

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Learning (supervised and unsupervised)

- From artificial neuron to deep learning
- Unsupervised learning: k -means and bandit algorithms
- See next devoted sections

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Linear artificial neuron

- Mc Cullogh & Pitts in 1943 : linear “neuron”. Emulate some logical functions ($\{0, 1\}$ -valued inputs and outputs).
- The first perceptron: 1957 by Rosenblatt ⁽⁹⁾ as a binary classifier.
- Learning by Hebb’s rule also known as reinforcement rule.
- Only able to classify linearly separable classes of data.

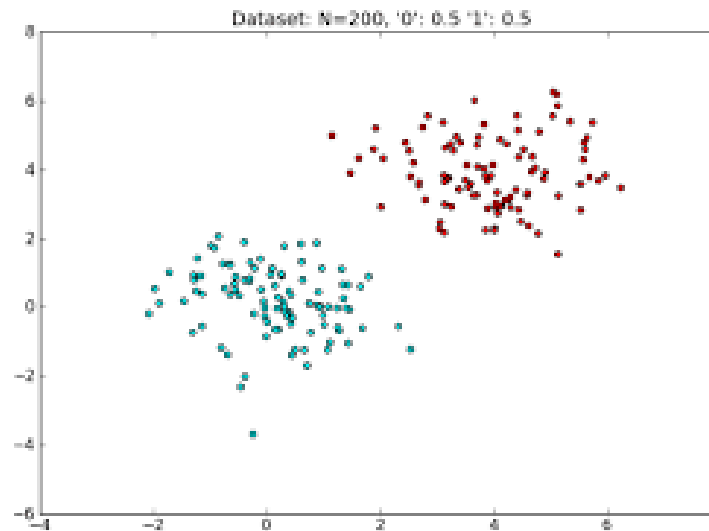


Figure: The Rosenblatt neural network performances.

⁹ Rosenblatt, Frank (1957), The Perceptron: a perceiving and recognizing automaton. Report 85-460-1, Cornell Aeronautical Laboratory.

No hidden layer: ADALINE (Adaptive Linear Neuron (B. Widrow & T. Hoff, 1960))

- Input data $x_1, \dots, x_N$ of the form $x_k = (1, x_k^1, \dots, x_k^d) \in \mathbb{R}^{d+1}$.
- Output (true) data are **real numbers** $y_1, \dots, y_N$.
- Data set: $z_k = (x_k, y_k)$, $k = 1 : N$ (supervised learning).
- Let $\Theta = \mathbb{R}^{d+1}$ the parameter set.
- The answer of the **ADALINE neuron/network** for $\theta \in \Theta$ “fed” with an input datum $x = (1, x^1, \dots, x^d)$ is

$$\theta^\top x = \sum_{i=0}^d \theta^i x^i = \theta^0 + \sum_{i=1}^d \theta^i x^i$$

where $\theta^\top$ denote the transpose of θ .

- **Local (convex) prediction/loss function**: $v(\theta, (x, y)) = \frac{1}{2}(y - \theta^\top x)^2$.

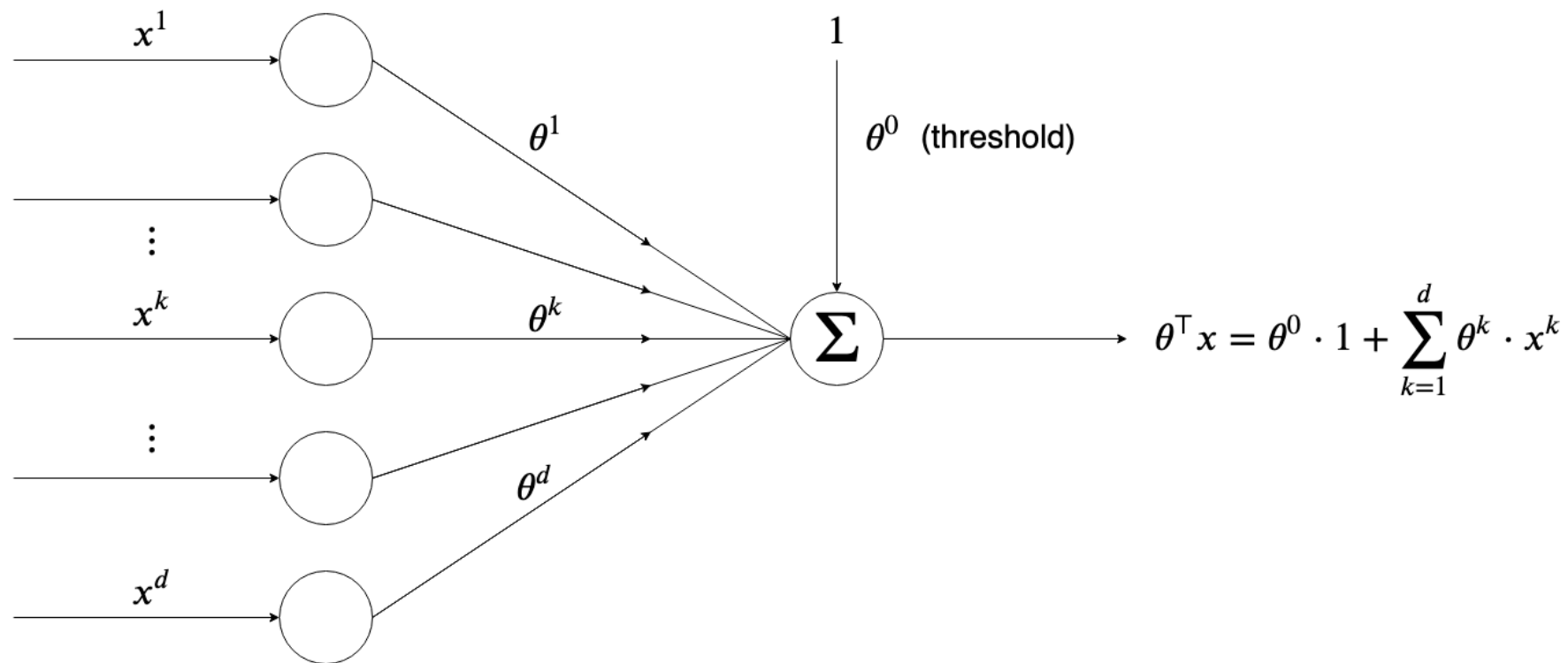


Figure: The ADALINE neural network.

Remark. When Σ is replaced by the indicator function $\mathbf{1}_{\mathbb{R}_+}$ i.e. the output is $\mathbf{1}_{\theta^T x \geq 0}$ this network becomes Rosenblatt's perceptron.

- $V(\theta) = \frac{1}{2} \int (\theta^\top x - y)^2 d\mu_N(x, y) = \mathbb{E} (\theta^\top x_I - y_I)^2$ (**convex!**) (where $I \stackrel{d}{=} \mathcal{U}(\{1, \dots, N\})$)
- i.e. **Global prediction/loss function = Ordinary Least squares !**
- $\nabla V(\theta) = \int (\theta^\top x - y)x d\mu_N(x, y) = \mathbb{E} (\theta^\top x_I - y_I)x_I.$
- The target is

$$\begin{aligned} \nabla V(\theta) = 0 &\iff \int \underbrace{(\theta^\top x)x}_{=(xx^\top)\theta} d\mu_N(x, y) = \int yx d\mu_N(x, y) \\ &\iff \theta_{opt} = \left(\int xx^\top d\mu_N(x, y) \right)^{-1} \left(\int yx d\mu_N(x, y) \right). \end{aligned}$$

- [Exercise: Prove that $(\theta^\top x)x = (xx^\top)\theta$].
- It is **linear regression** (as expected).

- Resulting (*SGD*):

$$\theta_{n+1} = \theta_n - \gamma_{n+1} (\theta_n^\top x_{l_{n+1}} - y_{l_{n+1}}) x_{l_{n+1}}$$

where $(l_n)_{n \geq 1}$ are i.i.d. with uniform distribution on $\{1, \dots, N\}$.

