

<https://introml.mit.edu/>

6.390 Intro to Machine Learning

Lecture 4: Linear Classification

Shen Shen

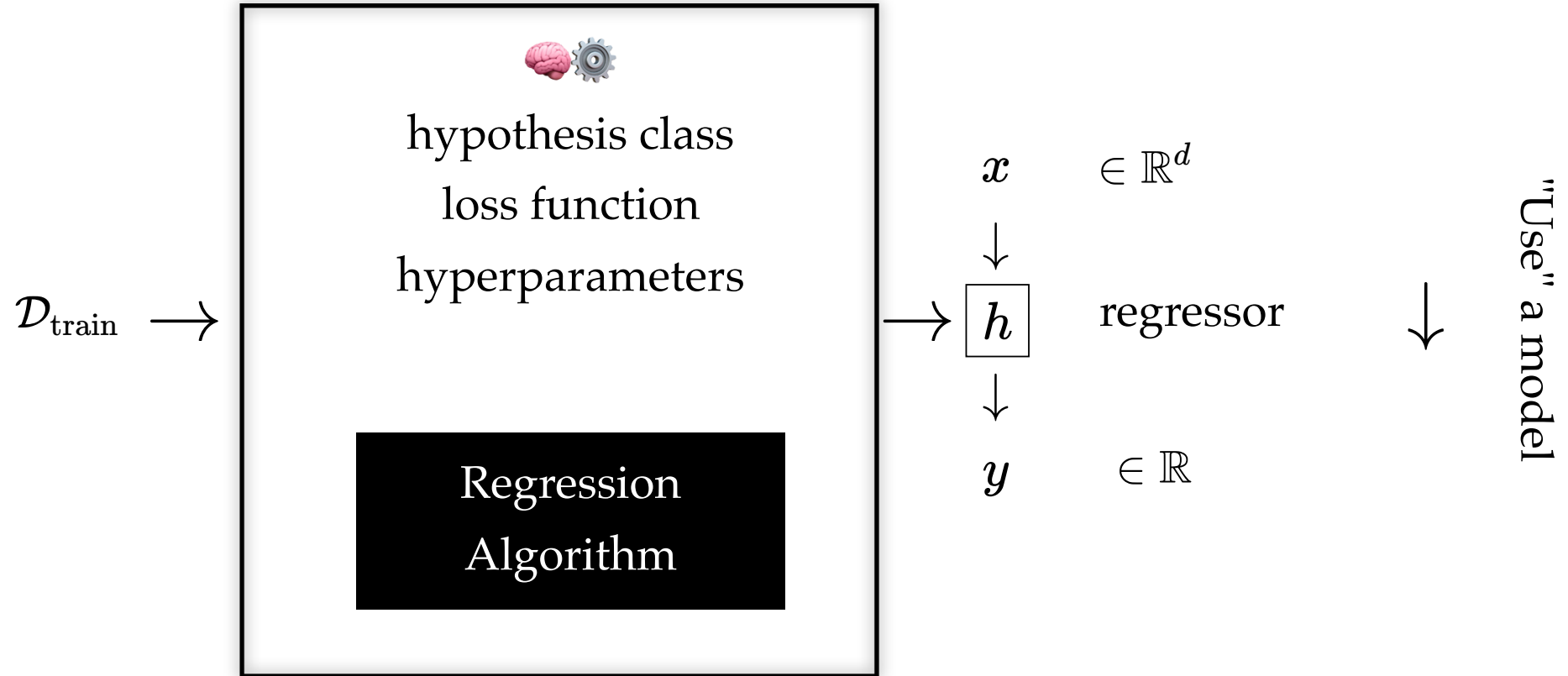
Sept 25, 2025

11am, Room 10-250

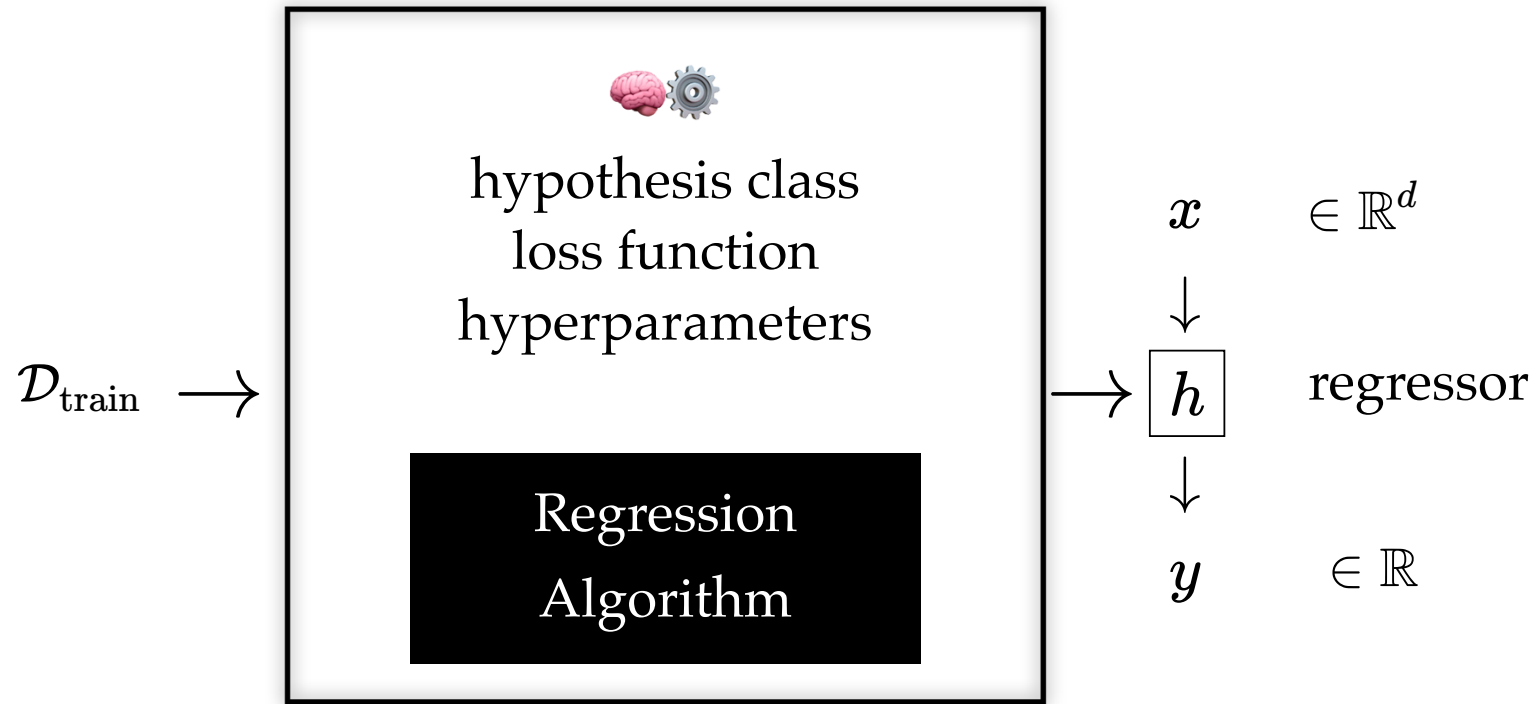
[Interactive Slides and Lecture Recording](#)

Recap:

"Learn" a model



Recap:



"Learn" a model $\rightarrow$

train, optimize, tune, adapt ...

