{x}_k + \eta_k \mathbf{d}_k) = \min_{\eta} f(\mathbf{x}_k + \eta \mathbf{d}_k). \quad (\text{A.3.1})$$

Pas 3. (Actualizare) Se calculează următorul punct

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k. \quad (\text{A.3.2})$$

Pas 4. (Determinarea următorului pas) Se calculează direcția $\mathbf{d}_{k+1}$ care este $\mathbf{Q}$ -conjugată direcțiilor $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_k$, adică:

$$\mathbf{d}_{k+1}^T \mathbf{Q} \mathbf{d}_j = 0, \quad j = 0, 1, \dots, k. \quad (\text{A.3.3})$$

Pas 5. (Ciclare) Se setează $k := k+1$ și se trece la Pasul 1.

Pas 6. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

Exemplul A.3.1. Pentru a ilustra aplicarea metodei bazate pe direcții conjugate considerăm forma pătratică precizată prin relația

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c, \quad (\text{A.3.4})$$

în care matricea $\mathbf{A} \in \mathbb{R}^{n \times n}$ este simetrică și pozitiv definită ($\mathbf{A}^T = \mathbf{A} > 0$), $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$, utilizată și în exemplele A.2.1 și A.2.2. Această funcție admite un punctul de minim unic, $\mathbf{x}^* = -\mathbf{A}^{-1}\mathbf{b}$. Pentru minimizarea funcției f vom utiliza direcții $\mathbf{A}$ -conjugate.

Pentru $k=0$ se alege $\mathbf{x}_0$, se calculează $\mathbf{g}_0 = \nabla f(\mathbf{x}_0) = \mathbf{A}\mathbf{x}_0 + \mathbf{b}$. Presupunând că $\mathbf{g}_0 \neq 0$ se alege direcția $\mathbf{d}_0$ astfel încât $\mathbf{d}_0^T \mathbf{g}_0 < 0$, ceea ce arată că $\mathbf{d}_0$ este o direcție de descreștere pentru funcția f pornind din $\mathbf{x}_0$.

La iterația $k \geq 0$ se calculează

$$\eta_k = \arg \min_{\eta} f(\mathbf{x}_k + \eta \mathbf{d}_k).$$

Pentru aceasta procedăm în același mod ca în exemplul A.2.1 și consideră funcția

$$\varphi_k(\eta) = f(\mathbf{x}_k + \eta \mathbf{d}_k) = \frac{1}{2} \mathbf{d}_k^T \mathbf{A} \mathbf{d}_k \eta^2 + \mathbf{g}_k^T \mathbf{d}_k \eta + f(\mathbf{x}_k),$$

unde $\mathbf{g}_k = \nabla f(\mathbf{x}_k) = \mathbf{A}\mathbf{x}_k + \mathbf{b}$. Funcția $\varphi_k(\eta)$ are gradul 2 în η și coeficientul termenului dominant pozitiv ($\frac{1}{2} \mathbf{d}_k^T \mathbf{A} \mathbf{d}_k > 0$ datorită faptului că matricea $\mathbf{A}$ este pozitiv definită), astfel că admite minim în

$$\eta_k = -\frac{\mathbf{g}_k^T \mathbf{d}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k}. \quad (\text{A.3.5})$$

Astfel, la iterația k a algoritmului direcțiilor conjugate se calculează punctul

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \frac{\mathbf{g}_k^T \mathbf{d}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k} \mathbf{d}_k. \quad (\text{A.3.6})$$

Se calculează apoi următoarea direcție de căutare $\mathbf{d}_{k+1}$ care se alege astfel încât să fie $\mathbf{A}$ -conjugată direcțiilor $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_k$, adică:

$$\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_j = 0, \quad j = 0, 1, \dots, k.$$

Arătăm în continuare că pornind dintr-un punct oarecare $\mathbf{x}_0$, se ajunge în punctul de minim în cel mult n iterații, adică $\mathbf{x}_n = \mathbf{x}^*$.

Deoarece direcțiile $\mathbf{A}$ -conjugate $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{n-1}$, sunt liniar independente, formează o bază în $\mathbb{R}^n$, în care vectorul $\mathbf{x}^* - \mathbf{x}_0$ admite descompunerea

$$\mathbf{x}^* - \mathbf{x}_0 = \alpha_0 \mathbf{d}_0 + \alpha_1 \mathbf{d}_1 + \dots + \alpha_{n-1} \mathbf{d}_{n-1} \quad (\text{A.3.7})$$

Deoarece $\mathbf{d}_k^T \mathbf{A} \mathbf{d}_j = 0$ pentru $j \neq k$, înmulțind la dreapta relația (A.3.3) cu $\mathbf{d}_k^T \mathbf{A}$ rezultă

$$\mathbf{d}_k^T \mathbf{A} (\mathbf{x}^* - \mathbf{x}_0) = \alpha_k \mathbf{d}_k^T \mathbf{A} \mathbf{d}_k. \quad (\text{A.3.8})$$

Pe de altă parte, din relația de recurență (A.3.2) scrisă în forma $\mathbf{x}_k = \mathbf{x}_{k-1} + \eta_{k-1} \mathbf{d}_{k-1}$, se obține

$$\mathbf{x}_k = \mathbf{x}_0 + \eta_0 \mathbf{d}_0 + \dots + \eta_{k-1} \mathbf{d}_{k-1}, \quad k = 1, 2, \dots,$$

relație ce arată că vectorul $\mathbf{x}_k - \mathbf{x}_0$ este conținut în subspațiul liniar generat de vectorii $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{k-1}$, astfel că $\mathbf{d}_k^T \mathbf{A} (\mathbf{x}_k - \mathbf{x}_0) = 0$. Scriind $\mathbf{x}^* - \mathbf{x}_0$ sub forma $\mathbf{x}^* - \mathbf{x}_0 = (\mathbf{x}^* - \mathbf{x}_k) + (\mathbf{x}_k - \mathbf{x}_0)$, se obține

$$\mathbf{d}_k^T \mathbf{A} (\mathbf{x}_k - \mathbf{x}_0) = \mathbf{d}_k^T \mathbf{A} (\mathbf{x}^* - \mathbf{x}_k) = -\mathbf{d}_k^T (\mathbf{A} \mathbf{x}_k + \mathbf{b}) = -\mathbf{d}_k^T \mathbf{g}_k \quad (\text{A.3.9})$$

deoarece $\mathbf{A} \mathbf{x}^* = -\mathbf{b}$ și $\mathbf{g}_k = \mathbf{A} \mathbf{x}_k + \mathbf{b}$. Din relațiile (A.3.8) și (A.3.9) rezultă

$$\alpha_k \mathbf{d}_k^T \mathbf{A} \mathbf{d}_k = -\mathbf{d}_k^T \mathbf{g}_k \Rightarrow \alpha_k = -\frac{\mathbf{d}_k^T \mathbf{g}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k},$$

care, prin comparație cu (A.3.5) revine la $\alpha_k = \eta_k$. Înlocuind în (A.3.7) se obține

$$\mathbf{x}^* = \mathbf{x}_0 + \eta_0 \mathbf{d}_0 + \eta_1 \mathbf{d}_1 + \dots + \eta_{n-1} \mathbf{d}_{n-1},$$

adică $\mathbf{x}^* = \mathbf{x}_n$. ■

Algoritmul general de minimizare a unei funcții $f: \mathbb{R}^n \rightarrow \mathbb{R}$ prin metoda direcțiilor conjugate este prezentat în continuare.

