

L01

Plan:

- ⑤
- lectures outline
 - prereqs: no IT today, but rev 5 min

15 min ① Motivation: tokenization, weight & act. quant.
+ classical name: lossy compression.

② Scalar quant: Panter-Dite
→ Smooth Quant error
→ dithering.

③ Non-MSE error:
→ motivation from nat net
→ L₁DLQ
→ YAKA?

① Lecture outline:

- "weights-only" {
- LO1: Scalar quantization
 - LO2: Vector Q. (VQ-VAE + IT)
 - LO3: Info-theory of VQ
- "wt + activation" {
- LO4: L&LM & Matmul Quantization
 - Open problems

① Motivation

→ Classic: signals → bits

→ Modern 1: L&LMs spend most time doing
thus $Y = WX$.

High Rate
 $R = 1 \dots 4$

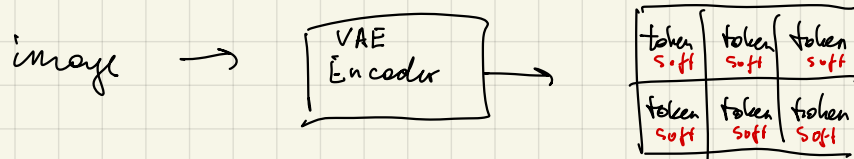
- GPU does 600-1000 FLOPS during the time it takes to load 1 byte to the multipliers!

↳ want to store w and x compactly
(FP32 → BF16 → int8 → int4, FP4, ...)

Energy: load FP32 from DRAM = 640 pJ
one FP add-mult = 4 pJ.
(int add-mult = 3 pJ)

→ Modern 2: tokenization of non-language modalities

text → $\langle \text{token 1} \rangle \langle \text{token 2} \rangle \dots$



Some arch's use soft tokens in MLM.

(DALL-E2+, Gemma3, Llama 3.2-Vision...,
Llava, ... Audio: OpenAI whisper

Others first convert them to discrete "hard" tokens:

Low Rate

$R < 1$

$R = 0.25$ VQ-GAN

$R = 0.4$ VQ-ViT

SoundStream?

(VQ-GAN, DALL-E1, ByteDance VAE
Audio: everyone except whisper
Meta Encoder
Google SoundStream
Boson.AI Migg52 ...)

→ Personal relevant theorem for LLMs.

(2) Scalar quantization.

(a) Simplest form.

$$x \rightarrow q_\varepsilon(x) = \varepsilon \left\lceil \frac{x}{\varepsilon} \right\rceil$$

----->

$$x = \text{int}() \quad \underline{\underline{RIN}}$$

Problem: ∞ number of integers

Classical: user's job is to make sure
 $x \in [-1/2, 1/2] \Rightarrow \# \text{ bits} = R = \log_2 2/\varepsilon$
Every input to your phone passes through ADC

Question: What is error?

If $X \sim P_X$ - dts density on $[-1, 1]$ $\Rightarrow$ unif. dts
 $\Rightarrow$ as $\varepsilon \rightarrow 0 \quad P_X(q_\varepsilon(x)) = \frac{1}{2\varepsilon} \rightarrow \text{Unif}[-1/2, 1/2]$

$$\Rightarrow \mathbb{E}(x - q_\varepsilon(x))^2 = \varepsilon^2 \frac{1 + o(1)}{12} \approx \frac{1}{12} 2^{-2R}$$

Remark: In sigproc quality is usually measured in dB:

$$10 \log_{10} \frac{1}{12} 2^{-2R} \approx -11 + (6 \text{ dB}) \times R$$

$$\text{I.e. } 1 \text{ bit} = 6 \text{ dB} \quad \left(\begin{array}{l} 4 \times \text{reduction} \\ \text{of MSE} \end{array} \right)$$

So 16 bit digitized audio is
often said to have SNR = 96 dB

Uniform Q. is still going strong today

(Tim Dettmers
LLM.int8()
 $\rightarrow$ bits and bytes ...)

6 Dithering:

Problem: • When analysing error we had to take $\epsilon \rightarrow 0$.
Is there a way to make analysis rigorous w/o $\epsilon \rightarrow 0$?

Yes! Dithering! Also helps with actual performance.

History: Invented by Royal Navy $\approx$ WW1 when navigation device was noticed to have poor perf. at stationary.

