

Logistic regression

A probabilistic classification model

$$\{(x_i, y_i)\}_{i=1}^N \quad x_i \in \mathbb{R}^D \quad y_i \in \{0, 1\}$$

We just saw generative models
e.g. $P(x|y=c, \theta) = \mathcal{N}(x|\mu_c, \Sigma_c)$

We then compute Posterior $P(y|x, \theta)$
 $\propto P(y) P(x|y, \theta)$
to classify.

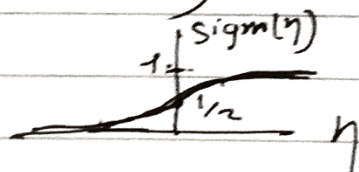
One might want to have
discriminative models directly
giving us the probability $P(y|x, \theta)$

We saw that two class Gaussian
model with tied covariance matrices
had a posterior class probability given
by a sigmoidal / logistic function
of a linear combination of ~~the~~ variable.
We could try to ~~learn~~ learn such ~~model~~ model
directly from labeled data.

$$P(y|x, w) = \text{Ber}(y | \text{Sigm}(w^T x))$$

remember

$$\text{Sigm}(\eta) = \frac{1}{1 + e^{-\eta}}$$



BTW, if we need to add a constant shift, say $P(y|x, \theta) = \frac{1}{1 + e^{-(w^T x + w_0)}}$

We could just define bigger vector

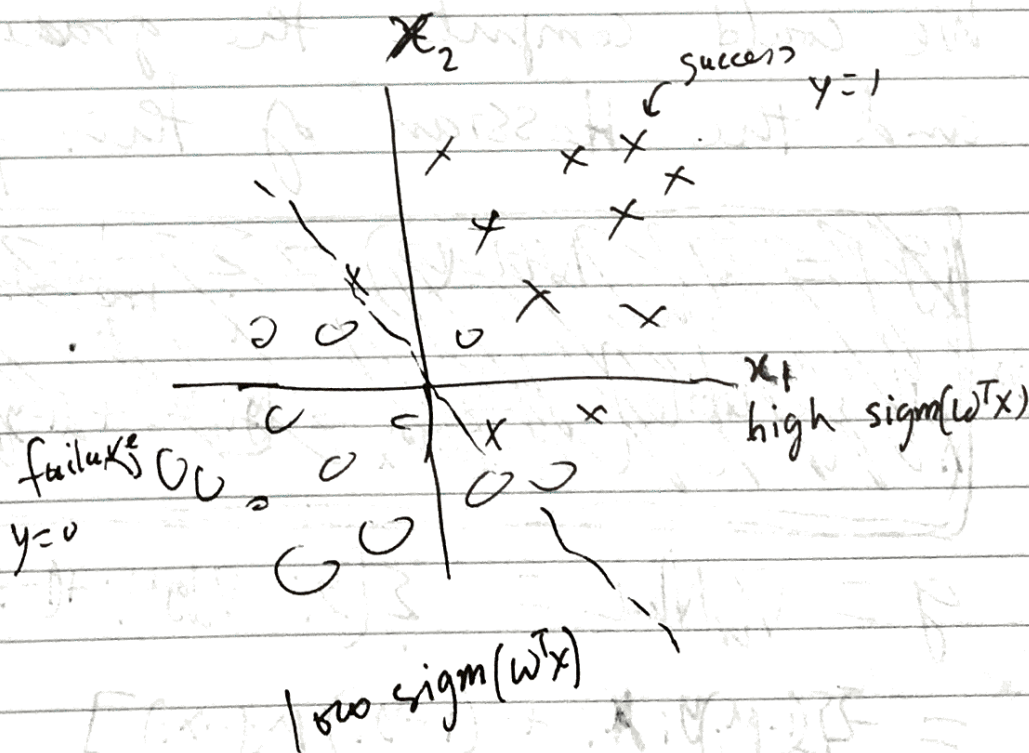
$$w \rightarrow (w_0, w) = \tilde{w}$$

$$x \rightarrow (1, x) = \tilde{x}$$

$$P(y|x, \theta) = \frac{1}{1 + e^{-\tilde{w}^T \tilde{x}}}$$

$$\prod_i P(y_i | x_i, w) = \prod_i \left[\mu_i^{I(y_i=1)} (1-\mu_i)^{I(y_i=0)} \right]$$

where $\mu_i = \frac{1}{1 + e^{-w^T x_i}}$



MLE Negative Log likelihood

$$NLL(w) = - \sum_i [y_i \log \mu_i + (1-y_i) \log(1-\mu_i)]$$

$$\text{By the way } 1-\mu = \left(1 - \frac{1}{1+e^{-w^T x}}\right) = \frac{e^{-w^T x}}{1+e^{-w^T x}}$$

$$\text{So, if we use } \tilde{y}_i = \pm 1 = 2y_i - 1 = \frac{1}{e^{w^T x} + 1}$$

$$NLL(w) = - \sum_i \log \frac{1}{1 + e^{\tilde{y}_i w^T x_i}}$$

$$= \sum_i \log(1 + e^{\tilde{y}_i w^T x_i})$$

Unfortunately we can't write down $\hat{w}$ in closed form
We could compute the gradient
and the Hessian of this function