adjusting / updating / finding θ

gradient based

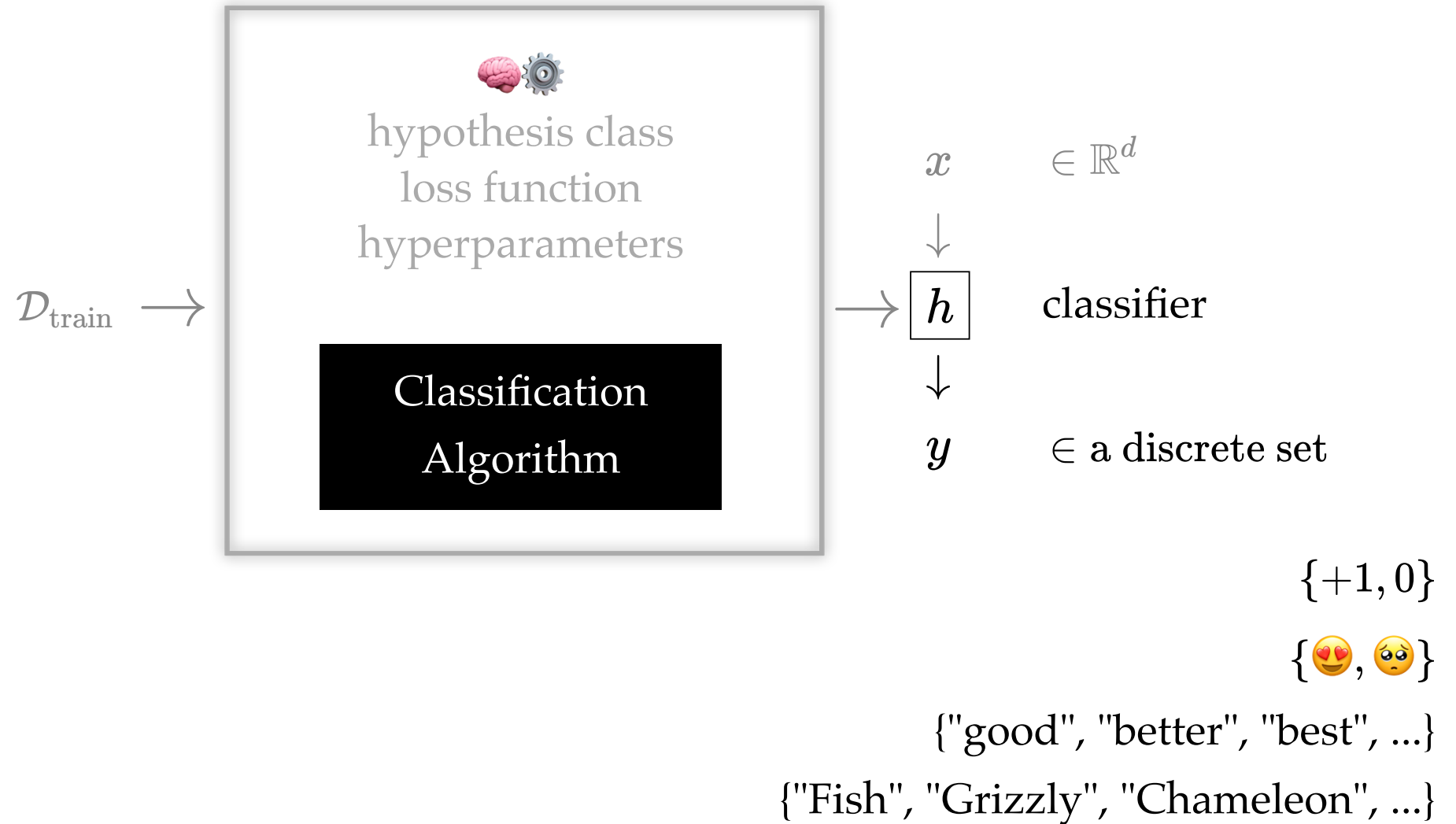
$\downarrow$ "Use" a model

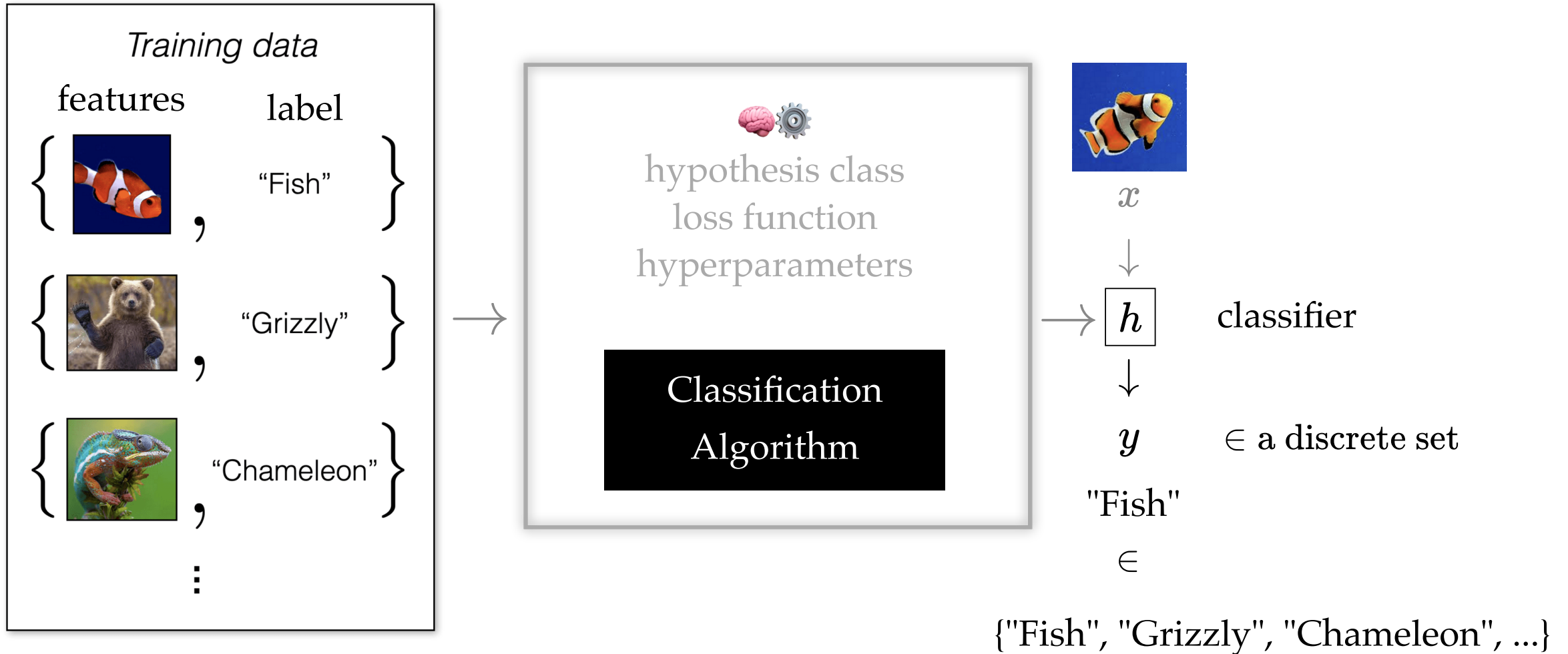
predict, test, evaluate, infer ...

plug in the θ found

no gradients involved

Today:





Outline

1. Linear (binary) classifiers

- to use: **separator, normal vector**
- to learn: difficult! won't do

2. Linear logistic (binary) classifiers

- to use: **sigmoid**
- to learn: **negative log-likelihood loss**

3. Linear multi-class classifiers

- to use: **softmax**
- to learn: **one-hot encoding, cross-entropy loss**

Outline

1. Linear (binary) classifiers

- to use: **separator, normal vector**
- to learn: difficult! won't do

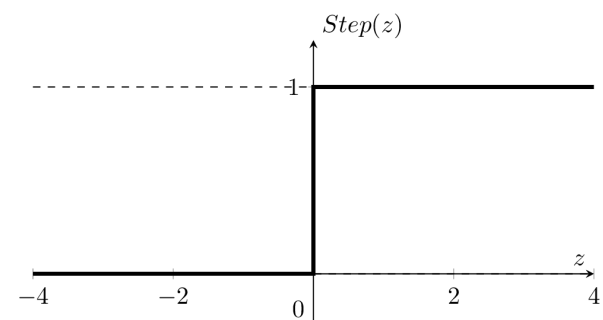
2. Linear logistic (binary) classifiers

- to use: sigmoid
- to learn: negative log-likelihood loss

3. Linear multi-class classifiers

- to use: softmax
- to learn: one-hot encoding, cross-entropy loss

	linear regressor	linear binary classifier
features	$x \in \mathbb{R}^d$	
label	$y \in \mathbb{R}$	$y \in \{0, 1\}$
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$	
linear combination	$\theta^T x + \theta_0 = z$	
predict	$g = z$	$g = \begin{cases} 1 & \text{if } z > 0 \\ 0 & \text{otherwise} \end{cases}$



today, we refer to $\theta^T x + \theta_0$ as z throughout.

Outline

1. Linear (binary) classifiers

- to use: separator, normal vector
- to learn: difficult! won't do

2. Linear logistic (binary) classifiers

- to use: sigmoid
- to learn: negative log-likelihood loss

3. Multi-class classifiers

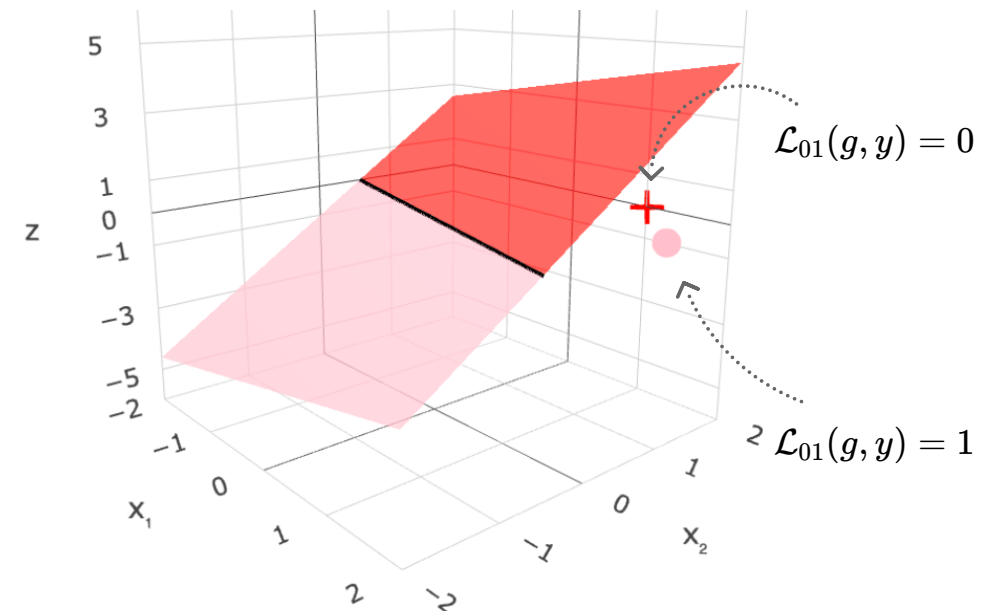
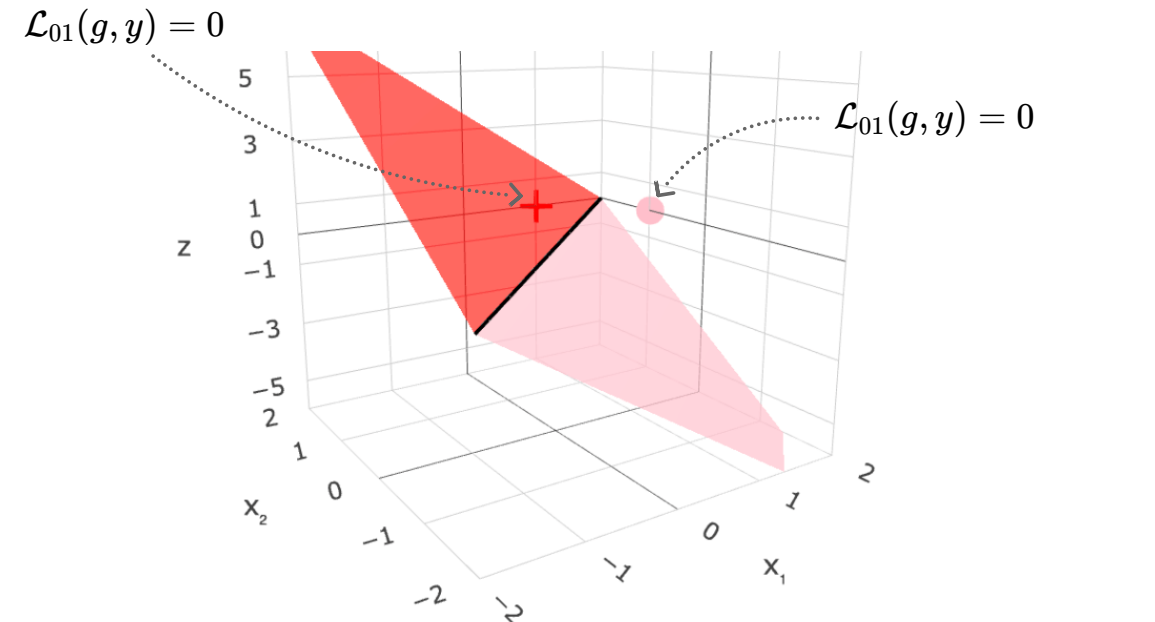
- to use: softmax
- to learn: one-hot encoding, cross-entropy loss

- To *learn* a model, need a loss function.
- One natural loss choice:

$$g = \text{step}(z) = \text{step}(\theta^\top x + \theta_0) \quad y$$

$$\mathcal{L}_{01}(g, y) = \begin{cases} 0 & \text{if guess} = \text{label} \\ 1 & \text{otherwise} \end{cases}$$

- Very intuitive, and easy to evaluate 🥰



<https://shenshen.mit.edu/demos/classification/01loss.html>

Very hard to optimize (NP-hard) 🙄

- "Flat" almost everywhere (zero gradient)
- "Jumps" elsewhere (no gradient)