• Given X let U be indep. dither $\text{Unif}[-\epsilon/2, \epsilon/2]$
(In practice, pseudo random $\rightarrow$ shared b/w compressor and decompressor)

$$\hat{X} = q(x+u) - u.$$

shared
gets stored/sent.

- Equiva to randomly shifting ϵ -grid
- helps when $x_1 = x_2 = \dots = x_{100} \approx \epsilon/3$

$$q(x+u) = \begin{cases} 1 & \text{if } x+u \in [0, \epsilon/2] \\ 0 & \text{if } x+u \in [\epsilon/2, \epsilon] \\ 0 & \text{if } x+u \in [-\epsilon/2, 0] \end{cases}$$

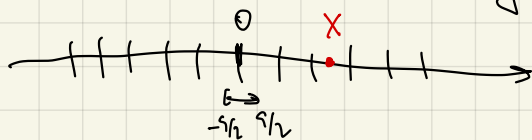
$$\mathbb{E} q(x+u) = x$$

Prop. (Crypto Lemma)

$$(\hat{X}, x) \stackrel{(d)}{=} (x + \tilde{U}, x) \quad \tilde{U} \perp x$$

Pf: $\hat{X} - x \mid x$ is $\text{Unif}[-\epsilon/2, \epsilon/2]$

Indeed, fix x then $\hat{X} = \text{n.v. on a } \epsilon\text{-grid}$ randomly displaced by $(\epsilon/2, \epsilon/2)$



$$q(x_0+u) = \begin{cases} 0 & , u < \epsilon/2 - x_0 \\ \epsilon & , u \geq \epsilon/2 - x_0 \end{cases}$$

$$\hat{x} - x = q(X+u) - u - x = . - (\text{rounding error of } x+u)$$

but $X \pm u$ given X is equally likely to appear anywhere in the bin.

- This is helpful, because now for Dithered Unif. Ques we can rigorously state

$$\mathbb{E}(x - \hat{x})^2 = \epsilon^2/12.$$

and do things like: what is error of x^2 ?

$$\begin{aligned} \mathbb{E}(x^2 - \hat{x}^2)^2 &= \mathbb{E}(x^2 - (x+u)^2)^2 = \\ &= \mathbb{E}(-2xu - u^2)^2 = (\mathbb{E}x^2)\epsilon^2/3 + \epsilon^4/80. \end{aligned}$$

we will see later: $\mathbb{E}(xy - \hat{x}\hat{y})^2 = \mathbb{E}(u_2x + u_1y + u_1u_2)^2$
 $= \frac{\epsilon^2}{12} (\mathbb{E}x^2 + \mathbb{E}y^2 + \frac{\epsilon^2}{12})$

- Empirical fact: DUQ analysis usually perfectly matches empirical perf. of UQ. on real world data., but is fully rigorous.

(This is important trick to do rigorous proofs in this field)

(C) Can we do better than uniform q.?

NR! Do not talk about Mx part

Example: 2024 MxFP4 standard (GPT-oss q.)

$$(x_1, \dots, x_{32}) = 2^{0 \dots 255} \times (\underbrace{\tilde{x}_1}_{\text{FP4}}, \dots, \underbrace{\tilde{x}_{32}}_{\text{FP4}})$$

$$R = 4.25$$

(vector q.)

$$\text{FP4} = 1S, 2E, 1M = \pm 2^{10, 1, 25} (1 + 2^{-1} \cdot \{0, 1\})$$

$$\pm (2^{-1} \{0, 1\})$$

$$\text{so we get: } \pm 0, \pm \frac{1}{2}, \pm 1, \pm \frac{3}{2}, \pm 2, \pm 3, \pm 4, \pm 6.$$

The above format was chosen for legacy & computational reasons.

Let us understand how can we do better?

$$D_{\text{scalar}}(R, P_x) = \inf_{|q(x)| \leq 2^R} \mathbb{E} |x - q(x)|^2$$

- Ugly non-convex problem.
- Solution unknown even for $P_x = \mathcal{U}(0, 1)$ and $2^R = 3$

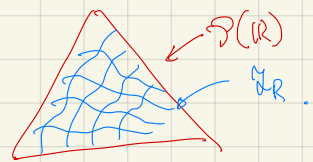
- Wasserstein projection.

