

LOGS 186 Genetic Algorithm

Intro ML lecture.

Discriminant model vs. Generative model.

↓
SVM, ...

↓
GAN, Transformer ...

Genetic Algorithm → An Optimization Algorithm

Naive Bayes.

↓
Bayesian Network

Assume some conditional dependence among some parameter

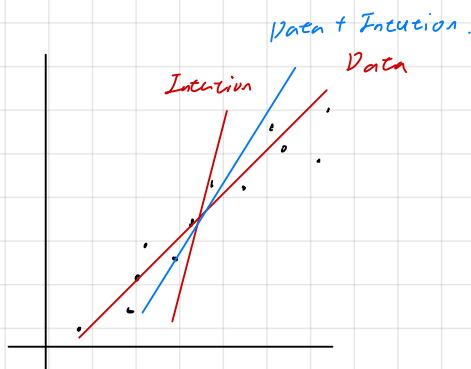
↓
as a network



Bayesian vs. Frequentist, different view of probability.

↓
some ground truth.
↓
prior belief.

Data + ground truth.



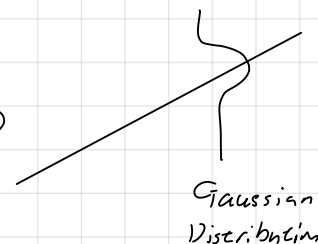
Linear Regression

- generative model independent (explanatory variable).

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$$

error term / variations

$$\varepsilon_i \sim N(0, \sigma^2)$$



Gaussian Distribution

$$y_i = \hat{\beta}_0 + \hat{\beta}_1 x_i + \hat{\varepsilon}_i$$

model: equation or function.

true value

estimated parameter

uncertainty

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_i$$

estimated value

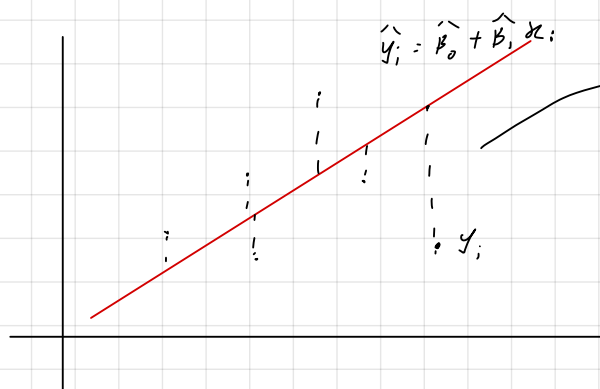
no random variation

linear regression

Perceptron

objective function.

Quantifying error / cost — the "goodness" of the line.



Distance.

$$\sum_{i=1}^n |\hat{y}_i - y_i| \rightarrow \text{Absolute loss}$$

$$= \sum_{i=1}^n |(m x_i + b) - y_i|$$

which parameters result in lowest "cost" / error.

$$\sum_{i=1}^n (\hat{y}_i - y_i)^2$$

→ sum of squared error.

least square objective function.

SSE:

$$\sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_i))^2 = \text{risk}$$

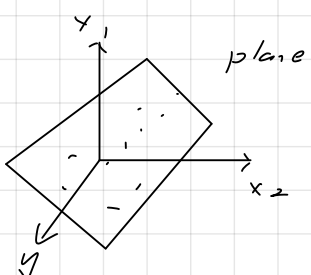
data points.

a function of the parameter

$$x_2 = \beta_1 x_1 + \beta_0$$

$$\beta_2 x_2 + \beta_1 x_1 + \beta_0 = 0$$

multiple features.

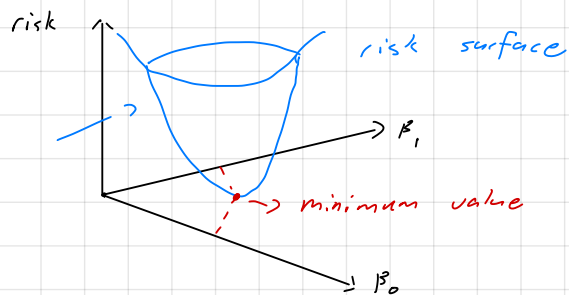


extend

$$\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n = 0$$

A hyper-plane.

bowl shape
b/c quadratic
function.



To find minimum, one way is to find derivative and set to 0.

$$SSE = \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2$$

$$2 \sum_{i=1}^n (\beta_0 + \beta_1 x_i - y_i) = 0$$

$$\frac{\partial SSE}{\partial \beta_0} = \frac{\partial}{\partial \beta_0} \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2$$

$\Rightarrow$

$$n \beta_0 + \beta_1 \sum_{i=1}^n x_i - \sum_{i=1}^n y_i = 0$$

$$= \sum_{i=1}^n 2(y_i - \beta_0 - \beta_1 x_i)(-1)$$

$$n \beta_0 + \beta_1 n \bar{x} - n \bar{y} = 0$$

$$\beta_0 + \beta_1 \bar{x} - \bar{y} = 0$$

$$\beta_0 = \bar{y} - \beta_1 \bar{x}$$

$$\frac{\partial SSE}{\partial \beta_1} = \frac{\partial}{\partial \beta_1} \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2$$

$\Rightarrow$

$$2 \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)(-x_i) = 0$$

$$= \sum_{i=1}^n 2(y_i - \beta_0 - \beta_1 x_i)(-x_i)$$

$$\sum (\beta_0 + \beta_1 x_i - y_i)(x_i) = 0$$

$$\sum \beta_0 x_i + \sum \beta_1 x_i^2 - \sum x_i y_i = 0$$

$$n \beta_0 \bar{x} + n \beta_1 \bar{x}^2 - n \bar{x} \bar{y} = 0$$

$\downarrow$

$$(\bar{y} - \beta_1 \bar{x}) \bar{x}$$

$\vdots$

$\Rightarrow$

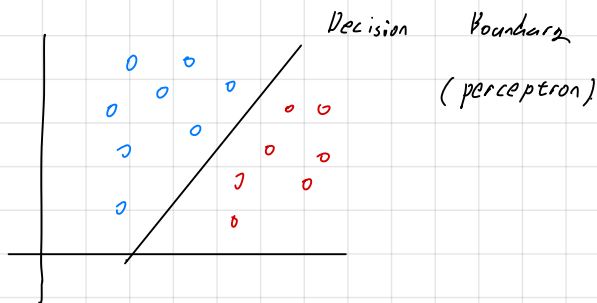
$$\beta_1 = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} = \frac{\text{Cov}(x, y)}{\text{Var}(x)}$$

$$\beta_0 = \bar{y} - \beta_1 \bar{x}$$

closed-form
solution! $\rightarrow$

Perceptron Algorithm

Classification problem.



Which line is best?



