

OPTIMIZATION METHODS FOR MACHINE LEARNING

BETHANY LUSCH

Asst. Computer Scientist

Argonne National Lab

Leadership Computing Facility

blusch@anl.gov



Argonne National Laboratory is a
U.S. Department of Energy laboratory
managed by UChicago Argonne, LLC.

August 9, 2019
ATPESC

WHAT IS OPTIMIZATION?

$$\begin{array}{ll} \underset{x}{\text{minimize}} & f(x) \\ \text{subject to} & \dots \text{constraints} \dots \end{array}$$

“objective” or “loss” function



- “best route” (Minimize cost of delivery subject to all mail is delivered)
- “best product” (Minimize -profit subject to safety)
- “best prediction model” (Minimize prediction error)

MACHINE LEARNING

minimize $f(x)$

subject to ... constraints...

- “best prediction model” (Minimize prediction error)
- “best recommendation” (Minimize # who don’t buy anything)
- “best clusters” (Minimize distance within clusters while maximizing distance between clusters)

Underneath most ML problems is an optimization problem



U.S. DEPARTMENT OF
ENERGY

Argonne National Laboratory is a
U.S. Department of Energy laboratory
managed by UChicago Argonne, LLC.

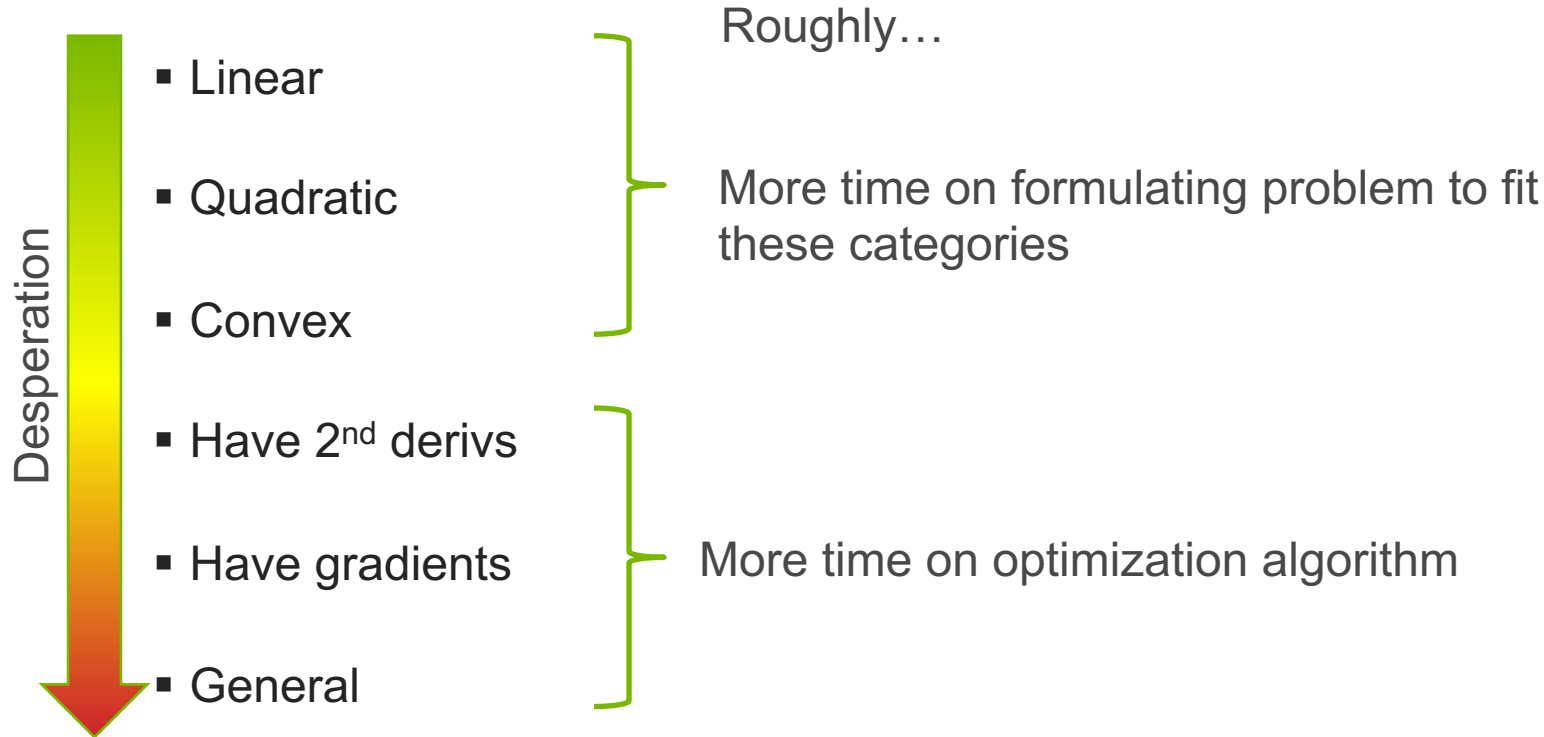
TYPES OF OPTIMIZATION



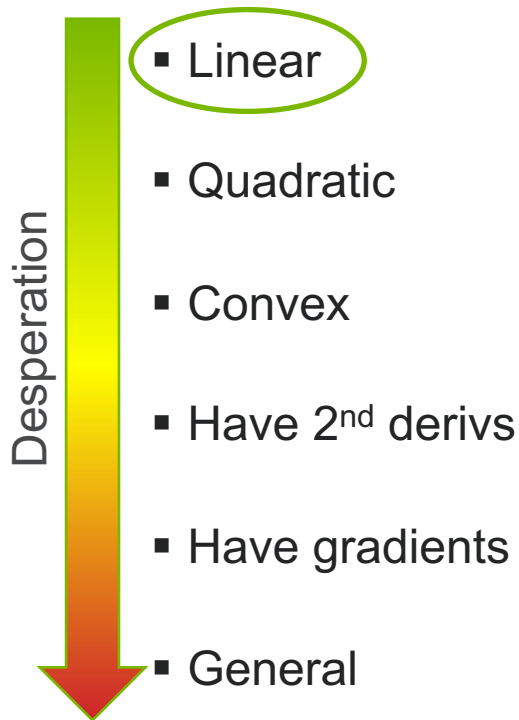
- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General

$$\begin{array}{ll} \underset{x}{\text{minimize}} & f_0(x) \\ \text{subject to} & f_i(x) \leq b_i, \quad i = 1, \dots, m. \end{array}$$

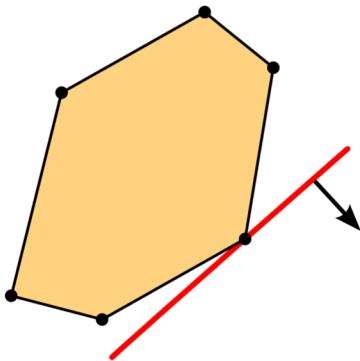
TYPES OF OPTIMIZATION



LINEAR PROGRAMMING



minimize
 x
subject to



linear

$$c^T x$$

$$f_0(x)$$

$$f_i(x) \leq b_i, \quad i = 1, \dots, m.$$

Convex polytope

$$Ax \leq b$$
$$x \geq 0.$$

Mature polynomial-time algorithms
Local minima are global optima

QUADRATIC PROGRAMMING



- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General

quadratic

$$\frac{1}{2}x^T Qx + c^T x$$

Convex polytope

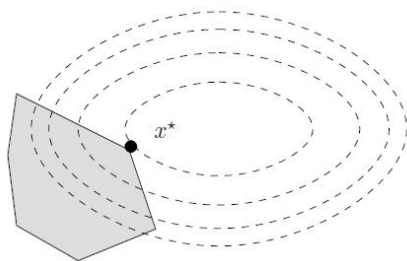
$$Ax \leq b$$
$$x \geq 0.$$

minimize
 x

$$f_0(x)$$

subject to

$$f_i(x) \leq b_i, \quad i = 1, \dots, m.$$



If Q is positive definite:
Weakly polynomial-time algorithms
Local minima are global optima

CONVEX OPTIMIZATION



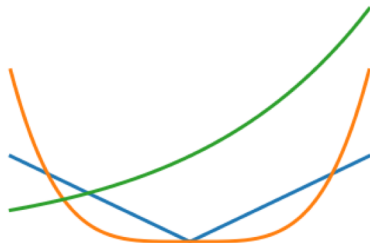
- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General

