

Inspired by evolution. Represent potential solutions to a specific problem on a data structure like a chromosome.

Implementation

- + Initial population of random chromosomes
- Allocate reproduction opportunities so that chromosomes which represent a better solution to the target problem have more chance to reproduce.

Broader Sense. Any population-based model that uses selection and recombination operators to generate new samples in a search space.

Optimize : $F(x_1, \dots, x_L)$
variables x_i are discrete
(usually a power of two).

require evaluating $F(x_1, x_2, \dots, x_L)$ to be fast (because we need to do this for each member of the population.)

Page 2.

Size of the search space is large $\sim 2^L$

Canonical Genetic Algorithm:

Initialize strings at random

Evaluate each string and assign a fitness function.

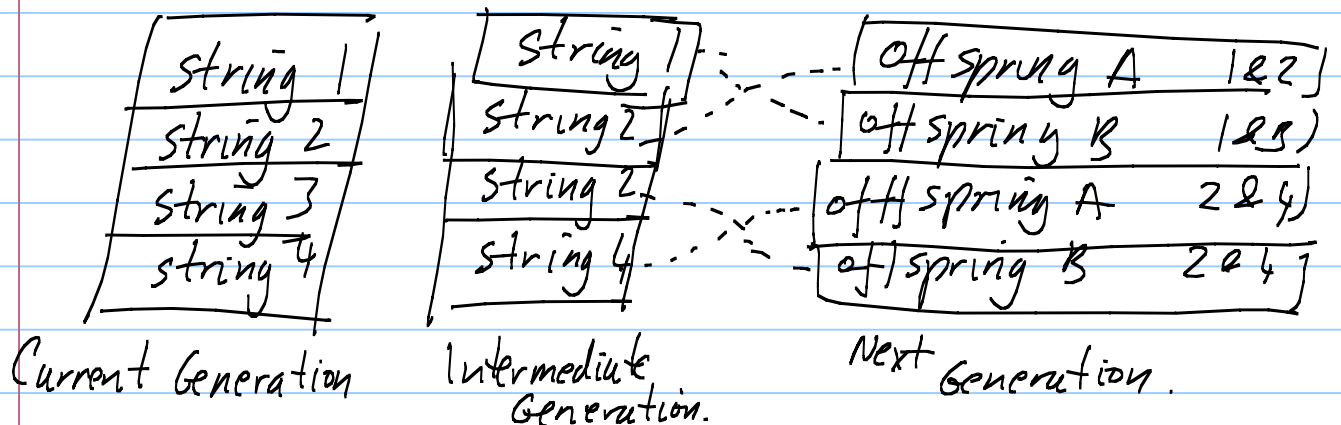
$$\text{Canonical fitness} = \frac{F(x)}{F}$$

/ Here F is the average for the population.

Execution in two stages:

(1) Current Population selection is applied to create an intermediate population.

(2) Recombination & Mutation are applied to the intermediate population to create the next population.



Page 3:

Probability that strings are selected (duplicated) is proportional to their current fitness $\frac{F(x)}{\bar{F}}$.

Several ways to do selection.

(1) Spin the Roulette Wheel

"stochastic sampling with replacement"

(2) "remainder stochastic sampling"

For each x s.t. $\frac{F(x)}{\bar{F}} > 1$

take the integer portion $\left[\frac{F(x)}{\bar{F}} \right]$ and place this number of copies in the intermediate population. All strings place additional copies with prob according to the fractional portion.

(e.g. $\frac{F(x)}{\bar{F}} = 1.36$ places one copy of x and another copy with probability 0.36)

Page 4.

Next recombination

Crossover is applied to randomly paired strings with probability P_c .

Recombine these strings to form two new strings.

E.G. $\begin{array}{c|c} 11010 & 01100101101 \\ \hline yxyyx & yxyyyxyxyxy \end{array}$

$\begin{array}{l} X=1, y=0 \\ \text{or} \\ X=0, y=1 \end{array}$

gives two offsprings

$\begin{array}{c} 11010yxyyyxyxy \\ yxyyx01100101101 \end{array}$

Mutation.

For each bit of the string,
mutate with probability p_m ($p_m \approx 0.01$)

Output next population.

Page 5.

Why does this work?

It encourages strings with good fitness to reproduce - and discourages those with poor fitness.