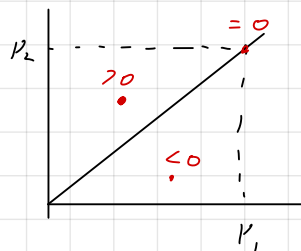
Objective function.



optimization problem

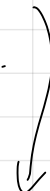
$$\beta_0 + \beta_1 x_1 + \beta_2 x_2 = 0$$

linear regression



if (p_1, p_2) on the line,

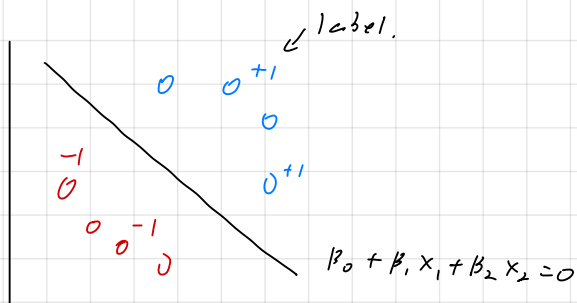
$$H(x) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 = 0.$$



Which side of the plane



classification!



$$\text{sign} \left(\beta_0 + \sum_{i=1}^d \beta_i x_i \right)$$

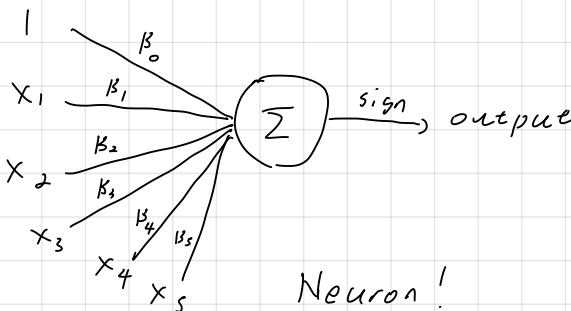


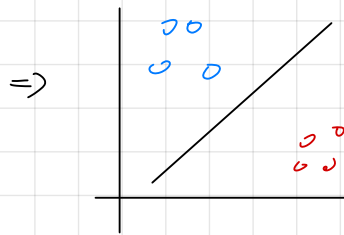
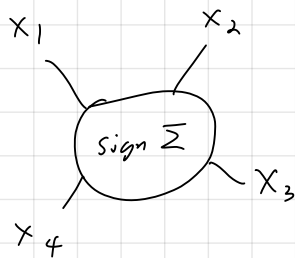
label +1 or -1.

sign function:

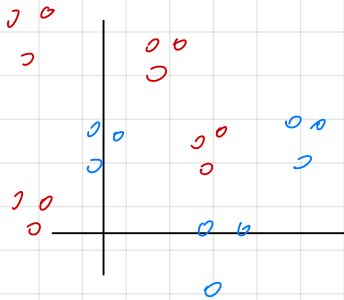


Activation function.





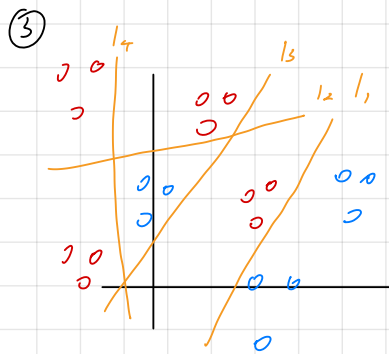
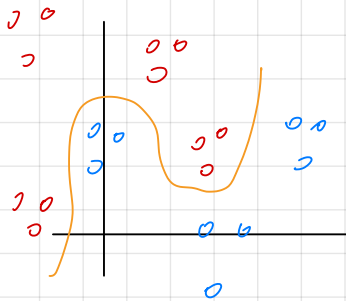
problem
not linearly separable.



Possible methods:

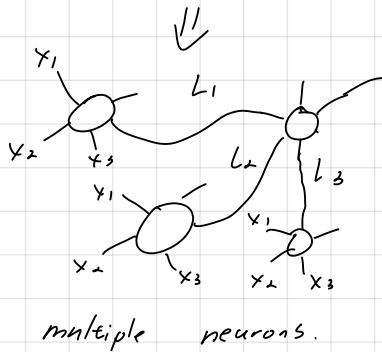
① bring up the dimension $\rightarrow$ feature mapping

② non-linear classifier.



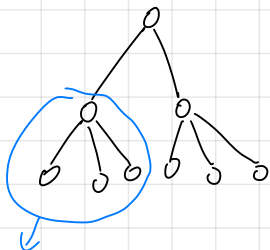
Perceptron

$\downarrow$
hyperplane



Neural Networks 1 or ANN

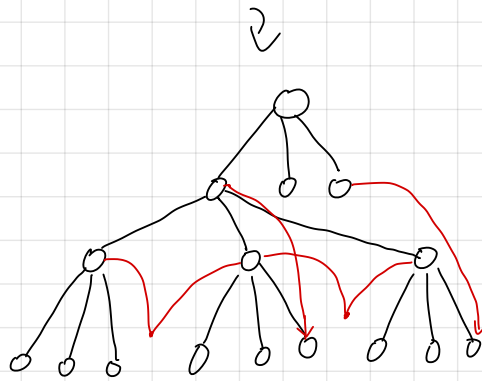
$\downarrow$
A collection of planes.
(perceptrons)



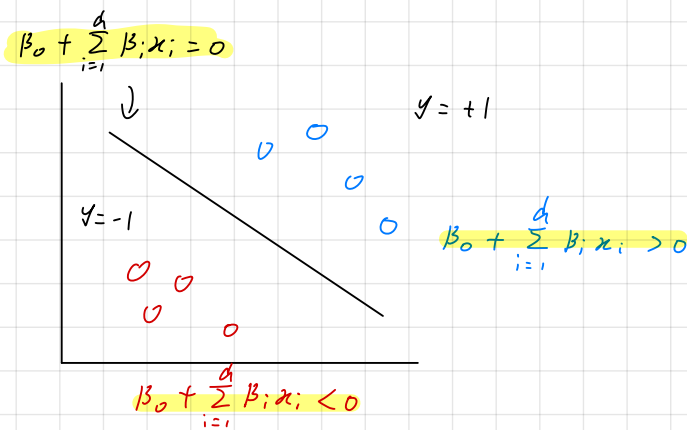
single layered NN
perceptron
separating hyperplane.



in deep learning, a convoluted neural networks



So, linear separating hyper-planes:



The Model:

$$\hat{y} = \mathcal{M}(x) = \text{sign} \left[\hat{\beta}_0 + (\hat{\beta}_1, \dots, \hat{\beta}_d) \overset{\vec{\beta}^T x^T}{x^T} \right]$$

where

$$\text{sign}(A) = \begin{cases} 1 & \text{if } A > 0 \\ -1 & \text{otherwise} \end{cases}$$

Decision boundary

$$\hat{\beta}_0 + (\hat{\beta}_1, \dots, \hat{\beta}_d) x^T = 0$$

Rosenblatt's Perceptron Learning Algorithm

- Starts with a random hyperplane $(\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_d)$
- Incrementally modify the hyperplane such that points that are misclassified move closer to the correct side.