- But some optimality claims can still be made.

- Given $q(R) = \{c_1, \dots, c_m\}$ - codebook.

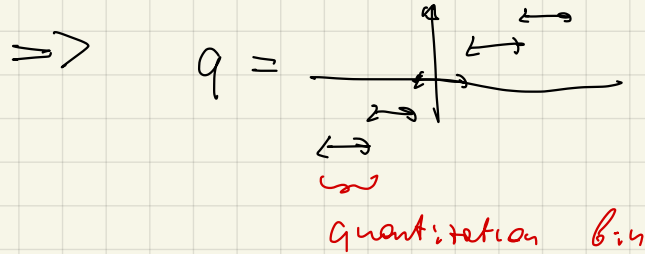
the map q is clear: $q(x) = \text{nearest nb in } C$

This is Wasserstein projection of P_x on $\mathcal{Q}_R = \{Q_x : \text{supp } Q_x = 2^R\}$



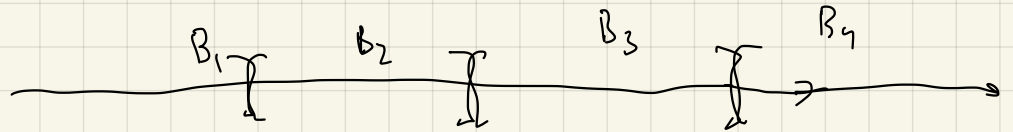
Prop

$$D_{\text{scalar}} = \inf_{P_x} \mathbb{E} |x - \pi|^2 : \text{supp } P_x \leq 2^R$$



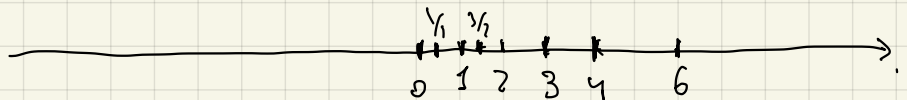
monotone,
piecewise constant

- Conversely if quant. bins are chosen



then $C_i = \underset{u}{\operatorname{argmin}} E[(x-u)^2 | x \in B_i] = E[x | B_i]$

- Immediate implication. Suppose you quantize X via FP4.



If $q(x) = 6$ Is it a good idea to
say $x \approx 6$ vs $x \approx E[X | X > 5]$?
(Note that block scaling tries
to move $P_X \subset [-6, 6]$)

I am waiting for a 6000 cit. paper
from DeepMind on this.

(d) High-rate quantization (Ponter-Dite).

[Th]. Suppose $X \sim P_x$, with cts bdd pdf f_x

then

$$D_{\text{scalar}}(R) = \frac{1}{12} \left(\int f_x^{1/3} dx \right)^3 2^{-2R} (1+o(1)), R \rightarrow \infty$$

Pf:

Introduce: $\lambda(x) \rightarrow$ point density. $\int \lambda(x) dx = 1$
 • put N points c_j at s.t. $\int_{\Delta_j} \lambda(x) dx = 1/N$.
 • $\Rightarrow$ in interval Δ $I \approx N \int_I \lambda(x) dx$ pts.

$$E |X - q(X)|^2 = \sum_j E [X - c_j]^2 | X \in \Delta_j] P[X \in \Delta_j]$$

$$\approx \sum_j \frac{\Delta_j^2}{12} P[X \in \Delta_j]$$

$$\approx \sum_j \frac{1}{12 N^2} \lambda(c_j)^{-2} P[X \in \Delta_j]$$

$$\approx \frac{1}{12 N^2} \int \frac{P_x(x)}{\lambda^2(x)} dx$$

$\Delta_j \approx \frac{1}{N}$
 $\int_{\Delta_j} \lambda(x) dx = 1/N$
 $\Delta_j \approx \lambda(c_j)^{-1} \frac{1}{N}$

$\lambda^* = \frac{\int P_x^{1/3}}{\int P_x^{1/3}}$

$$D_{\text{scalar}}(P_x, R) \approx \frac{1}{12} 2^{-2R} \left(\int P_x^{1/3} \right)^3, R \rightarrow \infty.$$

For G_{SN} : $P_x = \mathcal{N}(0, \sigma^2)$: $D \approx \frac{\pi \sqrt{3}}{2} \sigma^2 2^{-2R}$

By Reverse Hölder:

$$\int P_x \frac{1}{\lambda^2} \geq \|P_x\|_p \left\| \frac{1}{\lambda^2} \right\|_q$$

set $q = -1/2, p = 1/3$

$$\Rightarrow E |X - q(X)|^2 \geq \frac{1}{12 N^2} \left(\int P_x^{1/3} \right)^3$$

Mention naive argument:
 each bin $\approx 1/N$ prob-ty
 to maximize $H(\hat{X})$,
 (used in Smooth Quant etc)
 completely wrong!

