

Last time: $f(x) = |x| \Rightarrow \partial f(x) = \begin{cases} 1 & \text{if } x > 0 \\ [-1, 1] & \text{if } x = 0 \\ -1 & \text{if } x < 0 \end{cases}$

Ex: $f(x) = \|x\|_1, x \in \mathbb{R}^d$

$$\partial f(x) = \left\{ v : f(y) \geq f(x) + v^T(y-x) \right\}$$

$$= \left\{ v : \|y\|_1 \geq \|x\|_1 + v^T(y-x) \right\}$$

claim

$$= \left\{ \begin{bmatrix} v_1 \\ \vdots \\ v_d \end{bmatrix} : |y_i| \geq \underline{|x_i|} + \underline{v_i}(y_i - x_i) \right\}$$

$$\Rightarrow v \in \partial f(x) \text{ iff for } i=1, \dots, d$$

$$v_i = \begin{cases} 1 & \text{if } x_i > 0 \\ [-1, 1] & \text{if } x_i = 0 \\ -1 & \text{if } x_i < 0 \end{cases}$$

Recall our motivating ex for subgradient descent:

$$f(w) = \frac{1}{n} \sum_{i=1}^n (1 - y_i x_i^T w)_+ + \lambda \|w\|_1$$

for sparse SVM.

how to compute
subgrads
here?

we know how to
compute a
subgrad

Q: how to manipulate subgradient

Rules (for cvx functions)

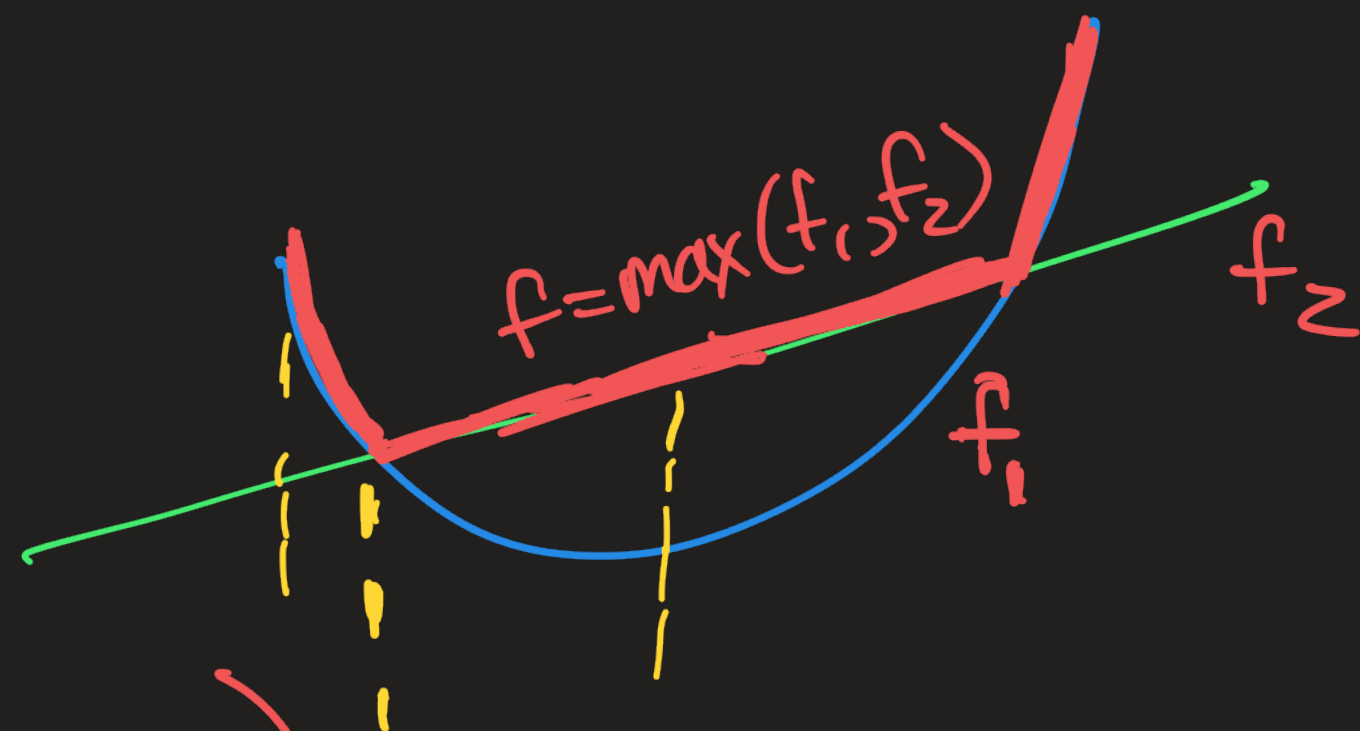
a) Scaling: $\partial(af) = a \partial(f)$ if $a \geq 0$