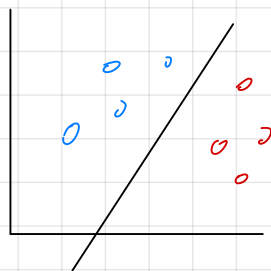
Gradient Descent Optimization Algorithm.

However, we can use genetic algorithm for optimization.

Optimization for perceptron

convex function

↓
always have one global minimum.



how far
is each point
from 0B

$$\beta_0 + \beta_1 x_1 + \beta_2 x_2 = 0$$

y. true

y. predicted.

$$\text{loss} = \sum_{i=1}^n -y_i \cdot (\beta_0 + \beta_1 x_{1i} + \beta_2 x_{2i})$$

↓
minimize using gradient descent.

$$\frac{\partial D}{\partial \beta_0} = - \sum_{i=1}^n \frac{\partial}{\partial \beta_0} (y_i (\beta_0 + \beta_1 x_{1i} + \beta_2 x_{2i}))$$

$$= - \sum_{i=1}^n y_i$$

$$\frac{\partial D}{\partial \beta_1} = - \sum_{i=1}^n y_i x_{1i}$$

More generally, we have the following derivative:

$$\frac{\partial D(\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_d)}{\partial \hat{\beta}_0} = - \sum_{i \in \mathcal{N}} y_i$$

$$\frac{\partial D(\hat{\beta}_0, \dots, \hat{\beta}_d)}{\partial \hat{\beta}_j} = - \sum_{i \in \mathcal{N}} y_i x_{ij}, \quad j = 1, \dots, d.$$

⇒ in vector form:

$$\frac{\partial}{\partial \beta} D = - \begin{bmatrix} y_1 \\ y_1 x_{11} \\ \vdots \\ y_i x_{ij} \end{bmatrix}$$

In gradient descent, we are updating:

$$\beta^{t+1} = \underbrace{\beta^t}_{\text{random guess}} - \alpha \frac{\partial}{\partial \beta} D = \begin{bmatrix} \beta_0^t \\ \vdots \\ \beta_d^t \end{bmatrix} - \alpha \begin{bmatrix} -y_i \\ -y_i x_{1i} \\ \vdots \\ -y_i x_{ij} \end{bmatrix}$$

random guess

↑
no summation b/c we are doing
each point individually

↓

Batch Perceptron:

While some (x_i, y_i) misclassified:

for i from 0 to n :

$$w = w + \alpha(y_i, x_i)$$

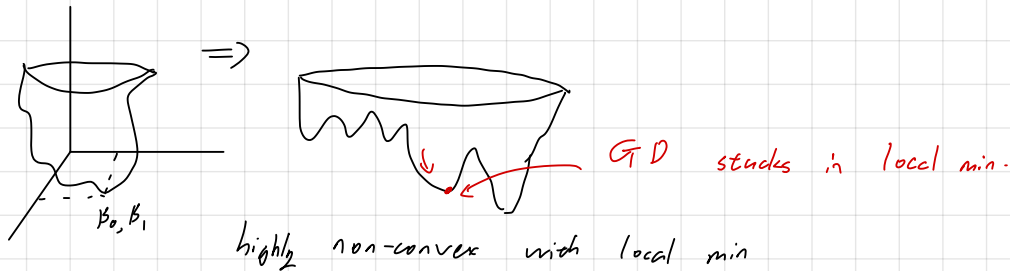
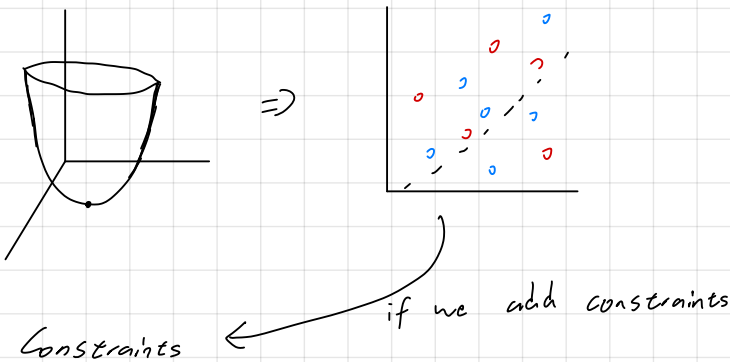
If points are linearly separable



perceptron learning with gradient descent guarantee to find the line that separates all points.

Optimization

Convex



Non-gradient Descent ("Hill-climbing")

idea ①: jumps to neighboring places.

if it is lower than previous, keep going.

if it is higher, continue jumping.

until finding the minimum.

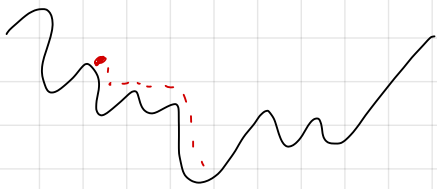
idea (2): Exploration phase

Exploration vs. Exploitation
↓ explore ↓ action.

Temperature. ↓
high exploring mindset
at first, meaning high probability to
take bad move. And lower exploration
later on.

(cooling)
Simulated Annealing

- Choose a number of "steps" (iterations).
- For each step
 - Randomly introduces a perturbation (a small change to the current combination).
 - Always accept the new alternative if it reduces the cost.
 - Randomly accept some alternatives that increases the cost (uphill change).
- Slowly decrease the uphill acceptance probability
 - later step are less likely to accept bad perturbations.



x_t

$$x_{t+1} = x_t + \Delta$$

if $f(x_{t+1}) < f(x_t)$,

$$x_t = x_{t+1}$$

else

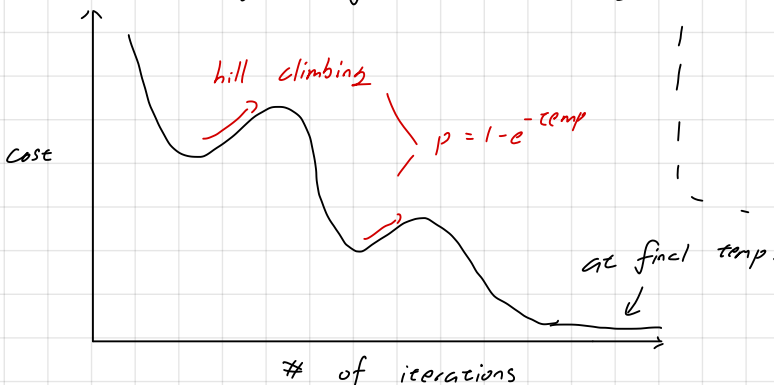
$p = 1 - e^{-temp}$ ← "cooling" temperature

$x_t = x_{t+1}$ with p probability.

else

$$x_t = x_t$$

Convergence of simulated annealing.



