

Antrenarea rețelelor neuronale

1. Modelul neuronului artificial
2. Rețele neuronale tip perceptron
3. Rețele neuronale MLP
 - 3.1. Metode de accelerare a învățării
 - 3.2. Considerente practice
4. Matlab Neural Network Toolbox
5. Aplicații

1. Modelul neuronului artificial

Rețelele neuronale 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 neuronale 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 neuronale, introdusă prin terminologie, se regăsește și în descrierea matematică a transferului intrare-ieșire. Rețelele neuronale 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 neuronale (feedforward sau recurente) îl constituie *neuronul artificial*. Pentru construirea unei rețele neuronale, neuronii artificiali se conectează în unul sau mai multe straturi, definind astfel, *arhitectura rețelei*.

În figura de mai jos este prezentată schema bloc a unui neuron cu mai multe intrări $\mathbf{x} = [x_1, \dots, x_n]^T \in \mathbb{R}^n$ și ieșirea scalară y . Constantele $\mathbf{w} = [w_1, \dots, w_n] \in \mathbb{R}^{1 \times n}$ se numesc *ponderi* (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 $\mathbf{x}$, valoarea cu care este alimentat blocul b este tot timpul 1. Referirea comună a constantelor $\mathbf{w}, b$ se face prin termenul de „*parametrii neuronului*”. În utilizarea neuronului, valorile parametrilor sunt *ajustabile*. Pe ajustarea adecvată a acestor valori se bazează capacitatea de învățare a neuronului artificial.

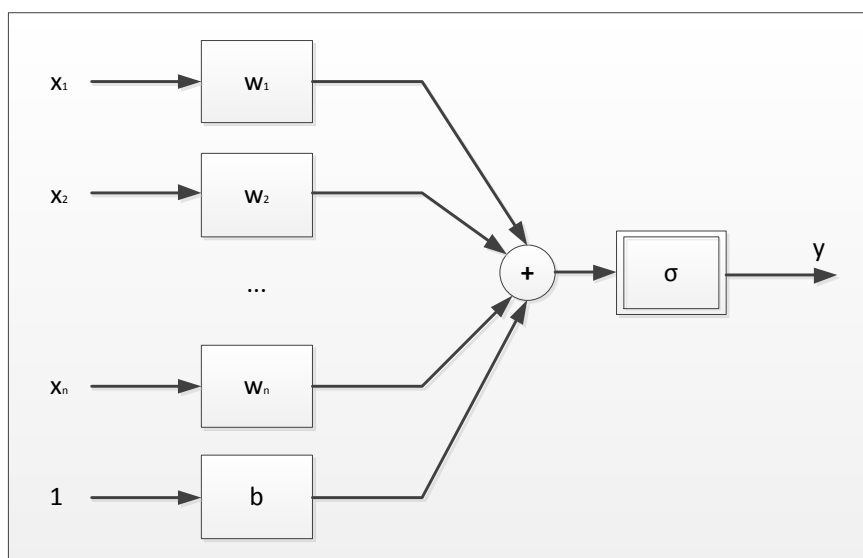
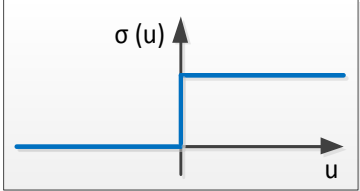
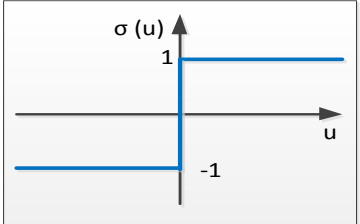
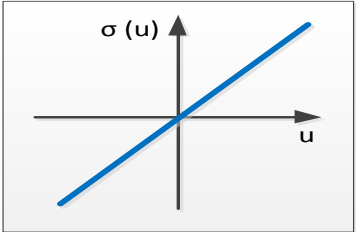
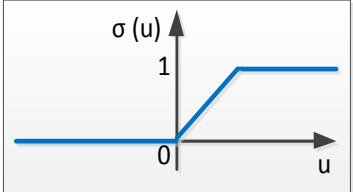
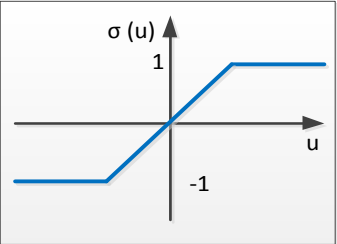
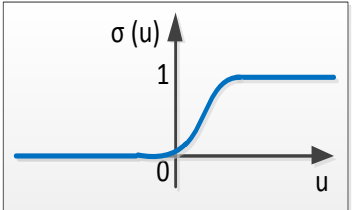
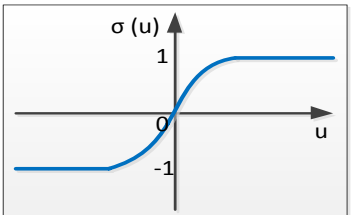


Figura 1. Neuronul artificial cu mai multe intrări

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

Neuronul artificial cu mai multe intrări 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:

$$f: \mathbb{R} \rightarrow \mathbb{R}$$

$$y = f(x) = \sigma(u(x)) = \sigma(r(x) + b) = \sigma(\mathbf{w}x + b) = \sigma\left(\sum_{i=1}^n w_i x_i + b\right)$$

Pentru cazul particular $n = 1$, se obține un neuron cu o singură intrare.