<https://shenshen.mit.edu/demos/classification/mse.html>

	linear regressor	linear binary classifier
features	$x \in \mathbb{R}^d$	
	$y \in \mathbb{R}$	$y \in \{0, 1\}$
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$	
linear combo	$\theta^T x + \theta_0 = z$	
predict	$g = z$	$g = \begin{cases} 1 & \text{if } z > 0 \\ 0 & \text{otherwise} \end{cases}$
loss	$(g - y)^2$	$\begin{cases} 0 & \text{if } g = y \\ 1 & \text{otherwise} \end{cases}$
optimize via	closed-form or gradient descent	both discrete NP-hard to learn

Outline

1. Linear (binary) classifiers

- to use: separator, normal vector
- to learn: difficult! won't do

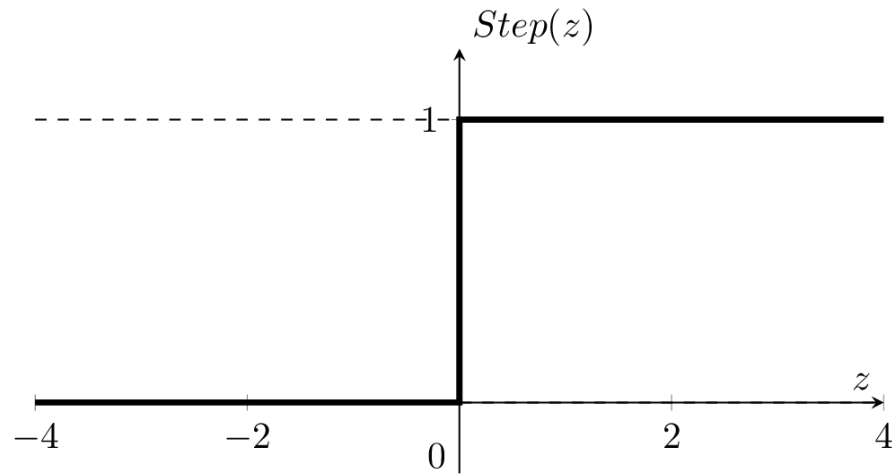
2. Linear logistic (binary) classifiers

- to use: **sigmoid**
- to learn: **negative log-likelihood loss**

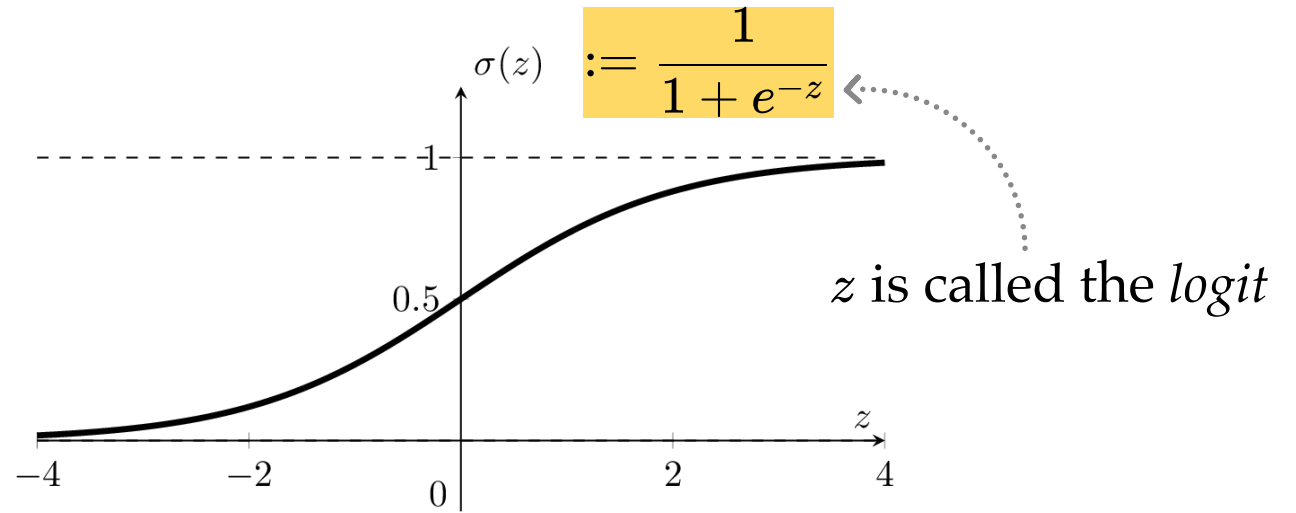
3. Linear multi-class classifiers

- to use: softmax
- to learn: one-hot encoding, cross-entropy loss

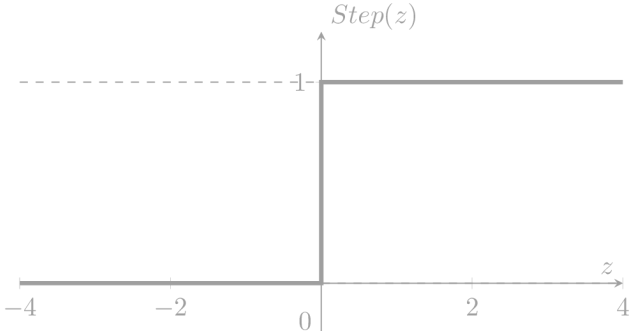
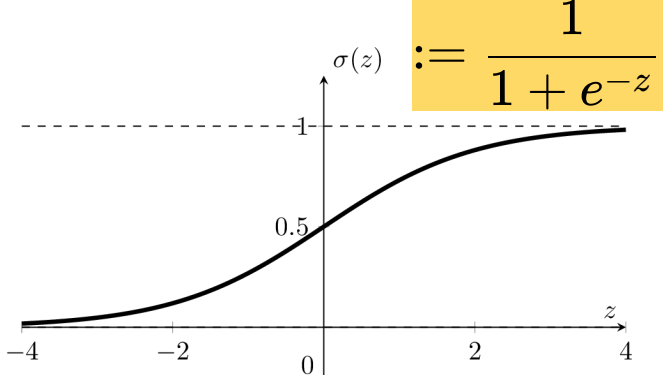
Sigmoid : a smooth step function

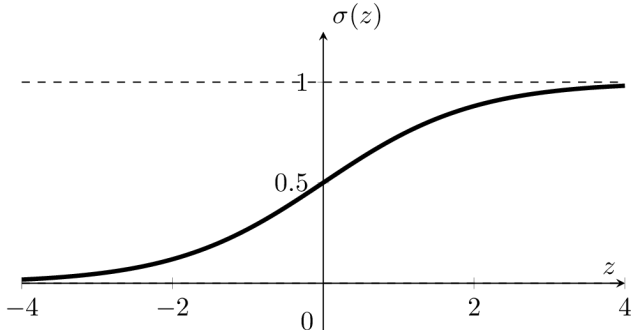


Predict $\begin{cases} 1 & \text{if } z > 0 \\ 0 & \text{otherwise} \end{cases}$



Predict $\begin{cases} 1 & \text{if } \sigma(z) > 0.5 \\ 0 & \text{otherwise} \end{cases}$

	linear binary classifier	linear <i>logistic</i> binary classifier
features	$x \in \mathbb{R}^d$	
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$	
linear combo	$\theta^T x + \theta_0 = z$	
predict	$\begin{cases} 1 & \text{if } z > 0 \\ 0 & \text{otherwise} \end{cases}$ 	$\begin{cases} 1 & \text{if } \sigma(z) > 0.5 \\ 0 & \text{otherwise} \end{cases}$ 

	linear <i>logistic</i> binary classifier
features	$x \in \mathbb{R}^d$
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$
linear combo	$\theta^T x + \theta_0 = z$
predict	$\begin{cases} 1 & \text{if } \sigma(z) > 0.5 \\ 0 & \text{otherwise} \end{cases}$ 

- The logit z is a linear combo of x via the parameters.
- Sigmoid squashes the logit z into a number in $(0, 1)$.
- Predict positive class label 1

$$\begin{aligned} \text{if } \sigma(z) &= \frac{1}{1 + e^{-z}} \\ &= \frac{1}{1 + e^{-(\theta^T x + \theta_0)}} > 0.5 \end{aligned}$$

$\sigma(\cdot)$: the model's *confidence* or *estimated likelihood* that feature x belongs to the positive class.

<https://shenshen.mit.edu/demos/classification/sigmoid.html>

https://shenshen.mit.edu/demos/classification/step_sigmoid.html

- θ, θ_0 can flip, squeeze, expand, or shift the $\sigma(x)$ graph *horizontally*
- $\sigma(\cdot)$ monotonic, very elegant gradient (see hw/lab)

linear logistic binary classifier

features

$$x \in \mathbb{R}^d$$

parameters

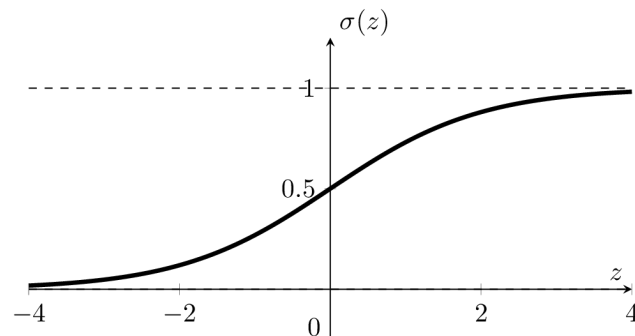
$$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$$

linear combo

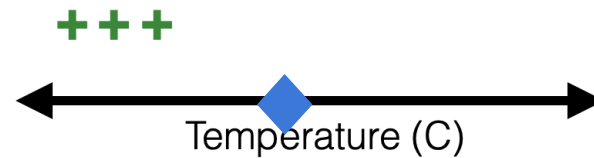
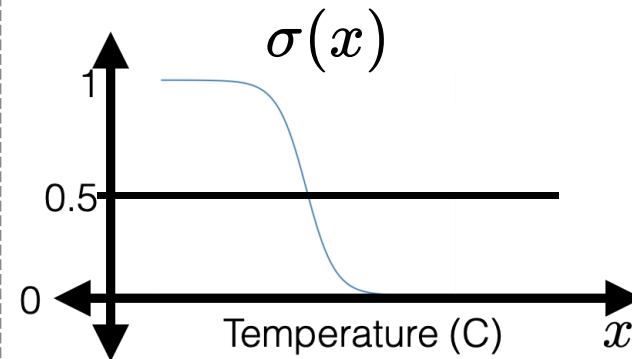
$$\theta^T x + \theta_0 = z$$

predict

$$\begin{cases} 1 & \text{if } \sigma(z) > 0.5 \\ 0 & \text{otherwise} \end{cases}$$



