

Deriving Δ rule - $\Delta w_i = -\epsilon \frac{\partial E}{\partial w_i}$ $y_N = \sum w_i x_{in}$ $E = \frac{1}{2} \sum_{n=1}^N (y_n - \hat{y}_n)^2$ $\frac{\partial E}{\partial w_i} = \sum (y_n - \hat{y}_n) \frac{\partial \hat{y}_n}{\partial w_i} = x_{in}$

CE $\frac{\partial E}{\partial w_k} = \partial - \sum t_k \ln h_k = \frac{\partial}{\partial w_k} [-t_k \ln(h_k) - (1-t_k) \ln(1-h_k)] = -t_k/h_k + (1-t_k)/(1-h_k) = (h_k - t_k)/h_k(1-h_k)$

Soft Max $\frac{\partial h_k}{\partial w_k} = \frac{\partial h_k}{\partial \exp(y_k)} \frac{\partial \exp(y_k)}{\partial y_k} = e^{y_k} [\sum_j e^{y_j}]^{-1} - e^{y_k} [\sum_j e^{y_j}]^{-2}] = \frac{e^{y_k}}{\sum_j e^{y_j}} [1 - \frac{e^{y_k}}{\sum_j e^{y_j}}] = h_k(1-h_k)$; $\partial - \sum_k \frac{\partial y_k}{\partial w_k} \log h_k = h_k - t_k > \text{SSE from Linear}$

Sigmoid $\frac{\partial h_j}{\partial y_j} = \frac{\partial (1/(1+\exp(-y_j)))}{\partial y_j} = -(1+\exp(-y_j))^{-2} (-\exp(-y_j)) = \exp(-y_j)/(1+\exp(-y_j))^2 = h_j(1-h_j) \rightarrow$ low (linear) capacity

linear $\frac{\partial y_j}{\partial w_{ij}} = \partial b + \sum_j x_j w_{ij} / \partial w_{ij} = y_j$

$Z = \sum x_i w_i, Z > 0 \Rightarrow y=1$ binary threshold
 $Z = \sum x_i w_i, -\theta < Z < \theta \Rightarrow y=0$ perceptron

Note: $\sigma(k) = \frac{1}{1+\exp(-k)}$

Softmax $y_i = e^{x_i} / \sum_j e^{x_j}$

Dist rep = Each concept rep by many neurons, each neuron partic. in rep of many concepts.

- Coarse Coding Saves #fine neurons / # Coarse = r^{k+1} $r = \text{Coarse grids, } N \text{ neurons } N \text{ vs. } N r^2 \rightarrow w_1, w_2 \geq \theta, w_1, w_2 \leq \theta \Rightarrow 0 \text{ N.L.S.}$

Replicate Features = Equiv., subsampling = Invar (local Avg) $\rightarrow$ set of weights $\vec{w}$: linearly separable

Convolut. NN $\rightarrow$ PERC. CONV $\rightarrow$ $\checkmark \Rightarrow$ leave weights alone $[\otimes \pm \text{subtract input}, 0 \Rightarrow \text{add input}] \Rightarrow (\text{weight} - \text{satisfactory } \vec{w})^2 \downarrow$

Speed Learning $\rightarrow$ Adaptive learning rate (global or separate for w_i) $\rightarrow$ Momentum $[\Delta w(t) = \alpha \Delta w(t-1) - \epsilon \frac{\partial E}{\partial w}(t)]$

$\rightarrow$ Stochastic GD $\rightarrow$ Conjugate gradient or $-\frac{\partial E}{\partial w_i} H^{-1} [\odot \rightarrow \odot]$ Prevent $\frac{\partial E}{\partial w_i} \rightarrow E[y - t]^2 = E[y + \sum w_i \epsilon_i - t]^2 = E[(y - t) + \sum w_i \epsilon_i]^2$