All assumptions are satisfied $\implies \theta_n \rightarrow \theta_{opt}$ *a.s.*

- Conclusion : the ADALINE network **performs linear regression without matrix inversion**.
- Extension to q -dimensional (true) output data y_k with $\Theta = \mathbb{M}(d, q)$ and

$$v(\Theta, (x, y)) = \frac{1}{2} |\Theta^\top x - y|^2.$$

“Non-linear” ADALINE

- Assume the output data y_k are e.g. $(0, 1)$ -valued.
- Let $\Psi \in C(\mathbb{R}, (0, 1))$, Ψ increasing C^2 -diffeomorphism from $\mathbb{R}$ onto $(0, 1)$.
- $V(\theta) = \frac{1}{2} \mathbb{E} (\Psi(\theta^\top x_I) - y_I)^2$ with $I \sim \mathcal{U}(\{1, \dots, N\})$.
- $\nabla V(\theta) = \mathbb{E} \left((\Psi(\theta^\top x_I) - y_I) \Psi'(\theta^\top x_I) x_I \right)$.
- Loss of convexity since

$$\nabla^2 V(\theta) = \int \left[\Psi'(\theta^\top x)^2 + \underbrace{(\Psi(\theta^\top x) - y)}_{\text{varying sign}} \underbrace{\Psi''(\theta^\top x)}_{\neq 0} \right] x x^\top d\mu_N(x, y)$$

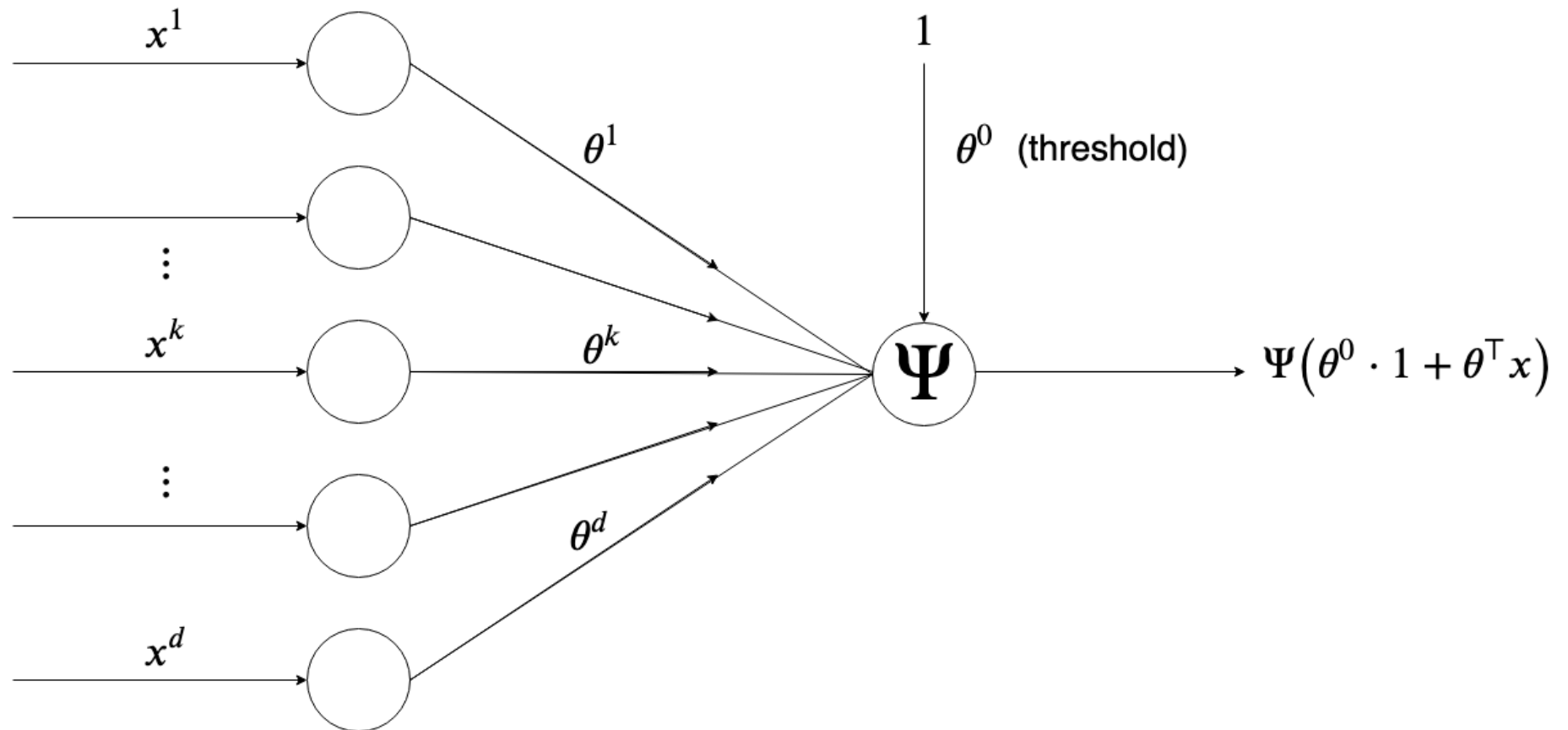


Figure: The **non-linear** ADALINE neural network.

- **Historical Rosenblatt's "linear perceptron"** (1957): $\Psi(u) = \mathbf{1}_{u \geq 0}$ is a linear classifier.