Preview: Shannon's breakthrough was to show
 that $D_{\text{vector}}(R; \mathcal{N}(0, \sigma^2)) = \frac{\pi \sqrt{3}}{2} \sigma^2 2^{-2R} < D_{\text{scalar}}$!

Note:

if domain of X is $[-1/2, 1/2]$ then.

$$\left(\int p_k^{1/3}\right)^3 \leq \left(\int p_k\right) = 1 ! \quad \text{by Jensen's}$$

$\Rightarrow$ Non-Uniform quant. provably wins.

3 What about quantizing a vector?

a Uniform case

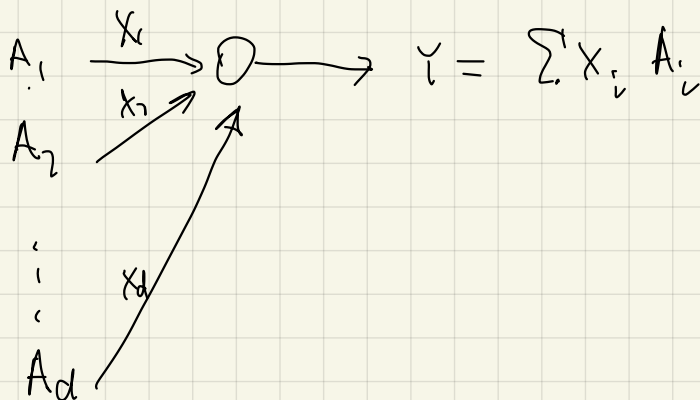
$$\begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix} \rightarrow \begin{bmatrix} q(x_1) \\ \vdots \\ q(x_d) \end{bmatrix}$$

Is that it?

b Modern case.

X - weights of linear layer.

A - activations from prev. layer.



$$\begin{aligned} \text{Mean Squared Error} \\ \mathbb{E}(\hat{y} - y)^2 \\ \mathbb{E}(\hat{y} - \sigma(XA))^2 \\ \mathbb{E}(\sigma(XA) - \sigma(\hat{X}A))^2 \\ \mathbb{E}(\sigma(XA) - \sigma(\hat{X}A))^2 \\ \mathbb{E}(\sigma(XA) - \sigma(\hat{X}A))^2 \\ \mathbb{E}(\sigma(XA) - \sigma(\hat{X}A))^2 \end{aligned}$$

We want to find $\hat{X} \in \mathbb{R}^{d \times d}$ s.t.

$$\min \mathbb{E} (X^T A - \hat{X}^T A)^2$$



$$\min_{\hat{X} \in \mathbb{R}^{d \times d}} (X - \hat{X})^T H (X - \hat{X})$$

$$H = \mathbb{E} A A^T$$

$$\text{or } \min \mathbb{E} (\sigma(XA) - \sigma(\hat{X}A))^2$$

$$\text{or } H = \mathbb{E} \sigma'(X^T A)^2 A A^T$$

$$\textcircled{C} \min_{\hat{x} \in \mathbb{Z}^d} (x - \hat{x})^T Q (x - \hat{x}).$$

Claim: This is post-quantum hard.

$$\text{Def:} \quad \min_{\hat{x}} (x - \hat{x})^T Q Q (x - \hat{x}) \Leftrightarrow \min_{\hat{y} \in \Lambda} \|y - \hat{y}\|^2$$

$\Lambda = \mathbb{Z}$ -span of col's of Q

Nearest nb. search in Lattice...