$\nabla_w \mu_i = \mu_i(1-\mu_i)x_i$

~~$$NLL(w) = \sum_i \log(1 + e^{\tilde{y}_i w^T x_i})$$~~
~~if we write $\frac{1}{1+e^{w^T x}} = \frac{e^{-w^T x}}{1+e^{-w^T x}} = \frac{e^{-w^T x}}{e^{-w^T x} + 1}$~~

$$g = \nabla_w NLL = - \sum_i [y_i \nabla_w \log \mu_i + (1-y_i) \nabla_w \log(1-\mu_i)]$$
$$= - \sum_i [(1-\mu_i)y_i x_i + (1-y_i)\mu_i(-x_i)]$$

$$= \sum_i (\mu_i - y_i) x_i = X^T (\mu - y)$$

$$H = \nabla_{\omega} g(\omega)^T = \sum_i (\nabla_{\omega} \mu_i) x_i^T$$

$$= \sum_i \mu_i (1 - \mu_i) x_i x_i^T$$

$$= X^T S X$$

$$\text{where } S \triangleq \text{diag}(\mu_1(1-\mu_1), \dots, \mu_N(1-\mu_N))$$

$$X = \begin{pmatrix} x_{11} & \dots & x_{1D} \\ \vdots & & \vdots \\ x_{N1} & \dots & x_{ND} \end{pmatrix}$$

$$X^T = \begin{pmatrix} x_{11} & \dots & x_{N1} \\ \vdots & & \vdots \\ x_{1D} & \dots & x_{ND} \end{pmatrix}$$

$$\mu - y = \begin{pmatrix} \mu_1 - y_1 \\ \vdots \\ \mu_N - y_N \end{pmatrix}$$

$$S = \begin{pmatrix} \mu_1(1-\mu_1) & & \\ & \mu_2(1-\mu_2) & \\ & & \ddots \\ & & & \mu_N(1-\mu_N) \end{pmatrix}$$

How do we minimize $NLL(w)$?

One thought: ~~Steepest~~ Steepest descent

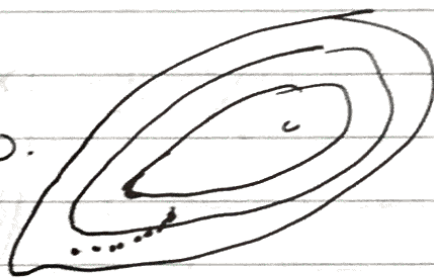
$$\theta_{k+1} = \theta_k - \eta_k g_k$$

Issues with step size

Big η
Potential oscillation



Small η
Maybe too slow.



There are many methods to deal with issues of stepsize selection. We describe ^{one of} the popular methods in the context of ~~the~~ logistic regression.

Newton's method

Initialize θ_1 :

for $k=1, 2, \dots$ until convergence do

Evaluate $g_k = \nabla f(\theta_k)$

Evaluate $H_k = \nabla^2 f(\theta_k)$;

Solve $H_k d_k = -g_k$;

Use line search to find step size η_k along d_k ;

$$\theta_{k+1} = \theta_k + \eta_k d_k;$$

Roughly
$$f(\theta + d) = f(\theta) + g \cdot d + \frac{1}{2} d^T H d + \dots$$

The second order approximation is not minimized when $g + H d = 0$ (assuming H_k is positive definite)

To avoid overstepping one uses line search in the direction of d .

For one dimension



For logistic regression with $\eta_k = 1$

$$\begin{aligned} W_{k+1} &= W_k - H_k^{-1} g_k \\ &= W_k + (X^T S_k X)^{-1} X^T (Y - M_k) \end{aligned}$$

$$\begin{aligned} \textcircled{2} &= (X^T S_k X)^{-1} [(X^T S_k X) W_k + X^T (Y - M_k)] \\ &= (X^T S_k X)^{-1} X^T S_k [X W_k + S_k^{-1} (Y - M_k)] = (X^T S_k X)^{-1} X^T S_k z_k \end{aligned}$$

where $z_k \triangleq X W_k + S_k^{-1}(y - \mu_k)$
which is called the working response

Note that W_{k+1} optimizes

$$\sum_{k_i} S_{k_i} (z_{k_i} - W^T x_i)^2$$

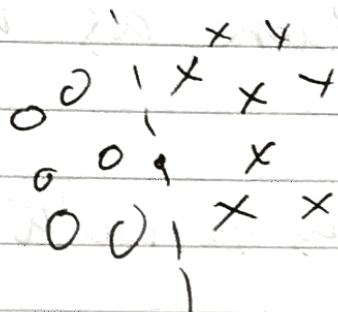
$$\nabla_W \sum S(\cdot)^2 = 0 \Rightarrow \sum_{k_i} \left[x_i^T S_{k_i} z_{k_i} - x_i^T S_{k_i} x_i W \right] = 0$$

$$z_{k_i} = W_k^T x_i + \frac{y_i - \mu_{k_i}}{\mu_{k_i}(1 - \mu_{k_i})}$$

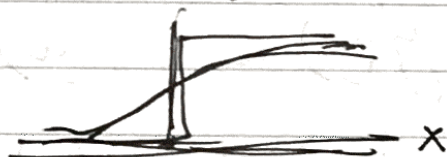
Iterating recomputation of z_{k_i}
with least square minimization
gives iteratively reweighted
least square (IRLS) algorithm.