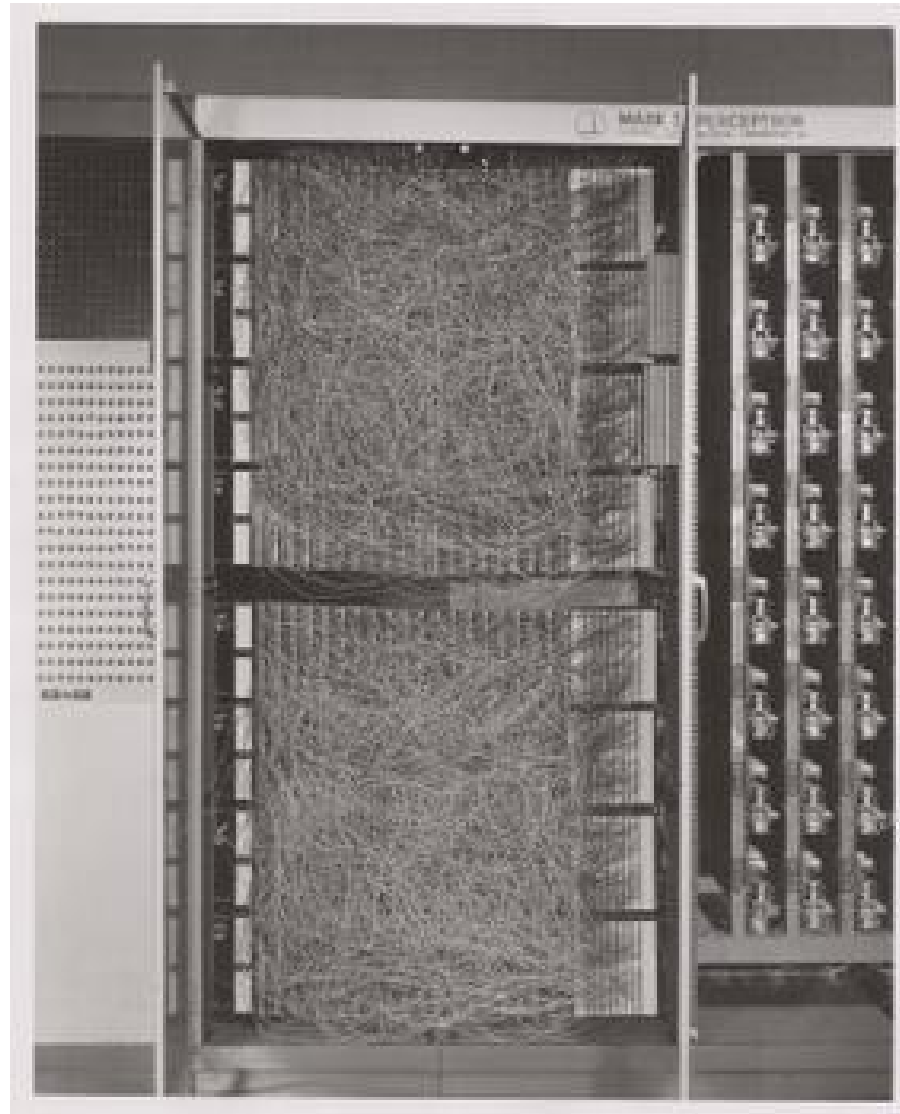


Figure: Historical Rosenblatt's perceptron: the hardware

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

One hidden layer: universal approximation property

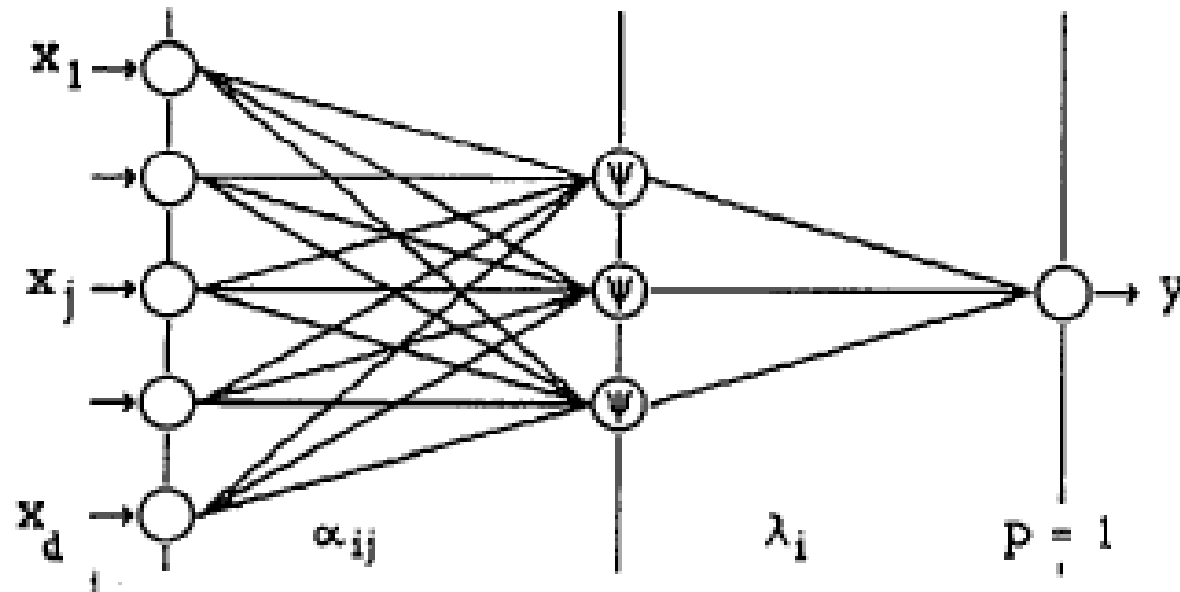


Figure: One hidden layer MLP (below: convention $\alpha_{ij} = w_{ij}$).

- Introduced by P. Werbos en 1974 (including backpropagation of errors) ⁽¹⁰⁾.

¹⁰Werbos, Paul J. (1994). The Roots of Backpropagation : From Ordered Derivatives to Neural Networks and Political Forecasting. New York: John Wiley & Sons.

One hidden layer: universal approximation property

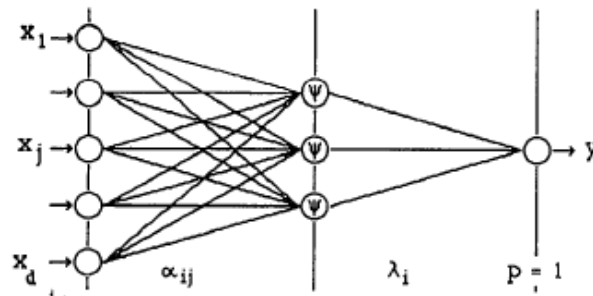


Figure: One hidden layer MLP (below: convention $\alpha_{ij} = w_{ij}$).

- $d = d_x$ -dimensional inputs. Switch to $x \rightsquigarrow \begin{pmatrix} 1 \\ x \end{pmatrix}$.
- L units $i = 1 : L$ on the hidden layer, with an *activation function*

$$\Psi \in C(\mathbb{R}, \mathbb{R}).$$

- Unit i of the hidden layer receives $x = (x^j)_{j=1:d}$ and emits

$$\Psi((w_i \cdot |x)) = \Psi\left(w_{i0} + \sum_{1 \leq j \leq d} w_{ij} x^j\right).$$

- The output layer receives $[\Psi((w_i \cdot |x))]_{i=1:L}$ and emits

$$\sum_{1 \leq i \leq L} \lambda_i \Psi((w_i \cdot |x)).$$

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Universal approximation property

Theorem (Cybenko, 1989)

(^a) If the *activation function* Ψ satisfies

$$(CL_\Psi) \equiv \Psi \in C_b(\mathbb{R}, \mathbb{R}), \quad \lim_{\xi \rightarrow -\infty} \Psi(\xi) = 0, \quad \lim_{\xi \rightarrow +\infty} \Psi(\xi) = 1,$$

then
$$F = \left\{ x \mapsto \sum_{i=1}^L \lambda_i \Psi((w_i \cdot x) + w_{i0}), \quad L \in \mathbb{N}^*, \lambda \in \mathbb{R}^L, w \in \mathbb{R}^{L \times (d+1)} \right\}$$

is $\|\cdot\|_{\text{sup}}$ -dense in $C([0, 1]^d, \mathbb{R})$.

^aG. Cybenko, Approximation by Superpositions of a Sigmoidal Function. Mathematics of Control, Signals, and Systems, 2:303-314, 1989.

- (CL_Ψ) can be slightly relaxed into Hornik et al's condition (¹¹, 1989)

$$(CS_\Psi) \equiv \Psi \in C_b(\mathbb{R}, \mathbb{R}), \quad \exists u_1, u_2 \in \mathbb{R} \text{ such that } \Psi(u_1) \neq \Psi(u_2).$$

- Extension to vector-valued outputs is straightforward.

¹¹Hornik K., Stinchcombe M. & White H., Multilayer feedforward networks are universal approximators. *Neural Networks*, 2, 359–366

Proof I

Assume $d = 1$ (for simplicity).

► STEP 1. If $\overline{F}^{\|\cdot\|_{\sup}} \subsetneq C([0, 1], \mathbb{R})$, $\exists$ a linear form $\Lambda \in C([0, 1], \mathbb{R})^*$ s.t.

$$\Lambda \neq 0 \quad \text{and} \quad \Lambda|_{\overline{F}} \equiv 0 \quad (\text{by Hahn-Banach's Theorem}).$$

By Riesz's representation Theorem, there exists a signed measure μ s.t.

$$\Lambda g = \int_{\mathbb{R}} g \, d\mu.$$

Hence, with for every weight $w = (\alpha, \beta)$,

$$\forall w = (\alpha, \beta) \in \mathbb{R}^2, \quad \int_{\mathbb{R}} \Psi(\alpha x + \beta) \mu(dx) = 0.$$

► STEP 2. Let $\mu = \mu^+ - \mu^-$, denote the/a decomposition of μ .

• Let $\alpha = 0, \beta \rightarrow +\infty$.

$$\int \mathbf{1} \mu(dx) = 0 \quad \text{so that} \quad \mu^+(\mathbb{R}) = \mu^-(\mathbb{R}).$$

Proof II

- Now setting $\beta = -\alpha u$ yields

$$\forall \alpha, u \in \mathbb{R}, \quad \int \Psi(\alpha(x - u)) \mu(dx) = 0.$$

Letting $\alpha \rightarrow +\infty$ implies, it follows from Lebesgue Dominated Cvgce theorem

$$\forall u \in \mathbb{R}, \quad \mu(\{x : x > u\}) + \Psi(0)\mu(\{u\}) = 0.$$

- The set $D = \{u : \mu^+(\{u\}) > 0 \text{ or } \mu^-(\{u\}) > 0\}$ is at most countable and, for every $u \in D^c$,