$\textcircled{d}$ Popular heuristic (GPVQ):

L₂D₂Q notes:

Target: $\min_{\hat{x} \in \mathbb{Z}^d} (x - \hat{x})^T H (x - \hat{x}), H > 0.$

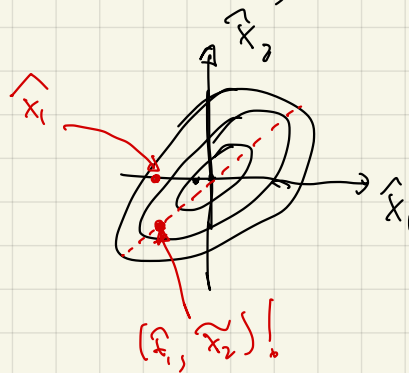
• Claim 1: NP-hard

Pf: Reduce to NN lattice quant.

• great heuristic: (OBC, GPTR, L₂D₂Q) arxiv: 2210.17323
arxiv: 2307.13304

1) $\hat{x}_1 = [x_1]$

2)



Now the minimizer $\arg\min_{\hat{x}_2^d} f(\hat{x}_1, \hat{x}_2, \dots, \hat{x}_d) \neq x_2^d!$

$$3) \arg\min_{\hat{x}_2^d} \begin{bmatrix} (x_1 - \hat{x}_1) \\ e \end{bmatrix}^T H \begin{bmatrix} x_1 - \hat{x}_1 \\ e \end{bmatrix} = \begin{bmatrix} H_{11} & H_{21}^T \\ H_{21} & H_{22} \end{bmatrix}$$

$$= \Delta_1^2 H_{11} + 2 \Delta_1 H_{21}^T e + e^T H_{22} e.$$

$$= \text{const} + (e - e^*)^T H_{22} (e - e^*)$$

$$e^* = -2 H_{22}^{-1} H_{21} \Delta_1, \quad 2 \Delta_1 H_{21} + H_{22} e \geq 0.$$

$$x_2^d - \hat{x}_2^d = -2 H_{22}^{-1} H_{21} \Delta_1$$

$$\hat{x}_2^d = x_2^d + 2 H_{22}^{-1} H_{21} \Delta_1.$$

4) From P.O.V. of $\hat{x}_2^d$ now we have the

$$f(x_1, \hat{x}_2, \dots, \hat{x}_d) = \text{const} + \underbrace{(\hat{x}_2^d - \hat{x}_2^d)^T H_{12} (\hat{x}_2^d - \hat{x}_2^d)}$$

5) thus
 $d \leftarrow d-1$
 $x \leftarrow \hat{x}_2^d$
 Repeat!

6) Does it matter? GPTQ paper on OPT-175B, 4-bit

	FP16	8.34
perp.w?	RTN	10.54
	GPTQ	8.37

98.5% reduction in gap!!

OPT-66B is even worse: RTN fails! GPTQ gives 9.55 vs 9.34.

7) Note: this time $\Delta = x - \hat{x}$ has non-diagonal cov. matrix. (errors are correlated). How to find this error dist.?

$$\begin{aligned} \Delta_k &= X_k - q\left(X_k + \sum_{j=1}^{k-1} a_{kj} \Delta_j\right) \\ &\approx X_k - \left(X_k + \sum_{j=1}^{k-1} a_{kj} \Delta_j + z_k\right) \\ &= -\left(\overset{\text{strictly low-triang.}}{A} \Delta\right)_k + z_k. \end{aligned}$$

$$\Delta \approx \left(\overset{\text{lower-triang.}}{I + A}\right)^{-1} z$$

$$\mathbb{E} \Delta^T H \Delta = \text{tr} H \Sigma_\Delta = \frac{1}{12} \text{tr} H \left(\overset{\text{lower-triang.}}{I + A}\right)^{-1} \left(\overset{\text{lower-triang.}}{I + A}\right)^{-T}$$

Note: If $H = (I + B)^T D (I + B)$

$$\text{tr} D \underbrace{(I + B)(I + B)^T}_{\mathbb{E}} E^T = \sum_{i,j} d_{ij} e_{ij}^2 \geq \sum_i d_{ii}$$

$$\begin{bmatrix} 0 & x \\ 1 & 1 \end{bmatrix}$$

$$U_{tr} \in \mathcal{U}_{tr}$$

$\Rightarrow$ so we see that if $M = (I + B)^T D (I + B)$
then $A = B$! this allows to skip
submatrix inversions. B is computed
via LDL decomposition.

History: QBC was via inversion
 $Q^T P Q / LDL^T Q \rightarrow$ via LDL / Cholesky