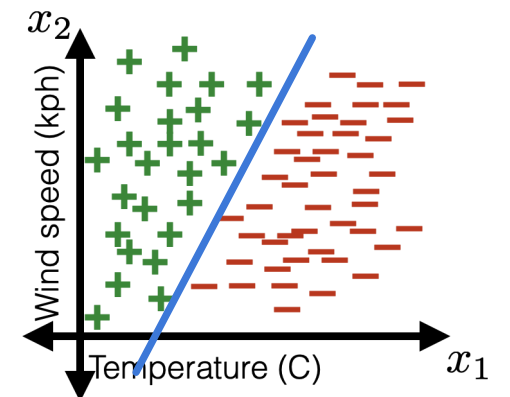
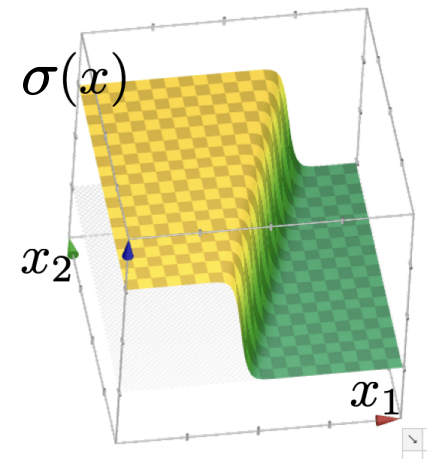
1d feature



$$z = \theta^T x + \theta_0 = 0$$

linear separator

2d feature



Outline

1. Linear (binary) classifiers

- to use: separator, normal vector
- to learn: difficult! won't do

2. Linear **logistic** (binary) classifiers

- to use: **sigmoid**
- to learn: **negative log-likelihood loss**

3. Linear multi-class classifiers

- to use: softmax
- to learn: one-hot encoding, cross-entropy loss

previously: $\mathcal{L}_{01}(g, y) = \begin{cases} 0 & \text{if } g = y \\ 1 & \text{otherwise} \end{cases}$

- if true label $y = 1$, we'd like the guess g to just *be* 1... a difficult ask

now:

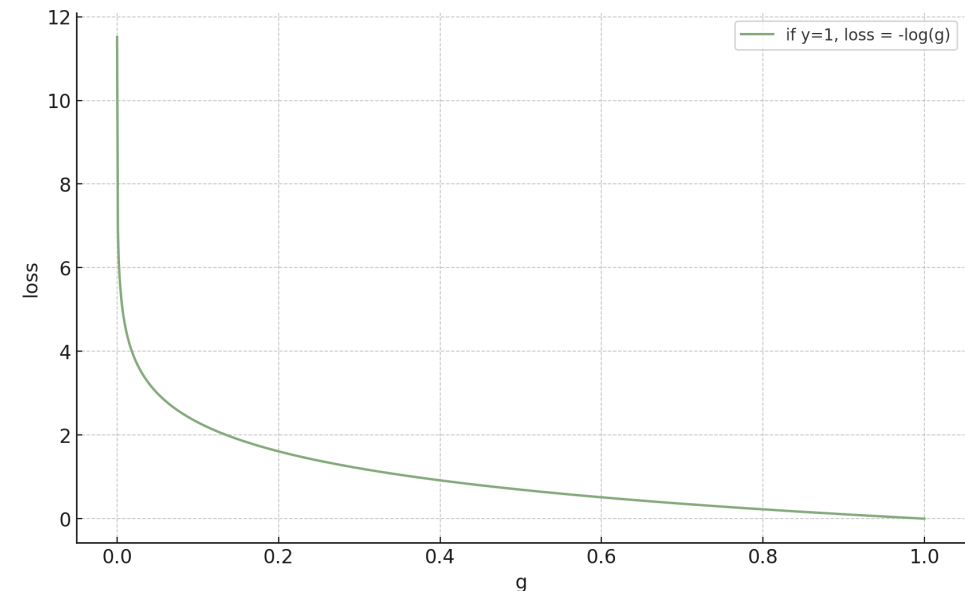
- if true label $y = 1$, we'd like the guess g to *gradually approach* 1



$$\mathcal{L}_{\text{nll}}(g, y) = \mathcal{L}_{\text{nll}}(g, 1) := -\log(g)$$

negative log likelihood

$g = \sigma(\cdot)$: the model's *confidence* or *estimated likelihood* that feature x belongs to the positive class.



training data:

+

← Temperature (C) →

true label y is 1



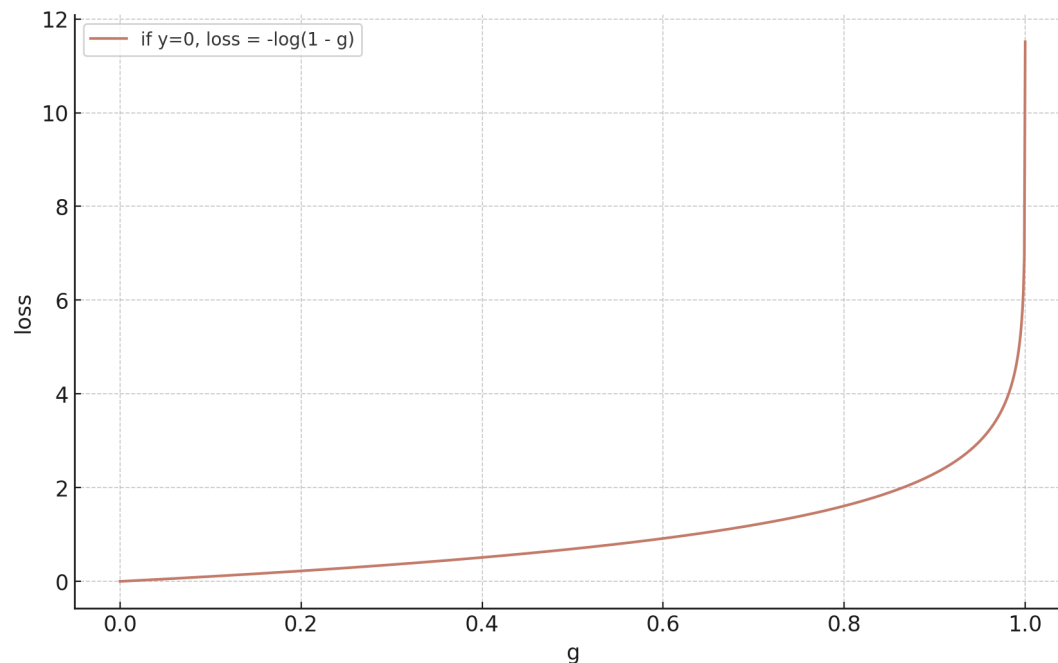
$$\mathcal{L}_{\text{nll}}(g, y) = \mathcal{L}_{\text{nll}}(g, 1) := -\log(g)$$

<https://shenshen.mit.edu/demos/nlloverfit.html>

- if true label $y = 0$, we'd like the guess g to *gradually approach* 0



$$\mathcal{L}_{\text{nll}}(g, y) = \mathcal{L}_{\text{nll}}(g, 0) := -\log(1 - g)$$



$g = \sigma(\cdot)$: the model's *confidence* or *estimated likelihood* that feature x belongs to the **positive** class.

$1 - g = 1 - \sigma(\cdot)$: the model's *confidence* or *estimated likelihood* that feature x belongs to the **negative** class.

training data:



true label y is 0

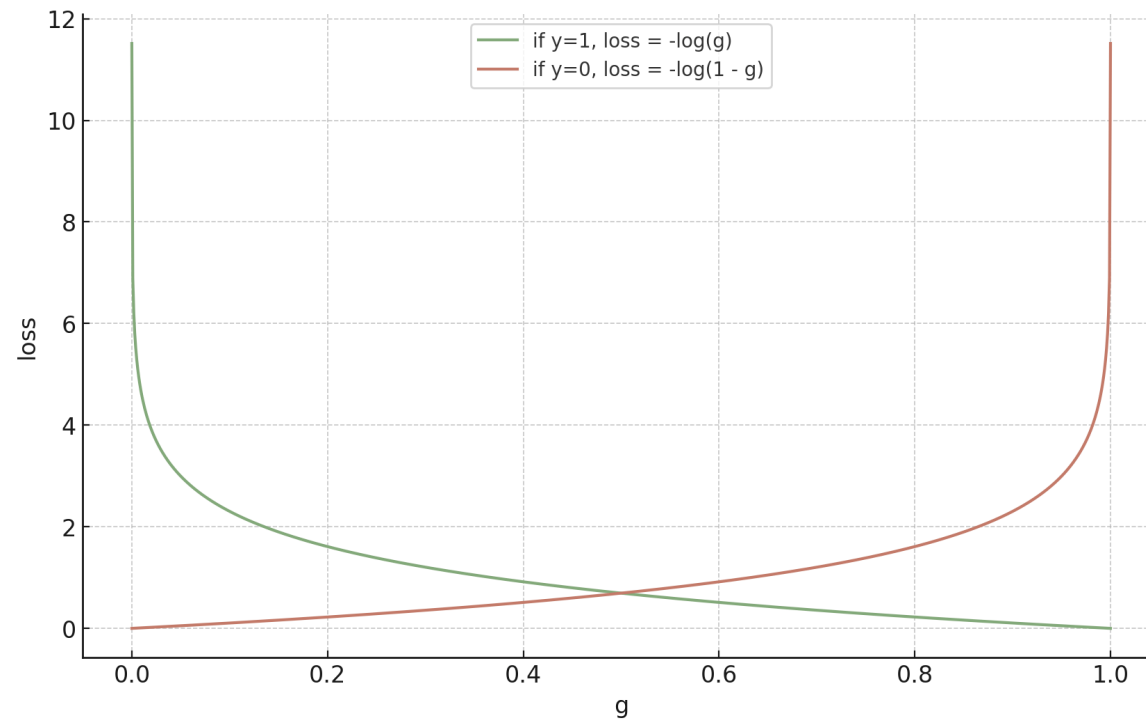


$$\mathcal{L}_{\text{nll}}(g, y) = \mathcal{L}_{\text{nll}}(g, 0) = -\log(1 - g)$$

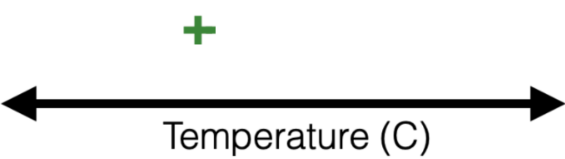
https://shenshen.mit.edu/demos/classification/nlloverfit_negative.html