Pe baza rezultatelor prezentate în exemplul A.3.1, algoritmul de minimizare a funcției pătratice (A.3.4)

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, \quad f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c,$$

cu $\mathbf{A} \in \mathbb{R}^{n \times n}$ matrice simetrică și pozitiv definită ($\mathbf{A}^T = \mathbf{A} > 0$), $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$, poate fi formulat în modul următor:

Algoritmul QCDM (metoda direcțiilor conjugate pentru forme pătratice).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$ și direcțiile

$\mathbf{A}$ -conjugate $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{n-1}$.

Pas 1. (Testare criteriu de stop). Dacă $k = n$, treci la Pasul 5.

Pas 2. (Determinarea lungimii pasului) Se calculează $\mathbf{g}_k = \nabla f(\mathbf{x}_k) = \mathbf{A} \mathbf{x}_k + \mathbf{b}$ și

$$\eta_k = -\frac{\mathbf{g}_k^T \mathbf{d}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k}.$$

Pas 3. (Actualizare) Se calculează următorul punct

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k.$$

Pas 4. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.

Pas 5. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_n$.

Aplicând algoritmul direcțiilor conjugate pentru minimizarea unei funcții pătratice este satisfăcută următoarea proprietate:

$$\mathbf{g}_{k+1}^T \mathbf{d}_i = 0, \quad (\text{A.3.10})$$

pentru orice $0 \leq k \leq n-1$ și $0 \leq i \leq k$. Această relație arată că $\mathbf{g}_{k+1}$ este ortogonal pe subspațiul generat de vectorii $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_k$. În plus, în afara faptului că $f(\mathbf{x}_{k+1}) = \min_{\eta} f(\mathbf{x}_k + \eta \mathbf{d}_k)$ (relație utilizată pentru determinarea punctului $\mathbf{x}_{k+1}$) se poate demonstra că

$$f(\mathbf{x}_{k+1}) = \min_{\alpha_0, \dots, \alpha_k} f\left(\mathbf{x}_0 + \sum_{i=0}^k \alpha_i \mathbf{d}_i\right).$$

A.3.2. Metoda gradientului conjugat

În algoritmul direcțiilor conjugate prezentat în secțiunea precedentă nu este precizată nici o regulă pentru alegerea direcțiilor conjugate, acestea putând fi specificate din start. Singura condiție impusă este aceea ca direcțiile să fie conjugate în raport cu o matrice dată.

Spre deosebire de acesta, în *algoritmul gradientului conjugat* direcțiile de căutare sunt calculate pe măsură ce algoritmul progresează. La fiecare pas, direcția de căutare este considerată ca o combinație liniară a direcțiilor precedente de căutare și a gradientului în punctul curent astfel încât toate direcțiile să fie conjugate. Metoda gradientului conjugat a fost propusă în anii 1950 de Hestenes și Stiefel pentru rezolvarea sistemelor lineare și extinsă apoi în anii 1960 de Fletcher și Reeves la cazul minimizării unei funcții generale. În prezent metodele de gradient conjugat sunt utilizate pentru rezolvarea problemelor de optimizare de dimensiuni mari.

Exemplul A.3.2. Ilustrăm în continuare aplicarea algoritmului gradientului conjugat pentru minimizarea formei pătratice

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, \quad f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c, \quad (\text{A.3.11})$$

în care matricea $\mathbf{A} \in \mathbb{R}^{n \times n}$ este simetrică și pozitiv definită ($\mathbf{A}^T = \mathbf{A} > 0$), $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$, utilizată și în exemplele A.2.1, A.2.2 și A.3.1.

Pentru $k=0$, se consideră punctul inițial $\mathbf{x}_0$ și se alege direcția $\mathbf{d}_0$ de descreștere maximă a lui f , adică

$$\mathbf{d}_0 = -\mathbf{g}_0,$$

unde $\mathbf{g}_0 = \nabla f(\mathbf{x}_0) = \mathbf{A} \mathbf{x}_0 + \mathbf{b}$.

Ținând cont de cele prezentate în exemplul A.3.1, se calculează

$$\eta_0 = \arg \min_{\eta} f(\mathbf{x}_0 + \eta \mathbf{d}_0) = -\frac{\mathbf{g}_0^T \mathbf{d}_0}{\mathbf{d}_0^T \mathbf{A} \mathbf{d}_0}$$

și se consideră

$$\mathbf{x}_1 = \mathbf{x}_0 + \eta_0 \mathbf{d}_0.$$

La pasul următor alegem o direcție $\mathbf{d}_1$, combinație liniară între $\mathbf{d}_0$ și $\mathbf{g}_1 = \nabla f(\mathbf{x}_1)$, de forma $\mathbf{d}_1 = -\mathbf{g}_1 + \beta_0 \mathbf{d}_0$, care să fie $\mathbf{A}$ -conjugată pentru $\mathbf{d}_0$, adică $\mathbf{d}_1^T \mathbf{A} \mathbf{d}_0 = 0$.

În general, la iterația k alegem direcția de căutare $\mathbf{d}_{k+1}$, ca o combinație liniară între $\mathbf{d}_k$ și $\mathbf{g}_{k+1}$ de forma

$$\mathbf{d}_{k+1} = -\mathbf{g}_{k+1} + \beta_k \mathbf{d}_k, \quad k = 0, 1, 2, \dots \quad (\text{A.3.12})$$

Coeficientul β_k este calculat astfel încât $\mathbf{d}_{k+1}$ să fie $\mathbf{A}$ -conjugată direcției $\mathbf{d}_k$. Aceasta conduce la

$$\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_k = -\mathbf{g}_{k+1}^T \mathbf{A} \mathbf{d}_k + \beta_k \mathbf{d}_k^T \mathbf{A} \mathbf{d}_k = 0,$$

de unde rezultă

$$\beta_k = \frac{\mathbf{g}_{k+1}^T \mathbf{A} \mathbf{d}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k}. \quad (\text{A.3.13})$$

Apoi se calculează noul punct utilizând relația

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k.$$

Se poate demonstra că direcțiile $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{n-1}$, obținute în modul prezentat mai sus sunt $\mathbf{A}$ -conjugate. În primul rând direcția $\mathbf{d}_1$ este selectată astfel încât să fie $\mathbf{A}$ -conjugată direcției $\mathbf{d}_0$. În continuare, presupunem că direcțiile $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_k$, $k < n-1$, sunt $\mathbf{A}$ -conjugate. Din relația (A.3.12) rezultă $\mathbf{g}_{k+1}^T \mathbf{d}_i = 0$, pentru orice $0 \leq i \leq k$. Cum fiecare $\mathbf{d}_i$ este ales de forma

$$\mathbf{d}_i = -\mathbf{g}_i + \beta_{i-1} \mathbf{d}_{i-1}, \quad i = 1, 2, \dots, k,$$

rezultă că

$$\mathbf{g}_{k+1}^T \mathbf{d}_i = -\mathbf{g}_{k+1}^T \mathbf{g}_i + \beta_{i-1} \mathbf{g}_{k+1}^T \mathbf{d}_{i-1},$$

astfel că

$$\mathbf{g}_{k+1}^T \mathbf{g}_i = 0, \quad i = 1, 2, \dots, k.$$

Cum $\mathbf{d}_0 = -\mathbf{g}_0$, are loc și $\mathbf{g}_{k+1}^T \mathbf{g}_0 = 0$.

În continuare demonstrăm că

$$\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_i = 0, \quad i = 0, 1, \dots, k-1.$$

Ținând cont de (A.3.12), avem $\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_i = (-\mathbf{g}_{k+1} + \beta_k \mathbf{d}_k)^T \mathbf{A} \mathbf{d}_i$.

