

Lecture 11: High-dimensional non-convex dynamics

①

SGD on NNs : if we are free to set architecture and initialization, then we can emulate any polynomial time algorithm.

↳ this is very different than what we do in practice

↳ there is some "hyperparameter tuning" but the space of search is relatively constrained : we only use some standard architectures and "generic" initializations
(fully connected / CNNs / transformers)

(iid initialization)
e.g. Gaussian
where we can tune the variance

To understand what SGD does on regular NNs, one needs to study the dynamics

↳ high-dimensional (many parameters)

+ non-convex

↳ this is hard!

While it is generically intractable to study these dynamics, there has been a lot of interesting work in the past few years.

Today: we will focus on a very simple setting that presents nonetheless a very interesting phenomenology + is highly instructive.

"Online SGD on non-convex losses from high-dim inference"
Ben Arous, Ghisleri, Jagannathan, 2021.

Setting: $x \sim N(0, I_d)$
 $y = f(\langle w_*, x \rangle)$ for some $\|w_*\|_2 = 1$

We learn this data using a "single neuron"

$\hat{f}(x) = f(\langle w, x \rangle)$ using SGD from a random initialization
 $w^0 \sim \text{Unif}(\mathbb{S}^{d-1})$

Test error:

$$L(w) = \frac{1}{2} \mathbb{E}[(f(\langle w_*, x \rangle) - f(\langle w, x \rangle))^2]$$

Hermite polynomials & the information exponent

$\{He_k\}_{k=0}^{\infty}$ the orthogonal basis of Hermite polynomials
(basis of $L^2(\mathbb{R}, N(0,1))$)

He_k is a degree- k polynomial such that

$$\mathbb{E}_{G \sim N(0,1)} [He_k(G) He_j(G)] = k! \mathbb{1}_{[j=k]}$$

$$He_0(x) = 1 \quad He_1(x) = x \quad He_2(x) = x^2 - 1 \quad He_3(x) = x^3 - 3x$$

Properties: (i) $\frac{d}{dx} He_k(x) = k He_{k-1}(x)$

↙ integration by part

(ii) $\mathbb{E}[He_k(G) g(G)] = \mathbb{E}[g^{(k)}(G)]$

(iii) $He_{k+1}(x) = x He_k(x) - k He_{k-1}(x)$

Any function $f: \mathbb{R} \rightarrow \mathbb{R}$ s.t. $\mathbb{E}[f(G)^2] < \infty$ can be decomposed as

$$f(x) = \sum_{k=0}^{\infty} \frac{\mu_k(f)}{k!} He_k(x) \quad \text{s.t.} \quad \mu_k(f) = \mathbb{E}[f(G) He_k(G)]$$

(if k times diff = $\mathbb{E}[f^{(k)}(G)]$)

(4)

Def: [Information exponent] We say that $f \in L^2(\mathbb{R}, N(0,1))$ has information exponent k_* if

$$k_* = \min \{ k \in \mathbb{N} : \mathbb{E}[f(G) He_k(G)] \neq 0 \}$$

Lemma: Let k_* be the info exponent of f . Then

$$L(\omega) = \frac{1}{2} \mathbb{E}[(f(\langle \omega_*, \alpha \rangle) - f(\langle \omega, \alpha \rangle))^2] = 1 - O(\langle \omega, \omega_* \rangle^k)$$

Proof: $\frac{1}{2} \|f(\langle \omega_*, \cdot \rangle) - f(\langle \omega, \cdot \rangle)\|_{L^2}^2 = 1 - \mathbb{E}[f(\langle \omega_*, \alpha \rangle) f(\langle \omega, \alpha \rangle)]$

$$= 1 - \sum_{k=k_*}^{\infty} \frac{\mu_k^2}{(k!)^2} \mathbb{E}[He_k(\langle \omega_*, \alpha \rangle) He_k(\langle \omega, \alpha \rangle)]$$