$$\forall u \in D^c, \quad \mu(\{x : x > u\}) = 0.$$

► STEP 3. Combined with $\mu^+(\mathbb{R}) = \mu^-(\mathbb{R})$, one has

$$\forall u \in \mathbb{R}, \quad \mu^+(\{x : x > u\}) = \mu^-(\{x : x > u\}).$$

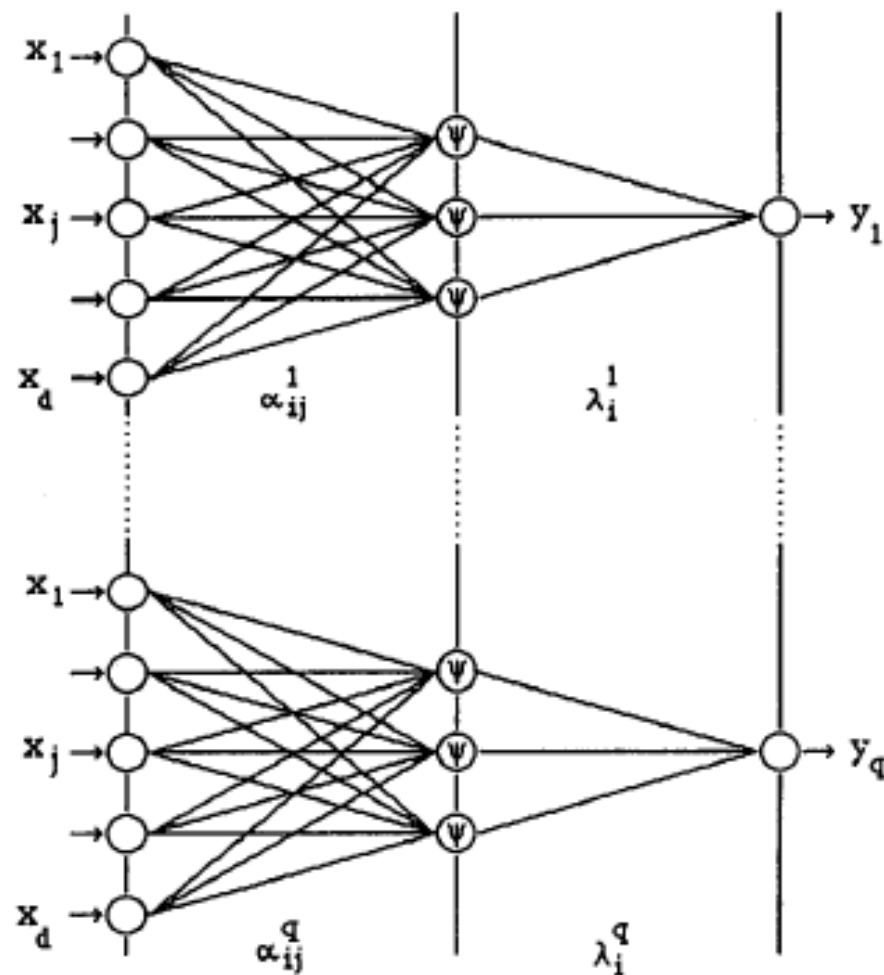
Hence $\mu^+ = \mu^-$ since D^c is everywhere dense in $\mathbb{R}$ i.e. $\mu \equiv 0$.

► STEP 4. **Contradiction! Since $\Lambda \neq 0$.**

□

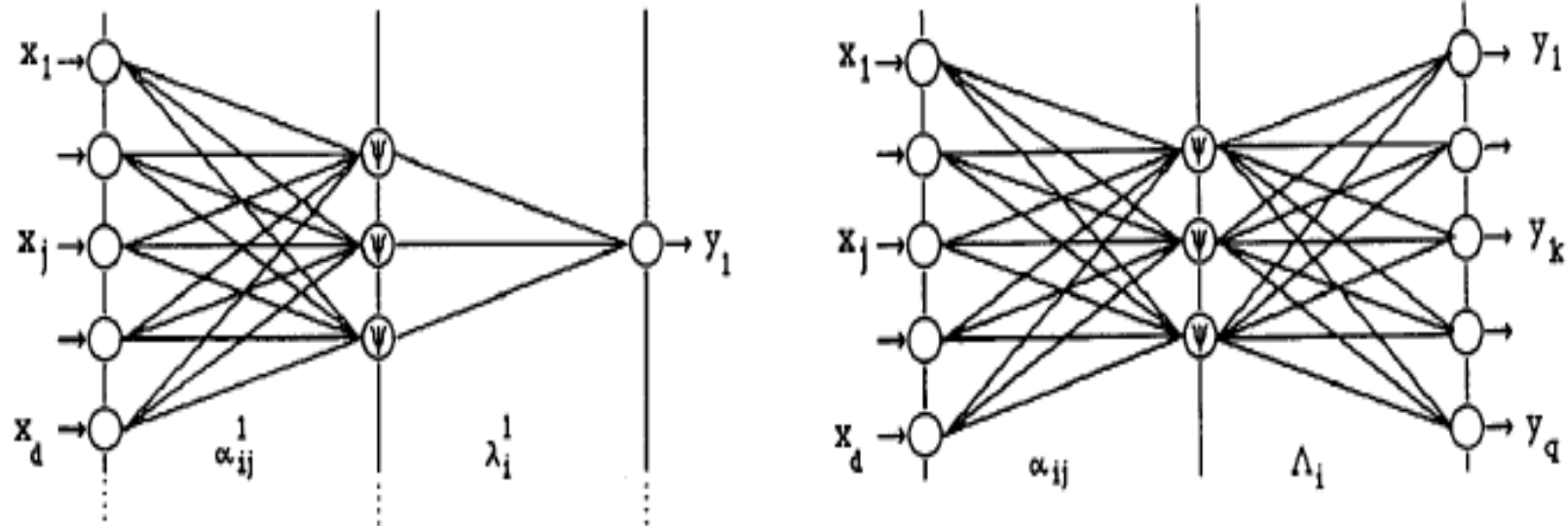
Higher dimensional outputs I

- Instant (theoretical) result by concatenating networks. . . (Warning! Unusual notations for indices)



Higher dimensional outputs II

- ... Although this is better in practice (Warning! same unusual notations for indices)



Constructive proof/smoothness

- Both proofs are **not constructive** (Hornik's relies on a distribution based argument).
- Constructive approach are possible including rates for higher order derivatives like. Among others...

Theorem (Attali-P., *Neural Networks*, 1997)

(^a) Assume the **activation function** $\psi \in C_b^r(\mathbb{R}, \mathbb{R})$ and satisfies

ψ is **non-polynomial on any open interval** and $\exists \xi_0$ s.t. $\psi^{(s)}(\xi_0) \neq 0$, $s = 0, \dots, r$.

Let $f \in C^\infty(\mathbb{R}^d, \mathbb{R})$. For every $L \in \mathbb{N}$, there exists weights $(w_{ij}^{(L)})_{i=1:L, j=1:d}$, $\lambda_{i=1:L}^{(L)}$ such that (with $\|g\|_{[0,1]^d} = \sup_{\xi \in [0,1]^d} |g(\xi)|$)

$$\max_{s=0:r} \left\| \partial^s \left(\sum_{i=1}^L \lambda_i^{(L)} \psi((w_i^{(L)} | x)) \right) - \partial^s f \right\|_{[0,1]^d} \leq C_f \cdot L^{-1/d}$$

^a J.-G. Attali, G. Pagès, Approximations of Functions by a Multilayer perceptron: a New Approach, *Neural Networks*, **10**(6):1069-1081, 1997.

Curse of dimensionality

- The proof heavily relies on multi-variable Bernstein polynomials and Vandermonde determinants.
- Obviously suffers from the **curse of dimensionality** through this $L^{-1/d}$ or even $L^{-(r+1-s)/d}$ for $\partial_s f$.
- A celebrated paper by A. Barron (seems to) break(s) this curse using **Maurey's Lemma** with a rate $O(L^{-1/2})$ for a function $f : [0, 1]^d \rightarrow \mathbb{R}$, C^1 , with an additional constraint on the Fourier coefficients (a series is converging) ⁽¹²⁾.
- Unfortunately this condition is more and more difficult to check as d grows: the space of such functions becomes sparse in $C^1([0, 1]^d, \mathbb{R})$.
- These papers came too late: **Vapnik's SVM** were coming at the front owing to a nice mathematical framework.
- Third winter of Neural networks begins.

¹²Universal approximation bounds for superpositions of a sigmoidal function, *IEEE Transactions on Information Theory*, Vol.39, Issue: 3, 930-945, May 1993.

By the way, what about ReLu ?

- What about using ReLu function ?
- Have in mind $\text{ReLu}(x) = x^+ = \max(x, 0)$.
- In fact

$$\psi(x) = \frac{1}{2}((x+1)^+ - (x-1)^+)$$

satisfies the assumptions of Cybenko's (and a fortiori Hornik et al. 's theorem(s))

- Hence, their theorem remains true by considering

$$\psi(x) = \text{Relu}(x) = x^+.$$

Renewal !

- G. Hinton, Y. Le Cun and Y. Bengio stroke back in 2012's with outstanding performances classification results at the 2012 edition of ImageNet challenge, using deep learning...
- New networks with more layers, not fully connected and new type of units (convolutive units, recurrent units).
- and even more recently Generative Adversarial Networks (2014) seem to get rid of curse of dimensionality,
- but, so far, no theoretical evidence.