b) $\partial(f_1 + f_2) = \partial f_1 + \partial f_2$ (set sum) ^{Minkowski sum}
 $= \{u + v : u \in \partial f_1, v \in \partial f_2\}$

c) if $g(x) = f(Ax + b)$ then $\partial g(x) = A^T \partial f(Ax + b)$
(chain rule)

d) Finite pointwise max:

$$f(x) = \max_{i=1, \dots, m} f_i(x)$$



$$\partial f(x) = \text{conv} \left(\bigcup_{i: f_i(x) = f(x)} \partial f_i(x) \right)$$

which is the convex hull of the union of the subdifferentials of all the active i :
the f_i s.t. $f_i(x) = f(x)$

e) Norm: Given $p \geq 1$ (maybe $p = \infty$)

p -norm
 ℓ_p -norm $\rightarrow \|x\|_p = \left(\sum |x_i|^p \right)^{1/p}$ $\left(\begin{array}{l} \text{ex: } p=2 \\ \|x\|_2 = \sqrt{\sum x_i^2} \\ p=1 \\ \|x\|_1 = \sum |x_i| \end{array} \right)$

and $p = \infty$,

$$\|x\|_p = \max_i |x_i| \quad \left(\begin{array}{l} \text{inf-norm, } \ell_\infty\text{-norm,} \\ \text{max norm} \end{array} \right)$$

Let q be such that $\frac{1}{q} + \frac{1}{p} = 1$ (note $q \geq 1$)

p & q are called a conjugate pair

Note $(1, \infty)$ is a conjugate pair

if $f(x) = \|x\|_p$

$$(x^{1/2})' = \frac{1}{2} x^{-1/2}$$

then

$$\partial f(x) = \{ y : \|y\|_q \leq 1 \text{ and } y^T x = \|x\|_p \}$$

Ex of this rule

1) $p=2 \Rightarrow q=2$

$$f(x) = \|x\|_2 = \sqrt{\sum x_i^2}$$



when $x \neq 0$: $\{\nabla f(x)\}_i = \frac{1}{2\sqrt{\sum x_j^2}} 2x_i = \frac{x_i}{\|x\|_2}$

$$\nabla f(x) = \frac{x}{\|x\|_2}, \quad x \neq 0$$

To evaluate $\partial f(0)$, use fact

$$\partial \|x\|_2 = \{y : \|y\|_2 \leq 1 \text{ and } x^T y = \|x\|_2\}$$

$$\Rightarrow \partial \|0\|_2 = \{y : \|y\|_2 \leq 1 \text{ and } 0 = 0\}$$

$$= \{y : \|y\|_2 \leq 1\}$$

Ex (using our subgradient rules)

$$f(w) = \frac{1}{n} \sum_{i=1}^n (1 - y_i w^T x_i)_+ + \lambda \|w\|_1$$

Then we have

$$\partial f(w) = \frac{1}{n} \sum_{i=1}^n \partial (1 - y_i w^T x_i)_+ + \lambda \partial \|w\|_1$$