$$\mathbb{E}[He_k(\langle \omega_*, \alpha \rangle) He_k(\langle \omega, \alpha \rangle)] \quad G_1, G_2 \sim N(0,1)$$

$$= \mathbb{E}[He_k(G_1) He_k(\langle \omega, \omega_* \rangle G_1 + \sqrt{1 - \langle \omega, \omega_* \rangle^2} G_2)]$$

$$= \mathbb{E}\left[\frac{d^k}{dG_1^k} He_k(\langle \omega, \omega_* \rangle G_1 + \sqrt{1 - \langle \omega, \omega_* \rangle^2} G_2)\right]$$

$$= \langle \omega, \omega_* \rangle^k \mathbb{E}[He_k^{(k)}(G)]$$

$$= k! \quad (\text{property (i)})$$

Int by part
on G_1
(property (ii))

(5)

Hence
$$L(\omega) = 1 - \sum_{k=k_0}^{\infty} \frac{\mu_k^2(b)}{k!} \langle \omega_*, \omega \rangle^k$$

□

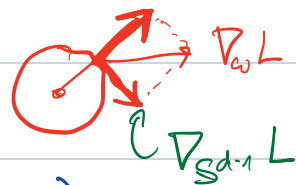
For simplicity, we will set $f(x) = \frac{\text{He}_k(x)}{\sqrt{k!}}$ from now on.

[results will hold for more general f of I.E. k]

Test error:
$$L(\omega) = 1 - \langle \omega, \omega_* \rangle^k = 1 - m^k \quad m = \langle \omega, \omega_* \rangle$$

We minimize over $\omega \in \mathbb{S}^{d-1}$. We will consider "spherical" gradient descent (gradient $\in$ tangent space of sphere manifold)

$$\begin{aligned} \nabla_{\mathbb{S}^{d-1}} L(\omega) &= (\text{Id} - \omega \omega^T) \nabla_{\omega} L(\omega) \\ &= (\text{Id} - \omega \omega^T) (-k \langle \omega, \omega_* \rangle^{k-1} \omega_*) \\ &= -k m^{k-1} (\omega_* - m \omega) \end{aligned}$$



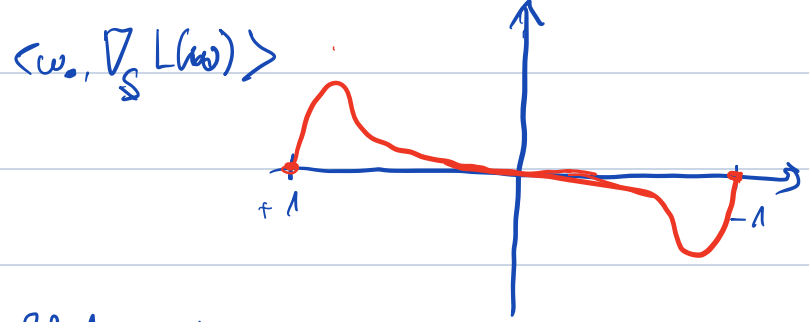
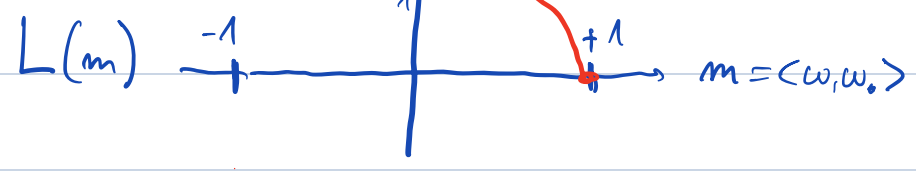
The gradient along the signal direction is

$$\langle \omega_*, \nabla_{\mathbb{S}^{d-1}} L(\omega) \rangle = -k m^{k-1} (1 - m^2)$$