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

From learning to deep learning

- Feedforward multilayer perceptron.
- **1998** (Y. Le Cun): MNIST database of handwritten figures.
 - **Input:** Image = 28×28 pixels $\times$ 256 grey levels $d_x \simeq 784$.
 - **Output:** probability y_i for each 10 figures : $[0, 1]^{10}$ i.e. $d_y \simeq 9$ or 10.
 - Predictor : Multilayer perceptron (MLP) with 1 hidden layers with 200 units $\implies 784 \times 200 \times 10 = 1.6 \cdot 10^6$ parameters
 - Size of the database: $N = 50\,000$ for learning, 10 000 for testing.
- **2010:** Same MNIST database
 - Predictor: MLP with convolutional units on *GPU* with 5 hidden layers
 $[2\,500, 2\,000, 1\,500, 1\,000, 500] \simeq 7,84 \cdot 10^9$ parameters
- **2013** (Google) : ImageNet database of $N = 162 \times 10^6$ images of 100×100 pixels and 256 grey levels: $d_x \simeq 10^4$.
 - $d_y = 21$ (classifier across 21 classes).
 - Predictor: Convolutional MLP network $\simeq 10^4 \times 1.72 \cdot 10^9$ parameters.
 - Error rate $< 1/1\,000$.

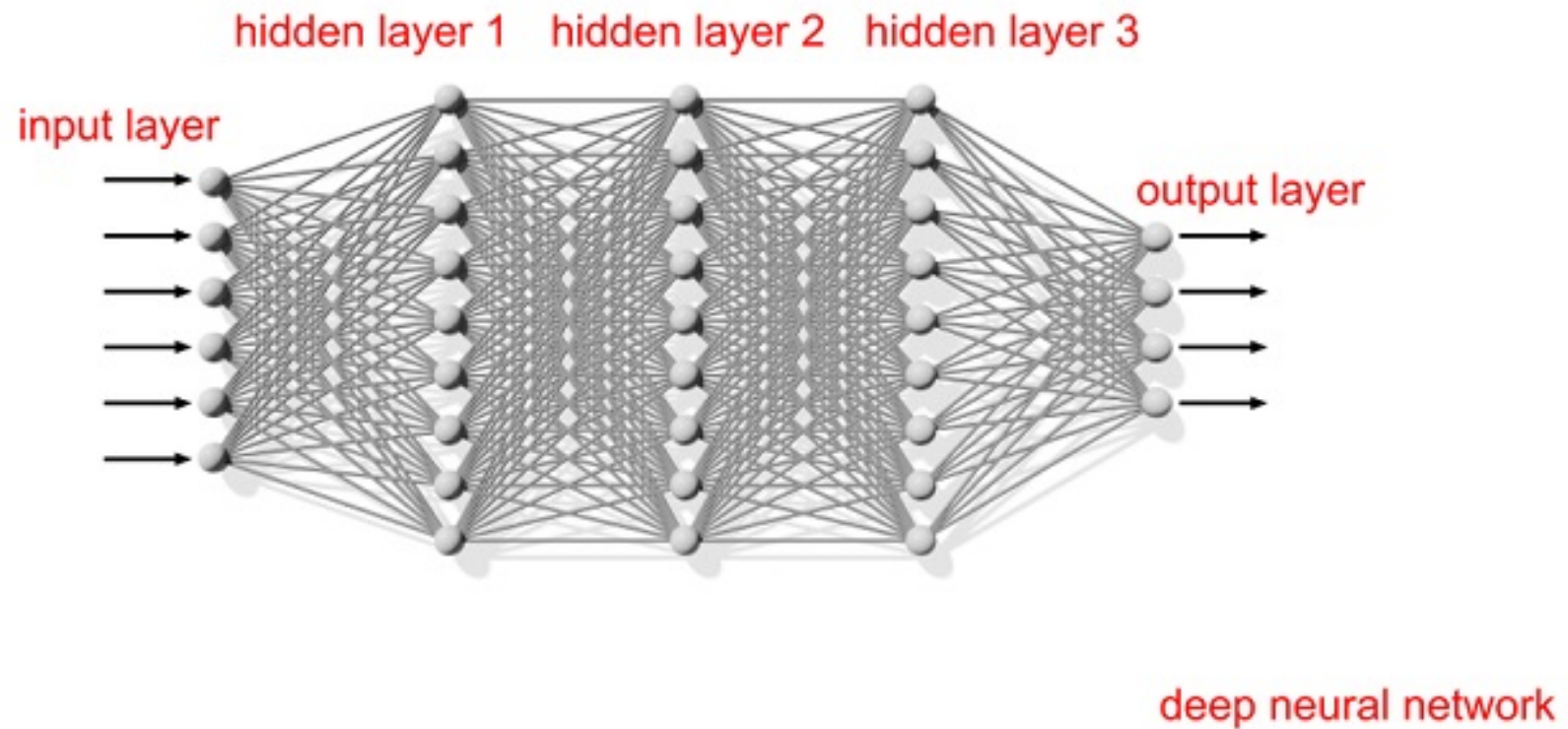


Figure: *A three-hidden layer feedforward perceptron.*

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Multilayer : toward Back propagation

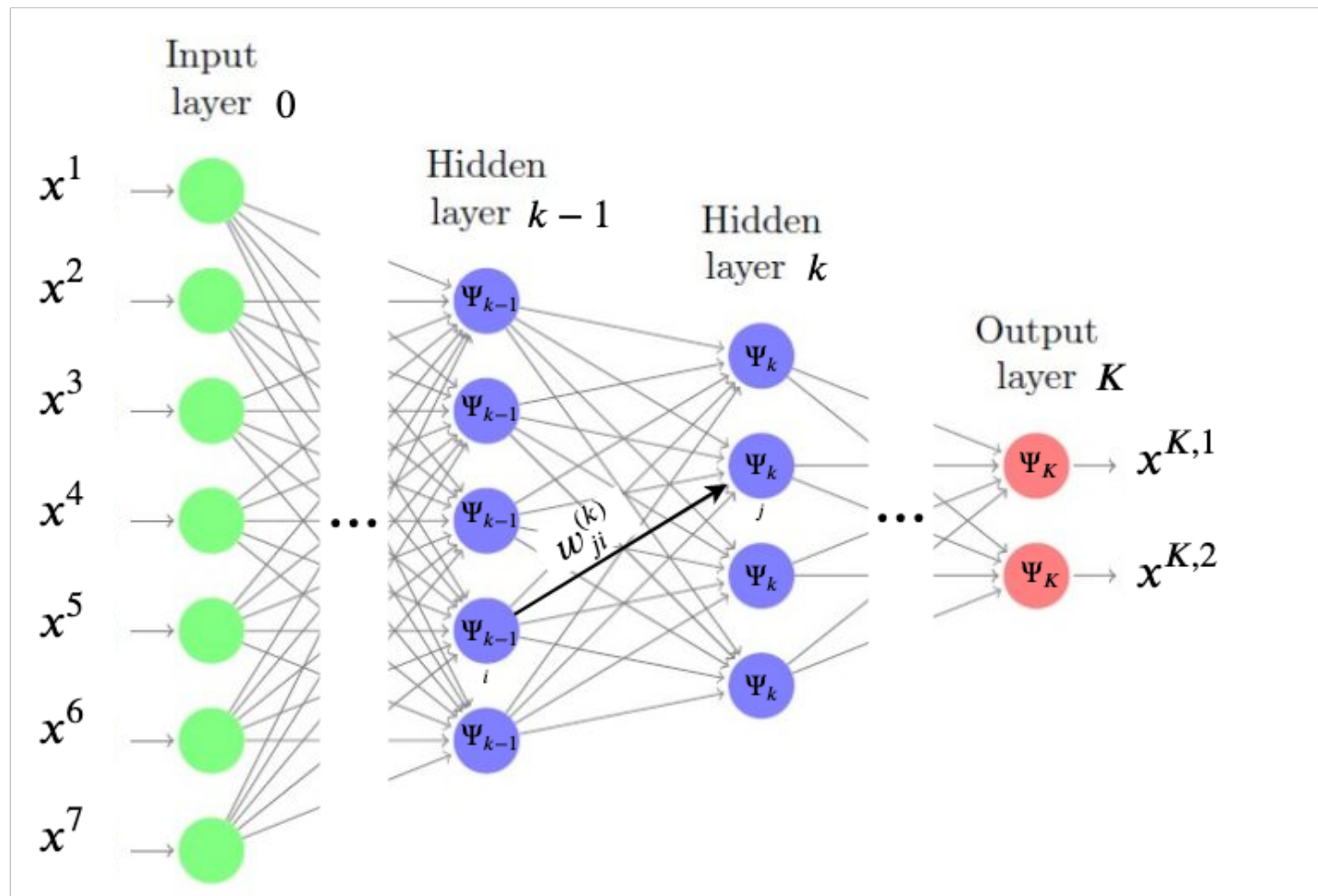


Figure: A $K - 1$ hidden layer feedforward perceptron.