minimize x
subject to

Convex function

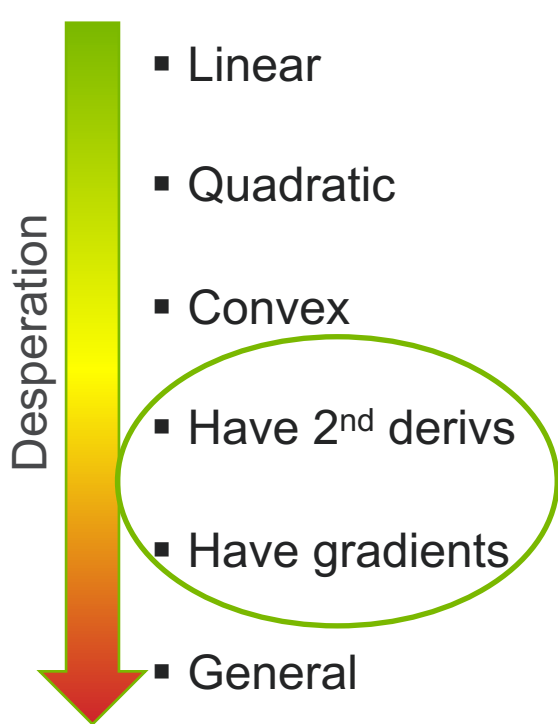
Convex set

$$f_0(x)$$
$$f_i(x) \leq b_i, \quad i = 1, \dots, m.$$



In some cases: poly-time
Local minima are global optima
Reliable methods

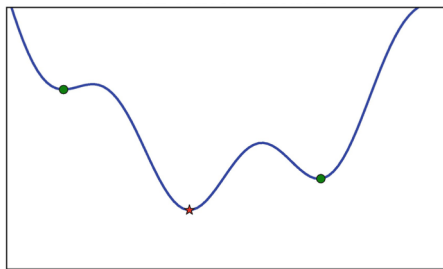
DIFFERENTIABLE OPTIMIZATION



Differentiable function

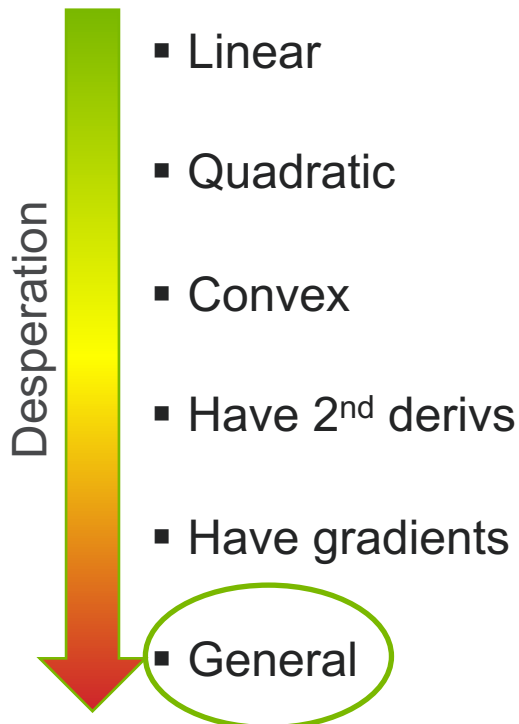
minimize $f_0(x)$

subject to $f_i(x) \leq b_i, i = 1, \dots, m.$

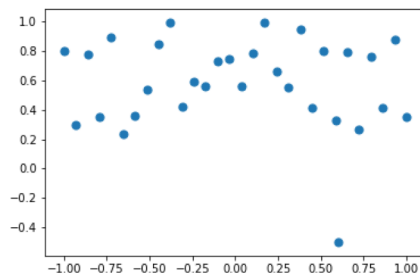


Generally NP-hard
Local minima problematic
Can use gradients and ideally
Hessians in algorithm

GENERAL OPTIMIZATION

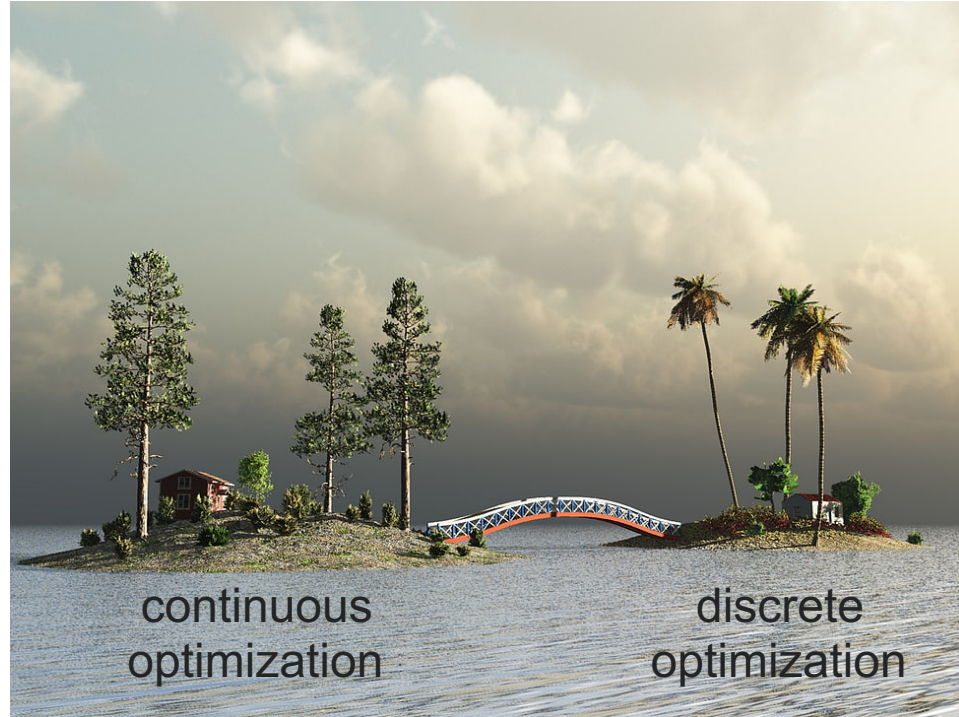
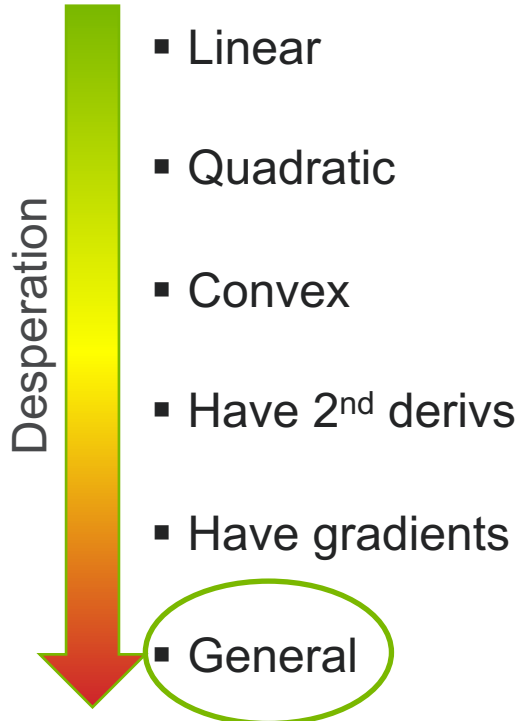