Now let us discuss issues of
regularization. What if training
data ~~was~~ for the two classes

were linearly separable, ~~by~~ by chance?

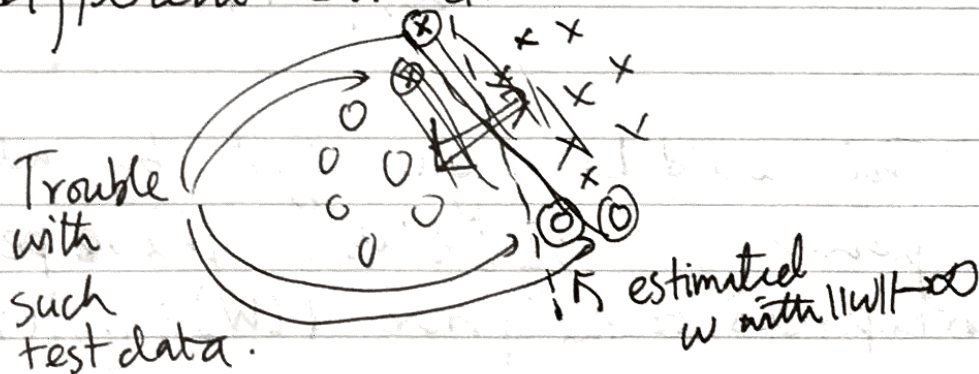


In that case,
once the direction
is established,
 $\|w\| \rightarrow \infty$



That way all $y=1$ have probability 1,
and all $y=0$ have probability 0,
maximizing the likelihood $\prod_{i:y=1} 1 \prod_{i:y=0} (1-0)$

This is brittle. Let ~~us~~ us say the
true model has finite w and a slightly
different direction



One option is to add an l_2 regularization as in ridge regression

$$f'(w) = NLL(w) + \lambda w^T w$$

$$g'(w) = g(w) + 2\lambda w$$

$$H'(w) = H(w) + 2\lambda I$$

Also, one could generalize logistic regression to multiclass ~~setting~~ setting.

$$P(y=c|x, w) = \frac{\exp(w_c^T x)}{\sum_{c'=1}^C \exp(w_{c'}^T x)}$$

This could be regularized by adding a penalty $\frac{1}{2} \sum_c w_c^T v_0 w_c$. This penalty corresponds to the prior $P(W) = \prod_c N(w_c | 0, v_0)$

As we will see, just regularizing may not be enough!

Bayesian logistic regression

As we discussed before, instead of using the mode of $p(w|D)$, namely MAP estimate, we might want to integrate over w .

Let us take a more general look at Bayesian integration in a particular approximation that becomes accurate in the large data size limit.

Think of $P(\theta|D) = \frac{e^{-E(\theta)}}{Z}$ with

$$Z = \int d\theta' e^{-E(\theta')}. \quad \text{If } E(\theta) = -\log P(\theta|D)$$

$$Z = P(D) \triangleq \text{evidence}.$$

Let us say θ^* is the MAP estimate (lowest energy $\triangleq$ highest posterior prob.)

$$E(\theta) = E(\theta^*) + (\theta - \theta^*)^T g + \frac{1}{2}(\theta - \theta^*)^T H (\theta - \theta^*) + \dots$$

If θ^* is the mode, $g = 0$

$$g \in \mathbb{R}^D \\ H \in \mathbb{R}^{D \times D}$$

$$Z = \int e^{-E(\theta)} d\theta = e^{-E(\theta^*)} \int d\theta e^{-\frac{1}{2}(\theta - \theta^*)^T H (\theta - \theta^*)}$$

$$= e^{-E(\theta^*)} (2\pi)^{D/2} |H|^{-1/2}$$

This approx is known as the Laplace approximation.

$$\log P(D) = -E(\theta^*) - \frac{1}{2} \log |H| + \text{constants}$$

$$= \boxed{\log P(D|\theta^*)} + \log P(\theta^*) - \frac{1}{2} \log |H| + \text{const.}$$

$$= \log P(D|\theta^*) + \log P(\theta^*) - \frac{1}{2} \log |H| + \text{const.}$$

Related
to MLE

Related to
MAP

Contribution
of Bayesian
integration in Laplace
approximation

Relation to Bayesian information criterion (BIC)

N pieces of data: $H = \sum_{i=1}^N H_i$, $H_i = \nabla \nabla \log p(D_i|\theta)$

$\boxed{\text{If each } H_i \approx \hat{H}, H \approx N \hat{H}}$

$$\log |H| \approx \log |N \hat{H}| = \log N^D |\hat{H}|$$

$$\log P(D) \sim \underbrace{\log P(D|\theta^*)}_{\sim N} + \underbrace{\log P(\theta^*)}_{\sim 1} - \frac{D}{2} \underbrace{\log N}_{\sim \log N} - \frac{1}{2} \underbrace{\log |\hat{H}|}_{\sim 1} + \dots$$

$$\sim \log P(D|\theta^*) - \frac{D}{2} \log N$$

For comparing models, the term $-\frac{D}{2} \log N$

~~is~~ penalizing more complex models -

(~~the~~ ^{higher} $n \log P(D|\theta^*)$ achieved by bigger D offset by $-\frac{D}{2} \log N$)

