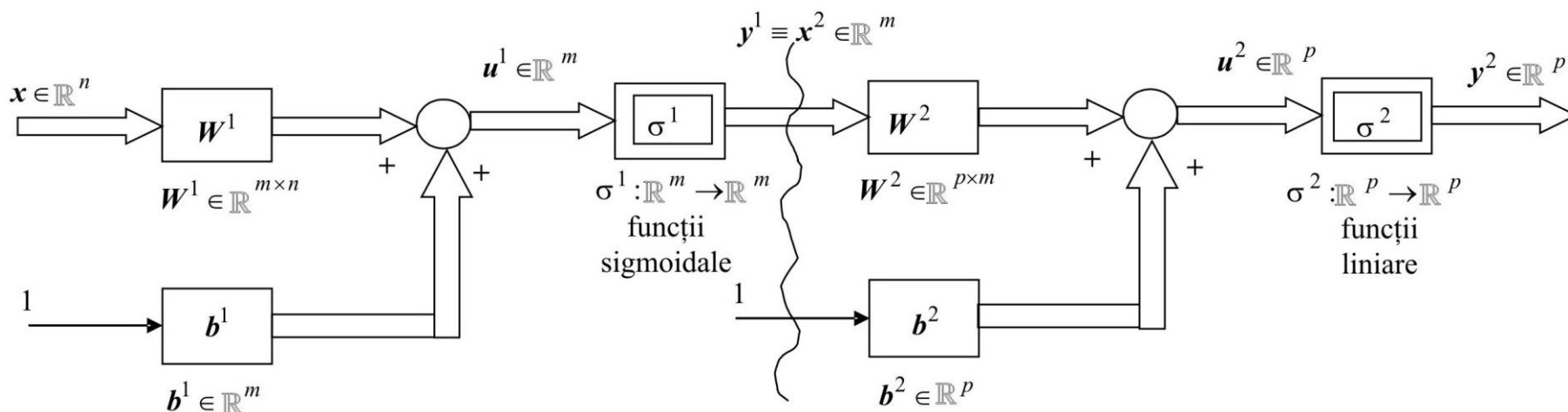


4.1. Proprietatea de aproximare a rețelelor feed-forward cu două straturi



$$f: \mathbb{R}^n \rightarrow \mathbb{R}^p, \quad y = f(x) = \sigma^2 W^2 \sigma^1 (W^1 x + b^1) + b^2 \quad (5.1.1)$$

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

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

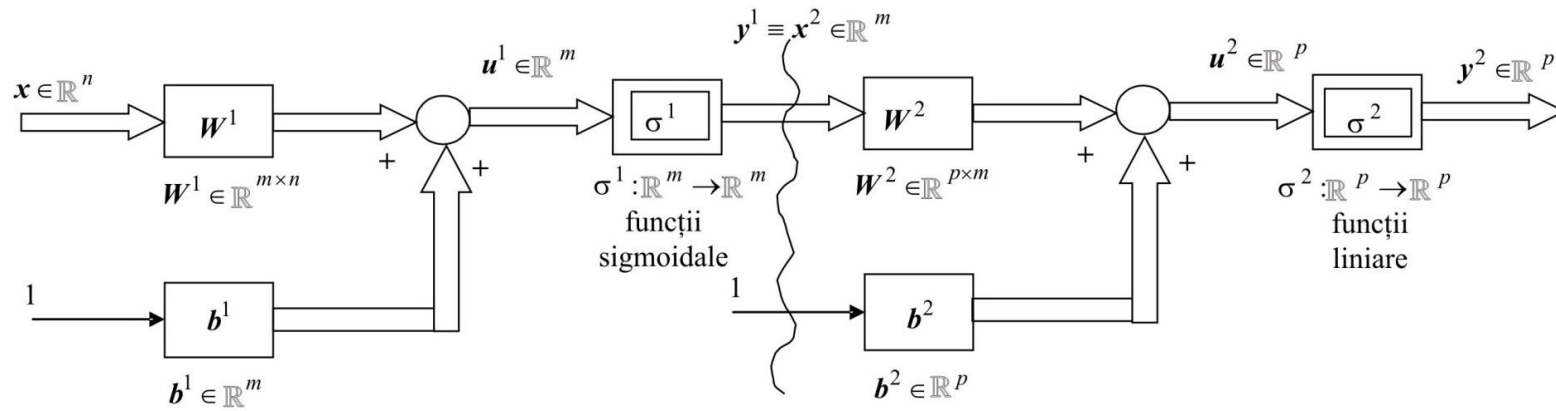
Parametrii rețelei:

$W^1 \in \mathbb{R}^{m \times n}$ – ponderi strat de intrare;

$b^1 \in \mathbb{R}^m$ – deplasări strat de intrare;

$W^2 \in \mathbb{R}^{p \times m}$ – ponderi strat de ieșire;

$b^2 \in \mathbb{R}^p$ – deplasări strat 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 de activare sigmoidală

$$\mathbf{u}^1 = \mathbf{W}^1 \mathbf{x} + \mathbf{b}^1 \stackrel{\text{not}}{=} [u_1^1 \dots u_m^1]^T, \text{ cu } \sigma^1(u) = \frac{1}{(1 + e^u)} \text{ sau } \sigma^1(u) = \frac{e^u - e^{-u}}{e^u + e^{-u}};$$

$\sigma^2: \mathbb{R}^p \rightarrow \mathbb{R}^p$, $\sigma^2(\mathbf{u}^2) = [\sigma^2(u_1^2) \dots \sigma^2(u_p^2)]^T = \mathbf{u}^2$, $\sigma^2(u) = u$ – funcția de activare liniară

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

Fie funcția netedă pe o mulțime compactă $S \subseteq \mathbb{R}^n$:

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

cunoscută în punctele $\mathbf{x}_i \in S \subseteq \mathbb{R}^n$:

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

Dacă rețeaua neuronală posedă *un număr adecvat de noduri* m în stratul de intrare, atunci parametrii rețelei pot fi determinați de așa manieră încât $\mathbf{f}$ (5.1.1) să aproximeze pe domeniul S funcția dată φ (5.1.2) cu o precizie impusă (în sensul de „oricât de bună”):

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

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

Rețeaua neurală MLP: Se alege numărul de neuroni m în stratul de intrare al rețelei
vectori prototip:

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

vectori țintă:

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

vectori de ieșire:

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

vectori eroare:

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

funcția obiectiv MSE:

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

$\forall m \Rightarrow \exists \mathbf{W}^{1*} \in \mathbb{R}^{m \times n}, \mathbf{b}^{1*} \in \mathbb{R}^m, \mathbf{W}^{2*} \in \mathbb{R}^{p \times m}, \mathbf{b}^{2*} \in \mathbb{R}^p$, astfel încât:

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

4.2. Antrenarea rețelelor feedforward cu două straturi

5.2.1. Obiectivul antrenării rețelei

- determinarea parametrilor $\mathbf{W}^{1*} \in \mathbb{R}^{m \times n}$, $\mathbf{b}^{1*} \in \mathbb{R}^m$, $\mathbf{W}^{2*} \in \mathbb{R}^{p \times m}$, $\mathbf{b}^{2*} \in \mathbb{R}^p$ care să asigure *minimizarea erorii pătratice globale*
- m va influența calitatea aproximării realizate de rețea pentru funcția φ

5.2.2. Algoritmul de antrenare prin propagare inversă

backpropagation: minimizarea erorii pătratice medii (MSE) prin metoda gradientului.