$$\begin{array}{ll} \underset{x}{\text{minimize}} & f_0(x) \\ \text{subject to} & f_i(x) \leq b_i, \quad i = 1, \dots, m. \end{array}$$



Generally NP-hard
Hopefully know *some* structure!

DISCRETE OPTIMIZATION



CLASSIFICATION EXAMPLE

- Problem: label each document x as related to politics or not (1 or -1).
- Hard to come up with rules by hand, so ML helps: learn function $h(x)$
- Really want to minimize expected risk of misclassification:

$$\underset{h}{\text{minimize}} \quad R(h) = \mathbb{P}[h(x) \neq y] = \mathbb{E}[\mathbb{1}[h(x) \neq y]]$$

- How do we pick family of functions to optimize over?
- How do we know which one is optimal?

REALITIES

- Really want to minimize expected risk of misclassification:

$$\underset{h}{\text{minimize}} \quad R(h) = \mathbb{P}[h(x) \neq y] = \mathbb{E}[\mathbb{1}[h(x) \neq y]]$$

- Don't know probability distribution, so minimize empirical risk:

$$\underset{h}{\text{minimize}} \quad R_n(h) = \frac{1}{n} \sum_{i=1}^n \mathbb{1}[h(x_i) \neq y_i]$$

- Easier if smooth loss and parameterized h :

$$\underset{w}{\text{minimize}} \quad R_n(h) = \frac{1}{n} \sum_{i=1}^n \ell(h(x_i; w), y_i)$$

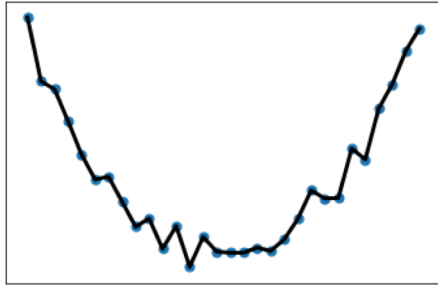
CHOOSING FUNCTION FAMILY

- Possibility of low empirical risk on training data
 - Expected risk and empirical risk don't have large gap
 - Can efficiently solve optimization
 - (convenient representation, smoothness,...)
- } Bias vs. variance
Overfitting vs. underfitting

Major themes of machine learning!

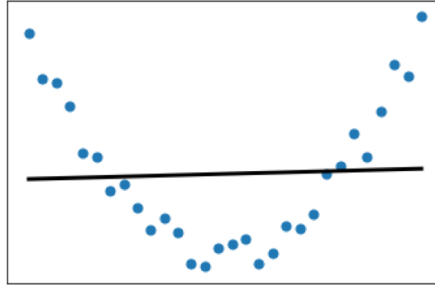
BIAS VS. VARIANCE

High variance



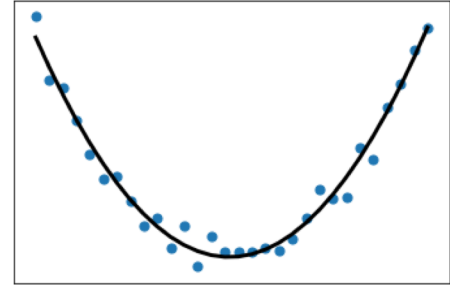
overfitting

High bias



underfitting

Low bias, low variance



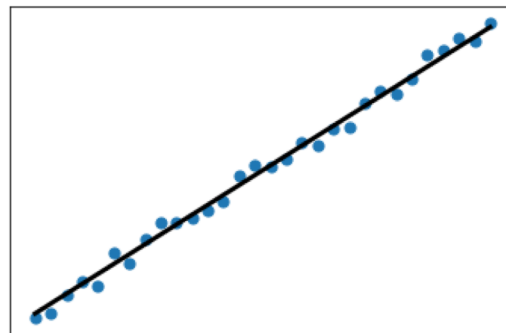
balanced

Picture similar to: Seema Singh, "Understanding the Bias-Variance Tradeoff"

LINEAR REGRESSION (LEAST-SQUARES)