Akaike information criterion (AIC)
considers $2D - 2 \log P(D|\theta^*)$
 $= -2(\log P(D|\theta^*) - D)$

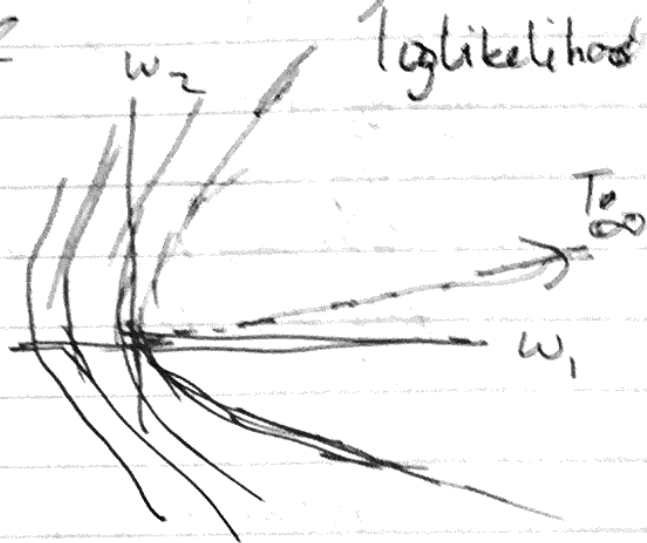
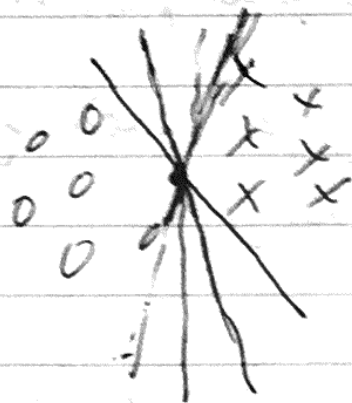
BIC depends on data size. Since $\frac{\log N}{2}$ is often greater than 1, BIC penalizes complex models more severely.

[Many prefer writing ~~$kD - 2 \log P(D|\theta^*)$~~ $kD - 2 \log P(D|\theta^*)$ and minimize it.

$k = \log N$ for BIC and $k = 2$ for AIC]

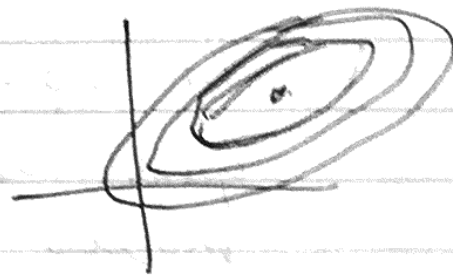
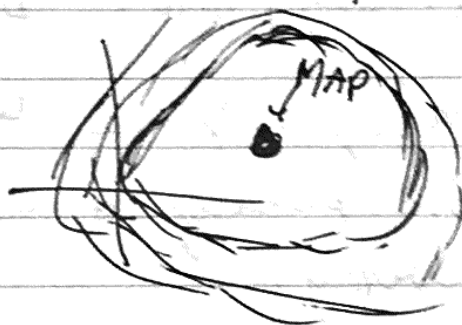
Back to the example of perfectly separable case. If data comes really

(a priori)
 from p distr + logistic regression calling
 the class (we will do a two class example,
 we will see this phenomenon only when
 N is not too large



log Posterior
 (with quadratic penalty)