2. Rețele neuronale de tip perceptron

Rețelele neuronale 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.

Se consideră un perceptron cu n intrări, cu funcție de activare de tip treaptă unipolară. Se observă că spațiul de intrare, poate fi separat în două clase în funcție de valorile ieșirii perceptronului:

$$\mathcal{C}_0 = \{\mathbf{x} \in \mathbb{R}^n \mid y(\mathbf{x}) = 0\}$$

$$\mathcal{C}_1 = \{\mathbf{x} \in \mathbb{R}^n \mid y(\mathbf{x}) = 1\}$$

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 cazurile 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.

O rețea neuronală cu un singur strat cu n intrări, p ieșiri (neuroni) și funcții de activare treaptă (unipolară sau bipolară) posedă proprietatea de a partiționa spațiul vectorilor de intrare $\mathbb{R}^n$ 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.

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

I. Etapa de inițializare

Se organizează vectorii prototip $\mathbf{x}_i \in \mathbb{R}^n$, $i = \overline{1, N}$ și cei țintă $\mathbf{z}_i \in \mathbb{R}^p$ sub forma matricelor: $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_N] \in \mathbb{R}^{n \times N}$, respective $\mathbf{Z} = [\mathbf{z}_1, \dots, \mathbf{z}_N] \in \mathbb{R}^{p \times N}$. Se inițializează ponderile și deplasările cu valori numerice aleatoare din intervalul $[-0.5; 0.5]$ (sau cu valori nule). 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 $\mathbf{W}_{old}$ pentru ponderi) și $\mathbf{b}_{old}$ (pentru deplasări).

II. Etapa de prezentare

Pentru valorile curente ale parametrilor rețelei se calculează ieșirile rețelei:

$$y_i = f(\mathbf{x}_i) = \sigma(\mathbf{W}_{old}\mathbf{x}_i + \mathbf{b}_{old}), i = \overline{1, N}$$

organizate sub forma matricei $\mathbf{Y} = [\mathbf{y}_1, \dots, \mathbf{y}_N] \in \mathbb{R}^{p \times N}$.

III. Etapa de verificare

Se calculează vectorii erorilor ca diferențe între vectorii țintă și vectorii ieșire la iterația curentă:

$$\mathbf{e}_i = \mathbf{z}_i - \mathbf{y}_i \in \mathbb{R}^p, i = \overline{1, N}$$

Algoritmul se oprește dacă este îndeplinită una din următoarele condiții:

- i) toți vectorii eroare sunt nuli $\mathbf{e}_i \equiv \mathbf{0}, i = \overline{1, N}$;
- ii) a fost atins numărul maxim de iterații.

La final, algoritmul va furniza valorile parametrilor rețelei rezultate din antrenare, care vor fi notate prin $\mathbf{W}_f$ (ponderi) și $\mathbf{b}_f$ (deplasări). Dacă nici una din condițiile de stop nu este îndeplinită, se continuă cu etapa următoare a iterației curente.

IV. Etapa de actualizare a parametrilor

Se construiește matricea erorilor:

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

cu ajutorul căreia se calculează matricele de actualizare a parametrilor:

- i) pentru ponderi: $\mathbf{D}_w = \mathbf{E} \cdot \mathbf{X}^T \in \mathbb{R}^{p \times n}$
- ii) pentru deplasări: $\mathbf{D}_b = \mathbf{E} \cdot [\mathbf{1} \dots \mathbf{1}]^T \in \mathbb{R}^p$.

Se calculează noile valori ale parametrilor:

- i) pentru ponderi: $\mathbf{W}_{new} = \mathbf{W}_{old} + \mathbf{D}_w$
- ii) pentru deplasări: $\mathbf{b}_{new} = \mathbf{b}_{old} + \mathbf{D}_b$.

Cu aceste calcule se încheie iterația curentă și se trece la o nouă iterație (Etapa II) efectuând mai întâi atribuirile: $\mathbf{W}_{old} \leftarrow \mathbf{W}_{new}$ și $\mathbf{b}_{old} \leftarrow \mathbf{b}_{new}$.

Exemplu

Ilustrăm aplicarea algoritmului de antrenare prezentat anterior în cazul determinării parametrilor unui perceptron ($p = 1$) cu două intrări ($n = 2$) ce implementează funcția logică SAU.

