

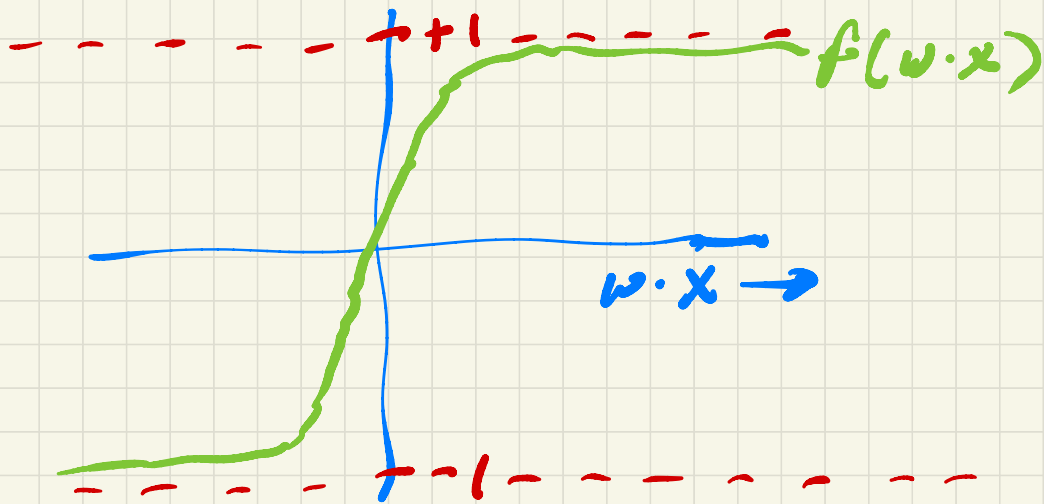

Ethical Algo Design in the Generative Era

Generative AI

- What is it?
- E.g. LLMs
- Next-word prediction
- Training & annotation
- History / precursors
- Embeddings
- Autoencoders / decoders

Neural Net Background

- single "neuron": on input $x \in \mathbb{R}^d$,
output = $f(w \cdot x)$
 $\approx f(w_1 x_1 + w_2 x_2 + \dots + w_d x_d)$
- typical choice for f : sigmoid

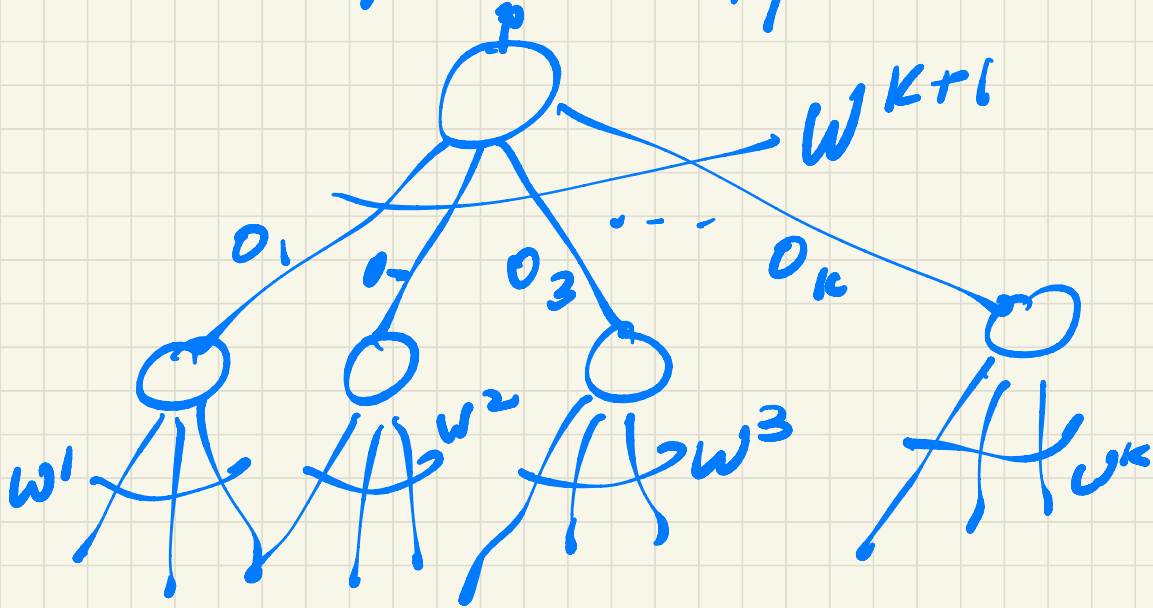


- training: weights w are parameters, loss fn is
$$L(w) = \sum_{x, y \in S} (f(w \cdot x) - y)^2$$