Laplace
 approximation



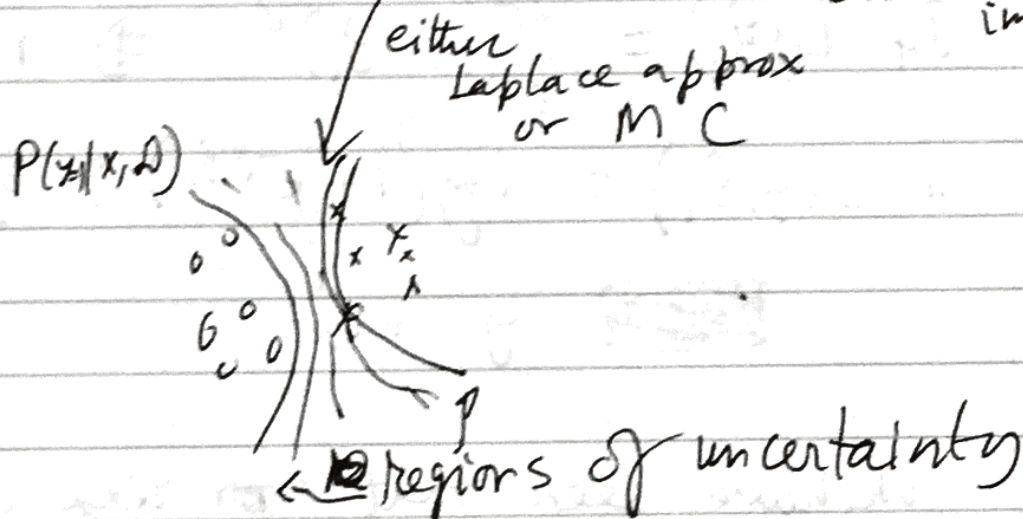
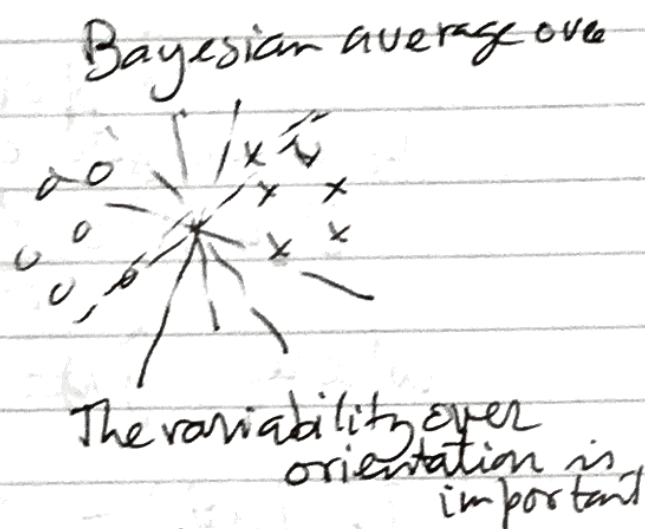
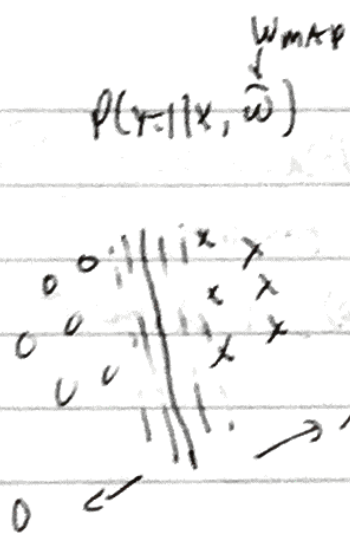
Approximating posterior predictive

$$P(y|x, \mathcal{D}) = P(y|x, w) P(w|\mathcal{D}) \approx P(y|x, \hat{w})$$

$\hat{w} \rightarrow w_{\text{MAP}} \text{ or } E[w]$

Alternative: Sample w from the posterior
 Monte Carlo (MC)

$$P(y|x, \mathcal{D}) \approx \frac{1}{S} \sum_{s=1}^S \text{sign}(w^{(s)} x)$$



Online logistic regression and perceptrons

For optimizing $f(\theta) = \frac{1}{N} \sum_{i=1}^N f(\theta, z_i)$ $\leftarrow -\log p(y_i|x_i, \theta)$

one could take an online approach: change θ in response to new data.

Grad descent $\theta_k = \theta_{k-1} - \eta_k g_k$

for logistic regⁿ this means

$$w_k = w_{k-1} - \eta_k (y_i - \hat{y}_i) x_i \quad y_i \in \{0, 1\}$$

If we consider using $y_i \in \{-1, 1\}$, and use $\hat{y}_i = \hat{y}_i - 1 \approx \pm 1$
and $\eta_k = \frac{1}{2} \eta_k$

$$w_k = w_{k-1} + \eta_k y_i x_i \quad \text{when } y_k \neq \hat{y}_k$$

$$= w_{k-1} \quad \text{otherwise}$$

Perceptron algorithm!

Support vector machines (SVM's)
 L_2 regularized empirical risk function

$$J(\omega, \lambda) = \sum_{i=1}^N L(y_i, \hat{y}_i(\omega)) + \lambda \|\omega\|^2$$

Where $\hat{y}_i(\omega) = \omega^T x_i + \omega_0$

If $L(y, \hat{y}_i) = \sum_i (y_i - \hat{y}_i)^2 \Rightarrow$ ridge regression

" $L(y, \hat{y}_i) = \log(1 + e^{-y \hat{y}_i}) \Rightarrow$ logistic regression

~~$y = +1, -1$~~ $y = +1, -1$

We have seen that for ridge regression

$$\hat{\omega} = (X^T X + \lambda I_N)^{-1} X^T Y$$

$$\hat{y} = \hat{\omega}^T x = Y^T X (X^T X + \lambda I_N)^{-1} x = Y^T (X X^T + \lambda I_N)^{-1} X x$$

This could be reexpressed as $\hat{y} = \sum_i \alpha_i K(x_i, x)$

$\alpha = A^{-1} Y$ $K(x, x') = x^T x'$ with $A = (X X^T + \lambda I_N)$

of course, here all x 's and y 's contribute

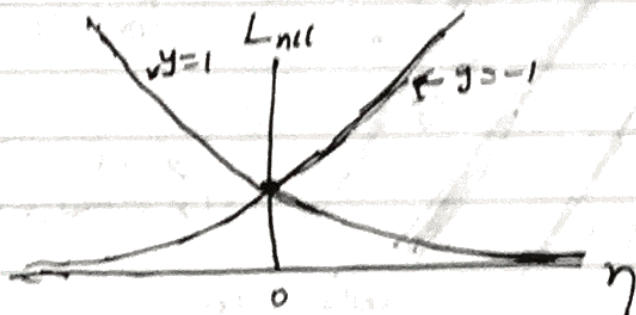
We will soon generalize the Kernel structure.