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 $\mathbf{Z} = [0 \ 1 \ 1 \ 1]$. Pentru simplitatea calculelor, inițializăm parametrii perceptronului cu valori nule: $\mathbf{W}_{old} = [0 \ 0]$, $\mathbf{b}_{old} = 0$

Iterația 1:

$$Y = [1 \ 1 \ 1 \ 1]$$

$$E = [-1 \ 0 \ 0 \ 0]$$

$$D_W = [0 \ 0], D_b = -1$$

$$W_{new} = [0 \ 0], b_{new} = -1.$$

Iterația 2:

$$W_{old} = [0 \ 0], b_{old} = -1$$

$$Y = [0 \ 0 \ 0 \ 0]$$

$$E = [0 \ 1 \ 1 \ 1]$$

$$D_W = [2 \ 2], D_b = 3$$

$$W_{new} = [2 \ 2], b_{new} = 2.$$

Iterația 3:

$$W_{old} = [2 \ 2], b_{old} = 2$$

$$Y = [1 \ 1 \ 1 \ 1]$$

$$E = [-1 \ 0 \ 0 \ 0]$$

$$D_W = [0 \ 0], D_b = -1$$

$$W_{new} = [2 \ 2], b_{new} = 1.$$

Iterația 4:

$$W_{old} = [2 \ 2], b_{old} = 1$$

$$Y = [1 \ 1 \ 1 \ 1]$$

$$E = [-1 \ 0 \ 0 \ 0]$$

$$D_W = [0 \ 0], D_b = -1$$

$$W_{new} = [2 \ 2], b_{new} = 0.$$

Iterația 5:

$$W_{old} = [2 \ 2], b_{old} = 0$$

$$Y = [1 \ 1 \ 1 \ 1]$$

$$E = [-1 \ 0 \ 0 \ 0]$$

$$D_W = [0 \ 0], D_b = -1$$

$$W_{new} = [2 \ 2], b_{new} = -1.$$

Iterația 6:

$$W_{old} = [2 \ 2], b_{old} = -1$$

$$Y = [0 \ 1 \ 1 \ 1]$$

$$E = [0 \ 0 \ 0 \ 0]$$

STOP.

Algoritmul se încheie furnizând valorile parametrilor rețelei rezultate din antrenare, $W_f = [2 \ 2], b_f = -1$. Ecuația dreptei de separație este $d(x) = 2x_1 + 2x_2 - 1 = 0$.

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. Rețele neuronale MLP

Perceptronul multistrat (MLP) este o rețea neuronală feedforward ce conține unul sau mai multe straturi ascunse. De obicei, rețeaua este formată dintr-un strat de intrare conținând neuronii sursă, cel puțin un strat ascuns conținând neuroni computaționali și un strat de ieșire conținând neuroni computaționali. Semnalele de intrare se propagă de la intrare spre ieșirea rețelei strat după strat.

Ca orice rețea neuronală, o rețea cu propagare înapoi este caracterizată de conexiunile dintre neuroni (care formează arhitectura rețelei), de funcțiile de activare folosite de neuroni și de algoritmul de învățare ce specifică procedura folosită pentru ajustarea ponderilor. De obicei o rețea neuronală cu propagare înapoi este o rețea multistrat ce conține trei sau patru straturi complet conectate. Fiecare neuron își calculează ieșirea similar cu perceptronul. Apoi valoarea de intrare este transmisă funcției de activare. Spre deosebire de perceptron, neuronii din rețelele cu propagare înapoi au funcții de activare de tip sigmoid. Derivata acesteia este foarte ușor de calculat și garantează ieșirea în intervalul $[0, 1]$.

Fiecare strat dintr-o rețea neuronală MLP îndeplinește o funcție specifică. Stratul de intrare acceptă semnale de intrare și rareori conține neuroni computaționali de aceea nu procesează tiparele de intrare. Stratul de ieșire acceptă semnale de ieșire (stimuli proveniți din stratul ascuns) și stabilește ieșirea rețelei în funcție de acesta. Neuronii din stratul ascuns detectează trăsăturile, iar

ponderile acestora reprezintă trăsăturile ascunse de tiparele de intrare. Aceste trăsături sunt folosite apoi de stratul de ieșire pentru a determina tiparul de ieșire.

Un strat ascuns ascunde ieșirea dorită. Neuronii din stratul ascuns nu pot fi observați prin intermediul comportamentului intrare-ieșire al rețelei. Nu există un mod evident de a cunoaște care trebuie să fie ieșirea dorită a stratului ascuns. Ieșirea dorită va fi determinată de însuși stratul ascuns.

Folosind un singur strat ascuns putem reprezenta orice funcție continuă a semnalelor de intrare, iar cu două straturi ascunse chiar și funcții cu discontinuități.

Rețelele neuronale au capacitatea de a învăța, însă modalitatea concretă în care se realizează acest proces este dată de algoritmul folosit pentru antrenare. O rețea se consideră antrenată dacă aplicarea unui vector de intrare conduce la obținerea unei ieșiri dorite, sau foarte apropiate de aceasta. Antrenarea constă în aplicarea secvențială a diferiți vectori de intrare și ajustarea ponderilor din rețea în raport cu o procedură predeterminată. În acest timp, ponderile conexiunilor converg gradual spre anumite valori pentru care fiecare vector de intrare produce vectorul de ieșire dorit.

Învățarea supervizată presupune aplicarea unei perechi vector de intrare – vector de ieșire dorit. După aplicarea unei intrări, se compară ieșirea calculată cu ieșirea dorită, după care diferența este folosită pentru modificarea ponderilor cu scopul minimizării erorii la un nivel acceptabil.

