

# **Eugenic Evolution for Combinatorial Optimization**

**John William Prior**

**Report AI98-268 May 1998**

`jprior@cs.utexas.edu`  
`http://www.cs.utexas.edu/users/jprior/`

Artificial Intelligence Laboratory  
The University of Texas at Austin  
Austin, TX 78712

Copyright

by

John William Prior

1998

# **Eugenic Evolution for Combinatorial Optimization**

by

**John William Prior, B.S. Applied Mathematics, Columbia University**

## **Thesis**

Presented to the Faculty of the Graduate School of

The University of Texas at Austin

in Partial Fulfillment

of the Requirements

for the Degree of

**Master of Arts**

**The University of Texas at Austin**

May 1998

# **Eugenic Evolution for Combinatorial Optimization**

**Approved by  
Supervising Committee:**

RISTO MIIKKULAINEN

RAYMOND MOONEY

# Acknowledgments

I would like to thank Professor Risto Miikkulainen for being unbelievably patient and encouraging during the creation of this thesis. I want to thank all my friends who put up with me saying "I'm going finish it soon ... really I am ..." thousands of times.

JOHN WILLIAM PRIOR

*The University of Texas at Austin*  
*May 1998*

# Eugenic Evolution for Combinatorial Optimization

John William Prior, M.A.  
The University of Texas at Austin, 1998

Supervisor: Risto Miikkulainen

In the past several years, evolutionary algorithms such as simulated annealing and the genetic algorithm have received increasing recognition for their ability to optimize arbitrary functions. These algorithms rely on the process of Darwinian evolution, which promotes highly successful solutions that result from random variation. This variation is produced by the *random* operators of mutation and/or recombination. These operators make no attempt to determine which alleles or combinations of alleles are most likely to yield overall fitness improvement. This thesis will explore the benefits that can be gained by utilizing a direct analysis of the correlations between fitness and alleles or allele combinations to intelligently and purposefully design new highly-fit solutions.

An algorithm is developed in this thesis that explicitly analyzes allele-fitness distributions and then uses the information gained from this analysis to purposefully construct new individuals “bit by bit”. Explicit measurements of “gene significance” (the effect of a particular gene upon fitness) allows the algorithm to adaptively decide when *conditional* allele-fitness distributions are necessary in order to correctly track important allele interactions. A new operator—the “restriction” operator—allows the algorithm to simply and quickly compute allele selection probabilities using these conditional fitness distributions. The resulting feedback from the evaluation of new individuals is used to update the statistics and therefore guide the creation of increasingly better individuals. Since explicit analysis and creation is used to guide this evolutionary process, it is not a form of Darwinian evolution. It is a pro-active, contrived process that attempts to intelligently create better individuals through the use of a detailed analysis of historical data. It is therefore a *eugenic* evolutionary process, and thus this algorithm is called the “Eugenic Algorithm” (EuA).

The EuA was tested on a number of benchmark problems (some of which are NP-complete) and compared to widely recognized evolutionary optimization techniques such as simulated annealing and genetic algorithms. The results of these tests are very promising, as the EuA optimized all the problems at a very high level of performance, and did so much

more *consistently* than the other algorithms. In addition, the operation of EuA was very helpful in illustrating the structure of the test problems. The development of the EuA is a very significant step to statistically justified combinatorial optimization, paving the way to the creation of optimization algorithms that make more intelligent use of the information that is available to them. This new evolutionary paradigm, eugenic evolution will lead to faster and more accurate combinatorial optimization and to a greater understanding of the structure of combinatorial optimization problems.