În ipoteza considerată, $\mathbf{d}_k^T \mathbf{A} \mathbf{d}_i = 0$, $i = 0, 1, \dots, k$, de unde rezultă $\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_i = -\mathbf{g}_{k+1}^T \mathbf{A} \mathbf{d}_i$, $i = 0, 1, \dots, k-1$. Dar $\mathbf{g}_{i+1} = \mathbf{A} \mathbf{x}_{i+1} + \mathbf{b} = \mathbf{A}(\mathbf{x}_i + \eta_i \mathbf{d}_i) + \mathbf{b} = \mathbf{g}_i + \eta_i \mathbf{A} \mathbf{d}_i$, astfel că $\mathbf{A} \mathbf{d}_i = \frac{1}{\eta_i} (\mathbf{g}_{i+1} - \mathbf{g}_i)$. Cum $\mathbf{g}_{k+1}^T \mathbf{g}_i = 0$, $i = 1, 2, \dots, k$, rezultă în continuare că $\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_i = 0$ pentru $i = 0, 1, \dots, k-1$. Pe de altă parte, $\mathbf{d}_{k+1}$ a fost calculată astfel încât să fie $\mathbf{A}$ -conjugată direcției $\mathbf{d}_k$, deci este satisfăcută și relația $\mathbf{d}_{k+1}^T \mathbf{A} \mathbf{d}_k = 0$.

Prin inducție matematică am demonstrat că dacă direcțiile $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_k$, $k < n-1$, sunt $\mathbf{A}$ -conjugate, atunci direcția $\mathbf{d}_{k+1}$ este $\mathbf{A}$ -conjugată acestora. Rezultă în final că direcțiile $\mathbf{d}_0, \mathbf{d}_1, \dots, \mathbf{d}_{n-1}$ sunt $\mathbf{A}$ -conjugate.

Pe baza celor prezentate pentru metoda generală a direcțiilor conjugate, algoritmul gradientului conjugat pentru minimizarea unei funcții pătratice se va termina în cel mult n iterații. Mai exact, se poate demonstra că algoritmul gradientului conjugat se va termina în cel mult m iterații, unde m reprezintă numărul de valori proprii distincte ale matricei A . ■

În continuare poate fi formulat algoritmul de minimizare a unei forme pătratice $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} + \mathbf{b}^T \mathbf{x} + c$, cu $\mathbf{A} \in \mathbb{R}^{n \times n}$, $\mathbf{A}^T = \mathbf{A} > 0$, $\mathbf{b} \in \mathbb{R}^n$ și $c \in \mathbb{R}$, prin metoda gradientului conjugat.

Algoritmul QCGM (metoda gradientului conjugat pentru forme pătratice).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$.

Pas 1. (Testare criteriu de stop). Se calculează $\mathbf{g}_0 = \nabla f(\mathbf{x}_0) = \mathbf{A} \mathbf{x}_0 + \mathbf{b}$.

Dacă $\mathbf{g}_0 = 0$, atunci se returnează $\mathbf{x}_f = \mathbf{x}_0$. STOP.

Altfel, se setează $\mathbf{d}_0 = -\mathbf{g}_0$.

Pas 2. (Determinarea lungimii pasului) Se calculează

$$\eta_k = -\frac{\mathbf{g}_k^T \mathbf{d}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k}.$$

Pas 3. (Actualizare) Se calculează următorul punct

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k.$$

Pas 4. Se calculează $\mathbf{g}_{k+1} = \nabla f(\mathbf{x}_{k+1}) = \mathbf{A} \mathbf{x}_{k+1} + \mathbf{b}$.

Dacă $\mathbf{g}_{k+1} = 0$, atunci se returnează $\mathbf{x}_f = \mathbf{x}_{k+1}$. STOP.

Pas 5. (Determinarea următoarei direcții de căutare) Se calculează

$$\beta_k = \frac{\mathbf{g}_{k+1}^T \mathbf{A} \mathbf{d}_k}{\mathbf{d}_k^T \mathbf{A} \mathbf{d}_k}$$

$$\mathbf{d}_{k+1} = -\mathbf{g}_{k+1} + \beta_k \mathbf{d}_k$$

Pas 6. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 2.

Pentru minimizarea unei funcții care nu este pătratică prin metoda gradientului conjugat se pornește de la ideea aproximării funcției de minimizat prin termenii de ordin cel mult 2 din dezvoltarea acesteia în serie Taylor în jurul punctului curent. Pentru a evita calcularea matricei Hessian a funcției respective (care, pentru o funcție pătratică este constantă) la fiecare iterație, trebuie modificate formulele care precizează valorile lui η_k și β_k astfel încât să nu mai apară matricea hessian.

Pentru calcularea lui η_k se poate utiliza o metodă de căutare liniară care să permită evaluarea

$$\eta_k = \arg \min_{\eta} f(\mathbf{x}_k + \eta \mathbf{d}_k).$$

În ceea ce privește calcularea factorului β_k , relația (A.3.13), este modificată astfel încât să nu mai apară matricea A . Printre formulele cele mai frecvent utilizate sunt:

$$\text{Fletcher-Reeves:} \quad \beta_k = \frac{\mathbf{g}_{k+1}^T \mathbf{g}_{k+1}}{\mathbf{g}_k^T \mathbf{g}_k} \quad (\text{A.3.14})$$

$$\text{Polak-Ribiere-Polyak:} \quad \beta_k = \frac{\mathbf{g}_{k+1}^T (\mathbf{g}_{k+1} - \mathbf{g}_k)}{\mathbf{g}_k^T \mathbf{g}_k} \quad (\text{A.3.15})$$

$$\text{Dixon:} \quad \beta_k = \frac{\mathbf{g}_{k+1}^T \mathbf{g}_{k+1}}{\mathbf{d}_k^T \mathbf{g}_k} \quad (\text{A.3.16})$$

$$\text{Dai-Yuan:} \quad \beta_k = \frac{\mathbf{g}_{k+1}^T \mathbf{g}_{k+1}}{\mathbf{d}_k^T (\mathbf{g}_{k+1} - \mathbf{g}_k)} \quad (\text{A.3.17})$$

Din punct de vedere istoric, prima formulă utilizată în acest sens este formula Hestenes-Stiefel:

$$\beta_k = \frac{\mathbf{g}_{k+1}^T (\mathbf{g}_{k+1} - \mathbf{g}_k)}{\mathbf{d}_k^T (\mathbf{g}_{k+1} - \mathbf{g}_k)}. \quad (\text{A.3.18})$$

Toate aceste formule sunt echivalente între ele în cazul în care funcția minimizată este pătratică.

Aplicarea acestor formule prezintă avantajul că nu necesită explicitarea matricei hessian la fiecare iterație, fiind nevoie numai de valorile gradientului. Pentru probleme nepătratice însă, algoritmul nu mai converge în n pași, iar proprietatea de „conjugare” a direcțiilor de căutare tinde să se piardă. De aceea o tehnică uzuală este de a reinițializa direcția de căutare cu direcția de descreștere maximă (opusă gradientului) la fiecare n (sau $n+1$) iterații. Algoritmul de va opri atunci când unul din criteriile de stop este satisfăcut.

Pentru simplitate prezentăm în continuare algoritmul Fletcher-Reeves cu reinițializare pentru minimizarea unei funcții $f: \mathbb{R}^n \rightarrow \mathbb{R}$.

Algoritmul RFRCGM (Restart Fletcher-Reeves Conjugate Gradient Method).

Pas 0. (Inițializare) Pentru $k = 0$ se precizează punctul inițial $\mathbf{x}_0 \in \mathbb{R}^n$, valoarea de prag $\varepsilon > 0$ și numărul maxim de iterații N . Se calculează $\mathbf{g}_0 = \nabla f(\mathbf{x}_0) = A$ și se setează $\mathbf{d}_0 = -\mathbf{g}_0$.

Pas 1. (Testare criterii de stop)
1.1 Dacă $\|\mathbf{g}_k\| < \varepsilon$, treci la Pasul 8.
1.2. Dacă $k = N$, treci la Pasul 8.