Combining both cases, since the actual label $\mathbf{y} \in \{+1, 0\}$

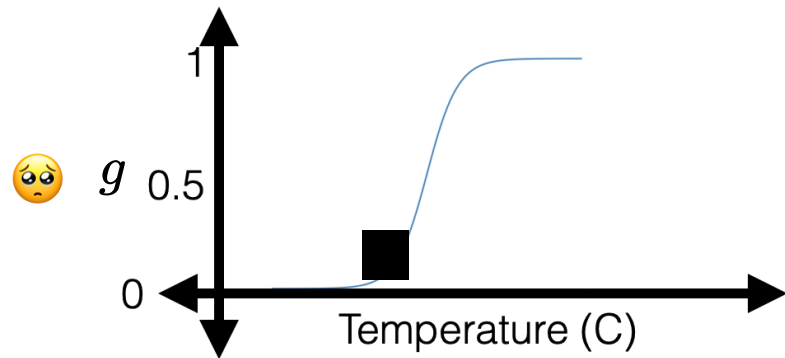
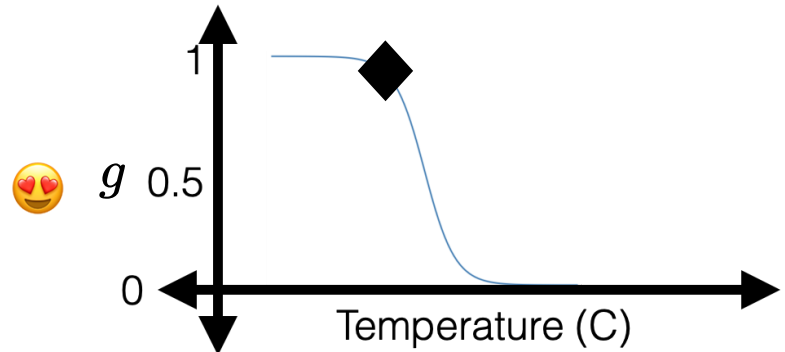
$$\mathcal{L}_{\text{nll}}(g, y) = -[y \log g + (1 - y) \log(1 - g)]$$



When $y = 1$

training
data: 

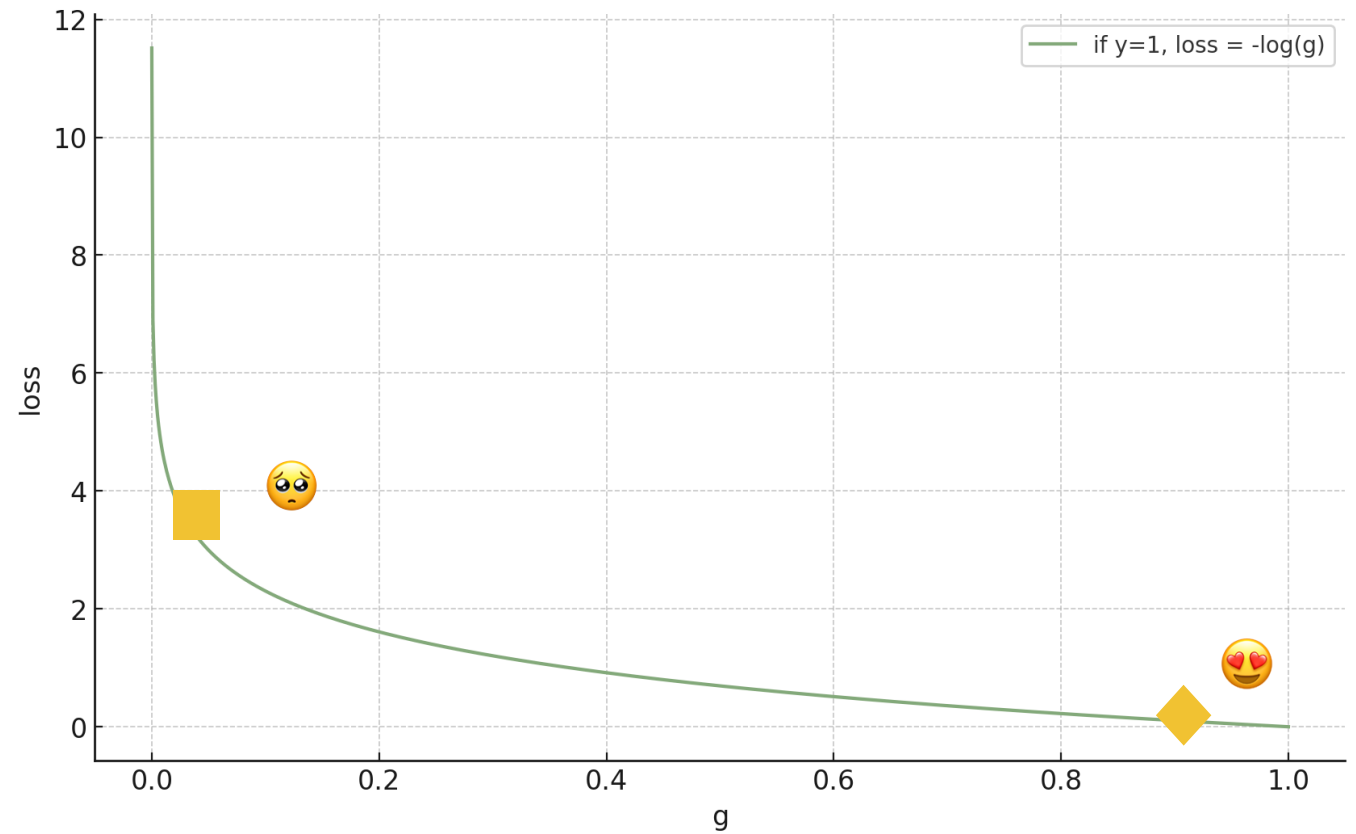
$$g(x) = \sigma(\theta x + \theta_0)$$



$$\mathcal{L}_{\text{nll}}(g, y)$$

$$= -[y \log g + (1 - y) \log (1 - g)]$$

$$= -\log g$$

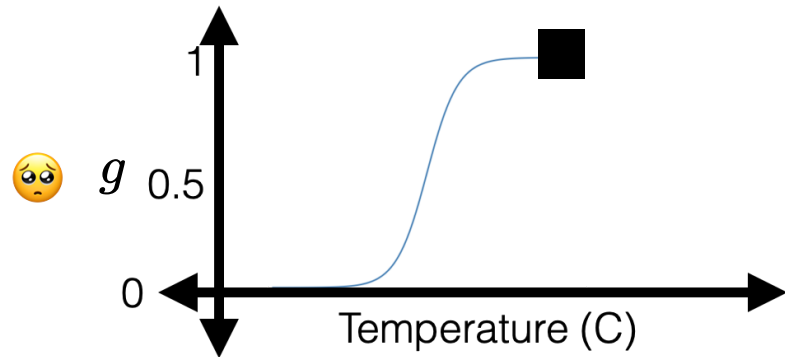
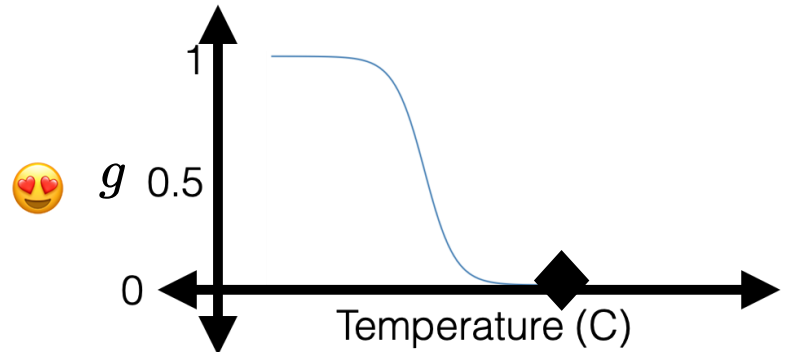


When $y = 0$

training
data:



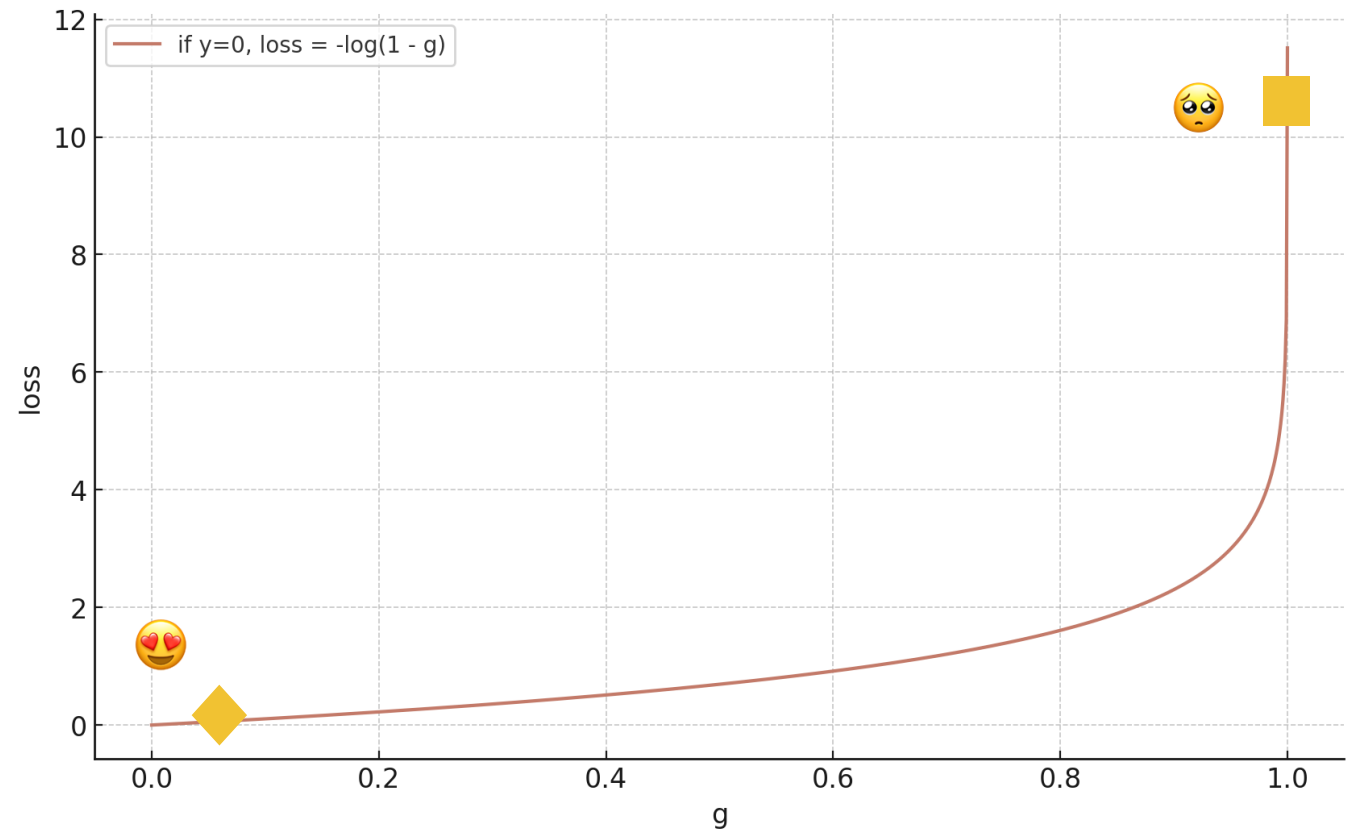
$$g(x) = \sigma(\theta x + \theta_0)$$



$$\mathcal{L}_{\text{nll}}(g, y)$$

$$= -[y \log g + (1 - y) \log (1 - g)]$$

$$= -[\log (1 - g)]$$



training data:

+

← Temperature (C) →

true label y is 1



$$\mathcal{L}_{\text{nll}}(g, y) = \mathcal{L}_{\text{nll}}(g, 1) := -\log(g)$$

<https://shenshen.mit.edu/demos/nlloverfit.html>

	linear regressor	linear binary classifier	linear <i>logistic</i> binary classifier
features	$x \in \mathbb{R}^d$		
label	$y \in \mathbb{R}$	$y \in \{0, 1\}$	
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$		
linear combo	$\theta^T x + \theta_0 = z$		
predict	$g = z$	$\begin{cases} 1 & \text{if } z > 0 \\ 0 & \text{otherwise} \end{cases}$	$\begin{cases} 1 & \text{if } g = \sigma(z) > 0.5 \\ 0 & \text{otherwise} \end{cases}$
loss	$(g - y)^2$	$\begin{cases} 0 & \text{if } g = y \\ 1 & \text{otherwise} \end{cases}$	$-[y \log g + (1 - y) \log (1 - g)]$
optimize via	closed-form or gradient descent	NP-hard to learn	• gradient descent • regularization still important

Outline

1. Linear (binary) classifiers

- to use: separator, normal vector
- to learn: difficult! won't do

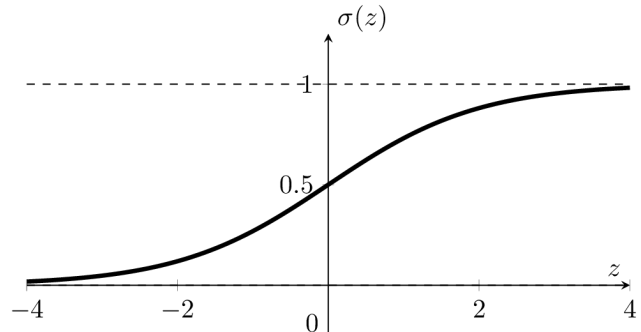
2. Linear logistic (binary) classifiers

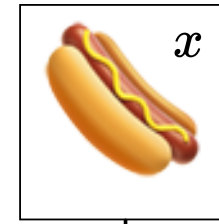
- to use: sigmoid
- to learn: negative log-likelihood loss

3. Multi-class classifiers

- to use: **softmax**
- to learn: **one-hot encoding, cross-entropy loss**



	linear logistic binary classifier
features	$x \in \mathbb{R}^d$
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$
linear combo	$\theta^T x + \theta_0 = z$
predict	$\begin{cases} 1 & \text{if } \sigma(z) > 0.5 \\ 0 & \text{otherwise} \end{cases}$ 



$$z \in \mathbb{R}$$

scalar logit
(raw hotdog-ness)

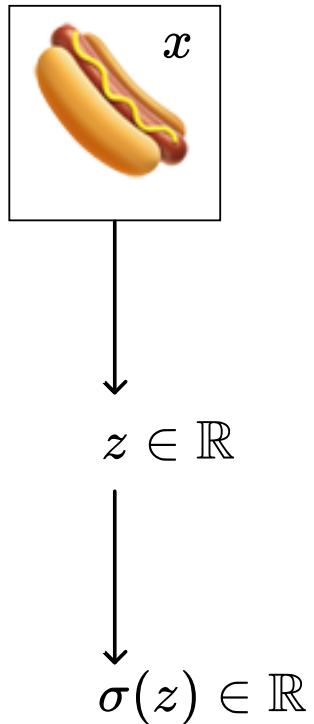


$$\sigma(z) \in \mathbb{R}$$

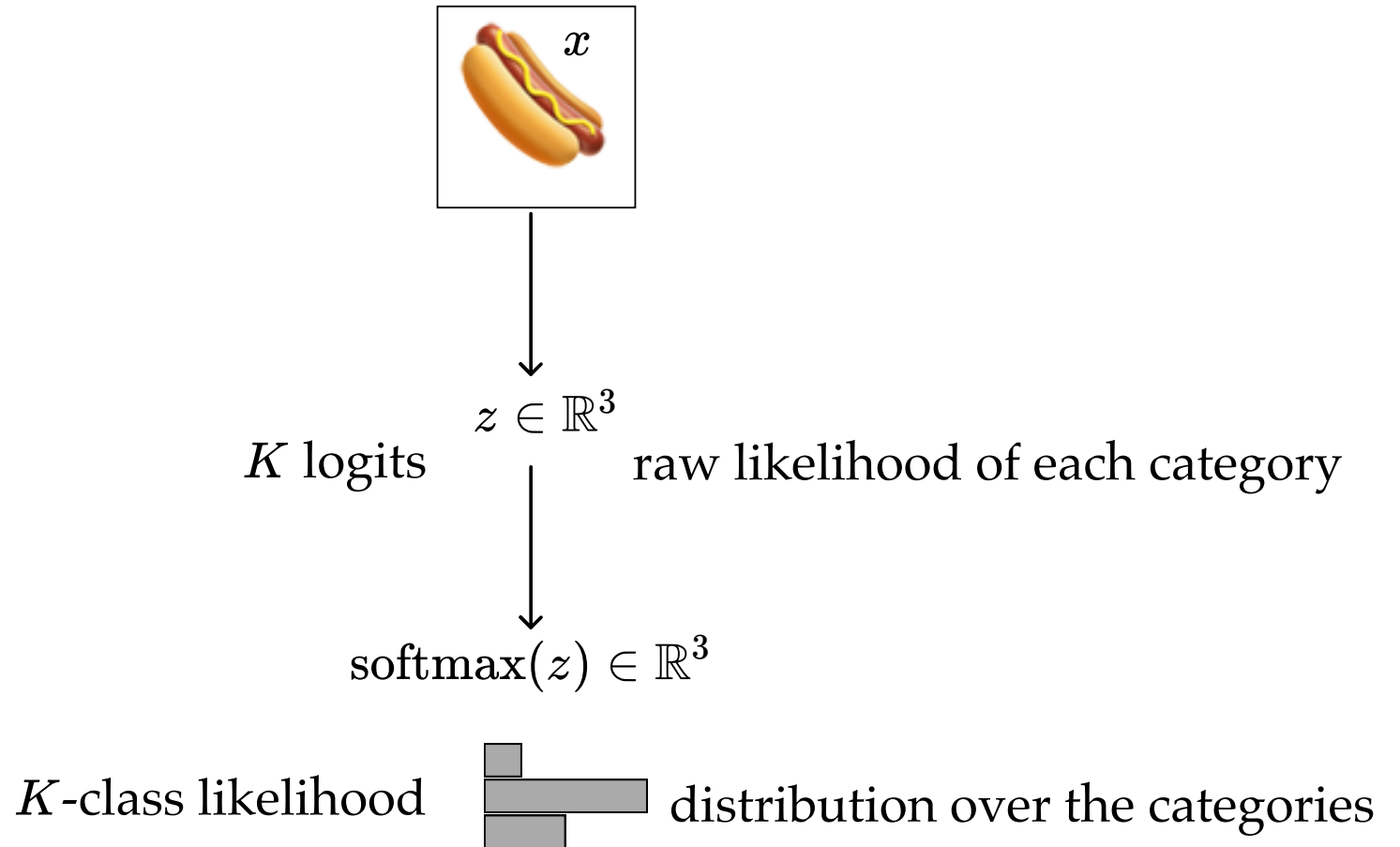
scalar likelihood
(normalized probability
of hotdog-ness)

- $\sigma(\cdot)$: the model's *confidence* or *estimated likelihood* that feature x belongs to the **hotdog** class;
- $1 - \sigma(\cdot)$: the **not-hotdog** class.

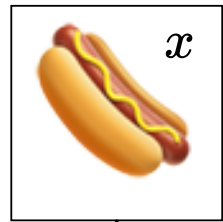
to predict hotdog or not,
a scalar logit suffices