Coming off discussions of logistic regression
 let us first take up support vector classification

Logistic regression: $\eta = f(x) = w_0 + w^T x$

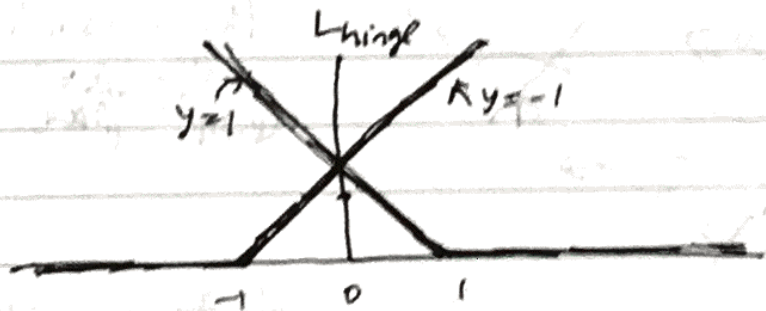
$$L_{nll}(y, \eta) = -\log p(y|x, w) = \log(1 + e^{-y\eta})$$

$y \in \{-1, +1\}$



Replace this loss function by hinge loss

$$L_{\text{hinge}}(y, \eta) = \max(0, 1 - y\eta) = (1 - y\eta)_+$$



Looks like
 door hinge

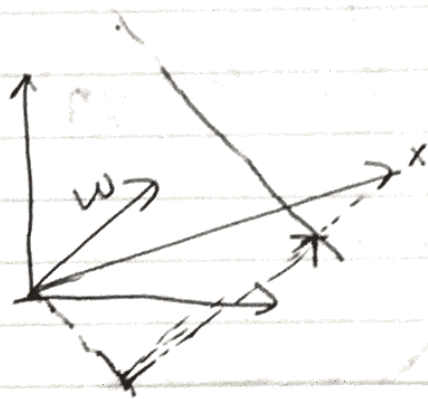
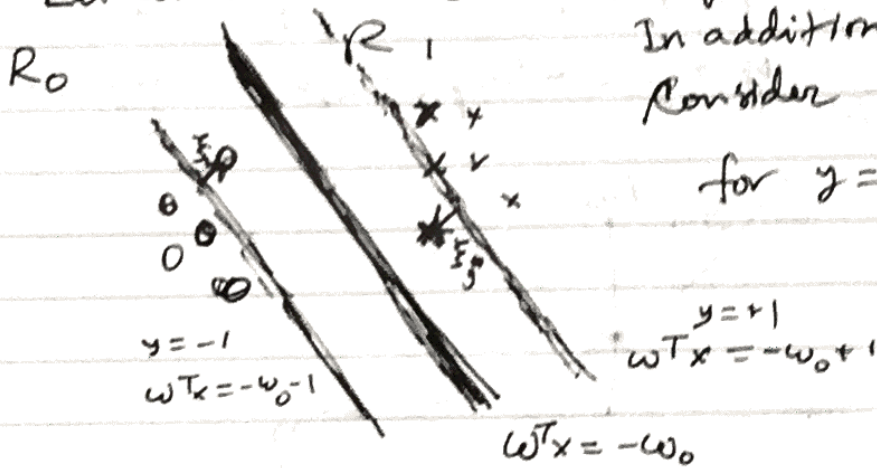
$$\min_{w_0, w} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N (1 - y_i f(x_i))_+$$

Introducing slack variables ξ_i , one
 could rewrite this optimization problem as

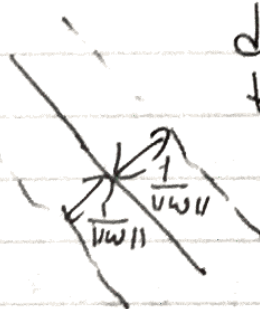
$$\min_{w_0, w, \xi} \frac{1}{2} \|w\|^2 + C \sum_i \xi_i \quad \text{st. } \xi_i \geq 0 \text{ and } y_i (w_0 + w^T x_i) \geq 1 - \xi_i \text{ for } i = 1, \dots, N$$

The points are classified by $f(x) \geq 0$.
 $f(x) = w_0 + w^T x$ is the discriminant function.

Let consider the geometry of relevant hyperplane
 In addition to $f(x) = 0$
 Consider $\eta f(x) = 1$
 for $y = \pm 1$

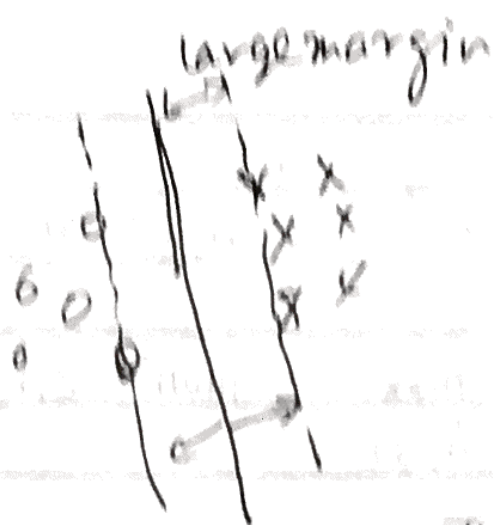


Distance from
 the decision boundary
 $\left(\frac{w}{\|w\|} \right)^T x - \frac{w_0}{\|w\|}$



distance between
 the decision boundary
 and the hyperplanes
 $\eta f(x) = 1$
 is $\frac{1}{\|w\|}$