Pas 2. (Determinarea lungimii pasului) Se determină factorul η_k astfel încât
$$f(\mathbf{x}_k + \eta_k \mathbf{d}_k) = \min_{\eta} f(\mathbf{x}_k + \eta \mathbf{d}_k).$$

Pas 3. (Actualizare) Se calculează următorul punct
$$\mathbf{x}_{k+1} = \mathbf{x}_k + \eta_k \mathbf{d}_k \text{ și } \mathbf{g}_{k+1} = \nabla f(\mathbf{x}_{k+1}).$$

- Pas 4. Dacă $k \equiv 0 \pmod{n}$ atunci $\beta_k = 0$,
 altfel $\beta_k = \frac{\mathbf{g}_{k+1}^T \mathbf{g}_{k+1}}{\mathbf{g}_k^T \mathbf{g}_k}$.
- Pas 5. (Determinarea următoarei direcții de căutare) Se calculează
 $\mathbf{d}_{k+1} = -\mathbf{g}_{k+1} + \beta_k \mathbf{d}_k$
- Pas 6. Dacă $\mathbf{d}_{k+1}^T \mathbf{g}_{k+1} > 0$, atunci $\mathbf{d}_{k+1} = -\mathbf{g}_{k+1}$.
- Pas 7. (Ciclare) Se setează $k := k + 1$ și se trece la Pasul 1.
- Pas 8. (Finalizare) Se returnează punctul final $\mathbf{x}_f = \mathbf{x}_k$.

Anexa B

PROBLEMATICA CLASIFICĂRII FORMELOR

B.1. Tipuri de clasificatori

Cum neuronul artificial a fost conceput pentru a modela (imita) neuronul biologic, rețelelor neurale artificiale le-au fost atribuite diverse sarcini specifice sistemului nervos.

Dintre acestea, operația de *recunoaștere a formelor* (eng. *pattern recognition*) sau *clasificare* (eng. *classification*) este caracteristică tuturor ființelor vii. Natura fizică a obiectelor clasificate poate fi foarte diversă (imagini, sunete, semnale electrice, etc) astfel încât recunoașterea formelor a devenit parte integrantă a sistemelor de decizie bazate pe inteligența artificială. În prezent, recunoașterea formelor are aplicații în domenii variate precum: vedere artificială, recunoașterea optică a caracterelor (litere sau cifre) (eng. *Optical Character Recognition – OCR*), recunoașterea sunetelor, diagnosticarea pacienților cu ajutorul computerului (eng. *computer aided diagnosis*) pe baza procesării unor date medicale variate (imagini obținute prin expunerea la raze X, tomografii, EKG, EEG), recunoașterea vorbirii (comanda verbală a sistemelor), procesarea de imagini în astronomie, geologie, aplicații militare (recunoașterea automată a țintelor).

Pentru a determina clasa căreia îi aparține un anumit obiect sunt evaluate mai întâi anumite *trăsături* (eng. *feature*) caracteristice ale acestuia; decizia privind clasificarea este luată apoi în funcție de valorile trăsăturilor. În general, un obiect poate fi caracterizat prin mai multe trăsături. În funcție de scopul clasificării, în această operație sunt utilizate numai unele dintre aceste trăsături, efectuându-se mai întâi operația de *extragere a trăsăturilor* (eng. *feature extraction*) (selectare a măsurătorilor relevante, evaluarea unor noi trăsături care vor fi utilizate efectiv în clasificare pe baza măsurătorilor existente).

Un concept de bază în problemele de clasificare a formelor este cel de *cluster* (eng. *cluster*). Un cluster constă dintr-un număr de obiecte (forme) similare care sunt grupate împreună. Operația de clasificare înseamnă de fapt împărțirea formelor în clustere și stabilirea pentru fiecare cluster în parte a centrului și a frontierei sale.

Cunoașterea numărului de clustere și a localizării lor (măcar aproximative) în etapa de proiectare a clasificatorului simplifică operația de clasificare. În acest caz se poate utiliza o procedură de *antrenare supervizată* a clasificatorului care presupune existența unui număr de *eșantioane* (*prototipuri* – eng. *patterns*) pentru care sunt cunoscute clasele cărora le aparțin. Clasificatorul trebuie, în primul rând, învățat să clasifice în mod corect aceste eşantioane. Apoi, la prezentarea unei noi forme, clasificatorul o va atribui uneia dintre clasele învățate, în funcție de trăsăturile sale.

Dacă asemenea cunoștințe nu sunt disponibile, se aplică un procedeu de *antrenare nesupervizată* a clasificatorului, care, în plus față de antrenarea supervizată, permite și determinarea numărului de clase în care pot fi clasificate eşantioanele considerate.

Următoarele metode de clasificare sunt frecvent utilizate în aplicații practice:

- **Funcțiile de decizie** (eng. *decision functions*) pot fi utilizate dacă numărul de clase este cunoscut apriori și există o separare geometrică între clase. De exemplu, pentru separarea a două clase $\mathcal{C}_1$ și $\mathcal{C}_2$ în $\mathbb{R}^n$ considerăm o funcție $d : \mathbb{R}^n \rightarrow \mathbb{R}$ și clasificăm o formă $\mathbf{x} \in \mathbb{R}^n$ după semnul funcției d : dacă $d(\mathbf{x}) > 0$ atunci $\mathbf{x} \in \mathcal{C}_1$ și dacă $d(\mathbf{x}) < 0$ atunci $\mathbf{x} \in \mathcal{C}_2$. Suprafața din $\mathbb{R}^n$ a cărei ecuație este $d(\mathbf{x}) = 0$ reprezintă *frontiera de decizie* (eng. *decision boundary*).
- Clasificarea bazată pe **distanța minimă** (eng. *minimum distance*) poate fi aplicată în cazul în care fiecare clasă poate fi reprezentată printr-un singur prototip, denumit *centrul clusterului*. Clasificatorul bazat pe distanța minimă (pe cel mai apropiat vecin – eng. *nearest neighbor*) clasifică o nouă formă măsurând distanțele de la aceasta la centrele clusterelor considerate și atribuind forma respectivă celei mai apropiate clase. Această metodă a fost extinsă la cazul în care este cunoscut numai numărul de clase, dar nu se cunoaște modul de clasificare a eşantioanelor din setul de antrenare. Centrele clusterelor pot fi găsite iterativ prin minimizarea unui criteriu de performanță (de exemplu algoritmi C-Means sau ISODATA).
- În situația în care există posibilitatea de suprapunere a unor clase, problema clasificării poate fi abordată **statistic**. Clasificatorul determină probabilitatea de apartenență a unei forme la fiecare din clasele considerate. De asemenea un clasificator statistic trebuie să evalueze gradul de risc asociat clasificării (probabilitatea de clasificare eronată a unei forme).
- **Rețele neurale** reprezintă un instrument frecvent utilizat în proiectarea clasificatorilor. Există atât metode de antrenare supervizată, cât și metode de antrenare nesupervizată, care permit ajustarea valorilor ponderilor și deplasărilor neuronilor până se obține o clasificare corectă (este minimizat un indicator de performanță).

- **Logica fuzzy** poate fi utilizată dacă există incertitudinii asupra modului de clasificare a prototipurilor (formele din setul de antrenare) sau o anumită formă poate aparține într-o măsură oarecare mai multor clase. Clasificatorul fuzzy determină în final gradul de apartenență a unei forme la fiecare clasă considerată.