Într-o rețea neuronală cu propagare înapoi, algoritmul de învățare cuprinde două etape: faza de prezentare a tiparelor de antrenare pentru stratul de intrare. Rețeaua propagă strat după strat tiparele de antrenare până la generarea tiparului de ieșire. Dacă acesta este diferit de tiparul țintă dorit, se va calcula eroarea și va fi propagată înapoi, de la ieșire spre intrare. Ponderile sunt actualizate simultan cu propagarea erorii.

Algoritmul back-propagation este cel mai cunoscut și utilizat algoritm de învățare supervizată. Numit și *algoritmul delta generalizat*, el se bazează pe minimizarea diferenței dintre ieșirea dorită și ieșirea reală, prin metoda gradientului descendent. Gradientul ne spune cum variază funcția în diferite direcții. Ideea algoritmului este găsirea minimului funcției de eroare în raport cu ponderile conexiunilor. Eroarea este dată bineînțeles de diferența dintre ieșirea dorită și ieșirea efectivă a rețelei. Cea mai utilizată funcție de eroare este eroarea pătratică medie. Dacă avem în setul de antrenare K modele iar rețeaua are N ieșiri, atunci:

$$E = \frac{1}{K \cdot N} \cdot \sum_k \sum_n (y_{kn} - o_{kn})^2,$$

unde y_{kn} este valoarea dorită pe ieșirea n a rețelei pentru modelul k , iar o_{kn} este ieșirea efectivă a rețelei.

O epocă reprezintă prezentarea tuturor exemplurilor din setul de antrenare. În majoritatea cazurilor, antrenarea rețelei presupune mai multe epoci de antrenare.

Pentru a păstra rigoarea matematică, ponderile vor fi ajustate numai după ce *toți* vectorii de test vor fi aplicați rețelei. În acest scop, gradientii ponderilor trebuie memorați și ajustați după fiecare model din setul de antrenare, iar la sfârșitul unei epoci de antrenare, se vor modifica ponderile *o singură dată* (există și varianta „on-line”, mai simplă, în care ponderile sunt actualizate direct; în acest caz, poate conta ordinea în care sunt prezentați rețelei vectorii de test).

Dacă s-a încheiat o epocă, se testează dacă s-a îndeplinit criteriul de terminare ($E < E_{\max}$ sau s-a atins un număr maxim de epoci de antrenare). Dacă nu, se trece la o nouă epocă. Dacă da, algoritmul se termină.

3.1. Metode de accelerare a învățării

Antrenarea unui perceptron multistrat este deseori destul de lentă, necesitând mii sau zeci de mii de epoci pentru probleme complexe. Cele mai cunoscute metode de accelerare a învățării sunt: metoda momentului și aplicarea unei rate de învățare variabile.

Metoda momentului (engl. „momentum”) propune adăugarea unui termen la ajustarea ponderilor. Acest termen este proporțional cu ultima modificare a ponderii, adică valorile cu care se ajustează ponderile sunt memorate și influențează în mod direct toate ajustările ulterioare.

Metoda ratei de învățare variabile este o tehnică simplă de accelerare a învățării și constă în utilizarea unei rate de învățare individuale pentru fiecare pondere și adaptarea acestor parametri în fiecare iterație, în funcție de semnele succesive ale gradientilor. Dacă pe parcursul antrenării eroarea începe să crească, în loc să scadă, ratele de învățare se resetează la valorile inițiale și apoi se continuă procesul.

Notă: mai multe detalii despre algoritmul back-propagation pot fi găsite în cursul 11 și laboratorul opțional 17.

3.2. Considerente practice

În implementarea algoritmului apar o serie de probleme practice, legate în special de alegerea parametrilor antrenării și a configurației rețelei.

În primul rând, o rată de învățare mică determină o convergență lentă a algoritmului, pe când una prea mare poate cauza eșecul (algoritmul va „sări” peste soluție). Pentru probleme relativ simple o rată de învățare $\eta = 0.7$ este acceptabilă, însă în general se recomandă rate de învățare în jur de 0.2. Pentru accelerarea prin metoda momentului, o valoare satisfăcătoare pentru α este 0.9. În cazul ratei de învățare variabile, valori tipice care funcționează bine în cele mai multe situații sunt: $u = 1.2$ și $d = 0.8$.

Alegerea funcției de activare pentru stratul de ieșire al rețelei depinde de natura problemei care trebuie rezolvată.

Pentru neuronii din straturile ascunse, se preferă funcții sigmoide, deoarece au avantajul atât al neliniarității cât și al diferențiabilității (condiție necesară pentru aplicarea algoritmului backpropagation). Influența cea mai mare a unei sigmoide asupra performanțelor algoritmului se pare că o are simetria față de origine. Sigmoida bipolară este simetrică față de origine, în timp ce sigmoida unipolară este simetrică față de punctul (0, 0.5), ceea ce scade viteza de convergență.

Pentru neuronii de ieșire se recomandă funcții de activare adaptate distribuției datelor de ieșire. Astfel, pentru probleme de categorizare binară (0/1), este potrivită sigmoida. Pentru clasificarea cu n clase, fiecare corespunde unei ieșiri binare a rețelei (de exemplu, într-o aplicație de recunoaștere optică a caracterelor).