Population landscape

non-convex

If k is odd

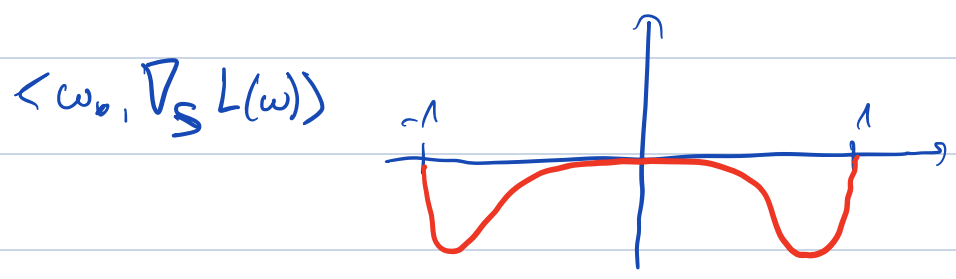
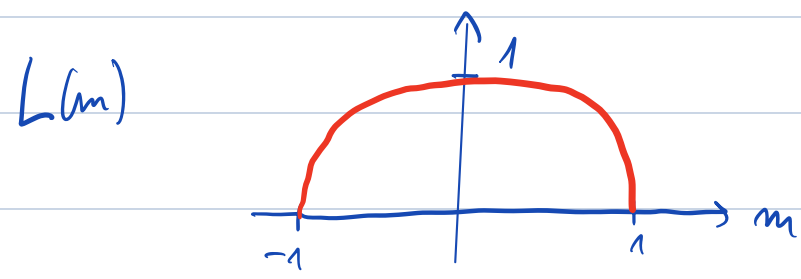


$\omega = \omega_*$ global minimum

$\langle \omega, \omega_* \rangle = 0$ stationary submanifold

$\omega = -\omega_*$ unstable stationary point (global maxime)

If k is even

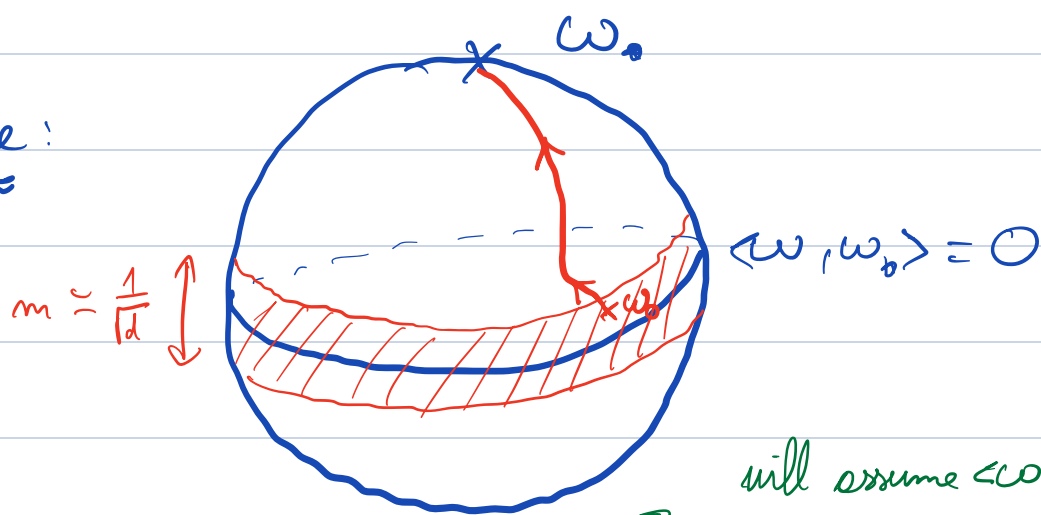


$\omega = \pm \omega_*$ 2 global minima

$\langle \omega, \omega_* \rangle = 0$ stationary submanifold

we will assume
 $\langle \omega^0, \omega_* \rangle > 0$
 (probe $1/2$)
 and $\omega^t \rightarrow +\omega_*$