If the ~~two~~ classes are linearly
 separable and we have a large C we want all $\xi_i = 0$



Not



We want $\frac{1}{\|w\|}$ to be as large s.t.

$$y_i (w_0 + w^T x_i) \geq 1 \text{ for all } i$$

That is equivalent to

$$\min \frac{1}{2} \|w\|^2, \text{ s.t. } y_i (w_0 + w^T x_i) \geq 1 \text{ for } i=1, \dots, N$$

In cases where perfect separation is not possible, we have the soft margin version.

$$\min_{w_0, w, \xi} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \xi_i \quad \text{s.t. } y_i (w_0 + w^T x_i) \geq 1 - \xi_i$$

For the sake of simplicity let us discuss perfectly separable case.

First $\min_w \frac{1}{2} \|w\|^2 \quad \text{s.t. } y_i (w_0 + w^T x_i) \geq 1$

Could be rewritten using Lagrange multipliers.

Consider

$$L(w_0, w, \lambda) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^N \lambda_i (y_i (w_0 + w^T x_i) - 1)$$

Turns out $\min_{w_0, w} \max_{\lambda_i \geq 0} \frac{1}{2} \|w\|^2 - \sum \lambda_i (y_i (w_0 + w^T x_i) - 1)$

The key difference is the Lagrange multipliers are sign restricted.

Think of it this way; If any constraint is not satisfied, i.e. $y_i (w_0 + w^T x_i) - 1 < 0$

We can send $\lambda_i \rightarrow +\infty$ and the obj. func. L would be ∞ . When minimizing over w , we will reject such solutions.

If a constraint is strictly satisfied

$$y_i (w_0 + w^T x_i) - 1 > 0$$

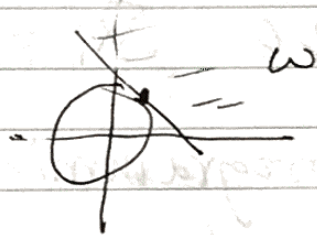
max over λ_i , sends $\lambda_i = 0$

So inactive constraints have corresponding $\lambda_i = 0$

What if $y_i (w_0 + w^T x_i) - 1 = 0$?

The marginal case is interesting

Now λ_i is determined by w variation being zero condition. It would usually be nonzero



These x_i 's are the support vectors.

Turns out that in this case, we can switch the order of min & max

$$\max_{\lambda_i \geq 0} \min_{w_0, w} \frac{1}{2} \|w\|^2 - \sum \lambda_i (y_i ((w_0 + w \cdot x_i) - 1))$$

Since w_0, w are ~~unrestricted~~ unrestricted, we could just differentiate.

$$w = \sum_i \lambda_i y_i x_i = \sum_i \alpha_i x_i$$

$$0 = \sum_i \lambda_i y_i x_i = \sum_i \alpha_i$$

$$\boxed{\alpha_i = \lambda_i y_i}$$

$$\tilde{L}(\alpha) = -\frac{1}{2} \sum_i \alpha_i x_i \cdot x_i + \sum_i \alpha_i y_i$$

$$\alpha_i y_i \geq 0 \text{ for all } i$$

$$\text{and } \sum \alpha_i = 0$$

$$\min_{\alpha} \frac{1}{2} \sum_i \alpha_i X_i^T X_i \alpha_i - \sum_i \alpha_i y_i$$

$$\text{s.t. } \alpha_i y_i \geq 0 \quad \sum_i \alpha_i = 0$$

Is a quadratic programming problem dual to the original one.

~~It is a quadratic programming problem~~

$$\text{The solution } \hat{w} = \sum_i \alpha_i X_i$$

$\hat{w}$ determined by picking an X_i
 s.t. $\alpha_i > 0$ or $y_i(\hat{w}_0 + \hat{w}^T X_i) = 1$ $\hat{w}_0 = y_i - \hat{w}^T X_i$

Primal

Dual

$$\hat{y}(x) = \text{sgn}(\hat{w}_0 + \hat{w}^T x)$$

Remember only certain α_i 's are non zero. Turns out that when only a small fraction of training data becomes support vectors, the model has good generalization properties.

How do we generalize this linear maximum margin classifier to problems which require non linear decision surfaces.

The trick is to realize that we could replace $x_i^T x_i$ by a general positive definite kernel $K(x_i, x_i)$.

$$\min_{\alpha} \frac{1}{2} \sum_i \alpha_i K(x_i, x_i) \alpha_i - \sum_i \alpha_i y_i$$

$$\text{s.t. } \alpha_i \geq 0 \quad \sum_i \alpha_i = 0$$

$$\hat{f}(x) = \text{sgn}(\hat{w}_0 + \sum_i \alpha_i K(x_i, x))$$

With $\hat{w}_0$ determined from a support vector.