- Desperation
- Linear
 - Quadratic
 - Convex
 - Have 2nd derivs
 - Have gradients
 - General

$$\underset{x}{\text{minimize}} \quad \|Ax - b\|_2^2$$



if multiply out:

$$\underset{x}{\text{minimize}} \quad x^T A^T A x - 2b^T A x + b^T b$$

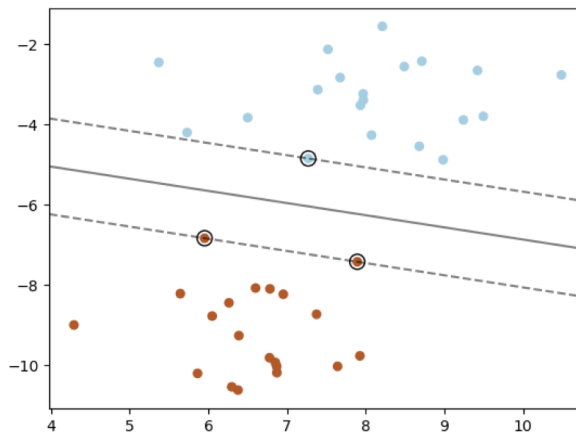
... quadratic program!

SUPPORT VECTOR MACHINE

Find linear classifier with maximum margin

Desperation

- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General



$$\underset{a,b,u,v}{\text{minimize}} \quad \|a\|_2 + \gamma(\mathbf{1}^T u + \mathbf{1}^T v)$$

$$\text{subject to} \quad a^T x_i - b \geq 1 - u_i, \quad i = 1, \dots, N$$

$$a^T y_i - b \leq -(1 - v_i), \quad i = 1, \dots, M$$

$$u \geq 0, v \geq 0$$

... quadratic program!



U.S. DEPARTMENT OF
ENERGY

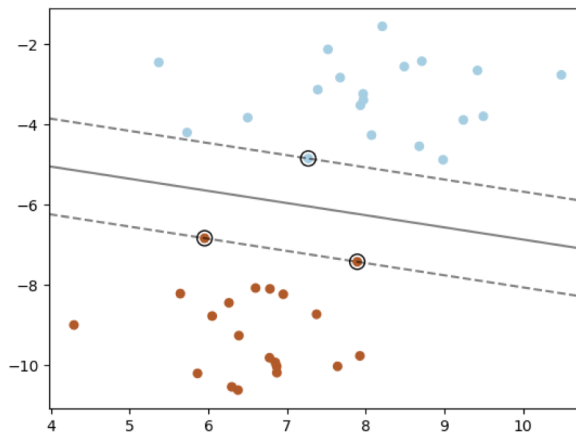
Argonne National Laboratory is a
U.S. Department of Energy laboratory
managed by UChicago Argonne, LLC.

SUPPORT VECTOR MACHINE

Find linear classifier with maximum margin

Desperation

- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General

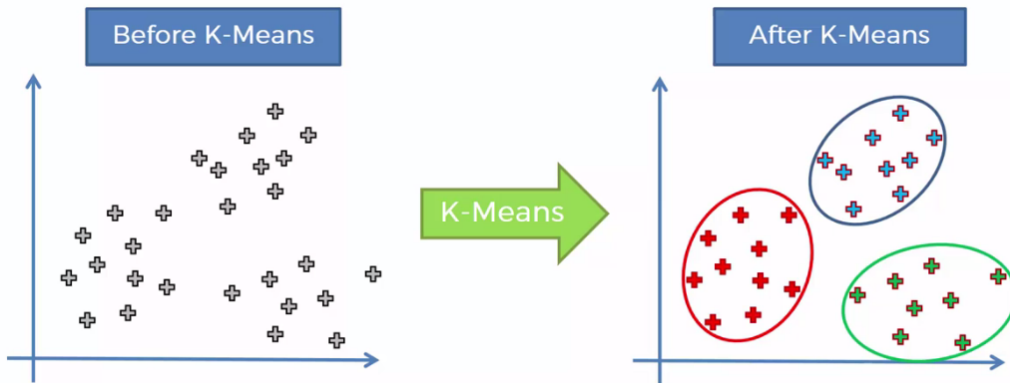


Kernel SVM can do nonlinear classification while remaining a quadratic program

K-MEANS CLUSTERING

Find clustering that minimizes distances within clusters

- Desperation
- Linear
 - Quadratic
 - Convex
 - Have 2nd derivs
 - Have gradients
 - General



$$\underset{S}{\text{minimize}} \quad \sum_{i=1}^k \sum_{x \in S_i} \|x - \mu\|_2^2$$