we know $\partial \|w\|_1$. How to compute $\partial (1 - y_i w^T x_i)_+$

$$\begin{aligned} \text{Let } g(w) &= (1 - y_i w^T x_i)_+ = (1 - y_i x_i^T w)_+ \\ &= f(a^T w + b) \end{aligned}$$

$$\text{where } f(x) = x_+$$

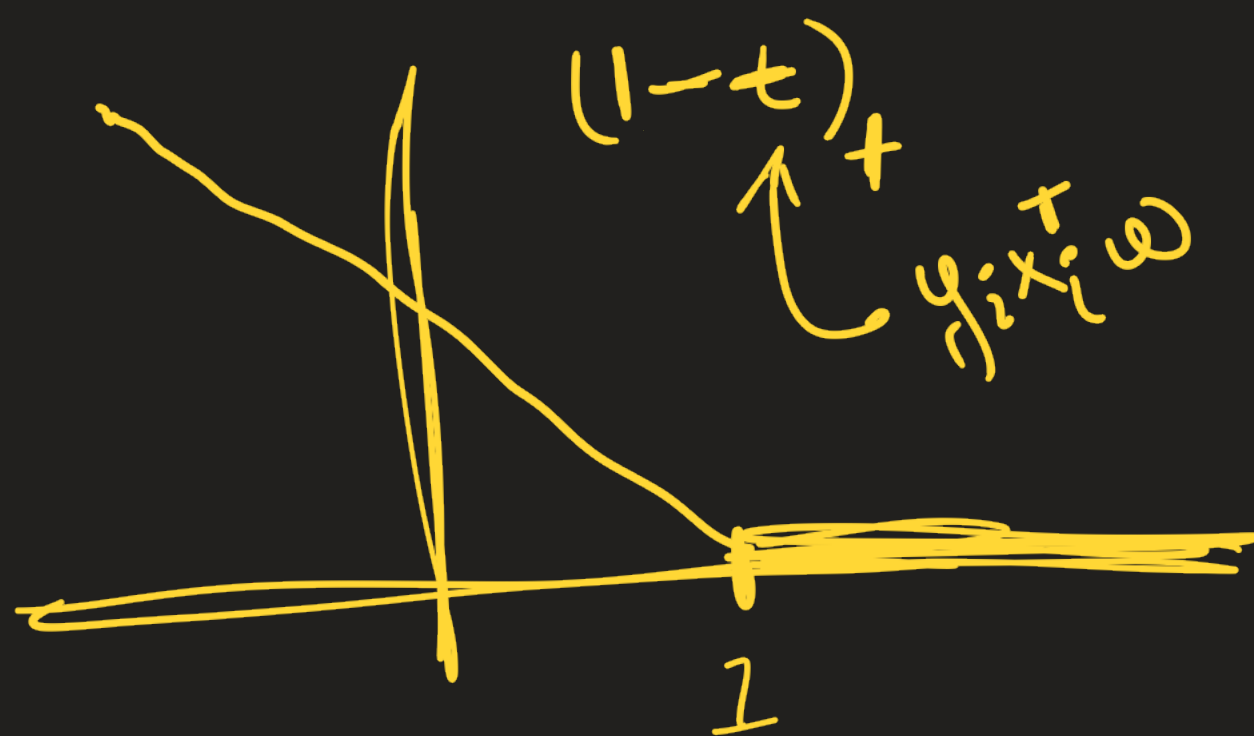
$$a = -y_i x_i$$

$$b = 1$$

So $\partial g(\omega) = a \partial f(a^T \omega + b)$ by chain rule

$$= a \begin{cases} 1 & \text{if } a^T \omega + b > 0 \\ [0, 1] & \text{if } a^T \omega + b = 0 \\ 0 & \text{if } a^T \omega + b < 0 \end{cases}$$

$$= \underline{-y_i x_i} \begin{cases} 1 & \text{if } 1 - y_i x_i^T \omega > 0 \\ [0, 1] & \text{if } 1 - y_i x_i^T \omega = 0 \\ \underline{0} & \text{if } \underline{1 - y_i x_i^T \omega} < 0 \end{cases}$$