Landscape:



w.p. $\geq \frac{1}{2} + \eta$

will assume $\langle \omega^0, \omega_* \rangle \gtrsim \frac{c}{\sqrt{d}}$

with $\omega_* \sim \text{Unif}(\mathbb{S}^{d-1})$ w.h.p. $\langle \omega^0, \omega_* \rangle \leq \frac{C}{\sqrt{d}}$

At initialization, $\|\nabla_S L(\omega^0)\|_2 = \frac{1}{d^{\frac{k-1}{2}}}$

$$\langle \omega_*, \nabla_S L(\omega^0) \rangle = \frac{1}{d^{\frac{k-1}{2}}}$$

very small gradient: dynamic initialized near a saddle-pt.
SGD will take a long time to escape this generic initialization

We want to study online SGD on this problem and in particular establish # of SGD steps to reach the global minimum.

Gradient flow limit

A standard approach is to approximate SGD by its continuous

(8)

limit (step size $\rightarrow 0$): $\dot{\omega}^t = -\nabla_{\omega} L(\omega^t)$

We only need to track $m_t = \langle \omega^t, \omega_* \rangle$

$$m_0 = \frac{c}{\sqrt{d}} \quad \dot{m}_t = k (m_t)^{k-1} (1 - m_t^2)$$

As long as m_t small $\dot{m}_t \approx k (m_t)^{k-1}$

$$\frac{dm_t}{m_t^{k-1}} = k dt$$

If $k=1$: $m_T \approx m_0 + kT \rightarrow \text{converge in } T = O_d(1)$

If $k=2$: $m_T \approx m_0 e^{kT} \rightarrow \text{converge in } T = O_d(\log d)$

If $k > 2$: $\frac{1}{k-2} \left[\frac{1}{m_0^{k-2}} - \frac{1}{m_T^{k-2}} \right] = kT \rightarrow m_T \approx \frac{1}{(d^{\frac{k-2}{2}} - kT)^{\frac{1}{k-2}}}$
 $\rightarrow \text{converge in } T = d^{\frac{k}{2}-1}$

Comparing k steps of SGD with step size η , we can show the propagation of error: with high proba $\sqrt{\frac{1}{k}} \times \eta$ ^{$\leftarrow$ steps}

$$| \underbrace{L(\omega^k)}_{\text{SGD}} - \underbrace{L(\omega^{k\eta})}_{\text{GF}} | \leq e^{c k \eta} \sqrt{d \eta}$$

(9)

For $k=1$, this shows that SGD with step size $\eta = \Theta_d(1/d)$ achieves small test error in $\Theta_d(d)$ steps/samples
 [note that $\Omega(d)$ samples is information-theoretic optimal]

For $k>1$, because of propagation of error, we cannot approximate discrete SGD by GF throughout the entire dynamic.

We need a different analysis!

Projected online SGD

We consider the following algorithm: initialize at ω^0

At each step t :

$$\begin{cases} \tilde{\omega}^{t+1} = \omega^t - \eta \nabla_S l(\omega^t, x^t) \\ \omega^{t+1} = \frac{\tilde{\omega}^{t+1}}{\|\tilde{\omega}^{t+1}\|_2} \end{cases} \quad \text{) project back on the sphere}$$

$$l(\omega^t, x^t) = \frac{1}{2} (f(\langle \omega, x_t \rangle) - f(\langle \omega, x_t \rangle))^2$$