differentiable

Multiple Layers

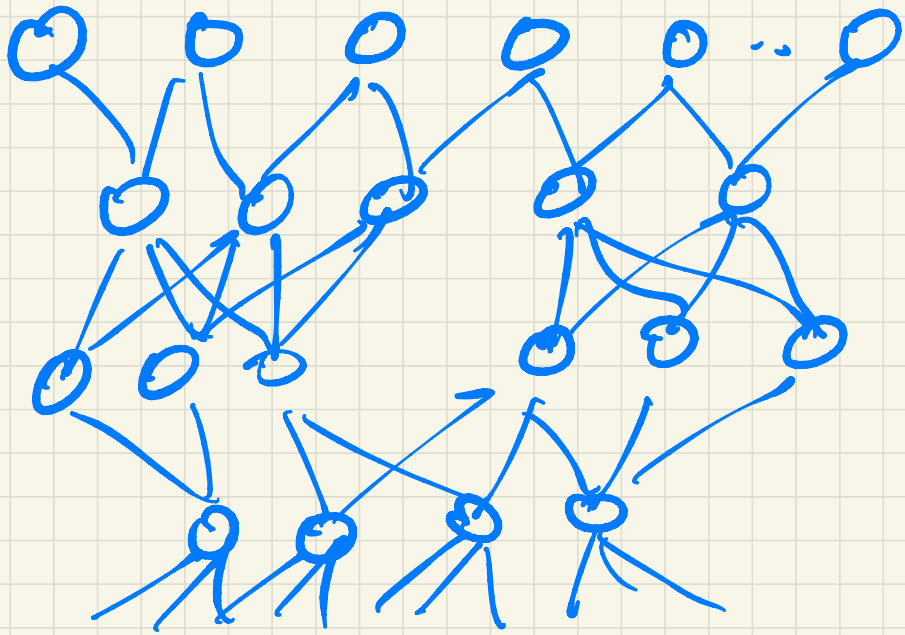
final output



- params now $w^1, w^2, \dots, w^{k+1}$
- still differentiable
(chain rule $\rightarrow$ backprop)