# Contents

|                                                                           |           |
|---------------------------------------------------------------------------|-----------|
| <b>Acknowledgments</b>                                                    | <b>v</b>  |
| <b>Abstract</b>                                                           | <b>vi</b> |
| <b>Chapter 1 Introduction</b>                                             | <b>1</b>  |
| <b>Chapter 2 Evolutionary Combinatorial Optimization</b>                  | <b>3</b>  |
| 2.1 Combinatorial Optimization and Evolution . . . . .                    | 3         |
| 2.2 Evolution . . . . .                                                   | 6         |
| 2.2.1 Darwinian Evolution . . . . .                                       | 6         |
| 2.2.2 Lamarckian Evolution . . . . .                                      | 7         |
| 2.2.3 Eugenic Evolution . . . . .                                         | 7         |
| 2.3 Current Methods of Evolutionary Combinatorial Approximation . . . . . | 8         |
| 2.3.1 Mutation-based and Pattern-based Evolutionary Algorithms . . . . .  | 8         |
| 2.3.2 The Population and Search Focus . . . . .                           | 11        |
| 2.3.3 Schemata Processing . . . . .                                       | 12        |
| 2.4 Eugenic Combinatorial Approximation . . . . .                         | 16        |
| <b>Chapter 3 Problem Structure and Difficulty</b>                         | <b>18</b> |
| 3.1 “Distance”: Genetic and Algorithmic . . . . .                         | 19        |
| 3.2 Attractors . . . . .                                                  | 20        |
| 3.3 Hill-Climbability . . . . .                                           | 21        |
| 3.4 Deception . . . . .                                                   | 23        |
| 3.5 Epistasis . . . . .                                                   | 24        |
| 3.6 Fitness Distribution . . . . .                                        | 30        |
| 3.6.1 ANOVA of Allele-Fitness . . . . .                                   | 31        |
| 3.6.2 Gene Significance and Fitness ANOVA . . . . .                       | 33        |

|                                                                         |           |
|-------------------------------------------------------------------------|-----------|
| <b>Chapter 4 The Eugenic Algorithm</b>                                  | <b>34</b> |
| 4.1 Example of Eugenic Algorithm Execution . . . . .                    | 34        |
| 4.1.1 Pseudo-Code for the Eugenic Algorithm . . . . .                   | 35        |
| 4.1.2 Example Execution . . . . .                                       | 36        |
| 4.2 Allele Selection . . . . .                                          | 40        |
| 4.2.1 Defining Allele Fitness: Average Fitness versus Maximum Fitness . | 41        |
| 4.2.2 Adaptive Allele Selection . . . . .                               | 42        |
| 4.2.3 Modifications to Fitness-Proportionate Selection . . . . .        | 43        |
| 4.2.4 Clustering and Allele Significance . . . . .                      | 44        |
| 4.3 Population Restriction . . . . .                                    | 45        |
| 4.3.1 The Restriction Operator . . . . .                                | 45        |
| 4.3.2 The Probability of Restriction . . . . .                          | 47        |
| 4.3.3 Consequences of Population Restriction . . . . .                  | 48        |
| 4.4 Steady-State Replacement . . . . .                                  | 51        |
| 4.5 Summary of the Eugenic Algorithm . . . . .                          | 52        |
| <b>Chapter 5 Experiments</b>                                            | <b>53</b> |
| 5.1 Quantifying Performance . . . . .                                   | 55        |
| 5.2 The Algorithms . . . . .                                            | 56        |
| 5.2.1 Increasing Mutation Hill Climbing (IMHC) . . . . .                | 57        |
| 5.2.2 Decreasing Mutation Hill Climbing (DMHC) . . . . .                | 58        |
| 5.2.3 Simulated Annealing (SA) . . . . .                                | 58        |
| 5.2.4 The Genetic Algorithm (GA) . . . . .                              | 59        |
| 5.2.5 The Eugenic Algorithm (EuA) . . . . .                             | 59        |
| 5.3 Combinatorial Optimization Experiments . . . . .                    | 60        |
| 5.3.1 F02: The Two-Dimensional Rosenbrock Function . . . . .            | 61        |
| 5.3.2 F30: Subset Sum Problem . . . . .                                 | 67        |
| 5.3.3 F31: Multiple-Knapsack Problem . . . . .                          | 70        |
| 5.3.4 F32: Maximum Cut Problem . . . . .                                | 76        |
| 5.3.5 F33: Minimum Tardy Task Problem . . . . .                         | 80        |
| 5.3.6 F38: Simulated Annealing Test Function . . . . .                  | 86        |
| 5.4 Summary . . . . .                                                   | 92        |
| <b>Chapter 6 Discussion and Future Work</b>                             | <b>95</b> |
| 6.1 The Eugenic Algorithm versus the Genetic Algorithm . . . . .        | 95        |
| 6.2 Problem Structure . . . . .                                         | 96        |
| 6.2.1 Fitness Variation . . . . .                                       | 97        |
| 6.2.2 Epistasis . . . . .                                               | 97        |