$\epsilon_{iid} \sim N(0, I_d)$

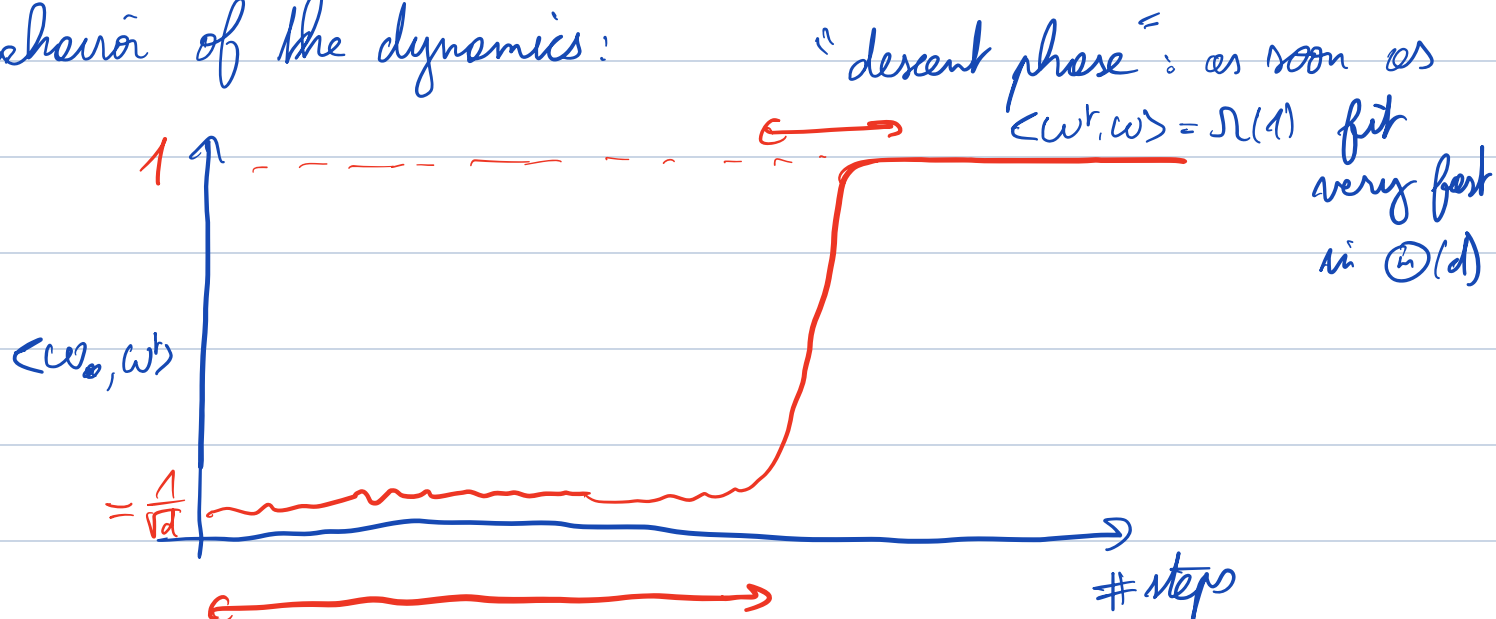
Thm: [Ben Arous, Ghahseri, Jagannathan, '21]

If we choose

$k=1$	$\eta = \tilde{\Theta}\left(\frac{1}{d}\right)$	$T = \tilde{\Theta}(d)$
$k=2$	$\eta = \tilde{\Theta}\left(\frac{1}{d}\right)$	$T = \tilde{\Theta}(d \log d)$
$k > 2$	$\eta = \tilde{\Theta}\left(d^{-\frac{k}{2}}\right)$	$T = \tilde{\Theta}(d^{k-1})$

Then online SGD with step size η and T steps will have $\langle w_*, w^T \rangle \approx 1$ w. h. p.

Behavior of the dynamics:



$\tilde{\Theta}(d^{k-1})$ samples spent to get correlated to the signal.

Main difficulty in this problem is leaving the high entropy region of the initialization. As soon as one has left this region, the “fitting” is fast.

Rule: Remember the case of learning k -parties. This is the same phenomenon!

Proof of the Theorem

Note that $\eta T = d^{\frac{T}{2}-1}$ which matches the continuous time guarantee from GF

However GF will not be a good approximation of SGD for $\eta T \gg 1$

Instead, we will (approximately) decompose the dynamics into a deterministic drift induced by the signal that biases dynamics to $\rightarrow \omega_*$ + a mean zero noise term.