Mutation is similar to standard MCMC - with mutation there is no interaction between strings.

Crossover & Recombination are new.

Schema Hypothesis

suppose a subpart of the problem can be encoded easily.

E.g. 010101*****

← wildcards

a string with wild cards is called a Schemata.

Low order hyperplanes are sampled by numerous strings in the population.

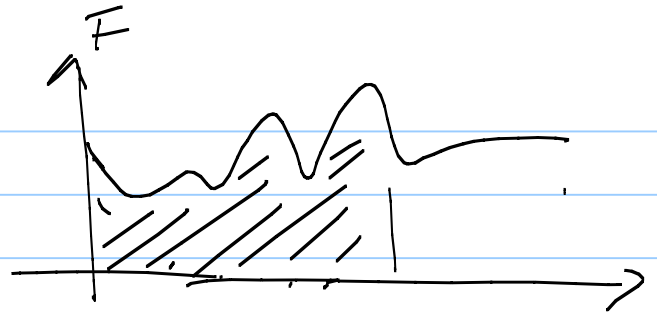
Implicit Parallelism - many hyperplanes computations are simultaneously solved in parallel.

Page 6.

Alternative View.

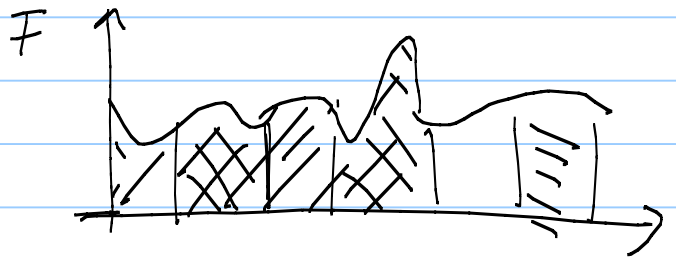
0***...* spans first half of space (shaded |||)

1***...* spans second half of space.



strings with 0***...* have on average better fitness than those 1***...*

So 0***...* will be encouraged.



The strings with
1***
are shaded \\\

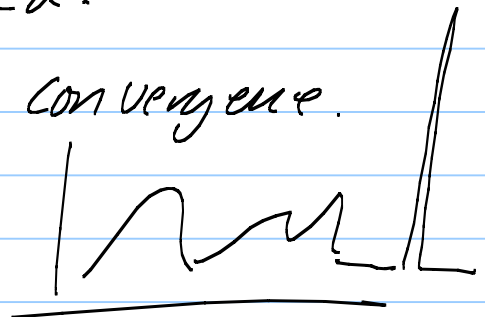
Shows that 0*1***** strings are good.

Local optima do not affect this.

Crossover & Recombination allows good partial solutions to be selected.

But doesn't guarantee convergence.

E.g. hard to find a narrow spike.



Page 7

What does recombination do?

Order-1 hyperplanes (e.g. $*y**.*$)
are unaffected.

Order-2 hyperplanes

for 1-point crossover

11**...*

→ prob of bits separated is $\frac{1}{L-1}$

1**...*1

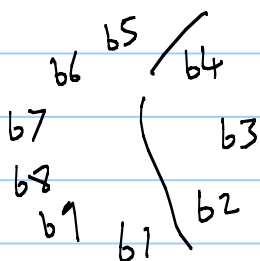
→ prob of bits separated is

$$\frac{L-1}{L-1} = 1.$$

Hence order is very important //

2-point crossover

— two randomly chosen crossover points



— exchange the segment that falls between these two points.

(Note: 1-point crossover is a special case, where one crossover point always occurs between the first and last bits).

Bits which are close together are less likely to be affected by crossover.

Favours compact representations.

Page 8.

Define the length of a schema to be the distance between the first and last bit. length
7
***1**0**10**

The Schema theorem provides a lower bound on the change in sampling rate for a single hyperplane.

$M(H, t)$ be no. of strings sampling hyperplane H at time t .

$$M(H, t + \text{inter}) = M(H, t) \frac{F(H, t)}{\bar{F}}$$

Here $F(H, t)$ is the average fitness of strings in H .

$$M(H, t+1) = (1 - p_c) M(H, t) \frac{F(H, t)}{\bar{F}} + p_c \left\{ M(H, t) \frac{F(H, t)}{\bar{F}} (1 - \text{losses}) + \text{gains} \right\}$$

Conservation Assumption

Crossover within the defining length of the schema is always disruptive.