to predict among K categories
say $K = 3$ categories: {hot-dog, pizza, salad}



two classes



$$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$$

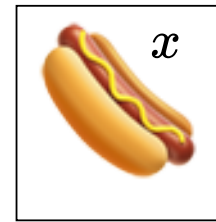


$$z \in \mathbb{R}$$



$$\sigma(z) \in \mathbb{R}$$

K classes



$$\theta \in \mathbb{R}^{d \times K}, \theta_0 \in \mathbb{R}^K$$



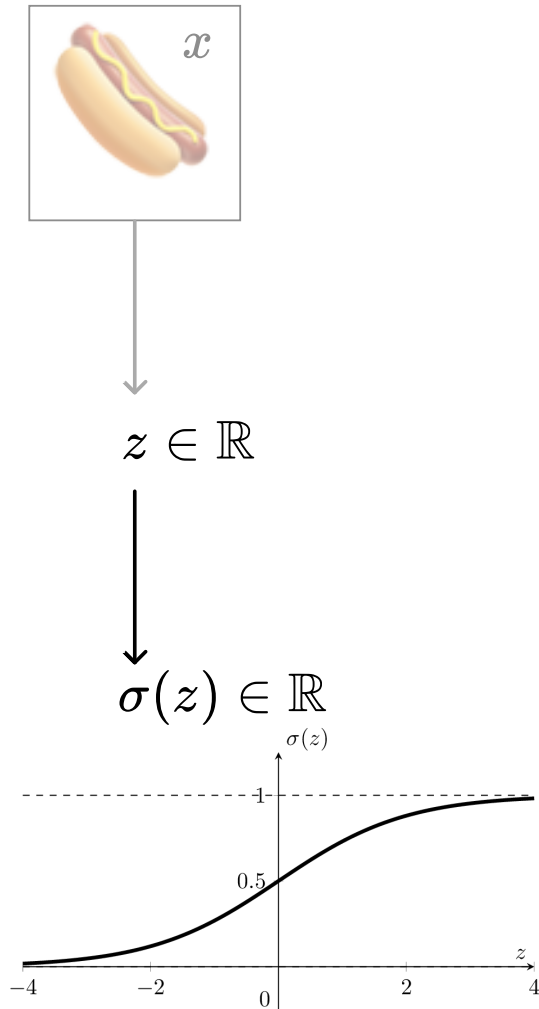
$$z \in \mathbb{R}^K$$

raw likelihood of each category

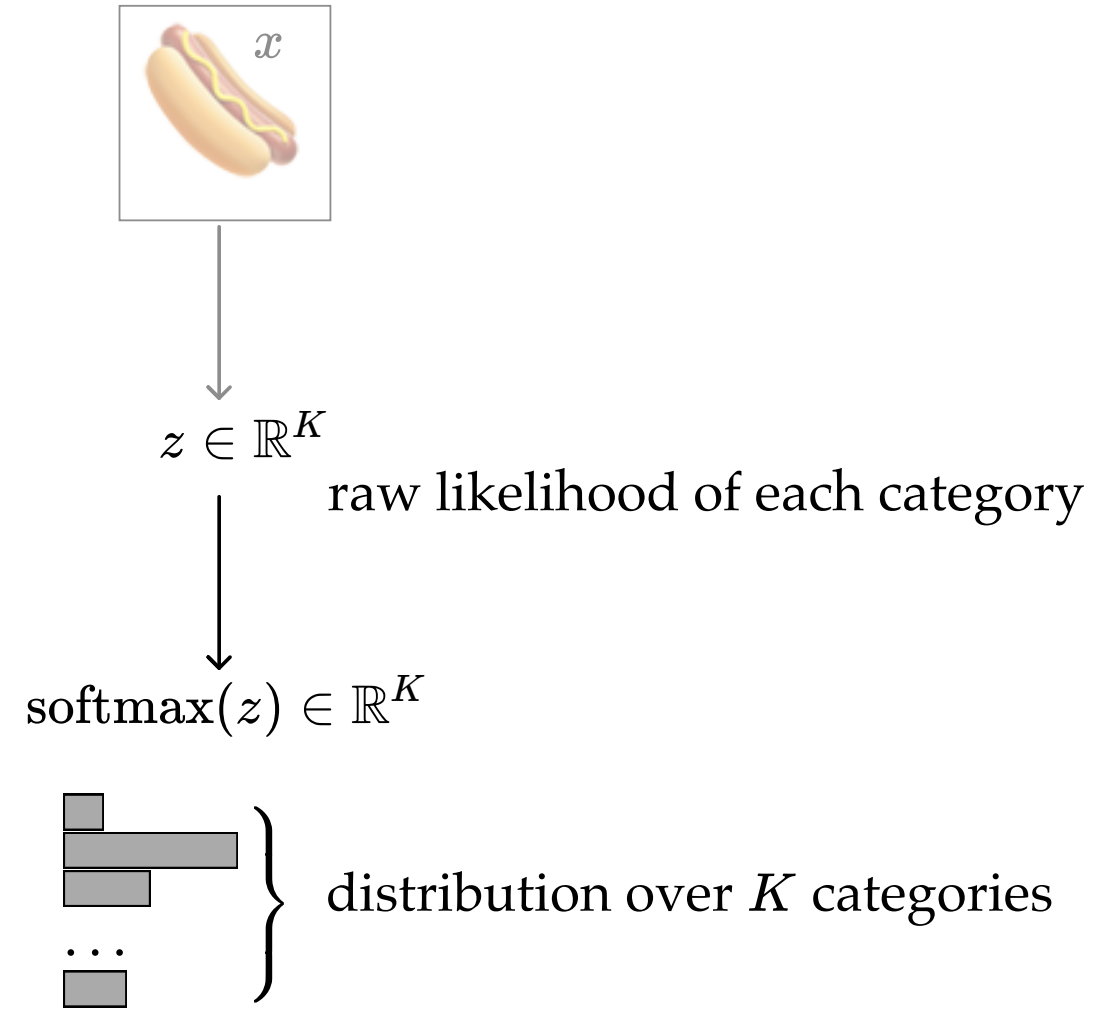


$$\text{softmax}(z) \in \mathbb{R}^K$$

two classes



K classes



softmax: $\mathbb{R}^K \rightarrow \mathbb{R}^K$

$$\text{softmax}(z) := \begin{bmatrix} \frac{\exp(z_1)}{\sum_{k=1}^K \exp(z_k)} \\ \vdots \\ \frac{\exp(z_K)}{\sum_{k=1}^K \exp(z_k)} \end{bmatrix}$$

each output entry is between 0 and 1, and their sum is 1

e.g.

$$\text{softmax}\left(\begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}\right) = \begin{bmatrix} \frac{e^1}{e^1 + e^2 + e^3} \\ \frac{e^2}{e^1 + e^2 + e^3} \\ \frac{e^3}{e^1 + e^2 + e^3} \end{bmatrix} = \begin{bmatrix} 0.0900 \\ 0.2447 \\ 0.6653 \end{bmatrix}$$

max in the input

"soft" max'd in the output

sigmoid $\mathbb{R} \rightarrow \mathbb{R}$

$$\begin{aligned}\sigma(z) &:= \frac{1}{1 + \exp(-z)} \\ &= \frac{\exp(z)}{\exp(z) + \exp(0)}\end{aligned}$$

implicit logit for the negative class

predict positive if $\sigma(z) > 0.5 = \sigma(0)$

equivalently, predicting the category with the largest raw logit.

softmax: $\mathbb{R}^K \rightarrow \mathbb{R}^K$

$$\text{softmax}(z) := \begin{bmatrix} \frac{\exp(z_1)}{\sum_{k=1}^K \exp(z_k)} \\ \vdots \\ \frac{\exp(z_K)}{\sum_{k=1}^K \exp(z_k)} \end{bmatrix}$$

predict the category with the highest softmax score

	linear logistic <i>binary</i> classifier	one-out-of- K classifier
features	$x \in \mathbb{R}^d$	
parameters	$\theta \in \mathbb{R}^d, \theta_0 \in \mathbb{R}$	$\theta \in \mathbb{R}^{d \times K}, \theta_0 \in \mathbb{R}^K$
linear combo	$\theta^T x + \theta_0 = z \in \mathbb{R}$	$\theta^T x + \theta_0 = z \in \mathbb{R}^K$
predict	$\sigma(z) = \frac{\exp(z)}{\exp(0) + \exp(z)}$ predict positive if $\sigma(z) > \sigma(0)$	$\text{softmax}(z) = \begin{bmatrix} \frac{\exp(z_1)}{\sum_{k=1}^K \exp(z_k)} \\ \vdots \\ \frac{\exp(z_K)}{\sum_{k=1}^K \exp(z_k)} \end{bmatrix}$ predict the category with the highest softmax score

Outline

1. Linear (binary) classifiers

- to use: separator, normal vector
- to learn: difficult! won't do

2. Linear logistic (binary) classifiers

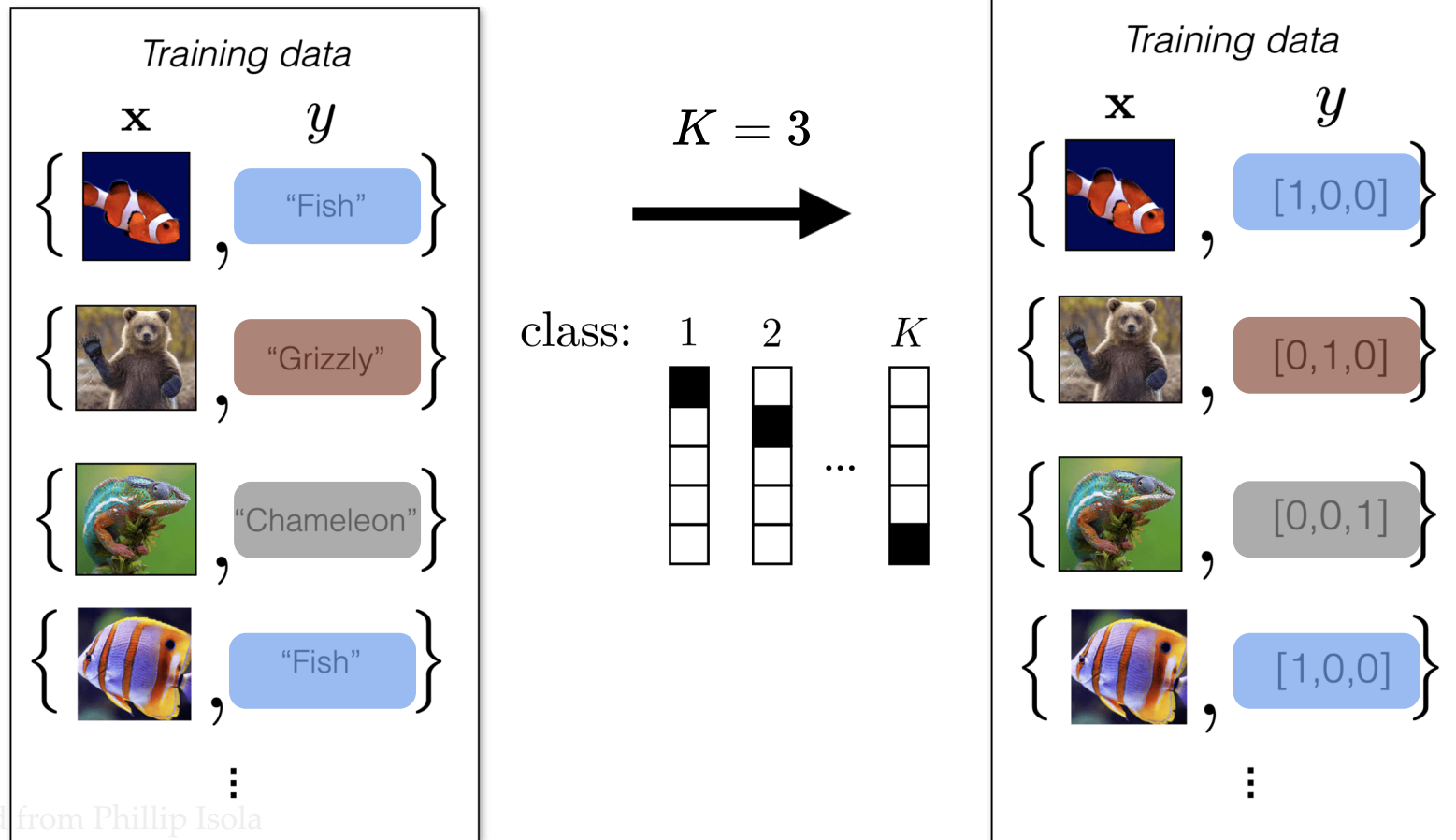
- to use: sigmoid
- to learn: negative log-likelihood loss

3. Multi-class classifiers

- to use: **softmax**
- to learn: **one-hot encoding, cross-entropy loss**

One-hot encoding:

- Generalizes from $\{0,1\}$ binary labels
- Encode the K classes as an $\mathbb{R}^K$ vector, with a single 1 (hot) and 0s elsewhere.