NP-hard discrete problem, so use approx. algorithm

DEEP LEARNING



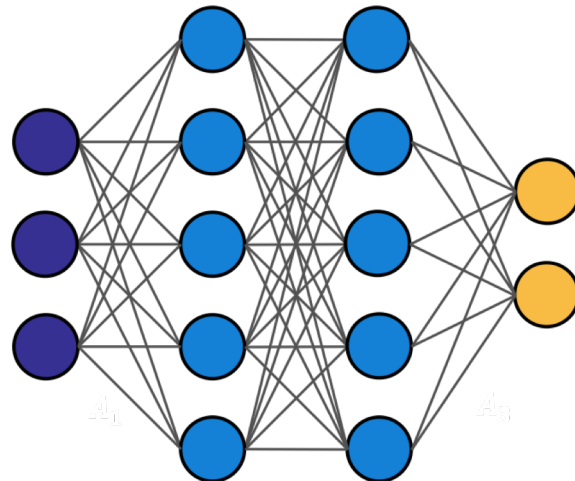
Desperation

- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General

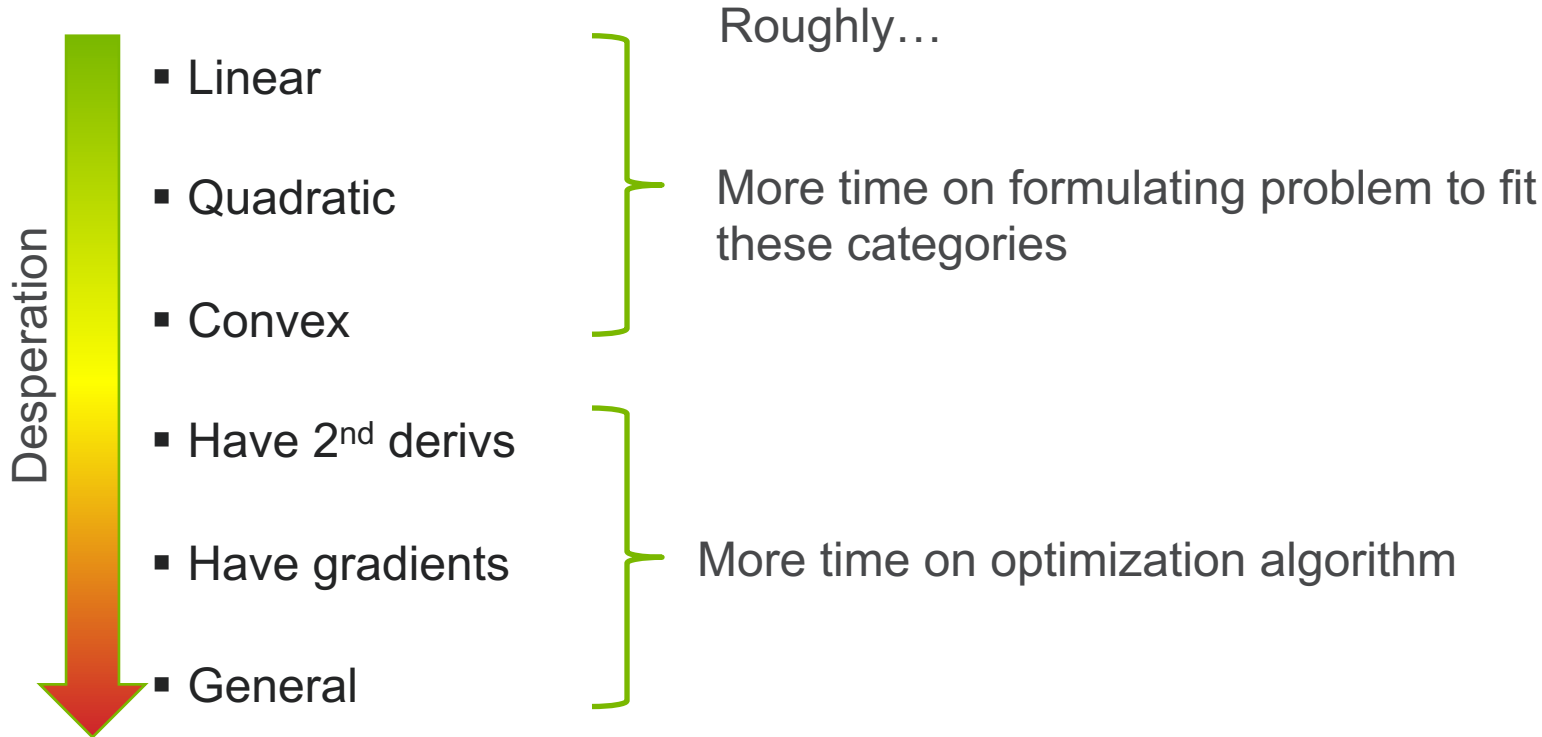
Too expensive

$$\underset{\theta}{\text{minimize}} \quad \|x - f(x)\|_2^2$$

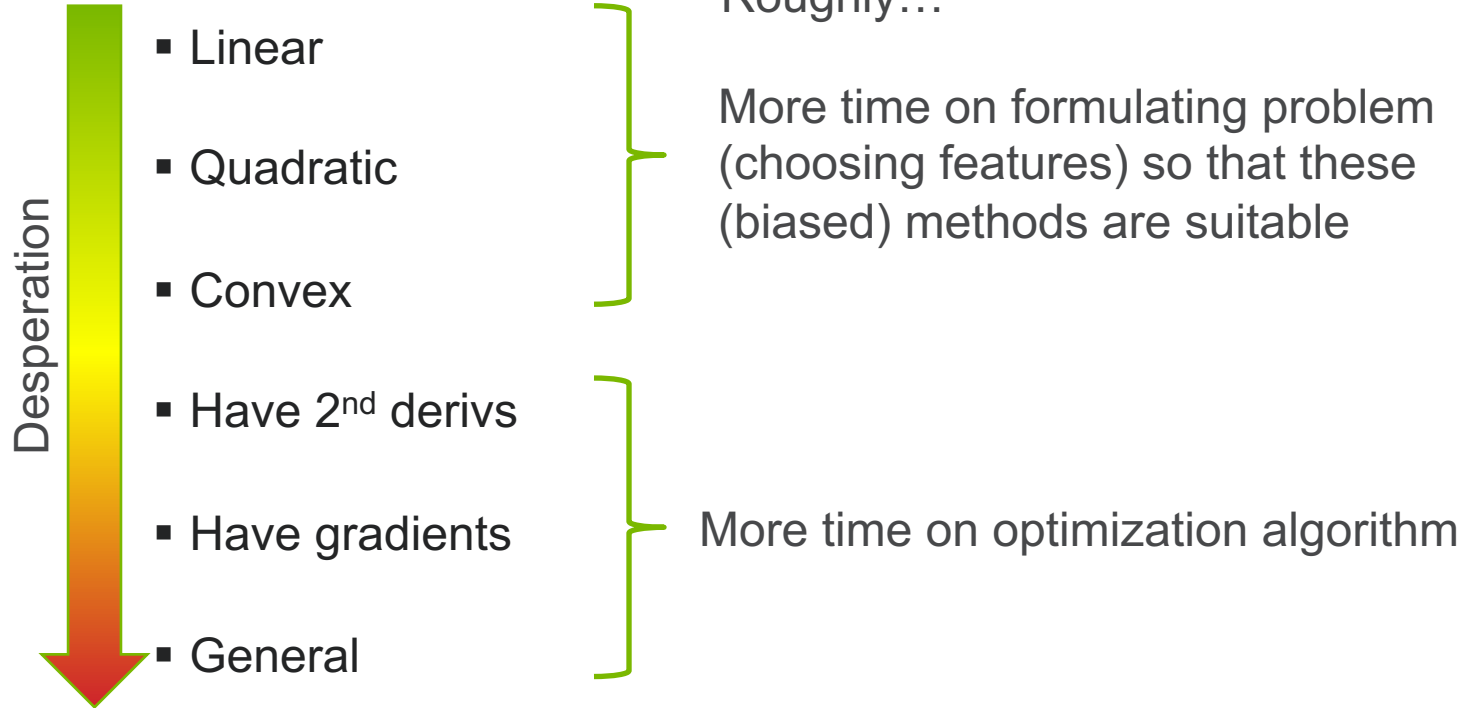
Getting desperate, plus want to be scale well,
so use gradient descent



RECALL: TYPES OF OPTIMIZATION



ANALOGOUSLY...