B.2. Clasificarea pe baza funcțiilor de decizie

B.2.1. Separarea în două clase

Pentru simplitate, considerăm situația în care se dorește separarea în două clase, notate $\mathcal{C}_1$ și $\mathcal{C}_2$, a unor forme caracterizate prin două trăsături. Presupunem că fiecare eșantion poate fi reprezentat printr-un vector bidimensional $\mathbf{x} = [x_1 \ x_2]^T \in \mathbb{R}^2$ și reprezentat grafic în planul $x_1 O x_2$ prin simbolul $\circ$ (dacă $x \in \mathcal{C}_1$) sau prin $\square$ (dacă $x \in \mathcal{C}_2$).

În cazul prezentat în figura B.2.1, cele două clase pot fi separate printr-o dreaptă, a cărei ecuație este:

$$d(\mathbf{x}) = w_1 x_1 + w_2 x_2 + b = 0. \quad (\text{B.2.1})$$

Semnele coeficienților $w_1, w_2, b \in \mathbb{R}$ pot fi alese astfel încât să fie satisfăcute condițiile:

$$\begin{cases} d(\mathbf{x}) > 0, & \forall \mathbf{x} \in \mathcal{C}_1, \\ d(\mathbf{x}) < 0, & \forall \mathbf{x} \in \mathcal{C}_2. \end{cases}$$

Pentru a clasifica o nouă formă $\tilde{\mathbf{x}} \in \mathbb{R}^2$ despre care se știe că aparține la una și numai una din clasele $\mathcal{C}_1$ și $\mathcal{C}_2$, se calculează valoarea funcției d în $\tilde{\mathbf{x}}$ și decizia asupra clasificării lui $\tilde{\mathbf{x}}$ se ia în funcție de semnul lui $d(\tilde{\mathbf{x}})$ astfel:

$$\begin{aligned} \tilde{\mathbf{x}} &\in \mathcal{C}_1, & \text{dacă } d(\tilde{\mathbf{x}}) > 0, \\ \tilde{\mathbf{x}} &\in \mathcal{C}_2, & \text{dacă } d(\tilde{\mathbf{x}}) < 0. \end{aligned} \quad (\text{B.2.2})$$

Funcția $d : \mathbb{R}^2 \rightarrow \mathbb{R}$, $d(\mathbf{x}) = w_1 x_1 + w_2 x_2 + b$, pe baza căreia se realizează clasificarea reprezintă o *funcție de decizie liniară*.

Acest concept poate fi extins la cazul general al clasificării unor forme caracterizate printr-un număr oarecare de trăsături. Presupunem că forme n -dimensionale reprezentate prin vectori $\mathbf{x} = [x_1 \ x_2 \ \dots \ x_n]^T \in \mathbb{R}^n$ pot fi separate în două clase $\mathcal{C}_1$ și $\mathcal{C}_2$, utilizând o *funcție de decizie liniară* de forma

$$d : \mathbb{R}^n \rightarrow \mathbb{R}, \quad d(\mathbf{x}) = w_1 x_1 + w_2 x_2 + \dots + w_n x_n + b = \mathbf{W}\mathbf{x} + b, \quad (\text{B.2.3})$$

unde $\mathbf{W} \triangleq [w_1 \ w_2 \ \dots \ w_n] \in \mathbb{R}^{1 \times n}$ și $b \in \mathbb{R}$, în conformitate cu (B.2.2). Geometric cele două clase din $\mathbb{R}^n$ pot fi separate prin hiperplanul de ecuație $d(\mathbf{x}) = \mathbf{W}\mathbf{x} + b = 0$.

Mai mult, funcția liniară (B.2.3) poate fi generalizată în cazul în care are forma unei combinații liniare de funcții

$$d : \mathbb{R}^n \rightarrow \mathbb{R}, \quad d(\mathbf{x}) = w_1 f_1(\mathbf{x}) + \dots + w_p f_p(\mathbf{x}) + b, \quad (\text{B.2.4})$$

cu $f_i: \mathbb{R}^n \rightarrow \mathbb{R}$, $i = \overline{1, p}$, în general funcții neliniare de n variabile, și $w_1, w_2, \dots, w_p, b \in \mathbb{R}$. Cu notațiile $W = [w_1 \dots w_p]^{1 \times p}$, $F: \mathbb{R}^n \rightarrow \mathbb{R}^p$, $F(x) = [f_1(x) \dots f_p(x)]^T$, funcția (B.2.4) poate fi adusă la forma $d(x) = WF(x) + b$.

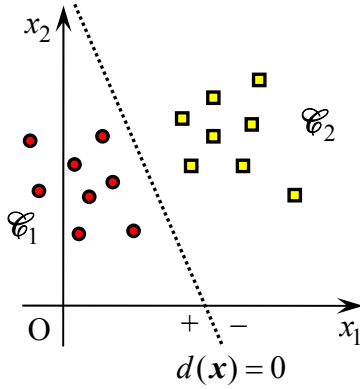


Figura B.2.1. Două clase liniar separabile.

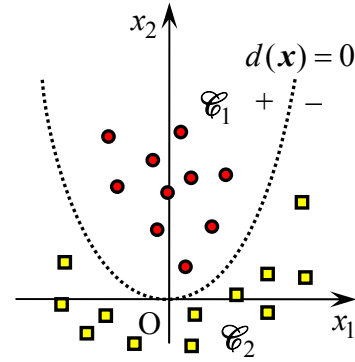


Figura B.2.2. Două clase neliniar separabile.

În cazul cel mai general funcția de decizie poate să nu fie liniară. Figura B.2.2 prezintă cazul a două clase din $\mathbb{R}^2$ ce pot fi separate utilizând funcția de decizie $d(x) = -x_1^2 + x_2$. Reprezentarea grafică a curbei de decizie $d(x) = 0$ este o parabolă. Clasificarea unei noi forme $\tilde{x} \in \mathbb{R}^2$ se face aplicând (B.2.2).

B.2.2. Separarea în m clase

În cazul cel mai general se dorește separarea unui număr m de clase, notate $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$, de forme din $\mathbb{R}^n$.

Dacă fiecărei clase de forme $\mathcal{C}_i$ îi poate fi atașată o funcție de decizie $d_i: \mathbb{R}^n \rightarrow \mathbb{R}$ care realizează separarea dintre clasa $\mathcal{C}_i$ și restul claselor $\mathcal{C}_j$, $j \neq i$, conform

$$\begin{cases} d_i(x) > 0, & \forall x \in \mathcal{C}_i, \\ d_i(x) < 0, & \forall x \in \mathcal{C}_j, j \neq i, \end{cases} \quad (\text{B.2.5})$$

atunci se spune că $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$ sunt *clase absolut separabile* (eng. *absolutely separable*).

Dacă $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$, sunt clase absolut separabile din $\mathbb{R}^n$ cu ajutorul funcțiilor de decizie $d_1, d_2, \dots, d_m$, mulțimea

$$\mathcal{D}_i = \{x \in \mathbb{R}^n \mid d_i(x) > 0, d_j(x) < 0, j \neq i\}, \quad i = 1, 2, \dots, m, \quad (\text{B.2.6})$$

se numește *regiune de decizie* corespunzătoare clasei $\mathcal{C}_i$. Utilizând criteriul (B.2.5), orice punct din $\mathcal{D}_i$ este clasificat ca aparținând clasei $\mathcal{C}_i$.

Figura B.2.3 prezintă cazul a trei clase $\mathcal{C}_1$, $\mathcal{C}_2$ și $\mathcal{C}_3$ absolut separabile în $\mathbb{R}^2$ prin funcții de decizie liniare. Regiunile de decizie corespunzătoare sunt prezentate în figura B.2.4. Se observă că regiunile de decizie ale celor trei clase nu acoperă întreg spațiul $\mathbb{R}^2$. O formă $\mathbf{x} \in \mathbb{R}^2 \setminus (\mathcal{D}_1 \cup \mathcal{D}_2 \cup \mathcal{D}_3)$ nu poate fi clasificată în nici una dintre clasele $\mathcal{C}_1$, $\mathcal{C}_2$, $\mathcal{C}_3$ utilizând funcțiile de decizie considerate.