key idea: intermediate
neurons can learn
new/better features

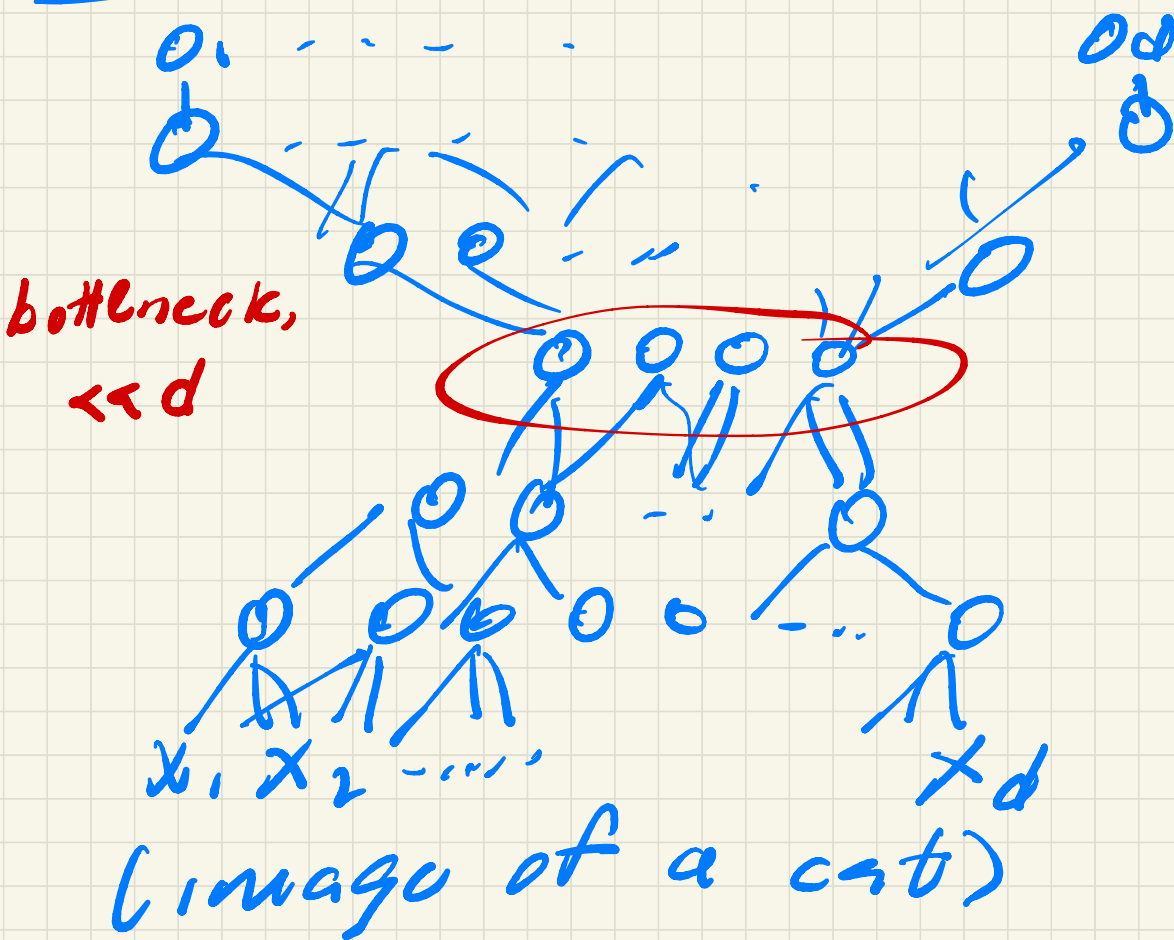
Multiple Outputs



$x_1, x_2, \dots, x_d$

- now outputs are high dimensional, still differentiable

Autoencoders



goal: minimize
$$\sum_{i=1}^d (o_i - x_i)^2$$

Fairness in GAI

- Compare e.g. group fairness in consumer lending
- Fairness in propensity?
- In tone?
- Biases vs. correlations, use cases

Toxicity

- Subjectivity
& context
- Quotations
& censorship
- Offensive
opinions
- Nuance &
indirection

Toxicity & Fairness

- Training data curation
- Must automate
 - human annotation (unpleasant)
- Train "guardrail" models
- Apply to training prompts & output

Privacy in GAI

- Regurgitation
- Virtual copying (e.g. code)
- Stylistic mimicry (more later)
- Could do DP model training, but impractical
- Robustness to membership inference attacks

Copyright & IP Concerns

- Stylistic mimicry
- Threat to creative livelihoods?

Copyright & IP

- Tech, policy & legal mechanisms
- Content attribution & compensation
- Machine unlearning & model disgorgement
 - DP
 - Sharding

Adversarial Attacks

- E.g. image manipulation, prompt injection, data poisoning, spam filter attacks, ...
- Robust adversarial training:
 - instead of minimizing e.g.
$$\sum_{\langle x_i, y_i \rangle} (h(x_i) - y_i)^2$$
 - minimize something like

$$\max_{\delta: \|\delta\| < \alpha} \sum_{\langle x_i, y_i \rangle} (h(x_i + \delta_i) - y_i)^2$$

Plagiarism & Cheating

- College essays, writing samples, etc.
- Tool vs. cheat?
- Author verification

Author Verification

- Train model to detect human/LLM (arms race)
- Redlist/greenlist (watermarking)

Hallucinations

- plausible but verifiably wrong
- E.g. citations

Hallucinations

- User education
- Independent verification
- External sources (e.g. citation DBs)
- Disclaimers
- Content attribution

A bit more on hallucinations...
(Kalai & Vempala, MM & MK, HW2)

- stylized facts, e.g.
IMDB records - unambiguous
true/false (also citations)
- schema or tuples:
<author, ..., title, year, journal>
- huge space of "factoids",
only a fraction are facts
- training data contains
only facts
- test: only generate
facts from training
- But in ML we want
to generalize...

Missing Mass Estimation

- How many facts are missing in data?
- Good-Turing: let $\zeta = \frac{\text{\#facts occurring only once in training}}{n}$

• Then

/ prob. next sample is a new fact $\approx \zeta / \zeta \frac{1}{\sqrt{n}}$

- Kato-I-Vempala: for any model

hallucination rate $\geq \zeta$ - model "miscalibration"
(one measure of overfitting)

Disruption to Work

- Passing bar exams, MBA courses, etc.
- IP concerns again

Job disruption/creation

- Prompt engineers
- "English is the new programming language"

Strongest
defense:

Use case
specialization