Problem of GP: Not all curves are continuous

ex) Traveling Salesman Problem

(Discrete Curve).

ex) Four color in n by n checkboard.

No same color is touching each other.

Cost is the number of color touching each other.

(Discrete Curve)

Learning via optimization / no domain knowledge.

Genetic Algorithm (Biology pov)



Biology



Evolution.

→ Darwin → Natural Selection.

can be used to "generate" objects
by using them to figure out the
rule that makes the object.

- Exponential growth in population

- Need $\rightarrow$ Not all member survive.
Which member survive is not random.

- Not two member are alike. Those variation are inherited.

- Limited Resources.



Not everyone will survive or reproduce.

Only the best suited member will survive.

- Mutation.

- Population evolve, not individuals.

- Evolution is continuous.

- Gene centric view of the evolution

Hamilton's Rule

$R \cdot B > C$, gene relatedness R , times the benefit to the group, B

↓ larger than the cost to the individual

the individual will demonstrate altruistic behavior
(individual sacrifice their own benefit to save others).

human's behavior is more complex.

How things evolve?

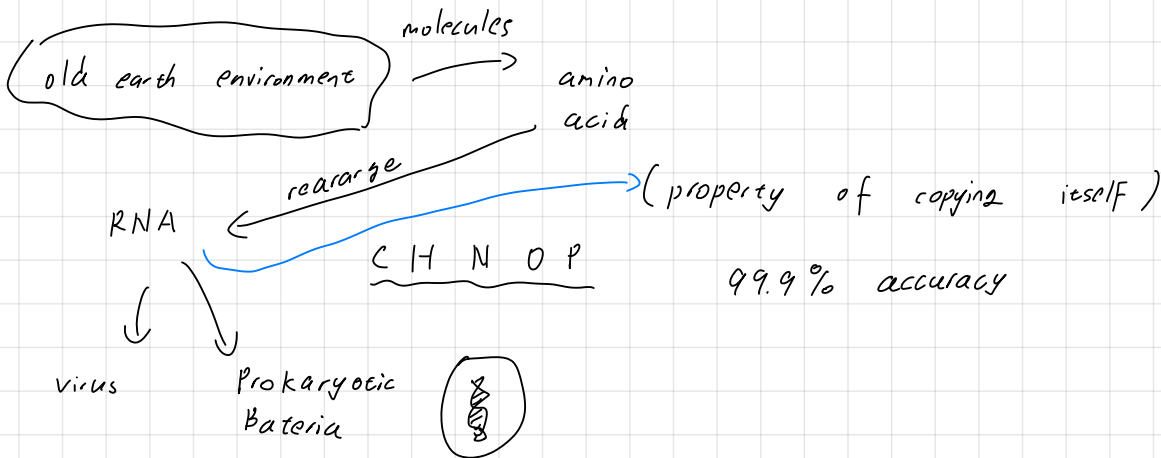
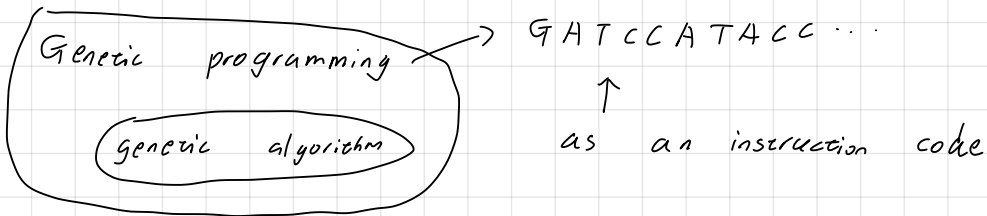
Genotype → Phenotype

(gene) (appearance).

Evolutionary Psychology

- Natural selection
- sexual
- group/kin
- language instinct.
- ...

Genetic Algorithm (Evolutionary Algorithm)



One replicator molecules RNA → DNA → Virus → Prokaryotic → Eukaryotic
Ant/plant ← ... ← Bacteria

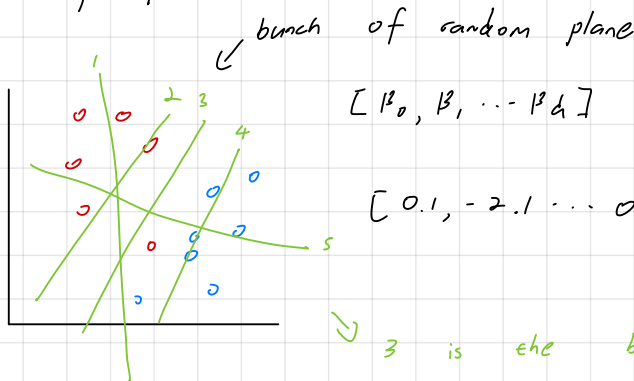
A G C T A G

arrange into something better fit

A G T C - - - - - ... optimize code such that

Survival reproduction (A G T G ...)
prob fitness/survival
this is high

ex) in perceptron



Genetic Algorithm:

- start with whole bunch of random solution
- keep the best one (natural selection).
- produce their mutated ones
- repeat the process over many generations

Fitness function (plane)

↓
classification accuracy
on the training set (plane)

↓
mutate & crossover the best member of the population.

mutate with parameters (with constraints)
ex) mutate 1~10% of "genes"
by 1~10%

idea.

$$f(x, y, z) = x^2 + \sin^2 x - \tan(x) + \frac{e^{-2x}}{10^x} + \dots$$

find (x, y, z) that minimizes/maximizes

initial population: random values of (x, y, z) .

10,000 $\left\{ \begin{array}{l} (1, 2.5, -3.5) \\ (0.5, 2.9, 4.9) \\ \vdots \end{array} \right\} \rightarrow$ evaluate at $f(\cdot)$
the "fitness" of individual

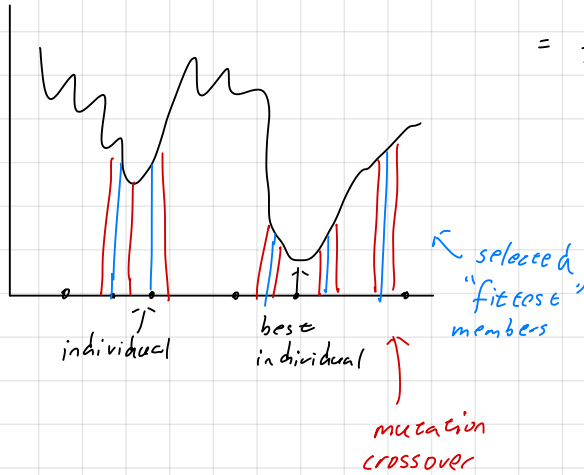
select 1,000 "fittest" members
 ↓
 recreate next generation of 10,000
 using mutation, cross over, breeding of the 1,000
 "fittest" members.

repeat

generation 1

$$X_1 = [X_{r0} \ X_{r1} \ X_{r2} \ \dots \ X_{rp}]$$

$$= f(X_{r0}) \ f(X_{r1}) \ f(X_{r2}) \ f(X_{rp})$$



X = genotype
 $f(X)$ = phenotype.