$= (y - t)^2 + \sum w_i \sigma_i^2 \rightarrow \frac{\partial E}{\partial w_i} = \frac{\partial E}{\partial w_i} + x w_i$ Bayes = $-\log P(W|D) = -\log P(D|W) - P(W) + P(D) = \frac{1}{2\sigma^2} \sum (y_c - t_c)^2 + \frac{1}{2\sigma^2} \sum w_i^2$

$= E + \frac{\sigma^2}{2} \sum w_i^2$ [Assume N prior + prediction] (MAKAY & DM) Guess σ_D^2, σ_W^2 - Learn - Set $\sigma_D^2 = \frac{1}{N} \sum (t_i - y_i)^2, \sigma_W^2 = \frac{1}{N} \sum w_{ij}^2$

- Repeat. Ass. $P(t_c | y_c) = \frac{1}{\sqrt{2\pi}\sigma} \exp(-\frac{1}{2\sigma^2}(t_c - y_c)^2), P(w) = \frac{1}{\sqrt{2\pi}\sigma} \exp(-\frac{w^2}{2\sigma_w^2})$ (Min Error = Max log prob.) $\frac{\partial}{\partial w_i} \frac{\partial E}{\partial w_i} = 0$

WMAP = $\arg \max_w [P(W|D)] = \arg \max_w [P(w)P(D|w)] = \arg \max_w [P(D|w)] = w_{ML}$; BAYES: $P(t | \text{input}) = \sum_g P(w_g | D) P(t | \text{input} + w_g)$

MIN Sq Err $\equiv$ MAX log prob $\sim N(y_c = \mu) [y_c = \text{models guess}]$ $\hat{t}_i^{(MAP)} = P(t_i | x_i, w_{MAP})$ (Likelihood term)

Combining Networks - use AVG of predictors $y_i(\vec{y})$. Data = $t_c y_c \rightarrow$ mix of "regimes", use Sep. Model for each.

$\rightarrow$ Cluster Cases into regimes, not inputs. Make y_i disagree: diff/ learning algorithms, bagging (resample w/ replacement), Boosting

$\rightarrow$ Fit next Model w/ data reweighted by how poorly prev. Model fit case.

M datapoints

ADABOOST $\rightarrow H_{\text{final}}$ (Strong Classifier) = $\text{Sgn}(\sum \alpha_t [\text{weight of rule of thumb}] h_t(x) [\text{Weak Classifier/perceptron}])$ ① Given $(x_i, y_i) \dots$

$y_i \in \{-1, 1\}$, for $t = [1, T]$, Make D_t on $[1, M]$: more weight on Misclass $|x_i$. Find $h_t(x) \rightarrow \{-1, 1\}$ w/ $\epsilon_t = P(h_t(x_i) \neq y_i) < 0.5$

Learning $\rightarrow$ Supervised $\rightarrow P(y|x)$ Unsup $\rightarrow P(x)$ $\rightarrow$ Give low Energy to data, $\downarrow R \Rightarrow \uparrow$ learning rate

PCA $\rightarrow X_n \rightarrow X_m$ ($m < n$) m in which data has most variance. Use mean for other $m-n$ Dim. (Autoenc.) $\rightarrow$ linear NN w_i span R^m ; or Nonlin NN $\rightarrow$ w/ encoding Hidunit, bottleneck, N input + output units.

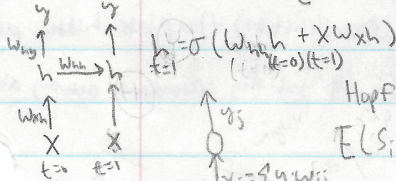
Clustering $\rightarrow$ Force Hidunit w/ $\vec{w}_i$ closest to $\vec{\text{input}}$ to have activity 1, $w_i = \text{Cluster } i$. $\frac{\partial E}{\partial w_{\text{input} \rightarrow \text{hidden}}} = 0$, Backprop w/ $\frac{\partial E}{\partial w_{\text{hidden} \rightarrow \text{output}}} \equiv \downarrow \text{SSE} \equiv \text{Centre of G.}$