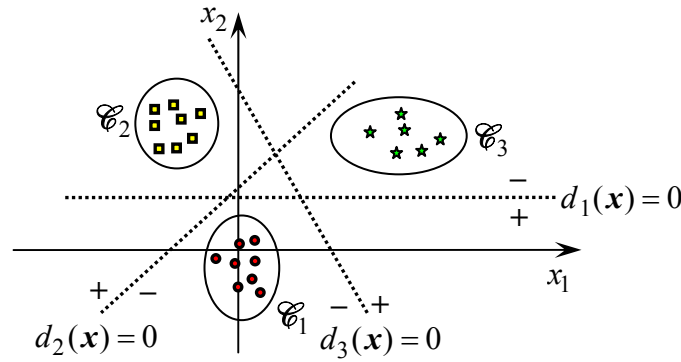


Figura B.2.3. Trei clase absolut separabile în $\mathbb{R}^2$ prin funcții de decizie liniare.

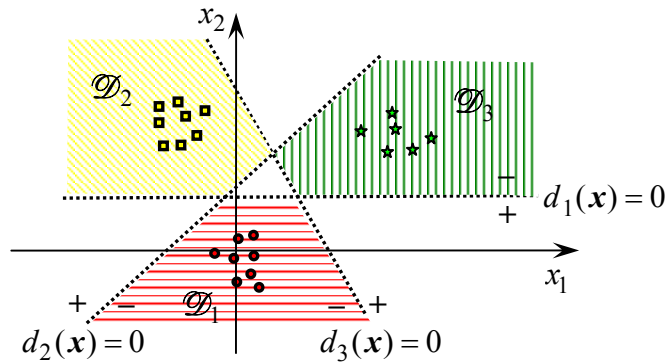


Figura B.2.4. Regiunile de decizie corespunzătoare claselor din figura B.2.3.

Fie clasele $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$ de forme din $\mathbb{R}^n$, care nu sunt absolut separabile. Dacă fiecărei perechi de clase $\mathcal{C}_i, \mathcal{C}_j, j \neq i$, îi poate fi asociată o funcție $d_{ij} : \mathbb{R}^n \rightarrow \mathbb{R}$ astfel încât

$$\begin{cases} d_{ij}(\mathbf{x}) > 0, & \forall \mathbf{x} \in \mathcal{C}_i, \\ d_{ij}(\mathbf{x}) < 0, & \forall \mathbf{x} \in \mathcal{C}_j, \end{cases} \quad (\text{B.2.7})$$

se spune că $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$ sunt clase *separabile două câte două* (eng. *pairwise separable*).

Dacă clasele $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$ sunt separabile două câte două, atunci pentru orice $j \neq i$ are loc $d_{ij}(\mathbf{x}) > 0, \forall \mathbf{x} \in \mathcal{C}_i$. De asemenea, funcțiile de decizie pot fi considerate astfel încât $d_{ji}(\mathbf{x}) = -d_{ij}(\mathbf{x}), \forall \mathbf{x} \in \mathbb{R}^n$.

Dacă $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$, sunt clase din $\mathbb{R}^n$ separabile două câte două cu ajutorul funcțiilor de decizie d_{ij} , mulțimea

$$\mathcal{D}_i = \left\{ \mathbf{x} \in \mathbb{R}^n \mid d_{ij}(\mathbf{x}) > 0, j \neq i \right\}, \quad (\text{B.2.8})$$

se numește *regiune de decizie* corespunzătoare clasei $\mathcal{C}_i$, $i = 1, 2, \dots, m$.

Figura B.2.5 prezintă cazul a trei clase $\mathcal{C}_1, \mathcal{C}_2, \mathcal{C}_3$ din $\mathbb{R}^2$ separabile două câte două prin funcții de decizie liniare. Așa cum se observă în figura B.2.6, regiunile de decizie ale celor trei clase nu acoperă întreg spațiul $\mathbb{R}^2$. O formă $\mathbf{x} \in \mathbb{R}^2 \setminus (\mathcal{D}_1 \cup \mathcal{D}_2 \cup \mathcal{D}_3)$ nu poate fi clasificată în nici una dintre clasele $\mathcal{C}_1, \mathcal{C}_2, \mathcal{C}_3$ utilizând funcțiile de decizie considerate.

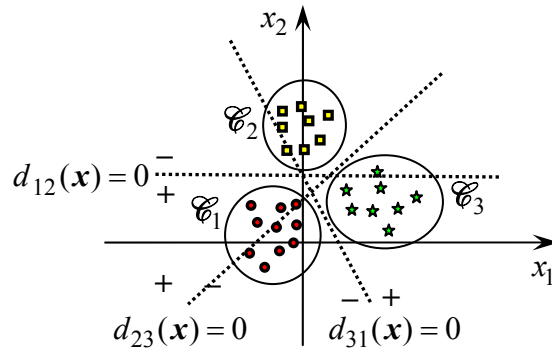


Figura B.2.5. Trei clase separabile două câte două prin funcții de decizie liniare.

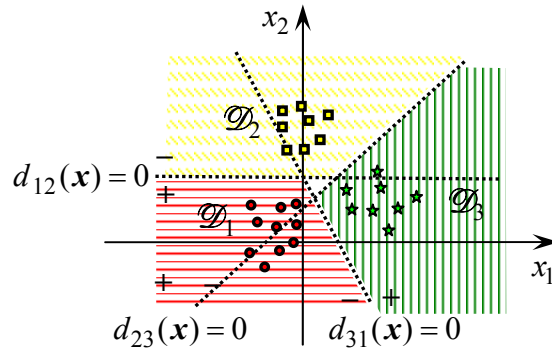


Figura B.2.6. Regiunile de decizie corespunzătoare claselor din fig. B.2.5.

B.3. Clasificarea pe baza distanței

B.3.1. Clase reprezentate printr-un singur prototip

Considerăm un număr m de clase, notate $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$, de forme din $\mathbb{R}^n$. Presupunem mai întâi că fiecare clasă $\mathcal{C}_i$ este reprezentată printr-un unic prototip $\mathbf{y}_i = [y_{i1} \ y_{i2} \ \dots \ y_{in}]^T \in \mathbb{R}^n$, $i = \overline{1, m}$. Distanța euclidiană de la o formă $\mathbf{x} \in \mathbb{R}^n$ la vectorul prototip $\mathbf{y}_i$ este

$$D_i = \|\mathbf{x} - \mathbf{y}_i\|_2 = \left[(\mathbf{x} - \mathbf{y}_i)^T \cdot (\mathbf{x} - \mathbf{y}_i) \right]^{1/2}, \quad i = \overline{1, m}. \quad (\text{B.3.1})$$

Un clasificator bazat pe distanța euclidiană atribuie forma $\mathbf{x} \in \mathbb{R}^n$ clasei $\mathcal{C}_j$ (sau echivalent prototipul $\mathbf{y}_j$) situată la distanță minimă de $\mathbf{x}$, adică

$$\mathbf{x} \in \mathcal{C}_j \Leftrightarrow D_j = \min_{1 \leq i \leq m} D_i = \min_{1 \leq i \leq m} \|\mathbf{x} - \mathbf{y}_i\|_2. \quad (\text{B.3.2})$$

Dacă valoarea minimă este atinsă pentru mai mulți indici j simultan, atunci forma $\mathbf{x}$ este fie clasificată în clasa corespunzătoare celui mai mic indice, fie nu este clasificată deloc.