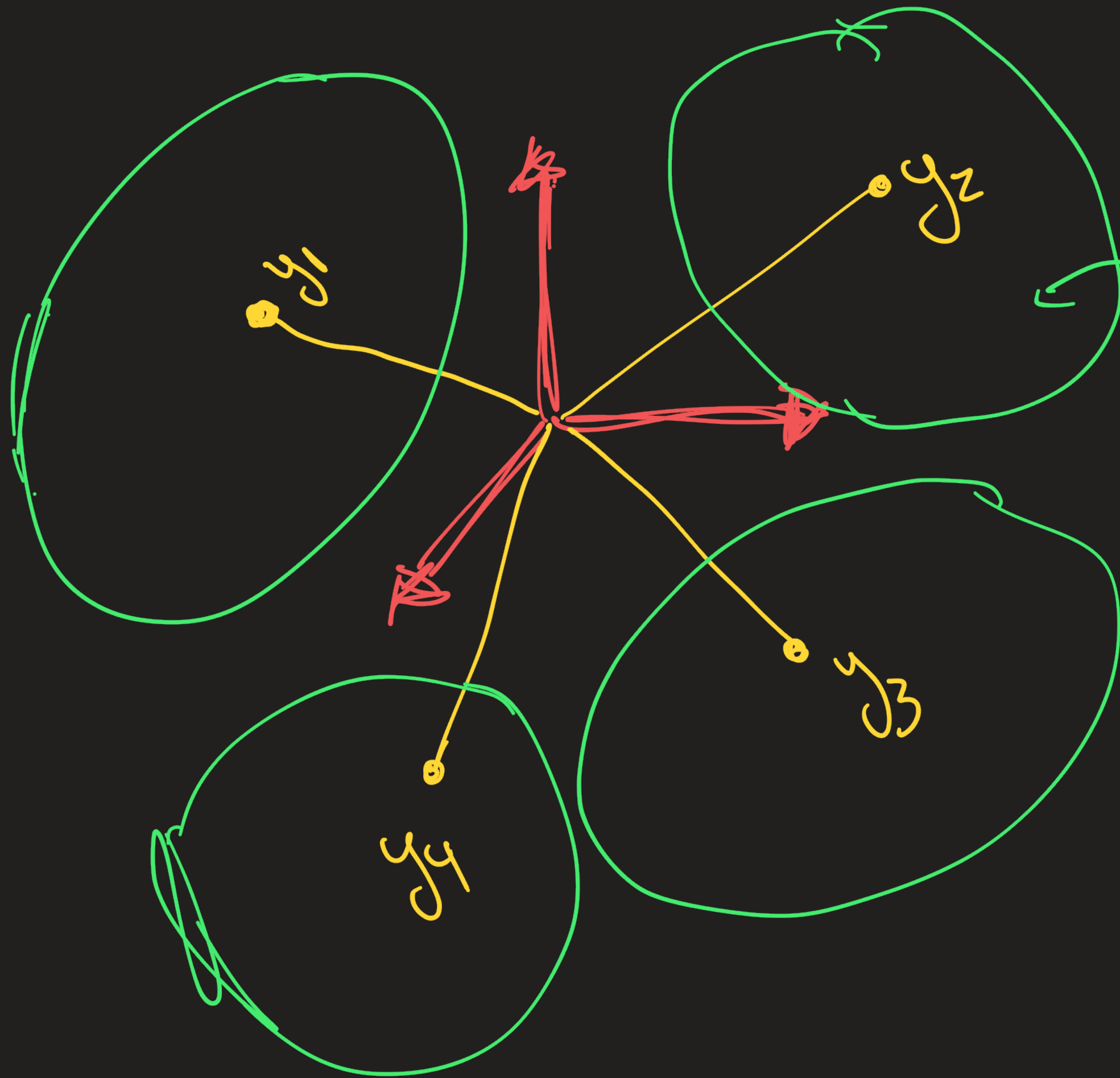


Ex ECOC (Error-correcting output coding)
and linear regression for classification (K-class)

Choose K different code words (vectors $\in \mathbb{R}^d$)
to represent our classes. So represent class i
with code-word $y_i \in \mathbb{R}^d$.