Fie $\mathbf{v} \in \mathbb{R}^{m(n+1)+p(m+1)} = \text{vect} \begin{bmatrix} \mathbf{W}^1 & \mathbf{b}^1 \end{bmatrix}; \begin{bmatrix} \mathbf{W}^2 & \mathbf{b}^2 \end{bmatrix}$, $\mathbf{W}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}^1 \in \mathbb{R}^m$, $\mathbf{b}^2 \in \mathbb{R}^p$

Funcția obiectiv:

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

Minimizarea prin metoda gradientului:

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

$\eta > 0$ - pasul metodei de optimizare,

$\nabla J(\mathbf{v}_{\text{old}}) = \left. \partial J / \partial \mathbf{v}^T \right|_{\mathbf{v}_{\text{old}}}$ - vectorul gradient (normal la suprafețele de nivel constant ale lui J

care trec prin $\mathbf{v}_{\text{old}}$).

Etapa 0 (Inițializare)

se fixează:

un *număr maxim de iterații* sau *epoci*

o *valoare de prag* $\varepsilon > 0$ care se dorește a fi atinsă pentru funcția criteriu

rată de învățare (eng. *learning rate*) = pasul metodei de optimizare $\eta > 0$

se construiesc:

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

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

Etapa I (Etapa de prezentare)

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

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

Etapa II (Etapa de verificare)

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

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

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:

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

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

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

pentru valorile curente $\mathbf{W}_{\text{old}}^1 \in \mathbb{R}^{m \times n}$, $\mathbf{W}_{\text{old}}^2 \in \mathbb{R}^{p \times m}$, $\mathbf{b}_{\text{old}}^1 \in \mathbb{R}^m$, $\mathbf{b}_{\text{old}}^2 \in \mathbb{R}^p$

pentru fiecare vector prototip $\mathbf{x}_i$, $i = 1, \dots, N$, se calculează *matricele Jacobian ale funcțiilor de activare* (vectoriale) ale celor două straturi:

- pentru stratul de intrare:

$$\mathbf{Q}_i^1 = \text{diag } q_i^1(1), \dots, q_i^1(m), \quad \mathbf{Q}_i^1 \in \mathbb{R}^{m \times m}, \quad i = 1, \dots, N, \quad (5.2.9-a)$$

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:

$$\mathbf{Q}_i^2 = \text{diag } q_i^2(1), \dots, q_i^2(m), \quad \mathbf{Q}_i^2 \in \mathbb{R}^{p \times p}, \quad i = 1, \dots, N, \quad (5.2.10-a)$$

unde:

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

funcțiile de activare pe al doilea strat sunt liniare $\Rightarrow \mathbf{Q}_i^2 = \mathbf{I}$ este *matricea unitate*, $i = 1, \dots, N$.

se definesc vectorii / matricele delta:

- pentru stratul de ieșire:

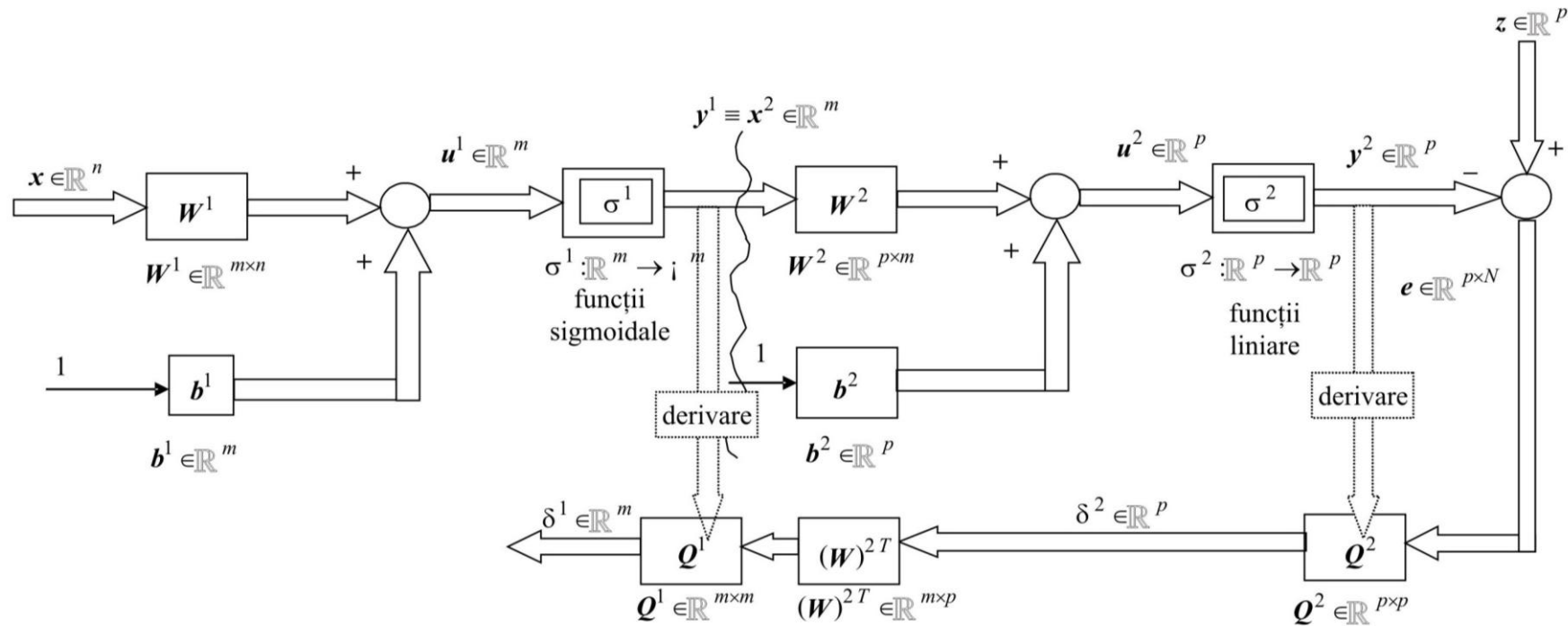
$$\delta_i^2 = Q_i^2 e_i = e_i \in \mathbb{R}^p, \quad i = 1, \dots, N; \quad (5.2.11)$$

$$\Delta^2 = [\delta_1^2 \dots \delta_N^2] \in \mathbb{R}^{p \times N}, \quad (5.2.13)$$

- pentru stratul de intrare:

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

$$\Delta^1 = [\delta_1^1 \dots \delta_N^1] \in \mathbb{R}^{m \times N}. \quad (5.2.14)$$



Etapa IV (Etapa de actualizare a parametrilor)

➤ se calculează *matricele de actualizare* a parametrilor aplicând *regula delta* pe fiecare strat:

- pentru stratul de ieșire

ponderi

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

deplasări

$$\mathbf{D}_{b^2} = \Delta^2 \underbrace{1 \dots 1}_N^T \in \mathbb{R}^p; \quad (5.2.15-b)$$

cu

$$\mathbf{x}_i^2 = \mathbf{y}_i^1 \in \mathbb{R}^m, \quad i=1, \dots, N \Rightarrow \mathbf{X}^2 = [\mathbf{x}_1^2 \dots \mathbf{x}_N^2] \in \mathbb{R}^{m \times N}. \quad (5.2.17)$$

- pentru stratul de intrare

ponderi

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

deplasări

$$\mathbf{D}_{b^1} = \Delta^1 \underbrace{1 \dots 1}_N^T \in \mathbb{R}^m. \quad (5.2.16-b)$$

➤ se determină noile valori ale parametrilor rețelei:

- pentru stratul de ieșire

$$\mathbf{W}_{\text{new}}^2 = \mathbf{W}_{\text{old}}^2 + \eta \mathbf{D}_{W^2}, \quad \mathbf{b}_{\text{new}}^2 = \mathbf{b}_{\text{old}}^2 + \eta \mathbf{D}_{b^2}; \quad (5.2.18)$$

- pentru stratul de intrare

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

➤ se trece la o nouă iterație cu:

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

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 senzitivitate 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 mărimea 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. ■