E.G. Schema 11**...*

If string 1110101 were combined with 100000 or 010000 then no disruptions will occur.

Also recombining 1000... with 010...0 would give a gain.

Ignore gains

Page 9

Conservative Assumption:

ignore gains, treat all disruptions
as losses.

$$\text{thus } M(H, t+1) \geq (1-p_c) M(H, t) \frac{F(H, t)}{\overline{F}} \\ + p_c \left\{ M(H, t) \frac{F(H, t)}{\overline{F}} (1 - \text{disruption}) \right\}$$

Description is $\frac{\Delta(H)}{L-1} (1-p(H, t))$

$\Delta(H)$ is the defining length associated with
1-point crossover.

($p(H, t)$ is the proportional representation of t obtained by dividing $M(H, t)$ by the population size)

Schema Theorem

$$p(H, t+1) \geq p(H, t) \frac{F(H, t)}{\overline{F}} \left\{ 1 - p_c \frac{\Delta(H)}{L-1} (1 - p(H, t) \frac{F(H, t)}{\overline{F}}) \right\}$$

Schema Theorem suggests that mutation is
less important than crossover.

But mutation is necessary to prevent the
system from getting stuck - e.g. 0.00.

Page 10.

Note: 1-point crossover is easy to model analytically.

But 2-point crossover is better for minimizing crossover disruption.

How effective are genetic algorithms?

Claim:

The standard genetic algorithm works poorly on most problems.

Non-standard genetic algorithms, which incorporate knowledge of specific problems work better.

Hybrids are good - combine existing algorithms with population methods.

Page 11,

Evolutionary Monte Carlo.

$$\pi(\underline{x}) \propto e^{-E(\underline{x})}, \quad \underline{x} \in \mathcal{X} \text{ sample space.}$$

$\underline{X} = \{\underline{x}_1, \underline{x}_2, \dots, \underline{x}_m\}$ population.
 m temperatures, $T = (t_1, t_2, \dots, t_m)$
ordered $t_1 > t_2 > \dots > t_{m-1} > t_m$.

Each chromosome $\underline{x}_i$ has a temperature t_i .

Gibbs for i^{th} chromosome is

$$\pi_i(\underline{x}_i) = \frac{1}{Z_i(t_i)} e^{-E(\underline{x}_i)/t_i}.$$

$$\pi(\underline{x}) = \frac{1}{Z[T]} e^{-\sum_{i=1}^m E(\underline{x}_i)/t_i} \quad Z[T] = \prod_i Z_i(t_i).$$

$$\underline{x}_i = (b_{i,1}, \dots, b_{i,d}) \quad b_{i,j} = 0 \text{ or } 1.$$

Mutation select $\underline{x}_k$ at random,
mutate $\underline{x}_k$ to $\underline{y}_k$.

New population $\underline{Y} = \{\underline{x}_1, \dots, \underline{y}_k, \dots, \underline{x}_n\}$

Accept with prob $\min(1, r_m)$
 $r_m = e^{-(E(\underline{y}_k) - E(\underline{x}_k))/t_k}$

Crossover: $\underline{x}_i$ & $\underline{x}_j$ are selected
to generate two offsprings $\underline{y}_i$ & $\underline{y}_j$

Page 12

New population $\underline{Y} = \{ \underline{x}_1, \dots, \underline{y}_i, \dots, \underline{y}_j, \dots, \underline{x}_m \}$,

accepted with probability $\min\{1, r_c\}$,

$$r_c = \exp \left\{ - \frac{(E(\underline{y}_i) - E(\underline{x}_i))}{t_i} - \frac{(E(\underline{y}_j) - E(\underline{x}_j))}{t_j} \right\} \frac{T(\underline{Y}|\underline{X})}{T(\underline{X}|\underline{Y})}$$