Multilayer and backpropagation in one slide

- Input data (layer 0): $x^0 = x$ i.e. $x_j^0 = x_j, j = 1 : d_0 = d + 1$,
- $K - 1$ hidden layers ($k = 1 : K - 1$) and $K + 1$ layers (0 and K)
- Output of unit i of layer $k - 1$: $x^{k-1,i}$.
- Input of unit j of layer k :

$$\sum_{1 \leq i \leq d_{k-1}} w_{ji}^{(k)} x^{k-1,i} = (w_{j\cdot}^{(k)} \mid x^{k-1,\cdot})$$

so that, after passing through the activation function

$$x^{k,j} = \psi_k((w_{j\cdot}^{(k)} \mid x^{k-1,\cdot})).$$

- Typical activation functions

$$\psi_k(\xi) = \frac{e^\xi}{e^\xi + 1}, \quad \psi_k(\xi) = \tanh(\xi), \quad \psi_k(\xi) = \xi_+ \text{ (ReLU)},$$

$$\text{and } \psi_k(\xi) = \log(1 + e^\xi) \text{ or swish}(\xi) = \xi \cdot \frac{e^\xi}{e^\xi + 1}.$$

Intermission: connection with AAD

- Automatic differentiation : Robert E. Wengert 1964 in his PhD thesis ⁽¹³⁾.
- Reverse/Adjoint AD: Bryson (1962); Werbos (1974) but also Le Cun (1985); Rumelhart, Hinton (the three for Backprop!) and Speelpenning (1980) for true “automatic”.
- **AAD** for Adjoint Automatic differentiation, proposes a model close to backpropagation but simpler.
- Let $g_k : \mathbb{R}^d \rightarrow \mathbb{R}^d$, $k = 1, \dots, n$ be n differentiable vector fields.
- Set $J_x(g) = \left[\frac{\partial g^i}{\partial x^j}(x) \right]$ the Jacobian of $g = (g^1, \dots, g^d)$ at $x = (x^1, \dots, x^d)$.
- Then
$$J_x(g_n \circ \dots \circ g_1)^\top = J_{y_0}(g_1)^\top \circ \dots \circ J_{y_n}(g_n)^\top$$
- where $y_k = g_{k-1}(y_{k-1}), \dots, y_0 = x$.

¹³Robert E. Wengert. A simple automatic derivative evaluation program. Communications of the ACM 7(8):463–4, Aug 1964

In fact two slides

- Parameter to be calibrated to perform the learning of the network:

$$\mathbf{w} = (w^{(0)}, w^{(1)}, \dots, w^{(K)})$$

with $w^{(0)} = Id$, $w^{(k)} \in \mathbb{M}_{d_k, d_{k+1}}(\mathbb{R})$.

- Dimension of $\mathbf{w} := D := d_0 \cdot d_1 + \dots + d_{K-1} \cdot d_K$.
- Local Predictor: x^0 input datum, $x^K(\mathbf{w})$ output of network, y true “label” /target output.

$$v(\mathbf{w}, (x^0, y)) = \frac{1}{2} |y - x^K(\mathbf{w})|^2$$

- Global Predictor: $V(\mathbf{w}) = \mathbb{E}_{\mu_N} v(\mathbf{w}, (x^0, y)) = \mathbb{E} v(\mathbf{w}, (x_I^0, y_I))$, $I \sim \text{Unif}(\{1, \dots, N\})$.
- Learning phase = Calibration

$$\min_{\mathbf{w} \in \mathbb{R}^D} V$$

- (SGD) ? Needs to differentiate V in $\mathbf{w}$ i.e. $v(\mathbf{w}, (x^0, y))$!!

In fact three slides

- Big Question: How to differentiate (and compute!) $v(\mathbf{w}, (x^0, y))$?
- Easy part: from $v(\mathbf{w}, (x^0, y)) = \frac{1}{2}|y - x^K(\mathbf{w})|^2$

($^\top$ for transpose) $\Downarrow$ (J for Jacobian matrix)

$$J_{\mathbf{w}} v(\mathbf{w}, (x^0, y)) = (J_{\mathbf{w}} x^K(\mathbf{w}))^\top (x^K(\mathbf{w}) - y).$$

- Use x^k as an **auxiliary variable** and the induction

$$x^k = \varphi_k(w^{(k)}, x^{k-1}) := [\Psi_k((w_j^{(k)} \mid x^{k-1}))]_{1 \leq j \leq d_k}, \quad k = 1 : K.$$

- Recursive formula: note that $J_{\mathbf{w}} x^k(\mathbf{w})^\top = J_{\mathbf{w}^{(1:k)}} x^k(\mathbf{w}^{(1:k)})^\top$ and

$$\begin{aligned} J_{\mathbf{w}} x^k(\mathbf{w})^\top &= J_{w^{(k)}} \varphi_k(w^{(k)}, x^{k-1}(\mathbf{w}))^\top + J_{\mathbf{w}} x^{k-1}(\mathbf{w})^\top J_x \varphi_k(w^{(k)}, x^{k-1}(\mathbf{w}))^\top \\ &= J_{w^{(k)}} \varphi_k(w^{(k)}, x^{k-1}(\mathbf{w}))^\top \\ &\quad + J_{\mathbf{w}^{(1:k-1)}} x^{k-1}(\mathbf{w}^{(1:k-1)})^\top J_x \varphi_k(w^{(k)}, x^{k-1}(\mathbf{w}))^\top. \end{aligned}$$

Finally ... in four slides

Lemma

In a ring $(A, +, \cdot)$ with regular multiplication, if $(b_k, c_k \text{ being known})$ the sequence $(a_k)_k$ recursively defined

$$a_k = b_k + a_{k-1}c_k, \quad k = 1 : K$$

reads $a_K = b_K + b_{K-1}c_K + b_{K-2}c_{K-1}c_K + \cdots + b_1c_2 \cdots c_K + a_0c_1 \cdots c_K$

which can be computed *recursively in a backward way*.

The backpropagation algorithm ⁽¹⁴⁾ is two-fold.

- A **forward step**: Compute the (**auxiliary!**) values x_j^k at each unit j of each layer k using the weights $w^{(k-1)}$.

- A **backward step**: Apply the lemma to the recursion to compute

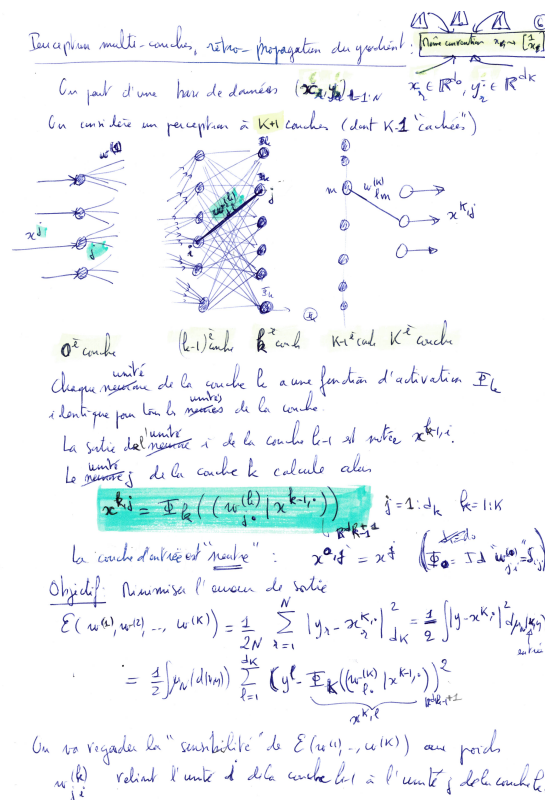
$$J_{w^{(k)}} \varphi_k(w^{(k)}, x^{k-1})^\top \quad \text{and} \quad J_{w^{(1:k-1)}} x^{k-1} (w^{(1:k-1)})^\top$$

- And that's (almost) it!! (see further on).