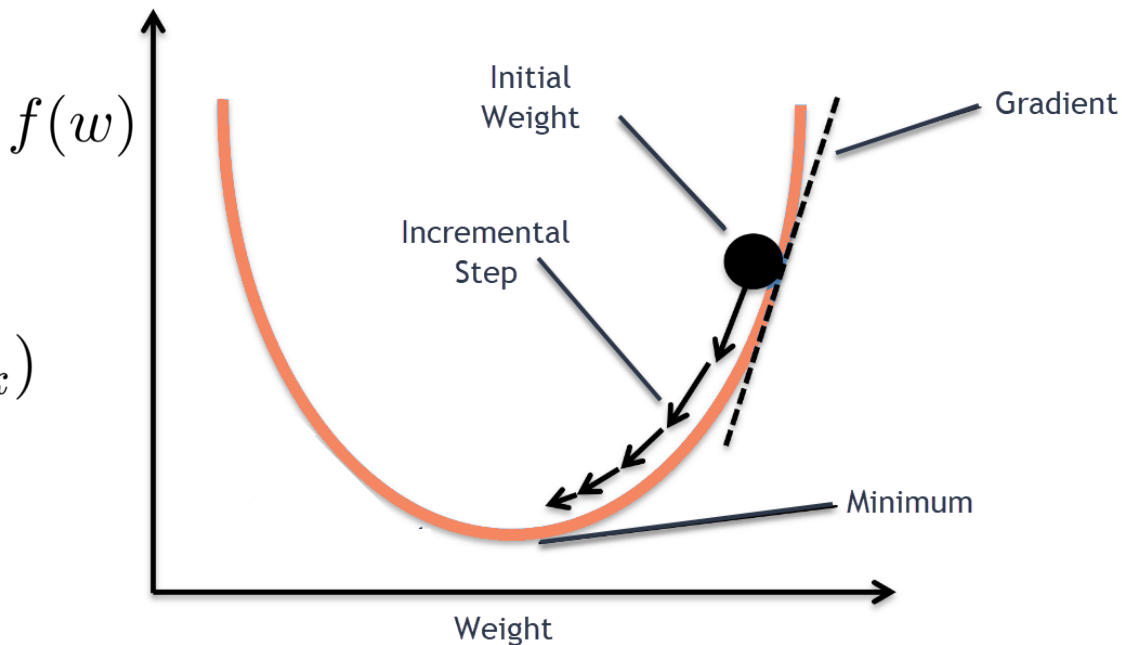
REMINDER

- Possibility of low empirical risk on training data
 - Expected risk and empirical risk don't have large gap
 - Can efficiently solve optimization
- } Bias vs. variance
Overfitting vs. underfitting
- Big neural networks can be very expressive (low bias)
 - So don't need to be as clever about input features
 - But then easy to overfit...
 - Optimization is tricky: optimization stalls, plus local minima or saddle points