O modalitate de clasificare echivalentă cu (B.3.2) este de a utiliza pătratele distanțelor D_i^2 în loc de D_i . Astfel o formă $\mathbf{x} \in \mathbb{R}^n$ este atribuită clasei $\mathcal{C}_j$ în modul următor :

$$\mathbf{x} \in \mathcal{C}_j \Leftrightarrow D_j^2 = \min_{1 \leq i \leq m} D_i^2 = \min_{1 \leq i \leq m} \|\mathbf{x} - \mathbf{y}_i\|_2^2. \quad (\text{B.3.3})$$

Datorită faptului că

$$D_i^2 = \|\mathbf{x} - \mathbf{y}_i\|_2^2 = \left[(\mathbf{x} - \mathbf{y}_i)^T \cdot (\mathbf{x} - \mathbf{y}_i) \right] = \mathbf{x}^T \cdot \mathbf{x} - 2\mathbf{y}_i^T \cdot \mathbf{x} + \mathbf{y}_i^T \cdot \mathbf{y}_i,$$

expresie în care $\mathbf{x}^T \cdot \mathbf{x}$ are valoare constantă pentru o formă $\mathbf{x}$ dată, relația (B.3.3) devine

$$\mathbf{x} \in \mathcal{C}_j \Leftrightarrow j = \arg \min_{1 \leq i \leq m} \left(-2\mathbf{y}_i^T \cdot \mathbf{x} + \mathbf{y}_i^T \cdot \mathbf{y}_i \right) = \arg \max_{1 \leq i \leq m} \left(\mathbf{y}_i^T \cdot \mathbf{x} - \frac{1}{2} \mathbf{y}_i^T \cdot \mathbf{y}_i \right).$$

Astfel se poate utiliza funcția

$$d_i : \mathbb{R}^n \rightarrow \mathbb{R}, \quad d_i(\mathbf{x}) = \mathbf{y}_i^T \cdot \mathbf{x} - \frac{1}{2} \mathbf{y}_i^T \cdot \mathbf{y}_i, \quad i = \overline{1, m}, \quad (\text{B.3.4})$$

drept funcție de decizie, clasificarea făcându-se astfel:

$$\mathbf{x} \in \mathcal{C}_j \Leftrightarrow d_j(\mathbf{x}) > d_i(\mathbf{x}), \quad j \neq i. \quad (\text{B.3.5})$$

Se observă că funcțiile de decizie (B.3.4) sunt liniare, putând fi puse sub forma $d_i(\mathbf{x}) = \mathbf{W}_i \mathbf{x} + b_i$ cu $\mathbf{W}_i = [w_{i1} \ w_{i2} \ \dots \ w_{in}] \in \mathbb{R}^{1 \times n}$, $w_{ik} = y_{ik}$, $k = \overline{1, n}$, și $b_i = -\frac{1}{2} \mathbf{y}_i^T \cdot \mathbf{y}_i$ pentru $i = \overline{1, m}$.

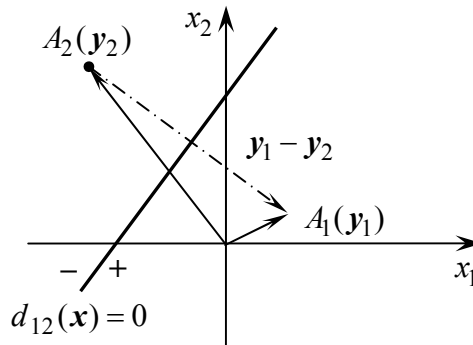


Figura B.3.1. Frontiera de separare a două clase din $\mathbb{R}^2$ pe baza distanței minime.

În particular, dacă se separă două clase $\mathcal{C}_1, \mathcal{C}_2$ în $\mathbb{R}^n$ cărora le corespund prototipurile y_1, y_2 , utilizând funcțiile $d_1(x)$ și $d_2(x)$ definite conform (B.3.4) se observă că frontiera de decizie are ecuația:

$$d_{12}(x) = d_1(x) - d_2(x) = (y_1 - y_2)^T \cdot x - \frac{1}{2}(y_1^T \cdot y_1 + y_2^T \cdot y_2) = 0,$$

fiind un hiperplan normal pe vectorul $y_1 - y_2$.

Acest raționament poate fi extins la cazul a m clase $\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_m$, din $\mathbb{R}^n$ ce pot fi separate două câte două cu ajutorul funcțiilor de decizie d_{ij} ,

$$d_{ij}(x) = d_i(x) - d_j(x) = (y_i - y_j)^T \cdot x - \frac{1}{2}(y_i^T \cdot y_i + y_j^T \cdot y_j), \quad i \neq j.$$

Figura B.3.1 prezintă cazul separării a două clase din $\mathbb{R}^2$ reprezentate prin prototipurile y_1, y_2 , cărora le corespund punctele A_1 și A_2 din plan. Frontiera de decizie, dreapta de ecuație $d_{12}(x) = 0$, reprezintă mediatoarea segmentului A_1A_2 . În figura B.3.2 este ilustrat cazul separării a trei clase din $\mathbb{R}^2$ reprezentate prin prototipurile y_1, y_2, y_3 , cărora le corespund punctele A_1, A_2 și A_3 din plan. Se observă că regiunile de decizie definite conform relației (B.2.8) pentru cele trei clase considerate acoperă întreg spațiul $\mathbb{R}^2$ (cu excepția punctelor situate pe semidreptele de separație a acestor regiuni).

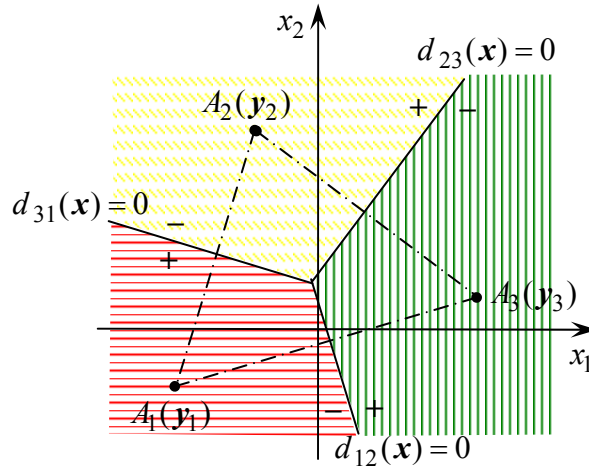


Figura B.3.2. Regiunile de decizie corespunzătoare separării a trei clase din $\mathbb{R}^2$ pe baza distanței euclidiene minime.

B.3.2. Clase reprezentate prin mai multe prototipuri

Sunt situații în care o clasă poate fi reprezentată prin mai multe prototipuri. Presupunem că un număr p_i de prototipuri notate $y_i^{(1)}, y_i^{(2)} \dots y_i^{(p_i)}$ reprezintă clasa $\mathcal{C}_i$, $i = \overline{1, m}$. În acest caz distanța de la o formă $x \in \mathbb{R}^n$ la clasa $\mathcal{C}_i$ este definită prin

$$D_i = \min_{1 \leq q \leq p_i} \|\mathbf{x} - \mathbf{y}_i^{(q)}\|, \quad i = \overline{1, m}. \quad (\text{B.3.6})$$

În mod similar cazului în care fiecare clasă este reprezentată printr-un singur prototip, un clasificator bazat pe distanța (B.3.6) atribuie forma $\mathbf{x} \in \mathbb{R}^n$ clasei $\mathcal{C}_j$ situată la distanță minimă de $\mathbf{x}$, adică

$$\mathbf{x} \in \mathcal{C}_j \Leftrightarrow D_j = \min_{1 \leq i \leq m} D_i. \quad (\text{B.3.7})$$

Dacă valoarea minimă este atinsă pentru mai mulți indici j simultan, atunci forma $\mathbf{x}$ este fie clasificată în clasa corespunzătoare celui mai mic indice, fie nu este clasificată deloc.