¹⁴ Rumelhart, D. E., Hinton, G. E., and Williams, R. J. (1986). Learning representations by back-propagating errors. *Nature*, 323, 533–536. Le Cun's PhD thesis: *Modèles Connexionnistes de l'Apprentissage* (Univ. Paris 6, 1987). But goes back to Paul Gilles PAGES (LPSM)

The truth... about Backpropagation of gradient

- **Backpropagation of gradient** is similar to the reverse mode of automatic differentiation developed independently in the seventies!!
- True formulas: very involved but some “invariant” **local errors** terms appear (in yellow) to alleviate computations (cf. Le Cun).



$$\frac{\partial E}{\partial w_{ji}^{(k)}} = \frac{1}{2N} \int p_N(d(y)) \sum_{i=1}^N \left(\Phi_K'(w_{ji}^{(K)} | x_i^{K-1}) - y_i \right) \Phi_K'(w_{ji}^{(K)} | x_i^{K-1}) \frac{\partial w_{ji}^{(K)}}{\partial w_{ji}^{(k)}} \quad (2)$$

$$\frac{\partial E}{\partial w_{ji}^{(k)}} = \sum_{m=1}^{d_K} w_{jm}^{(K)} \frac{\partial x_i^{K,m}}{\partial w_{ji}^{(k)}} = \int p_N(d(y)) \sum_{m=1}^{d_K} \frac{\partial x_i^{K,m}}{\partial w_{ji}^{(k)}} \sum_{l=1}^{d_K} \left(\Phi_K'(w_{jl}^{(K)} | x_i^{K-1}) - y_l \right) \Phi_K'(w_{jl}^{(K)} | x_i^{K-1}) w_{lm}^{(K)}$$

$$= \int p_N(d(y)) \sum_{m=1}^{d_K} \frac{\partial x_i^{K,m}}{\partial w_{ji}^{(k)}} e^{K,l} w_{lm}^{(K)}$$

$$= \int p_N(d(y)) \sum_{m=1}^{d_K} \frac{\partial x_i^{K,m}}{\partial w_{ji}^{(k)}} (e^{K,l} | w_{lm}^{(K)})_{l \in d_K}$$

On a $\frac{\partial x_i^{K-1,m}}{\partial w_{ji}^{(k)}} = \frac{\partial}{\partial w_{ji}^{(k)}} \Phi_{K-1} \left(\sum_{m=1}^{d_{K-1}} w_{jm}^{(K-1)} x_i^{K-2,m} \right)$

$$= \Phi_{K-1}' \left(w_{jm}^{(K-1)} | x_i^{K-2} \right) \sum_{m=1}^{d_{K-1}} w_{jm}^{(K-1)} \frac{\partial x_i^{K-2,m}}{\partial w_{ji}^{(k)}}$$

$$\frac{\partial E}{\partial w_{ji}^{(k)}} = \int p_N(d(y)) \sum_{m=1}^{d_K} \Phi_{K-1}' \left(w_{jm}^{(K-1)} | x_i^{K-2} \right) \sum_{m=1}^{d_{K-1}} w_{jm}^{(K-1)} \frac{\partial x_i^{K-2,m}}{\partial w_{ji}^{(k)}}$$

$$= \int p_N(d(y)) \sum_{m=1}^{d_K} \frac{\partial x_i^{K-2,m}}{\partial w_{ji}^{(k)}} \sum_{m=1}^{d_{K-1}} \Phi_{K-1}' \left(w_{jm}^{(K-1)} | x_i^{K-2} \right) (e^{K,l} | w_{lm}^{(K)})_{l \in d_K}$$

couche $K-1$ à $K-2$

$$\frac{\partial E}{\partial w_{ji}^{(k)}} = \int p_N(d(y)) \sum_{m=1}^{d_K} \frac{\partial x_i^{K-2,m}}{\partial w_{ji}^{(k)}} \sum_{m=1}^{d_{K-1}} \Phi_{K-1}' \left(w_{jm}^{(K-1)} | x_i^{K-2} \right) (e^{K,l} | w_{lm}^{(K)})_{l \in d_K}$$

donc $\rightarrow$ le aléa $\rightarrow$ m.

A une couche k intermédiaire on a la formule de propagation d'erreur

$$e^{K,l}(k-1, k) = \Phi_k' \left(\sum_{m=1}^{d_{k-1}} w_{jm}^{(k)} x_i^{k-1,m} \right) \times (e^{K,l} | e^{K,m})_{m \in d_k}$$

couche d'entrée $\rightarrow$ couche $k-1$ $\rightarrow$ couche k de la couche $k+1$

Lorsque l'on atteint la couche le aléa $\rightarrow$ m, il n'y a plus de couche $k-2$ dans Δ Faire $k = K-2$ dans Δ

$$\frac{\partial E}{\partial w_{ji}^{(k)}} = \int p_N(d(y)) \sum_{m=1}^{d_K} \frac{\partial x_i^{K-2,m}}{\partial w_{ji}^{(k)}} \times (w_{jm}^{(K-1)} | e^{K,m})_{m \in d_K}$$

$$= \Phi_K' \left(\sum_{m=1}^{d_{K-1}} w_{jm}^{(K-1)} x_i^{K-2,m} \right) \times \begin{cases} 0 & \text{si } l \neq j \\ x_i^{K-2,j} & \text{si } l = j \end{cases}$$

d'où finalement

$$\frac{\partial E}{\partial w_{ji}^{(k)}} = \int p_N(d(y)) \Phi_K' \left(\sum_{m=1}^{d_{K-1}} w_{jm}^{(K-1)} x_i^{K-2,m} \right) x_i^{K-2,j} (w_{jm}^{(K-1)} | e^{K,m})_{m \in d_K}$$

où les $e^{K,m}$ suivent la régression rétrograde (x)

À partir de ces formules, on met en œuvre une descente

1) la gradient (GD)

$$w_{ji}^{(k)} \leftarrow w_{ji}^{(k)} - \gamma \sum_{i=1}^N \frac{\partial E}{\partial w_{ji}^{(k)}} \quad \text{à partir de } (K) \text{ (aléa de } (K))$$

$$= w_{ji}^{(k+1)} = w_{ji}^{(k)} - \gamma \sum_{i=1}^N \frac{\partial E}{\partial w_{ji}^{(k)}} \rightarrow \text{calcul de } \frac{\partial E}{\partial w_{ji}^{(k)}} \text{ par la régression rétrograde à chaque pas.}$$

2) le gradient stochastique (SGD)

$$w_{ji}^{(k+1)} = w_{ji}^{(k)} - \gamma \sum_{i=1}^N \Phi_K' \left(\sum_{m=1}^{d_{K-1}} w_{jm}^{(K-1)} x_i^{K-2,m} \right) x_i^{K-2,j} (w_{jm}^{(K-1)} | e_{U_{i,j}}^{K,m})_{m \in d_K}$$

À partir de la propriété d'approximation universelle, Backprop à 1 couche (stat. scalaires)

Théorème (Hornik) Soit $\Phi: \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue, bornée et non constante (il y a $x_0 \neq x_1$ tq $\Phi(x_0) \neq \Phi(x_1)$). Alors

$$\left\{ x \mapsto \sum_{i=1}^N \lambda_i \Phi(w_i(x) + w_0) \right\}_{N \geq 1, \lambda_i \in \mathbb{R}, w_i \in \mathbb{R}^d}$$

est dense dans $C([0,1]^d, \mathbb{R})$. Le résultat reste vrai si l'on remplace $[0,1]^d$ par un compact quelconque de $\mathbb{R}^d$.

Résultat de K. Hornik en 1991 dans "Neural Networks": Approximation capabilities of Multilayer Feedforward Networks.

Tel qu'énoncé, ce théorème démontre la propriété d'approximation universelle

Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Back to paradigm of Learning

- Database $(z_k)_{k=1:N}$, parameters $\theta \in \Theta \subset \mathbb{R}^K$ and (local) loss function/predictor $v(\theta, z)$.
- Let $(I_k)_{k \geq 1}$ be an i.i.d. sequence $\mathcal{U}(\{1, \dots, N\})$ -distributed.
- The stochastic gradient descent reads

$$\theta_{n+1} = \theta_n - \gamma_{n+1} \nabla_{\theta} v(\theta_n, z_{I_{n+1}})$$

where $z_{I_{n+1}}$ means that a datum has been picked up at random in the database uniformly in $\{1, \dots, N\}$.