Pentru valori continue se poate realiza o preprocesare și o postprocesare a datelor, astfel încât rețeaua să opereze cu valori scalate, de exemplu în intervalul $[-0.9, 0.9]$ pentru sigmoida bipolară (tangenta hiperbolică).

De asemenea, pentru valori continue funcția de activare a neuronilor de ieșire poate fi liniară, mai ales dacă nu există limite cunoscute pentru intervalul în care se pot găsi acestea.

O altă problemă caracteristică acestui mod de antrenare este dată de minimele locale. Într-un minim local, gradientii erorii devin 0 și învățarea nu mai continuă. O soluție este încercarea mai multor antrenări independente, cu ponderi inițializate diferit la început, ceea ce crește probabilitatea găsirii minimului global. Pentru probleme mari, acest lucru poate fi greu de realizat și atunci pot fi acceptate și minime locale, cu condiția ca erorile să fie suficient de mici. De asemenea, se pot încerca diferite configurații ale rețelei, cu un număr mai mare de neuroni în stratul ascuns sau cu mai multe straturi ascunse, care în general conduc la minime locale mai mici. Totuși, deși minimele locale sunt într-adevăr o problemă, în practică nu reprezintă dificultăți de nesoluționat.

O chestiune importantă este alegerea unei configurații cât mai bune pentru rețea din punct de vedere al numărului de neuroni în straturile ascunse. În multe situații, un singur strat ascuns este suficient. Nu există niște reguli precise de alegere a numărului de neuroni. În general, rețeaua poate fi văzută ca un sistem în care numărul vectorilor de test înmulțit cu numărul de ieșiri reprezintă numărul de ecuații iar numărul ponderilor reprezintă numărul necunoscutelor. Ecuațiile sistemului sunt în general neliniare și foarte complexe și deci este foarte dificilă rezolvarea lor exactă prin

mijloace convenționale. Algoritmul de antrenare urmărește tocmai găsirea unor soluții aproximative care să minimizeze erorile.

O rețea neuronală trebuie să fie capabilă de generalizare. Dacă rețeaua aproximează bine setul de antrenare, aceasta nu este o garanție că va găsi soluții la fel de bune și pentru datele din alt set, cel de test. Generalizarea presupune existența în date a unor regularități, a unui model care *poate fi învățat*. În analogie cu sistemele liniare clasice, aceasta ar însemna niște ecuații redundante. Astfel, dacă numărul de ponderi este mai mic decât numărul de vectori de test, pentru o aproximare corectă rețeaua trebuie să se bazeze pe modelele intrinseci din date, modele care se vor regăsi și în datele de test. O regulă euristică afirmă că numărul de ponderi trebuie să fie în jur sau sub o zecime din produsul dintre numărul de vectori de antrenare și numărul de ieșiri. În unele situații însă (de exemplu, dacă datele de antrenare sunt relativ puține), numărul de ponderi poate fi chiar jumătate din produs.

Pentru un perceptron multistrat se consideră că numărul de neuroni dintr-un strat trebuie să fie suficient de mare pentru ca acest strat să furnizeze trei sau mai multe laturi pentru fiecare regiune convexă identificată de stratul următor. Deci numărul de neuroni dintr-un strat trebuie să fie de peste trei ori mai mare decât cel din stratul următor.

După cum am menționat, un număr insuficient de ponderi conduce la „sub-potrivire” (engl. “under-fitting”), în timp ce un număr prea mare de ponderi conduce la „supra-potrivire” (engl. “over-fitting”), fenomene prezentate în figura 1. Același lucru apare dacă numărul de epoci de antrenare este prea mic sau prea mare.

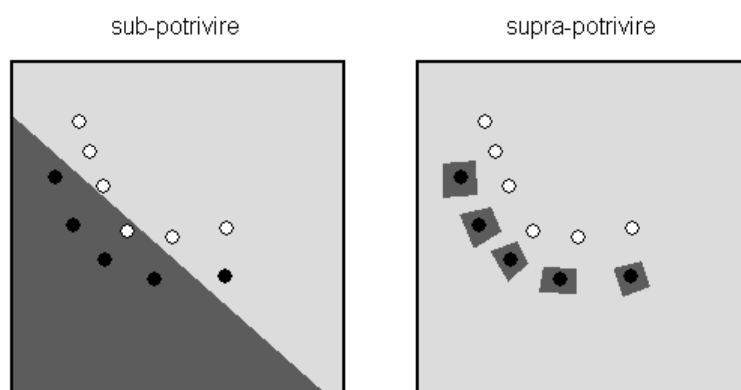


Figura 2. Capacitatea de aproximare a unei rețele neuronale în funcție de numărul de ponderi

O metodă de rezolvare a acestei probleme este oprirea antrenării în momentul în care se atinge cea mai bună generalizare. Pentru o rețea suficient de mare, s-a verificat experimental faptul că eroarea de antrenare scade în mod continuu pe măsură ce crește numărul epocilor de antrenare. Totuși, pentru date diferite de cele din setul de antrenare, se constată că eroarea scade la început până la un punct în care începe din nou să crească. De aceea, oprirea antrenării trebuie făcută în momentul când eroarea pentru setul de *validare* este minimă. Acest lucru se face împărțind datele de antrenare în două: între 67-90% din date vor fi utilizate pentru antrenarea propriu-zisă iar restul, numit set de validare este folosit pentru măsurarea erorii. Antrenarea se oprește atunci când începe să crească eroarea pentru setul de validare, moment numit „punct de maximă generalizare”.