For example: one could use kernels like

$$K(x, y) = (x \cdot y + 1)^p$$

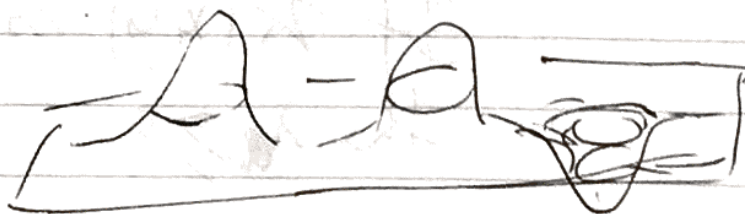
$$K(x, y) = e^{-\|x - y\|^2 / 2\sigma^2}$$

$$K(x, y) = \tanh(\kappa x \cdot y - \delta)$$

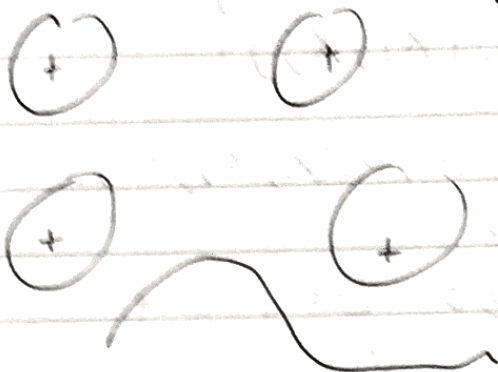
Gaussian
radial
basis functions
(RBF)

Decision Surface RBF

$$\sum_i \alpha_i e^{-\|x_i - x\|^2 / 2\sigma^2}$$



Level sets of RBF kernel
could be as complex as possible
with enough data.



With a small
enough σ^2 , it
can approximate
'any' decision
boundary.

Quite interestingly using $K(x, y)$
is equivalent to mapping x space
to an infinite dimensional feature
space by a nonlinear mapping.
Think of $K(x, y)$ as an infinite dimensional operator
$$\int K(x, y) \phi(y) dy = \psi(x)$$

If K is symmetric we have infinitely
many eigenvalues λ_a with
$$\int K(x, y) \phi_a(y) dy = \lambda_a \phi_a(x)$$

And

$$K(x, y) = \sum_a \lambda_a \phi_a(x) \phi_a(y)$$

If $K(x, y)$ is positive semi-definite,
the $\lambda_a \geq 0$.
 $\int \int K(x, y) \phi(x) \phi(y) dx dy \geq 0$

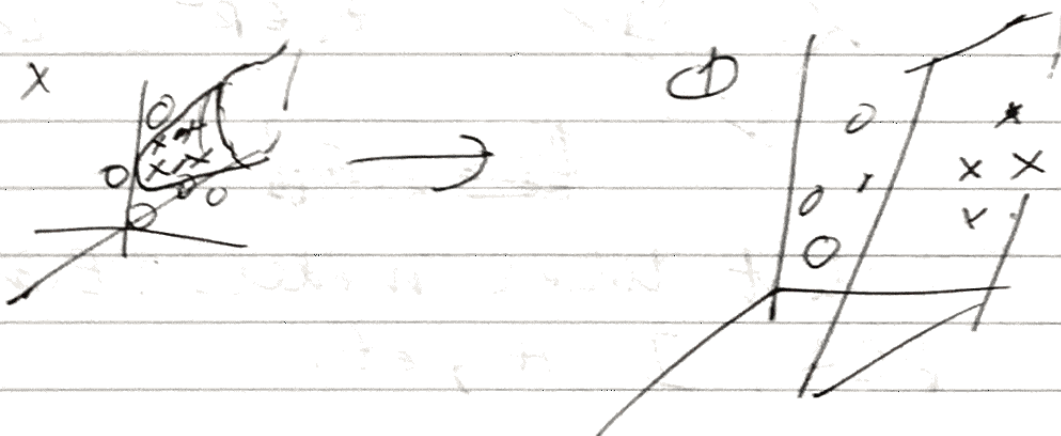
Define $\sqrt{x_i}: u_i = \phi_i$

$$K(x, y) = \sum_a \phi_a(x) \phi_a(y)$$

So $K(x_i, x_j) = \sum_a \phi_a(x_i) \phi_a(x_j)$

This mean $K(x_i, x_j) = \Phi(x_i)^T \Phi(x_j)$

where $\Phi(x) = (\phi_1(x), \phi_2(x), \phi_3(x), \dots)$



Similar idea's could be used for regression as well: $L(y, \hat{y}) \begin{cases} 0 & \text{if } y = \hat{y} \\ \infty & \text{otherwise} \end{cases}$
 $\hat{y} = w_0 + w^T x, \min_{w_0, w} L(y, \hat{y}) = \min_{w_0, w} \sum_i \max(0, 1 - y_i(w_0 + w^T x_i))^2$

As an aside, exponential of a hinge loss could be thought of as a mixture of many gaussian with different scales, allowing a probability interpretation of SVM's.

Trick $\int \frac{e^{-2a}}{e^{\frac{1}{2}(\frac{a}{\lambda} + \tilde{a})^2}} d\tilde{a} = \int_0^\infty e^{-\frac{1}{2} \frac{(a+\lambda)^2}{\lambda}} \frac{d\lambda}{\lambda} = e^{-2a} \int_0^\infty e^{-\frac{1}{2} (\frac{a}{\lambda} + \tilde{a})^2} \frac{\sqrt{\pi\lambda}}{\lambda} d\tilde{a}$