Fit a linear model $y = Wx$ that predicts y
given features x and take the class of a new
point x to be

$$\hat{i} = \underset{i}{\operatorname{argmin}} \|y_i - Wx\|_2$$



if Wx
closest to y_2
we say x should
be in class 2

Train: learn W to minimize the distance between x_i and the codewords for the class y_i

$$f(W) = \frac{1}{n} \sum_{i=1}^n \|c_{y_i} - Wx_i\|_2^2$$

$W \in \mathbb{R}^{l \times d}$ where c_{y_i} is the codeword for class y_i

then $f(W) = \frac{1}{n} \|Y - WX\|_F^2$

where $Y = \begin{bmatrix} c_{y_1} & \dots & c_{y_n} \end{bmatrix} \in \mathbb{R}^{l \times n}$

$$X = [x_1 \dots x_n] \in \mathbb{R}^{d \times n}$$

Aside $M \in \mathbb{R}^{n \times d}$

$$\|M\|_F^2 = \sum_{i=1}^n \sum_{j=1}^d m_{ij}^2 = \sum_{i=1}^n \left(\sum_{j=1}^d m_{ij}^2 \right)$$

$$M = [m_1 \dots m_d]$$

$$= \begin{bmatrix} \vdots \\ \vdots \\ \vdots \\ \vdots \\ \vdots \end{bmatrix}$$

$$= \sum_{i=1}^n \|r_i\|_2^2$$

$$= \sum_{j=1}^d \left(\sum_{i=1}^n m_{ij}^2 \right) = \sum_{j=1}^d \|m_j\|_2^2$$

$$\|M\|_F^2 = \sum_{j=1}^d \|m_j\|_2^2 \quad \text{where } M = [m_1 | \dots | m_d]$$

Consider $M = Y - WX = [c_{y_1} \dots c_{y_n}] - W[x_1 \dots x_n]$

$$= [c_{y_1} - Wx_1 | \dots | c_{y_n} - Wx_n]$$

$$\Rightarrow \|Y - WX\|_F^2 = \sum_{i=1}^n \|c_{y_i} - Wx_i\|_2^2 = n f(w)$$

$$\Rightarrow f(w) = \frac{1}{n} \|Y - WX\|_F^2$$

Fact $\|M\|_F^2 = \text{Tr}(M^T M) = \text{Tr}(M M^T)$

recall $\text{Tr}(A) = \sum_{i=1}^n A_{ii}$ if $A \in \mathbb{R}^{n \times n}$

Facts

1) $\text{Tr}(AB^T) = \text{Tr}(B^T A)$ if A & B have the same dimensions

2) $\text{Tr}(A(B+C)) = \text{Tr}(AB) + \text{Tr}(AC)$

3) $\text{Tr}(\alpha A) = \alpha \text{Tr}(A)$

These imply that if A & B have the same dimensions, then

$\langle A, B \rangle = \text{Tr}(AB^T)$ is a valid inner product

An inner-product generalizes the dot-product for two vectors $u, v \in \mathbb{R}^d$:

$$\langle u, v \rangle = u^T v$$

The dot product satisfies

1) $\langle u, v \rangle = \langle v, u \rangle$ symmetry

2) $\langle \alpha u + v, w \rangle = \alpha \langle u, w \rangle + \langle v, w \rangle$ for $\alpha \in \mathbb{R}$
linearity

3) $\langle u, u \rangle = \|u\|_2^2 \geq 0$ and $\langle u, u \rangle = 0$ iff $u = 0$
positive-definiteness

For matrices, $\langle A, B \rangle = \text{Tr}(AB^T)$ has the same properties: symmetry, linearity, and positive-definiteness.