- Evolution is not random:
- Reproduction is key to genetic algorithm for optimization.

perceptron classification problems:

Start with P

Parameters:

P = population size

E = elite size

M = mutation rates, how many elements

C = crossover type

Fitness function

T = selection type

G = generations

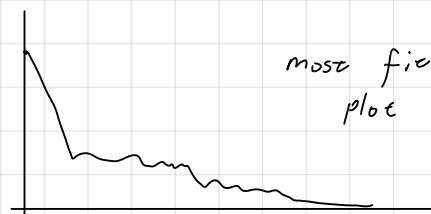
$$\text{Rosenblatt's} = \sum [Y \cdot (\vec{b} \cdot \vec{x})]$$

1. Initialize 1000 random planes

2. Evaluate the value of the fitness function for each one of these planes

3. Rank these planes by the "most to least fit"

4. Perform the crossover and mutation on the top 100.



ex) Most fit

2.7	-1.3	2.0	<u>3.1</u>	<u>2.5</u>	7.1
1.3	-0.5	7.0	<u>3.1</u>	<u>2.0</u>	8.0

Mutate by 1 : 2.8 -1.2 2.1 3.2 2.6 7.2

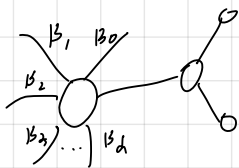
Crossover : 2.7 -1.3 2.0 3.1 2.0 7.1
 1.3 -0.5 7.0 3.1 2.5 8.0

Genetic algorithm :

is good for solving problems

when they are way too many potential solutions

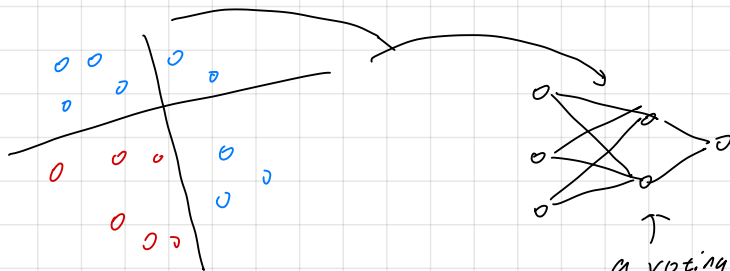
Neural Network + GA (Neuroevolution):



neurons learning
neurons firing.

Instance = recognition

idea of hidden layers in perceptron classification problem:



a voting mechanism on different planes.

Traveling Salesman Problem :

shortest path to go through all the cities one time.

- start with random tours
- reproduce from the fittest members
- crossing over

Selection Criteria :

- proportionate (roulette wheel selection)

GA will likely preserve section of chromosome.

- tournament selection

- Elitism.

Schema Theorem : $\rightarrow$ concept

Genes $\sim \rightarrow$ learned biological concept function

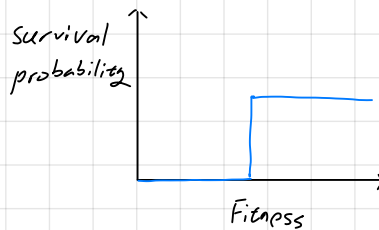
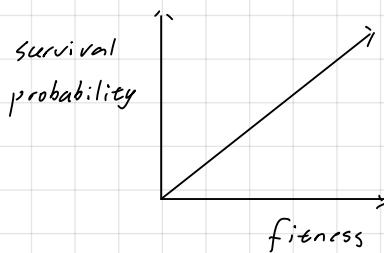


GA $\sim \rightarrow$ learning concepts.

Fitness of all individuals

ranking $\rightarrow$ in nature $\rightarrow$ modeling survival probability.

1 . . .
2 . . .
:
:
:



Selection Methods of GA:

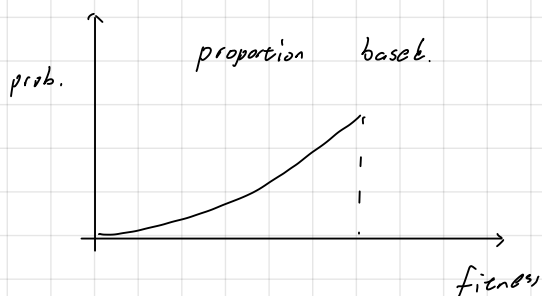
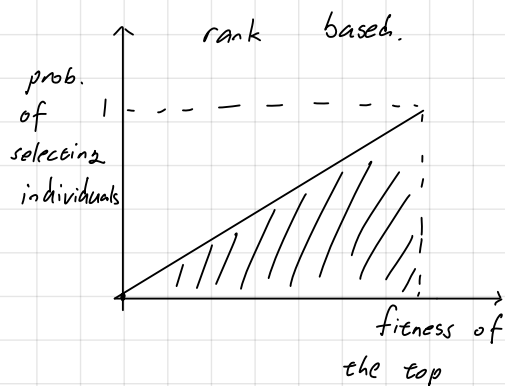
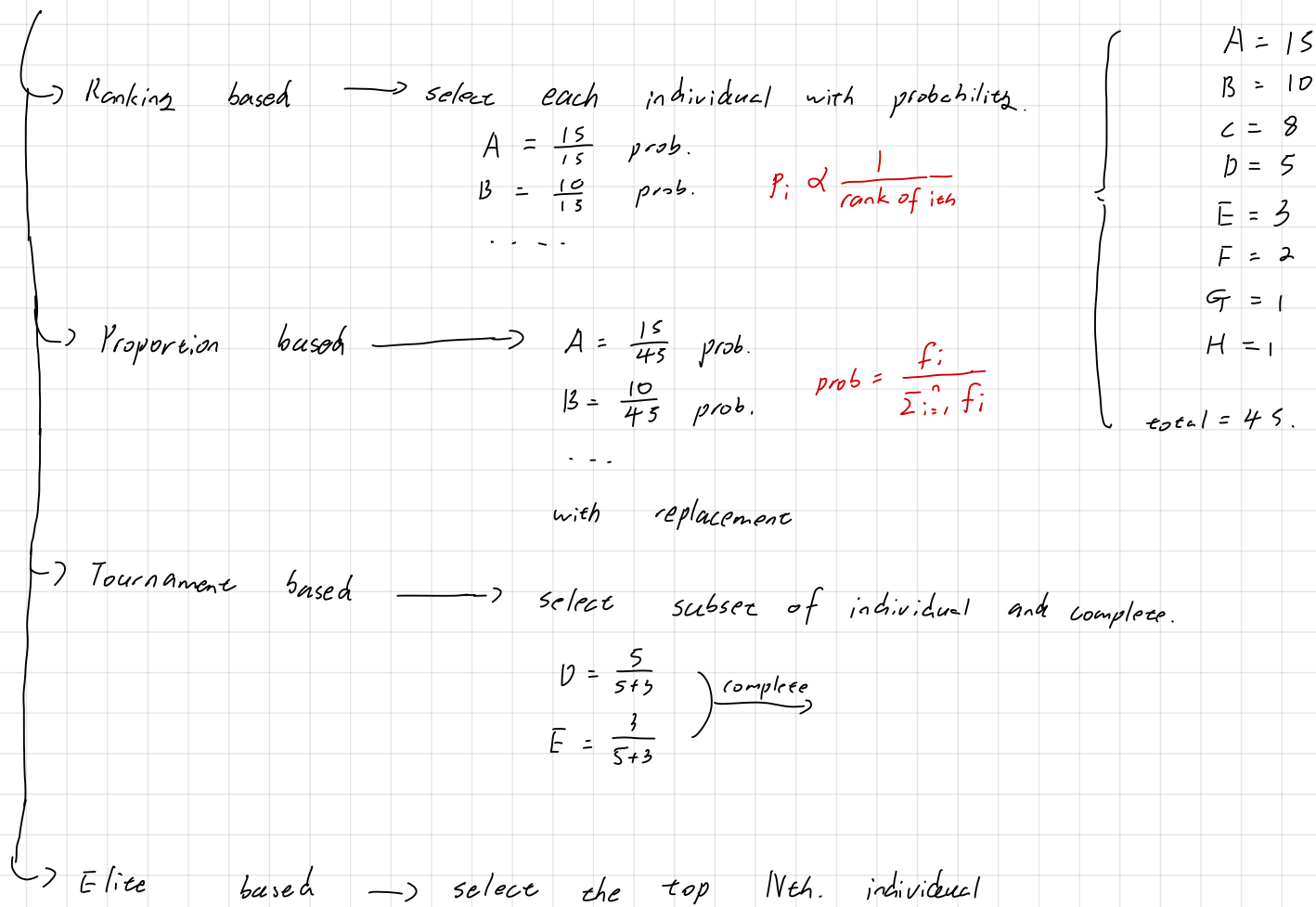
① Generate initial population $\begin{cases} \text{random} \\ \text{heuristics} \end{cases}$
(rough approx. way of getting a solution).