Transition

$$T(\underline{Y}|\underline{X}) = P[(\underline{x}_i, \underline{x}_j) | \underline{X}] P[(\underline{y}_i, \underline{y}_j) | (\underline{x}_i, \underline{x}_j)]$$

selection
probability of
 $(\underline{x}_i, \underline{x}_j)$ from $\underline{X}$

prob. of
generating
 $\underline{y}_i \& \underline{y}_j$

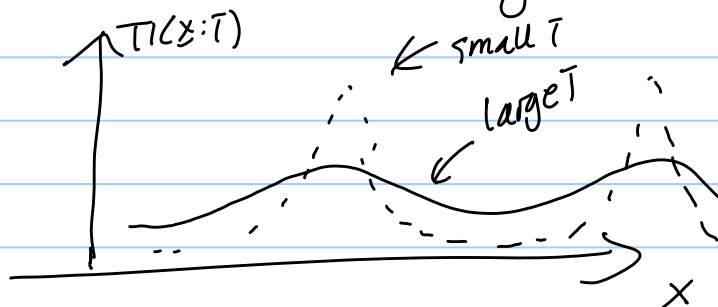
e.g. could have
 $P[(\underline{x}_i, \underline{x}_j) | \underline{X}] \propto \{ e^{-E(\underline{x}_i)/t_i} + e^{-E(\underline{x}_j)/t_j} \}$
weighted samples. $\underline{x}_i \neq \underline{x}_j$

Simulated Annealing

Want to sample from. $\pi(x) = \frac{1}{Z} e^{-E(x)}$...

Introduce a temperature T , create a family of distributions $\pi(x; T) = \frac{1}{Z(T)} e^{-E(x)/T}$

These all have the same peaks - at places where $E(x)$ is small - but distributions with larger T are smoother.



So it is usually quicker to sample $\pi(x; T)$ for large T . (easier to sample a smooth distribution)

Key idea of SA

Sample from $\pi(x; T)$ for large T

e.g. by Metropolis-Hastings

after M iterations, $x^1, \dots, x^M$

reduce $T \rightarrow T\lambda$ ($\lambda < 1$)

and sample $x^{M+1}, \dots, x^{M+1}$ from $\pi(x; T\lambda)$

repeat. until $T=1$.

i.e. $T\lambda^n = 1$, n number of times you decrease T

Question - what value for M ? what value of λ ?

Answer - Problem dependant - find by experiments

Page 14

You can also apply S.A. to find the lowest energy states $\rightarrow$ i.e. $\hat{x} = \underset{x}{\text{ARG MIN}} E(x)$

$$\pi(x; T) = \frac{1}{Z(T)} e^{-E(x)/T}$$

As $T \rightarrow 0$

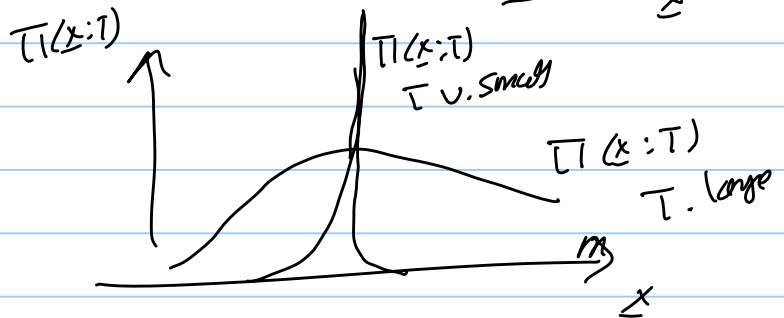
$$\pi(x; T) = 0, \quad x \neq \underset{x}{\text{ARG MIN}} E(x)$$

Sample from $\pi(x; T)$

start with T -large
gradually reduce T

$T \rightarrow \lambda T$

until $T \approx 0$.



For M.H. $\rightarrow$ Proposed $\pi(x^{t+1} | x^t)$

$$\text{Accept} \quad \min \left\{ 1, \frac{\pi(x^{t+1}; T) \cdot \pi(x^t | x^{t+1})}{\pi(x^t; T) \cdot \pi(x^{t+1} | x^t)} \right\}$$

$$\frac{\pi(x^{t+1}; T)}{\pi(x^t; T)} = e^{-\{E(x^{t+1}) - E(x^t)\}/T}$$

So as T decreases you alter the acceptance probability.

Note: S.A. is useful for some problems.

It is not as successful as originally hoped.

Theoretical proofs exist which guarantee convergence - but these require extremely slow smoothing rates (λ) - so not useful in practice.