În funcție de performanțele rețelei în acest moment, se pot încerca apoi diferite configurații, crescând sau micșorând numărul de neuroni din stratul (sau straturile) intermediar(e).

4. Matlab Neural Network Toolbox

newff – funcția construiește o rețea MLP (utilizată până la versiunea R2010a NNET 6.0.4)

Sintaxă


```
net = newff(P,T,S)
net = newff(P,T,S,TF,BTF,BLF,PF,IPF,OPF,DDF)
```

Descriere

newff(P,T,S) primește ca parametri:

P - matricea de dimensiuni $R \times Q1$ care conține cele mai reprezentative $Q1$ elemente ale celor R vectori prototip.

T - matricea de dimensiuni $SN \times Q2$ care conține cele mai reprezentative $Q2$ elemente ale celor SN vectori țintă.

Si - Dimensiunile celor $N-1$ straturi ascunse, $S1$ to $S(N-1)$, implicit []. Dimensiunea stratului de ieșire este determinată din **T**.

și returnează o rețea feed-forward ce folosește algoritmul backpropagation.

newff(P,T,S,TF,BTF,BLF,PF) primește următoarele intrări opționale,

TFi – Funcția de transfer a stratului i . Valoarea implicită este 'tansig' pentru straturile ascunse și 'purelin' pentru stratul de ieșire. Funcția de transfer $TF\{i\}$ poate fi orice funcție derivabilă, de exemplu: TANSIG, LOGSIG sau PURELIN.

BTF – funcția de antrenare a rețelei, cu valoarea implicită 'trainlm'. Funcția de antrenare poate fi una dintre funcțiile: TRAINLM, TRAINBFG, TRAINRP sau TRAINGD.

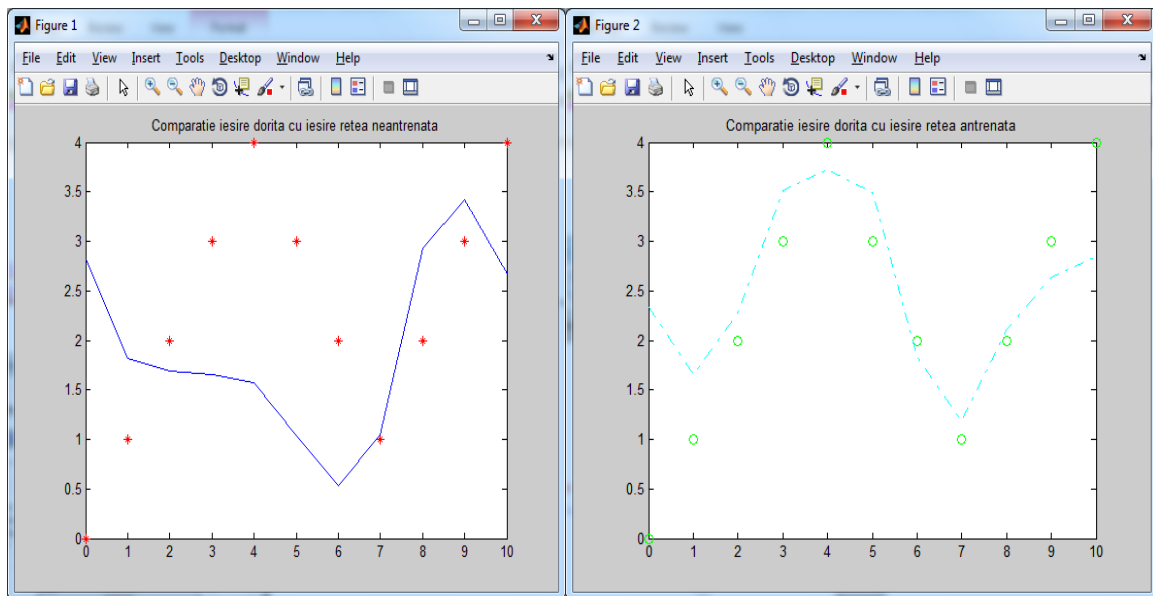
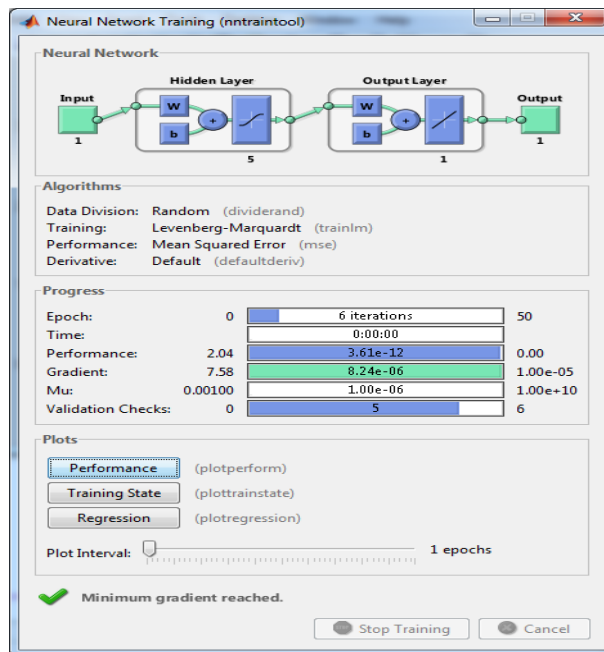
BLF – Funcția de învățare, cu valoarea implicită 'learngdm'. Funcția de învățare BLF poate fi una din următoarele funcții: LEARNGD sau LEARNGDM.