② Evaluate every individual fitness

③ Rank individuals based on the fitness

④ Select some individuals for reproduction & survive.

setting:



5) Mutate some individuals

ex)

$[0|1|0|1|1|0] \rightarrow [1|0|0|1|1|0]$

each bit has 1% of mutation.

ex) swap bits

$[A|B|C|D|E] \rightarrow [A|D|C|B|E]$

6) Crossover some individuals.



ex) single point ordered crossover

$$\begin{array}{cccccc|cccc} A & B & C & D & E & F & & & & & \\ G & H & I & J & K & L & & & & & \end{array} \xrightarrow{\text{crossover}} \begin{array}{cccccc|cccc} A & B & C & D & K & L & & & & & \\ G & H & I & J & E & F & & & & & \end{array}$$

disadvantage:

distance elements get broken up.

ex) double point crossover

$$\begin{array}{cccccc|cccc} A & B & C & D & E & F & & & & & \\ G & H & I & J & K & L & & & & & \end{array} \xRightarrow{\text{double point}} \begin{array}{cccccc|cccc} A & B & I & J & E & F & & & & & \\ G & H & C & D & K & L & & & & & \end{array}$$

ex) Average crossover

$$\begin{array}{ccc|c} 12 & 13 & 1.1 & \dots \\ \downarrow & \downarrow & \downarrow & \text{average} \\ 20 & 10 & 5 & \dots \end{array}$$

ex) unordered crossover

$$\begin{array}{cccccc|cccc} A & B & C & D & E & F & G & & & & \\ H & I & J & K & L & M & N & & & & \end{array} \xrightarrow{\text{unordered}} \begin{array}{cccccc|cccc} A & B & J & K & E & F & N & & & & \\ H & I & C & D & L & M & G & & & & \end{array}$$

ex) Constrained crossover

such as crossover in travelling salesman problem.

Schema Theorem:

pattern

"learned concept"

ex)

pattern

$$\underline{1} \quad * \quad \underline{1} \quad * \quad \underline{0}$$

3 fixed
2 varies

$$\begin{array}{ccccc} 1 & 0 & 1 & 0 & 1 \\ 1 & 1 & 1 & 0 & 0 \\ 1 & 0 & 1 & 1 & 0 \end{array}$$

all follows the pattern/schema.

• Order of Schema

of fixed bits

$O(H_1) = 3.$

$f(H)$ = average fitness of population member that contain H.

• Defining length of a schema.

$\delta(H_1) = 5 - 1 = 4.$

location of last fixed bit - location of first fixed bit

$M(H, t)$ = # of members in the population that contains H at generation t.

H is good if $f(H) > \frac{\sum f_i}{n}$ ← average fitness of the entire population.

Using proportional selection, where the probability of i th member getting selected is

of member with H schema that will get selected is

$$M(h, t+1) = M(h, t) \cdot \frac{f(H)}{\bar{f}}$$

$\bar{f}$ average of all

$$\frac{f_i}{\sum_{i=1}^n f_i}$$

if $f(H) > \bar{f}$, then $\frac{f(H)}{\bar{f}} > 1$,

$$M(H, t+1) > M(H, t)$$

Game Theory

- Social interaction games
- Prisoner's Dilemma. → Both people got caught.

		Other person	
		Stay quiet	Confess
You	stay quiet	-1, -1	-12, 0
	confess	0, -12	-8, -8

get interrogated separately

Nash Equilibrium. (If everyone is making a choice that is optimal given everyone else's choice stays the same).

best strategy is confessing regardless what other choose.

logical outcome to the result that both will confess

prisoner dilemma

ex) Coke, Pepsi

		Don't advertise	Advertise
Don't Advertise	Don't Advertise	500, 500	1000, 5
	Advertise	5, 1000	✓✓

ex)

		Opponent		
		A	B	C
You	left	2, 3	1, 4	7, 3
	middle	4, 1	3, 3	6, 2
	right	-1, 9	2, 9	8, -1

• opponent will never take C as B is better over all than C.

• you will never take right as middle is better

• Nash Equilibrium

ex)

	Deer	Rabbit
Deer	5, 5	0, 2
Rabbit	2, 0	1, 1

"rabbit hunting day"

Deer = 10 lbs, Rabbit = 1 lb, 2 rabbits, 1 deer

Contract Theory $\rightarrow$ constrain

Nash equilibrium shifted due to constrain

ex).

Stay Home Go Out

Stay Home	5, 5	2, 8
Go Out	8, 2	0, 0

Nash equilibrium tells you how others will play

or how others will keep a behavior without any enforcement

ex)

Dive Right Dive Left

kick right	0, 0	-1, 1
kick left	1, -1	0, 0

No Nash Equilibrium.