Dynamics:
$$\omega^{t+1} = \frac{\omega^t - \eta \nabla_S l(\omega^t, x^t)}{\|\omega^t - \eta \nabla_S l(\omega^t, x^t)\|_2}$$

Below we will focus on search phase, i.e., getting $\langle \omega^t, \omega_* \rangle = \Omega_d(1)$
(descent phase can be studied by GF)

For η small enough, $\omega^{t+1} \approx \omega^t - \eta \nabla_S l(\omega^t, \alpha^t)$ (*)

In these notes, we will ignore the projection step and directly consider (*).

$$\omega^{t+1} = \omega^t - \eta \nabla_S l(\omega^t, \alpha^t) = \omega_0 - \eta \sum_{\Delta=0}^t \nabla_S l(\omega^\Delta, \alpha^\Delta)$$

It is enough to track $m^t = \langle \omega^t, \omega_* \rangle$

$$m_{t+1} = m_0 + \eta \sum_{\Delta=0}^t g^\Delta$$

$$g^\Delta := -\langle \omega_*, \nabla_S l(\omega^\Delta, \alpha^\Delta) \rangle = (f(\langle \omega_*, \alpha^\Delta \rangle) - f(\langle \omega^\Delta, \alpha^\Delta \rangle)) f'(\langle \omega^\Delta, \alpha^\Delta \rangle) (\alpha^\Delta - \omega^\Delta \langle \omega^\Delta, \alpha^\Delta \rangle)$$

Denote $\bar{g}^\Delta = \mathbb{E}_{\alpha^\Delta} [g^\Delta] = -\langle \omega_*, \nabla_S L(\omega^\Delta) \rangle$

$$= k m_\Delta^{k-1} (1 - m_\Delta^2)$$

Then

$$m_{t+1} = \underbrace{m_0}_{= \frac{c}{1d}} + \underbrace{\eta \sum_{\Delta=0}^t \bar{g}^\Delta}_{D^t: \text{"deterministic drift"}} + \underbrace{\eta \sum_{\Delta=0}^t g^\Delta - \bar{g}^\Delta}_{M^t: \text{martingale}}$$

M^t is a martingale $(M^t = \sum_{s=0}^t X_s \quad \mathbb{E}[X_t | X_{\leq t-1}] = 0)$

① Controlling the martingale:

We can use Doob's maximal inequality: for every $p \geq 1$,

$$P(\max_{0 \leq t \leq T} |M_t| \geq \lambda) \leq \frac{p \mathbb{E}[|M_T|^p]}{(p-1) \lambda^p}$$

One can show $\left[\begin{array}{l} \text{using hypercontractivity of low-degree polynomials} \\ \text{on Gaussian data} \quad \|f\|_{L^1}^2 \leq (q-1)^k \|f\|_{L^2}^2 \\ f \text{ any degree } k \text{ polynomial} \end{array} \right]$

$$\mathbb{E}[|M_T|^p] \lesssim C_p \eta^p T^{\frac{p}{2}}$$

$$C \frac{\eta^2 T}{\lambda^2} = \delta$$

Hence, with prob $1 - \delta$,

$$\max_{0 \leq t \leq T} |M_t| \leq C \eta \sqrt{T} \underline{\delta^{-\frac{1}{2}}}$$

dependency can
be improved to
 $\log(\frac{1}{\delta})$

In particular if $C \frac{\eta \sqrt{T}}{\sqrt{\delta}} \leq \frac{m_0}{2}$

that is $\eta \sqrt{T} = O(\frac{1}{\sqrt{d}})$ (A1)

Then: $\frac{m_0}{2} + D^t \leq m_t \leq \frac{3}{2} m_0 + D^t$ for all $t \leq T$.

and we can neglect the Martingale/noise term

② Controlling the drift term

In the search phase (m^t small) $\bar{g}^\Delta \approx k m_\Delta^{k-1}$

Hence for all $t \leq T$ (under (A1))

$$m_T \geq \frac{m_0}{2} + \eta k \sum_{\Delta=0}^{T-1} m_\Delta^{k-1}$$

We can study this sequence and lower bound their growth. The Lemma below can be seen as a discrete lower bound version of Bihari-LaSalle inequality.