PF - funcția criteriu pentru evaluarea performanței, cu valoarea implicită 'mse'.
și returnează o rețea feed-forward ce folosește algoritmul backpropagation.

Exemplu

```
close all;
clear all;
clc;
% multimea vectorilor prototip
P = [0 1 2 3 4 5 6 7 8 9 10];

% multimea vectorilor tinta
T = [0 1 2 3 4 3 2 1 2 3 4];
% dimensiunile straturilor ascunse
% 1 strat ascuns cu 5 neuroni
S = [2 1];
net = newff(P, T, S);
net = newff([min(P) max(P)], S, {'tansig' 'purelin'}, 'traingd');
% evaluare retea neantrenata pe prototipurile din setul P
Y = sim(net, P);
figure(1);
plot(P, T, '*r', P, Y, '-b');
title('Comparatie iesire dorita cu iesire retea neantrenata');
% epoci de antrenare 100, 200, 500,...
net.trainParam.epochs = 50;
% eroarea 0.0001
net.trainParam.goal = 0.0001;
% rata de invatare 0.001, 0.1, 0.3
net.trainParam.lr = 0.01;
net = train(net, P, T);
Y = sim(net, P);
figure(2);
plot(P, T, 'og', P, Y, '-.c');
title('Comparatie iesire dorita cu iesire retea antrenata');
```



Parametrii rețelei antrenate:

- ponderile: $\text{net.IW}\{1\} = [-6.5669; -6.7676; -7.6076; 6.8881; -7.0563];$
- deplasările: $\text{net.b}\{1\} = [7.4293; 3.6983; 1.1361; 3.8130; -6.9153];$
 $\text{net.b}\{2\} = [0.6909];$

feedforwardnet – construiește o rețea neuronală feed-forward.

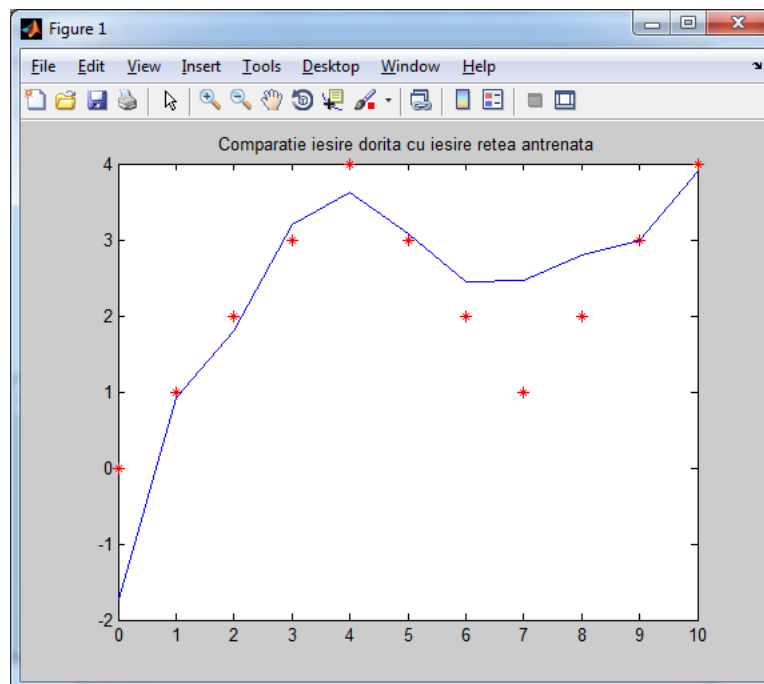
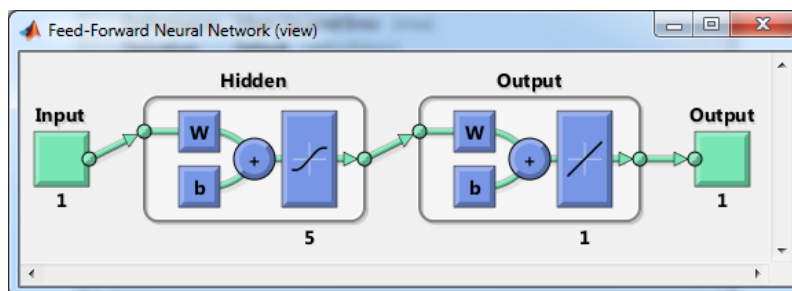
Sintaxă:

feedforwardnet(hiddenSizes,trainFcn) primește ca parametru vectorul linie cu dimensiunile celor N straturi ascunse și o funcție de antrenare specifică algoritmului backpropagation și returnează o rețea neuronală cu N+1 straturi. Intrările și ieșirile sunt setate pe 0, urmând a fi configurate cu datele specificate la apelarea funcției train sau manual prin apelul funcției configure. Argumentele implicite sunt (10, 'trainlm').