$\Rightarrow$ Random, $\frac{1}{2}$ right, $\frac{1}{2}$ left

②

	C	D
A	3, -3	-2, 2
B	-1, 1	0, 0

$\Rightarrow$

mixed strategy nash equilibrium

$\downarrow$

assume player B can choose either C or D with equal chance

Assume player ① pick

A with prob p ,

pick B with prob. $1-p$.

$\Rightarrow$ expected reward for picking C is $P_A(-3) + (1-P_A)(1)$
 $= -4P_A + 1$

reward for D is $P_A(2) + (1-P_A)(0)$

$= 2P_A$

• Pick P_A such that Reward C & Reward D are same.

$$\text{So, } P_A(-3) + (1 - P_A) = P_A(2) + 0$$

$$1 - 4P_A = 2P_A$$

$$P_A = 1/6$$

⇓

Player 1 picks A with prob. $1/6$, there is no strategy that player 2 can pick that increases the reward.

• Do the same for player 2.

$$\text{reward for A} = P_C(3) + (1 - P_C)(-2)$$

$$\text{reward for B} = P_C(-1) + (1 - P_C)(0)$$

$$\text{So, } 3P_C - 2 + 2P_C = -P_C$$

$$P_C = 1/3$$

Overtime, player 1 will pick A with $1/6$
player 2 will pick C with $1/3$

↓

Equilibrium.

maximizing the reward given that we don't know the opponent.

(A)

		(B)	
		keep quiet	confess
(A)	keep quiet	20, 20	0, <u>21</u>
	confess	<u>21</u> , 0	<u>5</u> , <u>5</u> (equilibrium)

maximize points

1. Given all the potential actions of my opponents, I pick my best move.

2. If my opponents know what I am going to do, they still can't hurt me.

Iterative Prisoner's Dilemma strategy → tip for tac

→ holding grudge (keep quiet, then confess)

Review of concepts:

(B)

N.E. if everyone is picking their best strategy

(A)

	L	M	R
U	<u>1</u> , 3	1, <u>4</u> ← best strategy for A	7, 3
M	4, 1	<u>3</u> , <u>3</u>	6, 2
D	-1, <u>9</u>	2, 8	<u>8</u> , -1 ↑ best strategy for A

ex) Two Nash Equilibrium.

ex) No Nash Equilibrium.

	L	R
U	<u>9</u> , <u>9</u>	4, 6
D	6, 5	<u>7</u> , <u>8</u>

	L	R
U	6, <u>8</u>	<u>4</u> , 7
D	<u>7</u> , 6	3, <u>7</u>

Battle of the Sexes

(F)

• people stick with the plan the other made.

(M)

	Boxing	Musical
Boxing	2, 1	0, 0
Musical	0, 0	1, 2

(S)

(G)

	L	R
L	0, 0	-1, +1
R	-1, +1	0, 0

⇒ mixed strategy nash equilibrium.

• Assume prob p and $1-p$.

• expected reward for (S) picking L if (G) pick L with p and R with $1-p$.

|

$$= (0)(p) + (1)(1-p)$$

$$= 1-p$$

$$\text{ER for K} = (1)(p) + (0)(1-p) = p$$

$$\Rightarrow \text{ER for L} = \text{ER for K}$$

$$1-p = p$$

$$p = \frac{1}{2}$$

ex)

	L	K
U	-3, 3	-2, 2
D	-1, 1	0, 0

mixed up picking strategy with probability.



genetic evolution!



run genetic algorithm that calculates evolving probability.



approach global min/max

nash equilibrium $\Leftrightarrow$ evolutionary stable strategy.



self-enforceable



• change the scale in the grid to approach some nash equilibrium for manipulation.

ex)

	U	D
U	-3, -3	1, -1
D	-1, 1	0, 0

• everyone is play D.

• Net gain in the population if one choose U.

{ Dove - Hawk Game }

	H	D
H	-1/4, -1/4	2, 0
D	0, 2	1, 1

• If only Doves, Hawk gets 2 points, Doves gets 1 or 0 points

• If only Hawk, Hawk gets -1/4 points, Doves get 0 pt.

Evolutionary Stable Strategy (ESS).

- In a population of a species, what will be the proportion of these strategy.
- With what probability will a species member will utilize one strategy over the other.

Different strategies (behavior)

Bully : Vore pretend to be hawk

Assessor: assess each other and play roles as needed.

- Indicators of fitness - Female choose and male compete.
- Conspicuous consumption - wasteful behavior to show indicator of fitness to be competitive.
(Biological Behavior).

(x)

17 D What is EES?

1- $0, 0$ $\frac{3}{2}, \frac{1}{2}$ • Expected pay offs for the Dove and Hawk have to be same.

D $\frac{1}{2}, \frac{3}{2}$ 1, 1 Suppose H with 6
D with 1 -

$$\text{expected}(H) = (0)(6) + (1-6)\left(\frac{3}{2}\right)$$

similar to Mixed strategy nash equilibrium.

$$\text{expected}(10) = \left(\frac{1}{2}\right)(6) + (1-6)(1)$$

expected 14 = expected 1)

$$\frac{3}{2} - \frac{3}{2}6 = \frac{1}{2}6 + 1 - 6$$

$$6 = \frac{1}{2}$$

		②	
		q	1-q
①	U	3, -3	-2, 2
	D	-1, 1	0, 0
		p	1-p

What should be the strategy for ① and ② for it to be stable mixed strategy nash equilibrium?

- ① should pick p such that payoff for player ② will be same if it plays L or R.

expected (L) if ① choose U with p.

$$E_2(L) = (-3)(p) + (1-p)(1)$$

$$E_2(R) = 2(p) + 0(1-p)$$

① plays U with $\frac{1}{6}$

$$-3p + 1 - p = 2p$$

$$p = \frac{1}{6}$$

$$E_1(U) = 3(q) + (-2)(1-q)$$

$$E_1(D) = (-1)(q) + 0(1-q)$$

② plays L with $\frac{1}{3}$.

$$3q - 2 + 2q = -q$$

$$q = \frac{1}{3}$$

Particle Swarm Optimization (PSO)

find the highest peak.



gradient descent simulated annealing. $\rightarrow$ start with one point

genetic algorithm $\rightarrow$ population spread across the func.

mutate, crossover, reproduce.