And

$$\|M\|_F^2 = \text{Tr}(MM^T) = \langle M, M \rangle$$

Just as in the vector case ($\nabla_u \|u\|_2^2 = 2u$),

$$\nabla_M \|M\|_F^2 = 2M$$

To compute $\nabla f(W)$, note

$$f(W) = \frac{1}{n} \|Y - WX\|_F^2 \Rightarrow$$

$$\nabla f(W) = -\frac{2}{n} (Y - WX)X^T \in \mathbb{R}^{d \times d}$$

With this expression we can use the gradient descent algorithm to fit an ECOC classification model.

How do we choose a stepsize? In this case, f is in fact β -smooth; let's determine β

$f: \mathbb{R}^d \rightarrow \mathbb{R}$
Recall that f is β -smooth iff for all $\omega, \omega_0 \in \mathbb{R}^d$

$$f(\omega) \leq f(\omega_0) + \langle \nabla f(\omega_0), \omega - \omega_0 \rangle + \beta \|\omega - \omega_0\|_2^2$$

For $f: \mathbb{R}^{l \times d} \rightarrow \mathbb{R}$, we need that for all $W, W_0 \in \mathbb{R}^{l \times d}$,

$$f(W) \leq f(W_0) + \langle \nabla f(W_0), W - W_0 \rangle + \beta \|W - W_0\|_F^2$$

To verify this, note

$$f(W) = \frac{1}{n} \|\varphi - WX\|_F^2 = \frac{1}{n} \|\varphi - W_0X + (W_0X - WX)\|_F^2$$

Now apply the identity

$$\begin{aligned}\|A+B\|_F^2 &= \langle A+B, A+B \rangle = \langle A, A \rangle + 2\langle A, B \rangle + \langle B, B \rangle \\ &= \|A\|_F^2 + \|B\|_F^2 + 2\langle A, B \rangle\end{aligned}$$

with $A = \psi - \omega_0 X$ and $B = (\omega_0 - \omega)X$ to see that

$$\begin{aligned}f(\omega) &= \frac{1}{n} \|\psi - \omega_0 X\|_F^2 + \frac{2}{n} \langle \psi - \omega_0 X, (\omega_0 - \omega)X \rangle \\ &\quad + \frac{1}{n} \|(\omega_0 - \omega)X\|_F^2\end{aligned}$$

Now apply the identity

$$\begin{aligned}\langle A, BC \rangle &= \text{Tr}(A(BC)^T) = \text{Tr}(BCA^T) = \text{Tr}(B(AC^T)^T) \\ &= \langle AC^T, B \rangle\end{aligned}$$

with

$$\begin{aligned}A &= \psi - \omega_0 X \\ B &= \omega_0 - \omega \\ C &= X\end{aligned}$$

to see that

$$f(W) = f(W_0) + \frac{2}{n} \langle (Y - W_0 X) X^T, W_0 - W \rangle \\ + \frac{1}{n} \| (W_0 - W) X \|_F^2$$

$$= f(W_0) + \langle \nabla f(W_0), W - W_0 \rangle + \frac{1}{n} \| (W_0 - W) X \|_F^2$$

Now note that

$$\frac{1}{n} \| (W_0 - W) X \|_F^2 \leq \frac{1}{n} \| X \|_2^2 \| W_0 - W \|_F^2$$

where $\|X\|_2$ is the spectral norm of X , so we have

$$f(W) \leq f(W_0) + \langle \nabla f(W_0), W - W_0 \rangle + \beta \| W_0 - W \|_F^2$$

where

$$\beta = \frac{1}{n} \| X \|_2^2$$

We therefore see that gradient descent will converge when $\alpha = \frac{1}{\beta} = \frac{n}{\|X\|_2^2}$ is used as the stepsize,

at rate

$$f(W_T) - f(W_*) = O\left(\frac{\beta}{\epsilon}\right)$$

Of course we can directly solve for W_* using OLS:

$$W_* = YX^T(XX^T)^{-1}$$

In practice we would do this as long as d is small enough that the $O(d^3)$ operation of inverting XX^T is feasible. Eg if $d = 50,000$ we would then use methods like gradient descent or, as we'll see next class, more likely stochastic gradient descent