Lemma: For a sequence $(u_t)_{t \in \mathbb{N}}$ that satisfies

$$u_t \geq a_0 + a_1 \sum_{s=0}^{t-1} u_s^{k-1}$$

Then u_t is lower bounded by

$k=2$: $u_t \geq a_0(1+a_1)^t$

$k>2$: $u_t \geq \min\left(1, \frac{1}{\left(a_0^{-(k-2)} - \frac{k-2}{(1+a_1)} a_1 t\right)^{\frac{1}{k-2}}}\right)$

For $k \geq 2$: $m_0^{-(k-2)} = d^{\frac{k}{2}-1}$

Hence $m_T = \Omega(1)$ as long as

$$qT = \Omega(d^{\frac{k}{2}-1}) \quad \text{(A2)}$$

In fact, if $t = o(d^{\frac{k}{2}-1})$ then $m_t \asymp \frac{1}{\sqrt{d}}$

$\hookrightarrow$ most of the search phase $m_t \asymp \frac{1}{\sqrt{d}}$

Combining (A1) and (A2)

$$\eta \sqrt{T} = \tilde{\Theta}\left(\frac{1}{\sqrt{d}}\right) \quad \eta T = \tilde{\Theta}\left(d^{\frac{k}{2}-1}\right)$$

$$\Rightarrow T = \tilde{\Theta}\left(d^{k-1}\right) \quad \eta = \tilde{\Theta}\left(d^{-\frac{k}{2}}\right)$$

for $k > 2$

For $k=2$, then $u_T \gtrsim m_0 e^{\eta T}$

$$\rightarrow \eta T = \tilde{\Theta}(\log d)$$

$$\Rightarrow T = \tilde{\Theta}(\log d) \quad \eta = \tilde{\Theta}(d^{-1})$$

Landscapes concentration

The above guarantees are for online SGD
that is $\# \text{ samples} = \# \text{ steps}$

What about if we reuse samples?

Batch 1 SGD choose one random sample from the empirical distribution $\hat{\mu}_n = \frac{1}{n} \sum_{i=1}^n \delta_{x_i}$

We can do the same drift + martingale decomposition but now

$$\begin{aligned} \bar{g}^s &\longrightarrow \hat{g}_n^s = \langle w_*, \nabla_S \hat{\mathbb{E}}_n[l(w^s, x)] \rangle \\ &= \frac{1}{n} \sum_{i=1}^n (y_i - f(\langle w^s, x_i \rangle)) f'(\langle w^s, x_i \rangle) (x_i - w^s \langle w^s, x_i \rangle) \end{aligned}$$

If we can show that $\hat{g}_n(w) \geq c \bar{g}(w)$

$$\text{for all } w \in \mathbb{S}^{d-1} \quad \langle w, w_* \rangle \geq \frac{c}{\sqrt{d}}$$

then the same analysis as for one-pass SGD goes through

For a given $w \in \mathbb{S}^{d-1}$, $\hat{g}_n(w)$ is an empirical average with mean $\bar{g}(w)$

For simplicity take β with $\|\beta\|_\infty, \|\beta'\|_\infty \leq C$.

$$P(|\hat{g}_m(\omega) - \bar{g}(\omega)| \geq \sqrt{\text{Var}(\hat{g}_m(\omega))}) \leq C e^{-c t^2}$$

Hence $P\left(\min_{\substack{\omega \in \mathbb{S}^{d-1} \\ \langle \omega, \omega_* \rangle \geq d^{-\gamma}}} \hat{g}_m(\omega) - \bar{g}(\omega) \geq \sqrt{\frac{C d \log d}{m}}\right) \leq \frac{C}{d^C}$