Kmeans $\rightarrow$ Ass. K clusters, random ① Assign each data point to closest μ [$\downarrow$ SSE] ② Move μ to Centre of G of assigned data ($\downarrow$ SSE) ③ Repeat until -- SSE. $\rightarrow$ Try non-local split-and-merge moves, Soft K

Means $\rightarrow$ Each Cluster has diff/ responsibility for each datapoint. $\rightarrow P(x^c \text{ is in } i\text{th cluster})$ $\rightarrow$ mixing proportion π_i

EM - use Validation Set or Kmeans to determine # Gaussians to use. (Estep) $\rightarrow$ Get $P(i|x^c) \equiv P(i)P(x^c|i)/P(x^c)$

(Mstep) - Use to update π_i, μ_i, σ_i^2 [ith N, Cth datapoint, $\in R^d$] $\pi_i^* = \sum_{c=1}^C P(i|x^c)/N$ $\mu_i^* = \sum_{c=1}^C P(i|x^c) x^c / \sum_{c=1}^C P(i|x^c)$
 $\sigma_i^{2*} = \sum_{c=1}^C P(i|x^c) \|x^c - \mu_i^*\|^2 / \sum_{c=1}^C P(i|x^c)$ Note: $P(x^c) = \sum_j P(j)P(x^c|j), [\sum_j P(j) = 1]; P(x^c|i) = \prod_{j=1}^D \frac{\exp(-\frac{1}{2\sigma_{ij}^2}(x_j^c - \mu_{ij})^2)}{\sqrt{2\pi}\sigma_{ij}}$



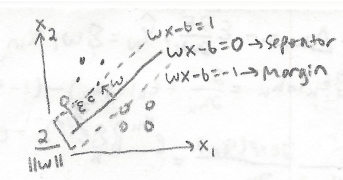
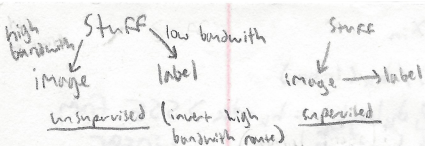
Hopfield Net: $E = - \sum_i S_i b_i - \sum_{i,j} S_i S_j w_{ij}$ [Note Symmetric weights] Change One $S_i, E(S_i=0) - E(S_i=1) = b_i - \sum_j S_j w_{ij}$ $\Delta w_{ij} = \eta (S_i - \frac{1}{2})(S_j - \frac{1}{2})$ or $w_{ij} = S_i S_j, S_i \in \{-1, 1\}$. Let Net settle then

Wlearn = $\downarrow$ spurious minima. Simulated Annealing: Start w/ lots of noise to $\downarrow$

E landscape, and $\uparrow$ cross E barriers, then $\downarrow$ Noise = Sharpen E landscape, $P(S_i=1) = 1/(1 + e^{-\Delta E_i/T})$

[Using Stochastic Binary Threshold units.] Thermal Equilibrium - prob dist of Configs settled down

At low T, prob of good states $\gg$ bad states. $P_A/P_B = \exp(-(E_A - E_B)/T)$.



Choose hidden -> get visible
 $P(\vec{s}^a) = \sum_{m \in \text{models}} \pi_m \prod_j (s_j^a p_j^m + (1-s_j^a)(1-p_j^m))$ [if $s_j^a = 1$ if $s_j^a = 0$]
 P_i = Prob of S_i^a being in training set

Energy based = Boltz Mixture Model = local = * Neurons exponential in $N_i = 1$ Model for each Combo -> want dist representation.