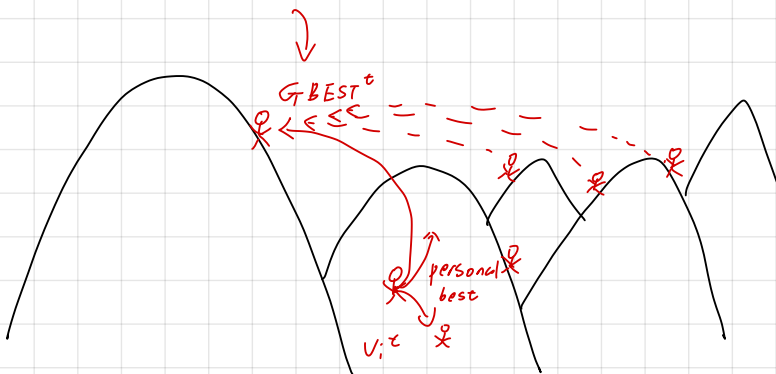
idea:

- Directional velocities

- Move in the direction of
 - global best individual
 - personal best
 - random or previously chosen direction

population fixed.

GBEST can change at every iteration.



velocity of i th individual at the t th step

X_i^t is the location of the i th person at the t th iteration.

$$\Rightarrow X_i^{t+1} = X_i^t + V_i^t$$

• How to update V_i^{t+1} ?

memory of better position.

$$V_i^{t+1} = \alpha V_i^t + \beta (\text{personal best direction})$$

+ γ (global best direction).

• α is the most exploring variable.

• Decrease the value of α every iteration.

$\downarrow$

advantage of idea:

individual can communicate with each other to share their information

• Guarantees to converge

$\downarrow$

global min?

local min?

decrease exploration mindsee
(pretty much like simulated annealing).

Knapsack problem - dynamic programming.

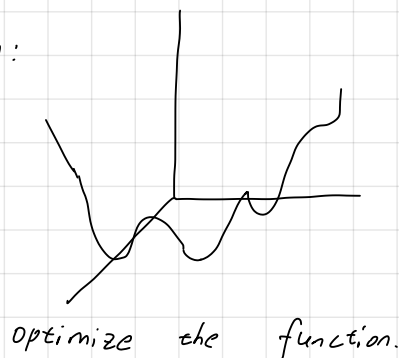
minimize risk & maximize reward.

- $\beta = \text{randomness} = r_1 c_1$ where r_1 is a random # between 0 & 1 and c_1 is maximum weight to the personal best.
- $\gamma = \text{randomness} = r_2 c_2$...

$$So, V_i^{t+1} = \underset{\substack{\uparrow \\ \text{decrease} \\ \text{with } t}}{\alpha} V_i^t + \underset{\substack{\uparrow \\ (0, c_1)}}{r_1 c_1} (\underset{\substack{\downarrow \\ \text{direction of} \\ \text{personal best}}}{P_i^t} - X_i^t) + \underset{\substack{\uparrow \\ \text{random \#} \\ (0, c_2)}}{r_2 c_2} (\underset{\substack{\uparrow \\ \text{direction of} \\ \text{global best}}}{G_i^t} - X_i^t)$$

$$X_i^{t+1} = X_i^t + V_i^{t+1}$$

Objective function:



$$7x^4 - 8y^4 + 3x^3y - \dots + 17$$

- Calculus - derivative / gradient - solve $\frac{df}{dx} = 0$, $\frac{df}{dy} = 0$

- gradient descent - $X_{t+1} = X_t + \alpha \frac{\lambda f}{\lambda x}(X_t)$
 $Y_{t+1} = Y_t + \alpha \frac{\lambda f}{\lambda y}(Y_t)$

- $(Ax + By + C)^4$ for some A, B, C ←

↓

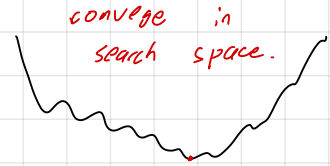
- $(Ax_1 + Bx_2 + Cx_3 + \dots + Zx_{26})^{38}$ ← 37 peaks on a 27 dimensional surface.

Search space is massive!

gradient descent not efficient
highly non-convex.

- Genetic Algorithm
 - PSO
- population based!
individual communicate!

↓
implementation.



Class particle:

position $(x_1, x_2, \dots, x_d)$

velocity $(v_1, v_2, \dots, v_d)$

current fitness $f(x_1, x_2, \dots, x_d)$

best-position $(x_{b1}, x_{b2}, \dots, x_{bd})$

best-fitness $(f(\text{best position}))$.

Class particle swarm:

start random particles

give them random initial velocity.

- compute everyone's fitness
- compute everyone's new position. ie. $x_{t+1} = x_t + v_t$
- update the velocity, $w = 0.99w$

"No Free Lunch Theorem" - No single optimization algorithm will be better than all the rest.

$$v_{t+1} = w v_t + r_1 c_1 (p_{best_i}^t - x_i^t) + r_2 c_2 (Global^t - x_i^t)$$

↑
eventually go to 0.

↑
only these the direction
exists and will converge to one position.

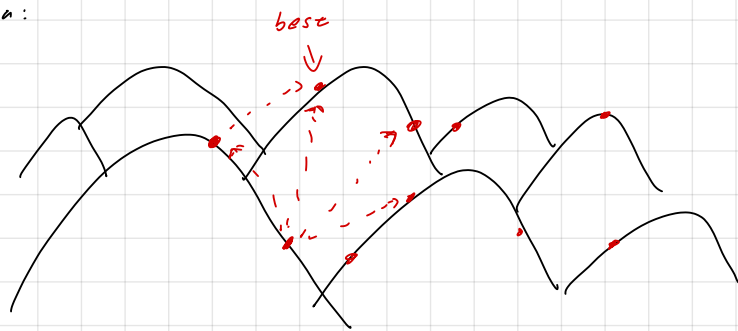
In PSO, No Crossover compared to genetic algorithm.
 No Mutation
 No Selection.

PSO does not guarantee global minimum/maximum.

Variation of PSO algorithm:

Firefly Algorithm:

idea:



idea: more than one particle will influence the velocity of the other.
 compared to PSO,

- Brightness of each firefly is the "Fitness".
- Fireflies are attracted to brightness of other fireflies.

For each firefly i

For each firefly j

if i^{th} brightness $<$ j^{th} brightness:

move $\frac{(x_j - x_i)}{r^2}$ towards j

$\beta_0 e^{-\gamma r_{ij}^2}$ ← brightness that i^{th} firefly sees from j^{th} firefly.

In convergence, we will get cluster of fireflies.

Update: $x_i^{t+1} = x_i^t + \underbrace{\beta_0 e^{-\gamma r_{ij}^2}}_{\text{some exponential decay}} (x_j^t - x_i^t) + \alpha \varepsilon_i^t$

- β_0 is the brightness level at distance 0. generally set to 1.

moving in the direction of j^{th}
scaled by some factor

- γ is the light scattering factor. Larger it is. The less light (visibility) there is.
- r_{ij} is the euclidean distance between x_i & x_j .

$$\sqrt{(x_{i1} - x_{j1})^2 + \dots + (x_{id} - x_{jd})^2}$$

- d is a random exploring tendency.
- Σ_i^t is a random velocity.

Firefly is good at finding multiple peaks.

- nash equilibrium in game theory.