- Check that

$$\begin{aligned} \mathbb{E} \nabla_{\theta} v(\theta_n, z_I) &= \frac{1}{N} \sum_{k=1}^N \nabla_{\theta} v(\theta_n, z_k) \\ &= \int \nabla_{\theta} v(\theta_n, z) \mu_N(dz) = \nabla V(\theta_n) \end{aligned}$$

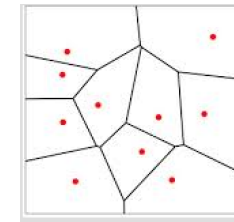
Table of Contents

- 1 Zero search, Optimization (deterministic, the origins)
- 2 Examples from Finance
 - Implicitation
 - Minimization
- 3 Learning procedures
 - Abstract Learning
 - Supervised Learning
 - Unsupervised Learning (clustering)
- 4 Stochastic algorithms/Approximation
 - Paradigm of stochastic approximation
 - From Robbins-Monro to Robbins-Siegmund
 - Application: Stochastic Gradient Descent (SGD) and $S(\text{pseudo-})GD$
- 5 Examples revisited by SGD
 - Numerical Probability
 - Learning theory
- 6 Application to Neural Networks and deep learning
 - Linear neural network
 - One hidden layer feedforward perceptron
 - Universal approximation property
 - Toward deep learning
 - Multilayer feedforward perceptron and Backpropagation
- 7 Unsupervised learning

Unsupervised learning (example of clustering)

- Only input $z_k = x_k \in \mathbb{R}^d$, $k = 1 : N$.
- **Prototype parameter** set: $\theta := (\theta^1, \dots, \theta^r) \in \Theta = (\mathbb{R}^d)^r$, $r \in \mathbb{N}$.
- (An example of) **Local loss function**: **nearest neighbor** among the prototypes: $x \in \mathbb{R}^d$, $\theta \in \Theta$.

$$v(\theta, x) = \frac{1}{2} \min_{i=1:r} |\theta^i - x|^2 = \frac{1}{2} \text{dist}(x, \{\theta^1, \dots, \theta^r\})^2$$



(minimal distance to prototypes).

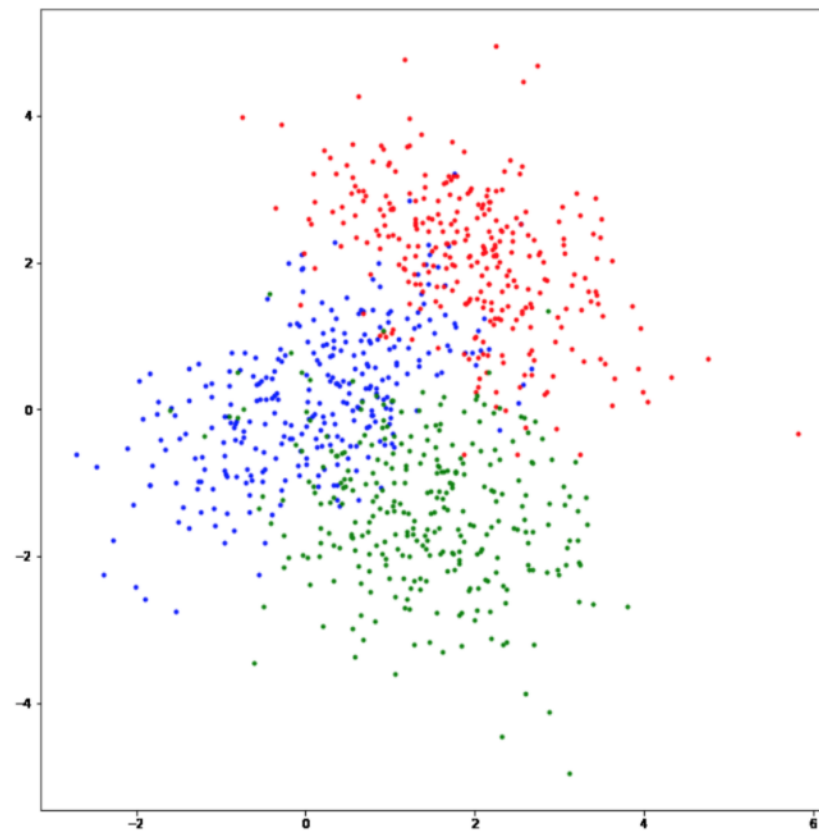
- $v(\theta, x)$ is **not convex** in θ !
- Global loss function (**Distortion**):

$$V(\theta) = \frac{1}{2N} \sum_{k=1}^N \min_{i=1:r} |\theta^i - x_k|^2 \quad (\text{mean minimal distance to prototypes}).$$

- Searching for the **best prototypes**: $\min_{\theta \in (\mathbb{R}^d)^r} V(\theta)$

A data set

The true data set is ...



Batch k -means/Forgey's algorithm

- Gradient at θ s.t. $\theta^i \neq \theta^j$:
$$\nabla V(\theta) = \frac{1}{2N} \sum_{k=1}^N \nabla_{\theta} v(\theta, x_k)$$

with,

$$\forall i = 1 : r, \quad \partial_{\theta^i} v(\theta, x_k) = (\theta^i - x_k) \mathbf{1}_{\{|x_k - \theta^i| < \min_{j \neq i} |x_k - \theta^j|\}} \in \mathbb{R}^d.$$

- Compute the vector of $(\mathbb{R}^d)^r$: $\mathbf{1}_{\{|x_k - \theta^i| < \min_{j \neq i} |x_k - \theta^j|\}}$ = nearest neighbour search.
- Compute $\nabla V(\theta) = \frac{1}{2N} \sum_{k=1}^N \nabla_{\theta} v(\theta, x_k)$.
- $\implies N \times$ nearest neighbour searches among r prototypes of dim d !
- Forgey's algorithm = GD algorithm (or batch GD algorithm):

$$\theta_{n+1} = \theta_n - \gamma_{n+1} \nabla V(\theta_n).$$

CLVQ/ k -means (unsupervised learning)

- Aim:

$$\min_{(\theta^j)_{j=1:r}} \left[V(\theta) = \frac{1}{2} \sum_{k=1}^N \min_{i=1:r} |\theta^i - x_k|^2 \right]$$

(mean minimal distance to prototypes).

- Competitive Learning Vector Quantization: let $\gamma_n \in (0, 1)$, $n \geq 1$, then

$$\theta_{n+1}^i = \begin{cases} \theta_n^i - \gamma_{n+1}(\theta_n^i - x_{n+1}) & \text{if } |x_{n+1} - \theta_n^i| < \min_{j \neq i} |x_{n+1} - \theta_n^j| \\ \theta_n^i & \text{otherwise} \end{cases}$$

- In other words: $\rightarrow n + 1$ reads
 - Nearest neighbour search to the datum among r prototypes of dimension d .
 - Moving the winner by a dilatation centered at the datum with ratio $1 - \gamma_{n+1} > 0$.

The true data set is ...

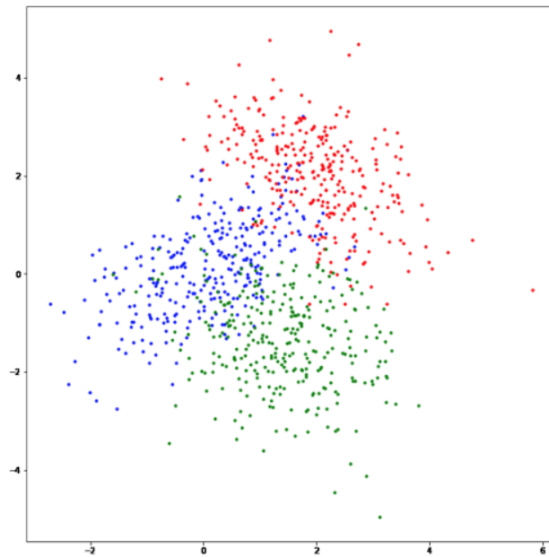


Figure: A mixture of 3 normal distributions.

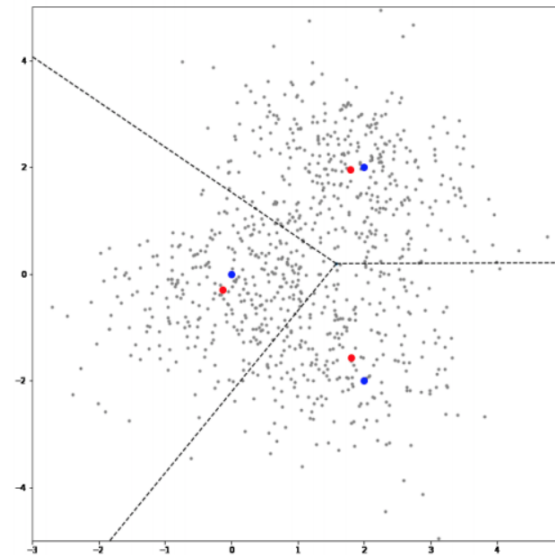


Figure: The red points are an optimal quantizer of size 3 of this data set. The blue points are the three centers of the normal distributions.

Figure: The CLVQ algorithm