$P(V) = \sum_h P(h)P(V|h)$ Causal, generative model, (MOG) Pick π_i then generate image. Boltzmann Mach - E of joint config $S = \langle S_1, \dots, S_k, S_{k+1}, \dots, S_n \rangle$
 prob joint config $P(V, h) = e^{-E_{vh}} / \sum_{h,v} e^{-E_{vh}}$, $P(V) = \sum_h e^{-E_{vh}} / \sum_{h,v} e^{-E_{vh}}$ Cannot compute so use MCMC Sampling.

for training vecs $\{V_1, \dots, V_n\}$ Maximize $\prod P(V_i) = \sum \log(P_i) = \log \sum e^{-E_{vh}} - \log \sum_{h,v} e^{-E_{vh}}$ and $\frac{\partial \log P(V)}{\partial w_{ij}} = \langle S_i S_j \rangle_v - \langle S_i S_j \rangle_{\text{free}}$
 $P(V) = \sum_h e^{-E_{vh}} / \sum_{h,v} e^{-E_{vh}}$
 $\frac{\partial \log P(V)}{\partial w_{ij}} = \frac{1}{\sum_{h,v} e^{-E_{vh}}} \sum_{h,v} \frac{\partial \exp(-E_{vh})}{\partial w_{ij}} = \frac{1}{\sum_{h,v} e^{-E_{vh}}} \sum_{h,v} e^{-E_{vh}} \frac{\partial (-E_{vh})}{\partial w_{ij}} = \frac{1}{\sum_{h,v} e^{-E_{vh}}} \sum_{h,v} e^{-E_{vh}} (-S_i S_j) = \langle S_i S_j \rangle_v - \langle S_i S_j \rangle_{\text{free}}$
 Note $E(V, h) = - \sum_i S_i b_i - \sum_{i,j} S_i S_j w_{ij}$
 $\frac{\partial \log P(V)}{\partial w_{ij}} = \langle S_i S_j \rangle_v - \langle S_i S_j \rangle_{\text{free}}$ unlearning, $\downarrow$ prob of non data vecs.
 $\downarrow$ E (prob of states out TE when vector clamped to vis units)

Learning = +ve phase clamp V to V_{vis} , let hid reach TE, Sample $S_i S_j$. Repeat $\forall V_i \in \text{Training Set}$. -ve phase Do not clamp, get to TE, Gibbs

Sample $S_i S_j$. Sampling -> Start w/ training vec -> update hidden in parallel -> the visible (reconstruction) -> update hidden Continuous divergence
 $\Delta w_{ij} = E(\langle S_i S_j \rangle_v - \langle S_i S_j \rangle_{\text{free}}) = \frac{1}{N} \sum_{i,j} [V_i^h \hat{V}_j^h - \hat{V}_i^h \hat{V}_j^h]$ Wake-Sleep algorithm wake: use input data to get $h_j = 1/(1 + \exp(-\sum_i V_i w_{ij} - b_j))$
 the $h_j \leftarrow 1$ if $h_j > \text{rand}(0,1)$, 0 o.w. Sleep use h_j , get -ve data $\hat{V}_i = 1/(1 + \exp(-\sum_j w_{ij} h_j - b_i))$; use $\hat{V}_i$ get -ve hid activities
 $\hat{h}_j = 1/(1 + \exp(-\sum_i \hat{V}_i w_{ij} - b_j))$. $\Delta b_i = \frac{1}{N} \sum_{i,j} [V_i^h - \hat{V}_i^h]$, $\Delta b_j = \frac{1}{N} \sum_{i,j} [h_j^h - \hat{h}_j^h]$. $P(V) = \sum_h P(V|h)P(h)$ Train RBM to learn meaningful

Patterns of activation. Training case info = $\log(\# \text{ possible labels}) \Rightarrow$ Need lots of labelled data. $\uparrow$ layers $\Rightarrow$ Tor $\downarrow \frac{\partial E}{\partial w_{ij}}$ Soln $\Rightarrow$ Greedy unsupervised learning

the Fine-tune backprop. Equivariance -> Change in view = Change Neural Activity. 1) Place Coded = Discrete change in property of
 Deep autoencoder. Visual Entity = Discrete change in which Neuron used for Encoding 2) Rate Coded = Real valued change in Out-put, but not in which Neurons used. -> Capsule: Info = prob $\exists$ Object, Pose of Object $\in \mathbb{R}^n$ Steps