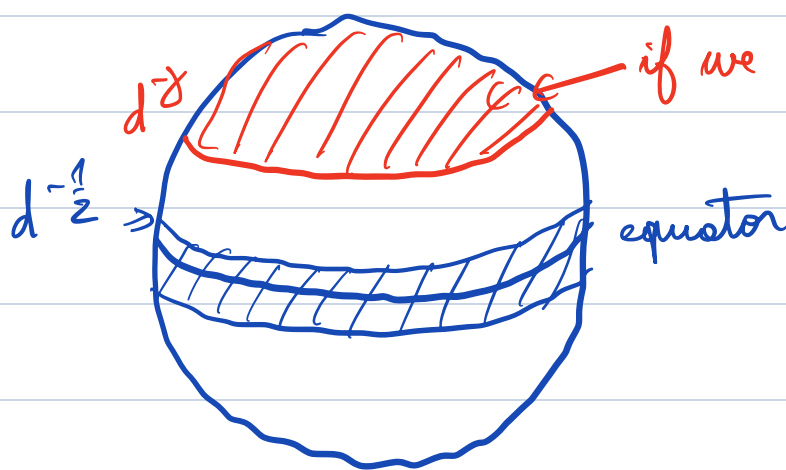
[using a standard uniform convergence bound]

Hence for $\frac{1}{2} \leq \gamma \leq 0$

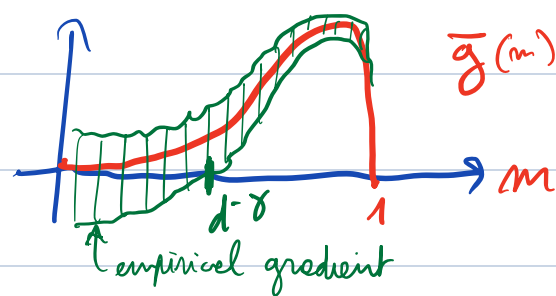
then if we take $\sqrt{\frac{d}{m}} \leq \frac{1}{d^{\gamma(h-1)}}$

$$\Rightarrow m \geq d^{2\gamma(h-1)+1}$$

then the landscape is "benign" for any $\langle \omega, \omega_* \rangle \geq d^{-\gamma}$



if we initialize on this cap, then $\omega^t \rightarrow \omega_*$



If $\gamma = \frac{1}{2}$, $n \gtrsim d^k$ random initialization falls in this cap w.h.p.

and $w_t \rightarrow w_*$ in d^{k-1}

↳ so in fact see each sample only once during the dynamics w.h.p.
and we didn't gain anything compared to online SGD

Open question: can we show that batch-1 multi-pass SGD still succeeds with high-probability for $n \ll d^{k-1}$ despite the landscape being not benign?

How tight are these guarantees?

Online SGD requires $n = T \gtrsim d^{k-1}$ samples

and runtime $Td \gtrsim d^k$ (each step compute a d -dim gradient)

* for parity fcts: SQ lower bound is $\Omega(d^k)$
 $\rightarrow$ match best SQ runtime

* for Hermite-k: CSQ lower bound is $\Omega(d^{k/2})$

$\rightarrow$ modifying the dynamics, online SGD succeeds
 with $n = T \asymp d^{k/2}$ (runtime $d^{\frac{k}{2}+1}$)

[Damian, Nichani, Ge, Lee, '23]

* for Hermite-k: SQ lower bound is much smaller
 $\Omega(d)$

$\rightarrow$ taking L^1 loss $\ell(y, \hat{y}) = |y - \hat{y}|$

online SGD succeeds in $n = T \asymp d$

$\rightarrow$ reusing samples with L^2 loss

full-batch GD with $n = \tilde{\omega}(d)$ and $T = \tilde{O}_d(1)$
 succeeds

* There exists single index functions that seem to have a fundamental statistical/computational trade-off

$n = \Omega(d)$ is info-theoretically optimal

$n = \Omega(d^{1/2})$ samples is necessary (and sufficient) for a polynomial-time algo to exist.

(under low-degree poly conjecture)

[See Korts and Lede presentation !!!]