O formă echivalentă a acestui mod de clasificare a formelor din $\mathbb{R}^n$ este de a determina distanțele de la $\mathbf{x}$ la toți vectorii prototip $\mathbf{y}_1^{(1)}, \dots, \mathbf{y}_1^{(p_1)}, \mathbf{y}_2^{(1)}, \dots, \mathbf{y}_m^{(p_m)}$, și de a atribui $\mathbf{x}$ clasei căreia îi aparține prototipul situat la distanța minimă de $\mathbf{x}$. Acest algoritm poartă numele de *clasificare pe baza celui mai apropiat vecin* (eng. *nearest-neighbor classification*).

În afară de norma euclidiană (norma 2), în probleme de clasificare mai sunt frecvent utilizate și normele Hölder 1 (denumită și distanța Manhattan) și ∞ (sau distanța Chebychev) definite pentru doi vectori oarecare $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, $\mathbf{x} = [x_1 \dots x_n]^T$, $\mathbf{y} = [y_1 \dots y_n]^T$, prin:

$$d_1(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_1 = \sum_{i=1}^n |x_i - y_i|, \quad (\text{B.3.8})$$

respectiv,

$$d_\infty(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_\infty = \max_{i=\overline{1, n}} (|x_i - y_i|). \quad (\text{B.3.9})$$

De asemenea, este frecvent utilizată distanța Mahalanobis, definită prin

$$d_A(\mathbf{x}, \mathbf{y}) = \left[(\mathbf{x} - \mathbf{y})^T \mathbf{A} (\mathbf{x} - \mathbf{y}) \right]^{\frac{1}{2}}, \quad (\text{B.3.10})$$

unde matricea $\mathbf{A} \in \mathbb{R}^{n \times n}$ este simetrică și pozitiv definită.

Bibliografie

- Agarwal, M., A systematic classification of neural-network-based control, *IEEE Control Systems Magazine*, vol. 17, no. 2, pp. 75-93, 1997.
- Chong, E.K.P. and Żak, S.H., *An Introduction to Optimization*, Second Edition, John Wiley & Sons, Inc., New York, 2001.
- Cybenko, G., 1989. "Approximation by superpositions of a sigmoidal function", *Mathematics of Control, Signals and Systems*, vol 2, pp. 303 – 314.
- Demuth, H., Beale, M. and Hagan, M., *Neural Networks Toolbox 5 – User's Guide*, The MathWorks Inc., 2007.
- Duch, W. and Jankowski, N., Survey of Neural Transfer Functions, *Neural Computing Surveys*, vol. 2, pp. 163-212, 1999.
- Dumitrache, I., Constantin, N., Drăgoicea, M., *Rețele neurale. Identificarea și Conducerea Proceselor*, MatrixRom, 1999.
- Eberhart, R.C. and Yuhui Shi, *Computational Intelligence: Concepts to Implementations*, Morgan Kaufmann; 2007.
- Graupe, D., *Principles of Artificial Neural Networks*, 2nd Edition, World Scientific Publishing Co. Pte. Ltd., 2007.
- Gupta, M.M., Liang Jin and Homma, N., *Static and Dynamic Neural Networks: From Fundamentals to Advanced Theory*, John Wiley & Sons, Inc., Hoboken, NJ, USA, 2003
- Hagan, M.T., Demuth, H.B. and Beale, M.H., *Neural Network Design*, PWS Pub. Co., 1996.
- Haykin, S., *Neural Networks: A Comprehensive Foundation, 2nd Ed.*, Prentice-Hall, Pearsons Education, India, 2005.
- Hu, Y.H. and Hwang, J.N. (Eds.), *Handbook of Neural Network Signal Processing*, CRC Press LLC, Boca Raton, FL, USA, 2002.
- Hunt, K.J., Sbarbaro, D., Zbikowski, R. and Gawthrop, P.J., Neural Networks for Control System - A Survey, *Automatica*, vol. 28, pp. 1083-1112, 1992.
- Jang, J.S.R., Sun, C.T. and Mizutani, E., *Neuro-Fuzzy and Soft Computing: A Computational Approach to Learning and Machine Intelligence*, Prentice-Hall Inc., Upper Saddle River, NJ, USA, 1997.

- Kecman, V., *Learning And Soft Computing. Support Vector Machines, Neural Networks and Fuzzy Logic Models*, The MIT Press, Cambridge, MA, 2000.
- Kloetzer, M., Ardelean, D., Păstrăvanu, O., Developing Simulink tools for teaching neural-network-based identification, *9th IEEE Mediterranean Conference on Control and Automation, MED 01*, June 27-29, 2001, Dubrovnik, Croatia.
- Kloetzer, M., Păstrăvanu, O., Simulink blockset for neural predictive control, *Buletinul Științific al Universității Politehnica din Timișoara, Periodica Politechnica, Transactions On Automatic Control And Computer Science*, Vol. 49 (63), No.2, pp. 19-22, 2004.
- Lazăr, C., *Conducerea predictivă a proceselor cu model cunoscut*, Ed. Matrix Rom, București, 1999.
- Levin, A.U. and Narendra, K.S, Control of nonlinear dynamical systems using neural networks. I. Controllability and stabilization, *IEEE Trans. Neural Networks*, vol. 4, iss. 3, pp. 192-206 1993.
- Levin, A.U. and Narendra, K.S, Control of nonlinear dynamical systems using neural networks. II. Observability, identification, and control, *IEEE Trans. Neural Networks*, vol. 7, iss. 1, pp. 30-42, 1996.
- Lewis, F.L., Jagannathan, S. and Yeildirek, A., *Neural Network Control of Robot Manipulators and Nonlinear Systems*, Taylor & Francis, New York, 1999.
- Ljung, L., Sjöberg, J., Hjalmarson, H., "On neural network model structures in system identification" in Bittanti, S., Picci, G. (Eds.) *Identification, Adaptation, Learning*, NATO ASI Series, Springer, 1996.
- Liu, G.P., Kadiramanathan, V. and Billings, S.A., Predictive Control for Non-Linear Systems Using Neural Networks. *International Journal of Control*, nr. 71, 1998, pp. 1119-1132.
- Masoud, M., Sarker, R.A. and Yao, X., *Computational Intelligence in Control*, Idea Group Publishing, London, UK, 2003.
- Narendra, K.S. and Lewis, F.L. (Eds.), Special issue on neural network feedback control, *Automatica*, vol. 37, no. 8, 2001.
- K.S. Narendra, K. Parthasarathy, "Identification and control of dynamical systems using neural networks", *IEEE Trans. on Neural Networks*, vol. 1, issue 1, pp. 4-27, 1990.
- Nguyen, H.T., Prasad, N.R., Walker, C.L. and Walker, E.A., *A First Course in Fuzzy and Neural Control*, Chapman & Hall/CRC Press, Boca Raton, FL, 2003.
- Rosenblatt, F., *Principles of Neurodynamics*, Spartan Press, Washington D.C., 1961.
- Sarangapani, J., *Neural Network Control of Nonlinear Discrete-Time Systems*, CRC Press, Taylor & Francis Group, Boca Raton, FL, USA, 2006.
- Sima, V. și Varga, A., *Practica optimizării asistate de calculator*, Editura Tehnică. București, 1986.
- Sun, W. and Ya-Xiang Yuan, *Optimization Theory and Methods: Nonlinear Programming*, Springer, USA, 2006.
- *** *Using Simulink*, The MathWorks Inc., 2007.
- Widrow, B. and Lehr, M.A., 30 Years of adaptive neural networks: Perceptron, Madaline, and Backpropagation, *Proceedings of the IEEE*, vol. 78 no.9, September 1990.