Exemplu

```
close all;
clear all;
clc;
% multimea vectorilor prototip
P = [0 1 2 3 4 5 6 7 8 9 10];
% multimea vectorilor tinta
T = [0 1 2 3 4 3 2 1 2 3 4];
% dimensiunile straturilor ascunse
% 1 strat ascuns cu 5 neuroni
S = [5];

net = feedforwardnet(S, 'trainlm');
net = train(net, P, T);
view(net)
Y = net(P);
perf = perform(net, P, Y);
figure(1);
plot(P, T, '*r', P, Y, '-b');
title('Comparatie iesire dorita cu iesire retea antrenata');
```



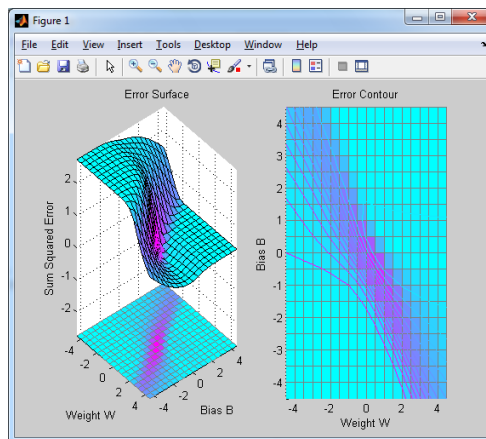
```
perf = 14.4846
net.IW{1} = [-7.0298; -6.9963; -6.9970; 6.8387; -6.7973 ];
net.b{1} = [6.9627; 3.4981;-0.0030; 3.7237; -7.2267];
net.b{2} = [ -1.0255];
```

Utilizare plotes

```
close all;
clear all;
clc;
P = 0:pi/12:pi/2;
T = sin(P);
```

```
w_range = -4.5:0.5:4.5;
b_range = -4.5:0.5:4.5;

figure(1);
ES = errsurf(P, T, w_range, b_range, 'tansig')/length(T);
plotes(w_range, b_range, ES, [50 25]);
```



5. Aplicații

5.1. Implementați algoritmul back-propagation pentru o rețea neuronală de tip perceptron multistrat, cu un singur strat ascuns și cu un număr variabil de unități în stratul de intrare și în stratul ascuns pentru a aproxima funcția neliniară φ definită prin vectorii prototip

$$x_i = \frac{(i-1)\pi}{12}, i = 1, \dots, 7 \text{ și vectorii țintă: } z_i = \sin(x_i), i = 1, \dots, 7.$$

Se vor preciza:

- valorile inițiale ale parametrilor, reprezentarea grafică a aproximării realizate de rețea pe baza acestor valori și eroarea corespunzătoare;
- valorile finale ale parametrilor (rezultate din antrenare), reprezentarea grafică a aproximării realizate de rețea pe baza acestor valori și eroarea corespunzătoare;
- numărul total de iterații cât a durat antrenarea și reprezentarea grafică a evoluției erorii ca funcție de numărul de iterații;
- 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.

5.2. Proiectați o rețea neuronală care să aproximeze funcția:

$$z_i = \varphi(x_i) = 1 + e^{-x_i} \cdot \sin\left(2\pi x_i - \frac{\pi}{2}\right), i = 1, \dots, N. \text{ Tiparele de intrare din setul de antrenare se aleg astfel: } u_i = \frac{i-1}{N-1}, i = 1, \dots, N. \text{ Se vor considera cazurile: } N = 51 \text{ sau } 101. \text{ Se va urmări ieșirea rețelei antrenate pe setul de date de testare } u = 0:\frac{0.5}{N}:1. \text{ Se vor alege valori convenabile pentru numărul de neuroni ascunși (2, 4 sau 8), rata de învățare (0.01, 0.1 sau 0.3) și numărul de epoci de antrenare (100, 250 sau 500). Se vor nota rezultatele obținute.}$$

Se vor considera cazurile: $N = 51$ sau 101 . Se va urmări ieșirea rețelei antrenate pe setul de date de testare $u = 0:\frac{0.5}{N}:1$. Se vor alege valori convenabile pentru numărul de neuroni ascunși (2, 4 sau 8), rata de învățare (0.01, 0.1 sau 0.3) și numărul de epoci de antrenare (100, 250 sau 500). Se vor nota rezultatele obținute.

5.3. Să se reia problema 5.2. pentru aproximarea funcției:

$$\varphi(x_i) = 1 + e^{-x_i} \sin\left(8\pi x_i - \frac{\pi}{2}\right), N = 51.$$