GRADIENT DESCENT

minimize $f(w)$

$$w_{k+1} \leftarrow w_k - \alpha_k \nabla f(w_k)$$



Picture source: Divakar Kapil in "Stochastic vs Batch Gradient Descent"

TYPES OF GRADIENT DESCENT

i.e. for empirical risk, explicitly summing over data points

$$R_n(w) = \frac{1}{n} \sum_{i=1}^n f_i(w)$$

$$w_{k+1} \leftarrow w_k - \frac{\alpha_k}{n} \sum_{i=1}^n \nabla f_{ik}(w_k)$$

Batch GD: use all points every step

$$w_{k+1} \leftarrow w_k - \alpha_k \nabla f_{ik}(w_k)$$

Stochastic GD: use one point per step

$$w_{k+1} \leftarrow w_k - \frac{\alpha_k}{|S_k|} \sum_{i \in S_k} \nabla f_{ik}(w_k)$$

Mini-batch GD: use a subset each step

TYPES OF GRADIENT DESCENT

Batch GD: use all points every step

Each step is accurate but expensive

Stochastic GD: use one point per step

Each step is noisy but fast

Mini-batch GD: use a subset each step

Happy medium?



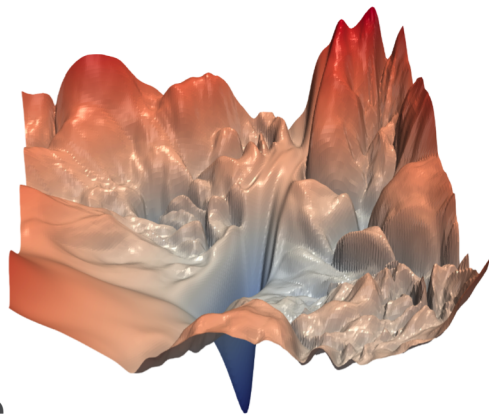
Very common in deep learning, but often call it SGD

GRADIENT DESCENT CONSIDERATIONS

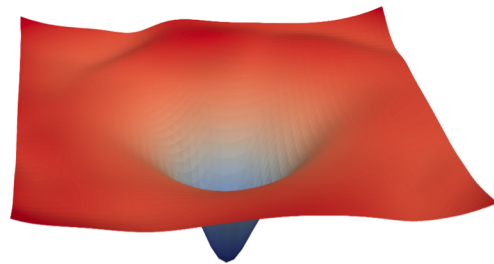
$$\underset{w}{\text{minimize}} \quad f(w)$$

$$w_{k+1} \leftarrow w_k - \alpha_k \nabla f(w_k)$$

- Step size α_k
 - Too big: overshoot
 - Too small: very slow
 - (But can be good to escape local minima)
- Initialization
- Can you make the problem easier?



(a) without skip connections



(b) with skip connections

Li, et al. "Visualizing the Loss Landscape of Neural Nets" NeurIPS 2018

VARIANT: ADAM

Popular improvement on GD: Adam optimizer

- Separate learning rate for each weight
- Momentum: uses moving average of the gradient
- Also incorporates squared gradients

Cool exploration/visualization of momentum: <https://distill.pub/2017/momentum/>

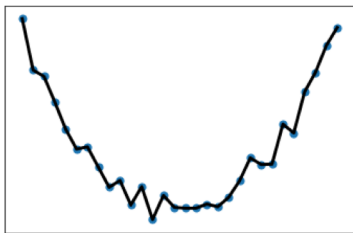
(For those familiar: combines the best properties of
AdaGrad, momentum, and RMSProp)

REGULARIZATION

- Common way to avoid overfitting: regularization
- Most common: L2 regularization

$$\underset{w}{\text{minimize}} \quad \underbrace{\frac{1}{n} \sum_{i=1}^n (h(x_i; w) - y_i)^2}_{\text{error}} + \underbrace{\lambda \|w\|_2^2}_{\text{regularization}}$$

balance
↓



Roughly: big coefficients/weights correspond to large variation

OVERFITTING CAUTION

- Can you generalize outside of your training set to validation/testing set?
- What about interpolating to data you haven't collected?
- Extrapolation extra unlikely to work

SUMMARY



- Linear
- Quadratic
- Convex
- Have 2nd derivs
- Have gradients
- General

Major themes in machine learning:

- Overfitting vs. underfitting
- Ability to efficiently solve optimization problem

For more, see:

SIAM REVIEW
Vol. 60, No. 2, pp. 223–311

© 2018 Society for Industrial and Applied Mathematics

Optimization Methods for Large-Scale Machine Learning*

Léon Bottou[†]
Frank E. Curtis[‡]
Jorge Nocedal[§]

ANY QUESTIONS?

Thinking ahead to next talk: how would you parallelize gradient descent?