1) Compute Capsule Outputs for Image 2) Apply transformation to output of Capsule 3) Predict transformed image from transformed Output from Capsule. Invariance -> Removes Spatial relationships. SVM - Choose Model, Structural risk Minimization
 Error = train E + f(Size of train set, VC, prob bound fails); VC = How many data can learn $\forall$ label permutations. hyper-plane $\in \mathbb{R}^K$ VC = K+1
 $E_{\text{test}} < E_{\text{train}} + ((h + h \log(N/h) - \log(P/4))/N)^{1/2}$ TRAIN $w \cdot x + b > 1 - \epsilon^c$ (true cases), $< -1 + \epsilon^c$ (false cases) $\|w\|^2/2 + \lambda \sum_i \epsilon^c$ small, $\epsilon^c > 0$

Slack Variable. [Finding MaxMargin Separator Using SVs] If No Separating plane = $\uparrow$ Set of Features, Slack vars. - Can put plane closer to data then the margin = Min dist to any data, $\uparrow \Rightarrow \downarrow$ VC. Input vecs -> high D space -> Kernel trick to avoid dot prod.
 $b + \sum_{s \in SV} w_s K(x^{\text{test}}, x^s) > 0 \Rightarrow$ Select: SV to max margin, Compute w for each SV, Choose good Kernel function. * You Only do this
 $K(x_1, x_2) = \phi(x_1) \cdot \phi(x_2)$. SVs = Samples on margin.

MOG Generative Model -> pick Gaussian w/ prob π_i , generate random point from Gaussian.
 Stacked RBMs -> 1) learn layer of hidden features 2) treat feature activation as data, learn 2nd layer.

weights in RBM define $P(V_i=1|h) = 1/(1 + e^{-\sum_j w_{ij} h_j})$ sample based on this dist.

Generate data from RBM -> 1) top RBM -> TE 2) Single top-down pass to get states of lower layers.

CD -> 1) Start w/ Train vec on Vis units. 2) update hidden in parallel 3) update visible (reconstruction) 4) update hidden again

Wake Sleep -> output [wakes] Modified to $\uparrow$ prob they reconstruct activity below
 input data -> hid acts -> -ve data -> -ve hid acts.
 output [sleep] Modified to $\uparrow$ prob they reconstruct activity above
 input -> $\hat{h}_k = \sigma(\sum_j w_{jk} \hat{h}_j + b_k)$ $\hat{h}_j = \sigma(\sum_k w_{jk} \hat{h}_k + b_j)$ $\hat{V}_i = \sigma(\sum_j w_{ij} \hat{h}_j + b_i)$ $\Delta w_{ij} = E \hat{V}_i (\hat{h}_j - \sigma(\sum_k w_{ik} \hat{h}_k + b_i))$
 Wake = generate $h_j = \sigma(\sum_i w_{ij} V_i + b_j) > \text{rand}(0,1)$ $\Delta w_{ij} = E h_j (V_i - \sigma(\sum_k w_{ik} h_k + b_i))$, $h_k = \sigma(\sum_j w_{jk} h_j + b_k) > \text{rand}(0,1)$ $\Delta w_{kj} = E h_k (h_j - \sigma(\sum_i w_{ij} \hat{h}_j + b_j))$
 Sleep = $\hat{h}_k = \sigma(\sum_j w_{jk} \hat{h}_j + b_k)$ $\hat{h}_j = \sigma(\sum_k w_{jk} \hat{h}_k + b_j)$ $\hat{V}_i = \sigma(\sum_j w_{ij} \hat{h}_j + b_i)$ $\Delta w_{ij} = E \hat{V}_i (\hat{h}_j - \sigma(\sum_k w_{ik} \hat{h}_k + b_i))$
 MoG = κ . Soft K-means.