Negative log-likelihood K – classes loss (aka, cross-entropy)

g : softmax output

g_k : probability or confidence of belonging in class k

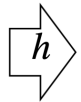
$$\mathcal{L}_{\text{nllm}}(g, y) = - \sum_{k=1}^K y_k \cdot \log(g_k)$$

y : one-hot encoding label

y_k : k th entry in y , either 0 or 1

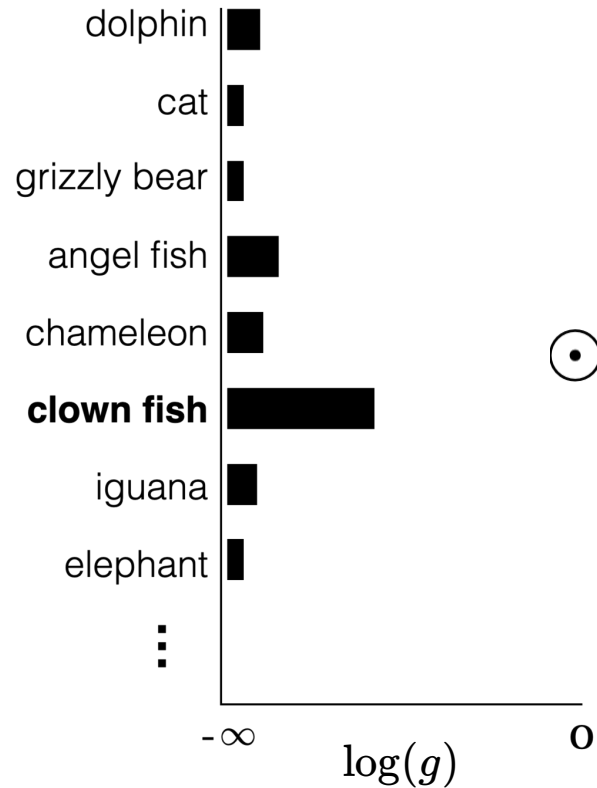
- Generalizes negative log likelihood loss $\mathcal{L}_{\text{nll}}(g, y) = - [y \log g + (1 - y) \log (1 - g)]$
- Although this is written as a sum over K terms, for a given training data point, only the term corresponding to its true class label contributes, since all other $y_k = 0$

feature x



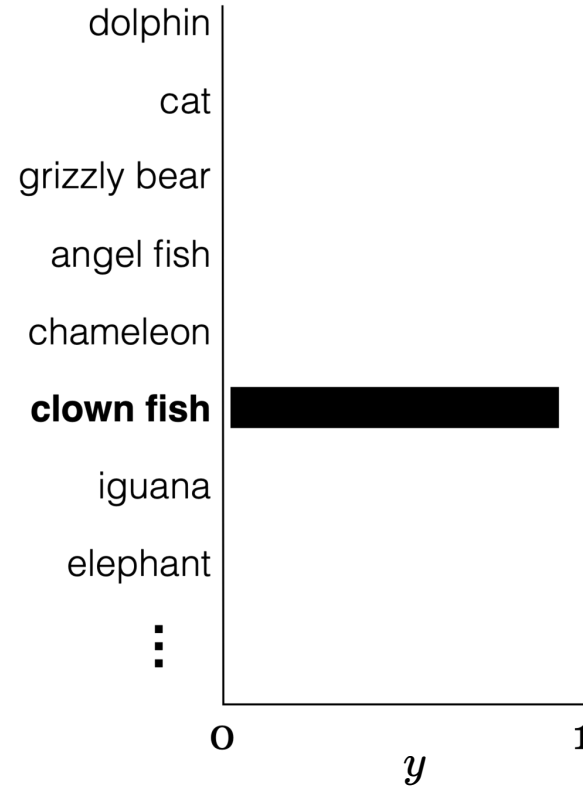
current prediction

$$g = \text{softmax}(\cdot)$$

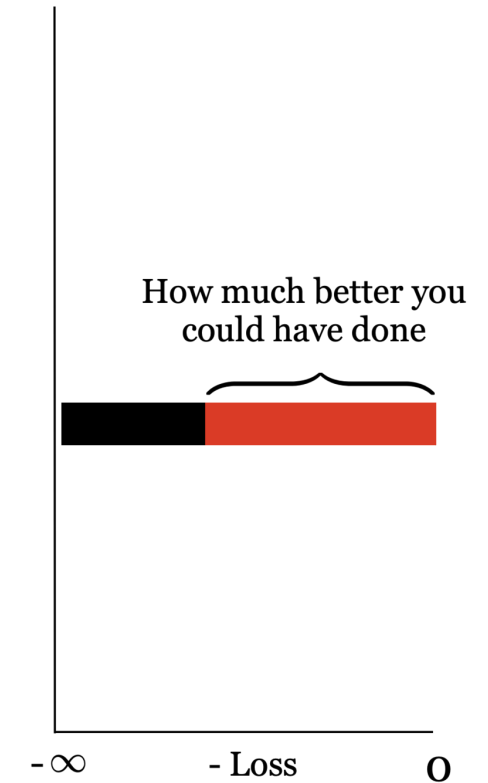


true label y

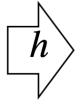
$$[0, 0, 0, 0, 0, 1, 0, 0, \dots]$$



$$\text{loss } \mathcal{L}_{\text{nllm}}(g, y) = - \sum_{k=1}^K y_k \cdot \log(g_k)$$



feature x



current prediction
 $g = \text{softmax}(\cdot)$



true label y
 $[0, 0, 1, 0, 0, 0, 0, 0, \dots]$

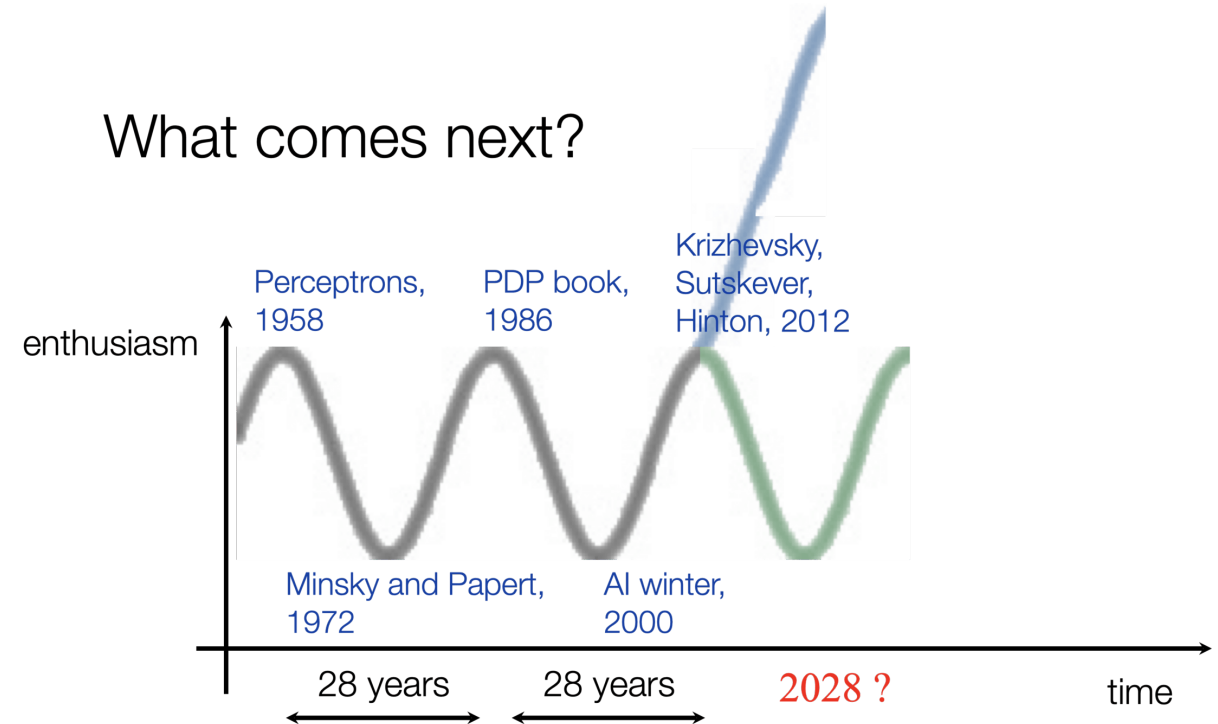
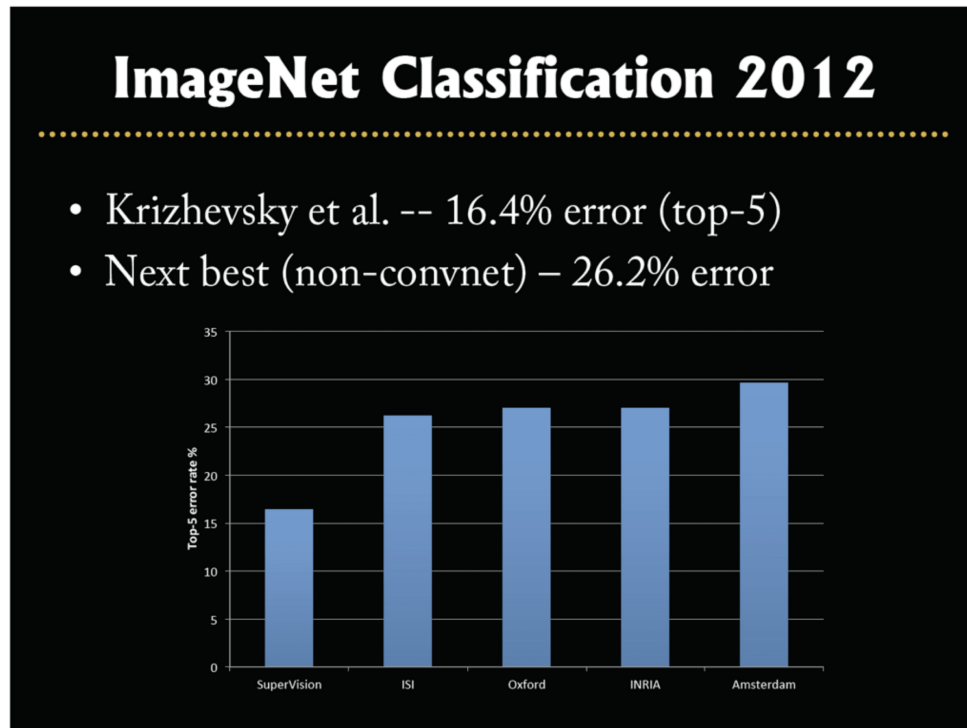


loss $\mathcal{L}_{\text{nllm}}(g, y)$
 $= - \sum_{k=1}^K y_k \cdot \log(g_k)$



Classification

Image classification played a pivotal role in kicking off the current wave of AI enthusiasm.



Summary

- Classification: a supervised learning problem, similar to regression, but where the output/label is in a discrete set.
- Binary classification: only two possible label values.
- Linear binary classification: think of θ and θ_0 as defining a hyperplane that **cuts** the d -dimensional feature space into two half-spaces.
- 0-1 loss: a natural loss function for classification, BUT, hard to optimize.
- Sigmoid function: smoother and probabilistic step function, motivation and properties.
- NLL loss: smooth loss and has nice probabilistic motivations. Can optimize via (S)GD.
- Regularization is still important.
- The generalization to multi-class via (one-hot encoding, and softmax mechanism)
- Other ways to generalize to multi-class (see hw/lab)

[https://docs.google.com/forms/d/e/1FAIpQLSfG1vnfaOvy8jugeVHrJWJQB-_15IWBq683-XI8z1AJf6YZNg/viewform?
embedded=true](https://docs.google.com/forms/d/e/1FAIpQLSfG1vnfaOvy8jugeVHrJWJQB-_15IWBq683-XI8z1AJf6YZNg/viewform?embedded=true)

We'd love to hear
your **thoughts**.

Thanks!