|                             |                                                   |            |
|-----------------------------|---------------------------------------------------|------------|
| 6.2.3                       | Local Optima and Deception . . . . .              | 98         |
| 6.3                         | Eugenic Algorithm Parameters . . . . .            | 99         |
| 6.3.1                       | Population Size . . . . .                         | 99         |
| 6.3.2                       | Selection Noise and Creation Rate . . . . .       | 102        |
| 6.3.3                       | The Restriction Operator . . . . .                | 102        |
| 6.4                         | Extensions and Future Work . . . . .              | 103        |
| 6.4.1                       | Adaptive Parameters and the “Multi-EuA” . . . . . | 103        |
| 6.4.2                       | The Statistics used by the EuA . . . . .          | 105        |
| 6.4.3                       | Continuous Optimization . . . . .                 | 108        |
| <b>Chapter 7 Conclusion</b> |                                                   | <b>109</b> |
| <b>Bibliography</b>         |                                                   | <b>113</b> |
| <b>Vita</b>                 |                                                   | <b>114</b> |

# Chapter 1

## Introduction

The field of combinatorial optimization deals with problems whose inputs are discrete arrangements or permutations. These problems are typically NP-complete, and therefore no polynomial time algorithms exist which can solve them. However, some algorithms forsake the goal of finding *optimal* solutions and instead attempt to find *good* solutions in polynomial time. Since there is no guarantee that an optimal solution will be found, these algorithms are called “approximation” algorithms; it is hoped that solutions almost as good as the optimal solution will be found. Because most of these algorithms guide their quest for better solutions by making small modifications to existing solutions, they rely on the heuristic of “neighborhood search” (Garey and Johnson, 1979). Both genetic algorithms (“GA”s) (Goldberg, 1989; Holland, 1975) and simulated annealing (“SA”) (Kirkpatrick and Sherrington, 1988) are examples of neighborhood search algorithms. In past years, both GAs and SA have received increasing recognition for their ability to efficiently approximate arbitrary functions. Both of these algorithms are “evolutionary” algorithms, as they are characterized by survival-of-the-fittest search pruning and the inheritance of information from one iteration (“generation”) of solutions to the next.

There are many methods by which evolutionary algorithms pass information from generation to generation, but most modern algorithms rely on *random* choices to decide which information shall be inherited. For example, genetic algorithms perform a type of beam search, where a simulated process of “natural selection” prunes the event space (“population”) and selects solutions from the population to be modified by the random state-space operators of “recombination” and “mutation”. Standard recombination operators *randomly* combine parts of highly fit solutions to form complete new solutions. No attempt is made to determine which parts might perform best with each other. SA is a probabilistic optimization algorithm that can be viewed as a simplification of the GA; it uses a population of only a single solution, and relies solely on random mutation of this solution to create new

solutions. The standard mutation operators used by both GAs and SA make small random changes to existing solutions. No attempt is made to determine what changes would be most likely to yield improvements in fitness.

What benefit might come from “smart” recombination operators that analyze interactions and dependencies among parts of solutions and solution fitness, and then *intelligently* assemble these parts into new, hopefully higher fitness, solutions? What benefit might come from “smart” mutation operators that only perform those mutations most likely to increase the fitness of selected solutions? This thesis explores these questions through the development of an algorithm that explicitly analyzes the distribution of fitness with respect to specially selected allele combinations, and *intelligently* creates solutions using the information gained from this analysis. Instead of relying on *random* recombinations or mutations of parent solutions to generate offspring, this new algorithm—the “Eugenic Algorithm” (EuA)—constructs new solutions “bit by bit”, probabilistically including parts of solutions based upon statistical measures of their historical performance. The major innovation of this algorithm, the “restriction operator”, will be shown to be paramount to effectively tracking and exploiting dependencies among solution parts (“epistasis”). In particular, the EuA studied in thesis will measure the benefit of genotypic parts by computing the average fitness of solutions in which these parts are found, and then selecting these parts for use in new individuals proportionately to these average fitnesses. The EuA will be compared to standard evolutionary techniques on a variety of combinatorial optimization problems, and will be demonstrated to be superior to these techniques in many situations. The EuA’s success at selecting the right combinations of alleles through the use of statistical measures of performance of both single alleles and dependent groups of alleles represents a major step forward in approximating very difficult combinatorial optimization problems.

## Chapter 2

# Evolutionary Combinatorial Optimization

Evolution is the adaptation of a population to its environment. This adaptation causes the creation of individuals of increasingly higher “fitness”; in environments where the definition of fitness remains static, evolution drives the population to better and better individuals. This process is similar to approximation—the search for good solutions to a particular problem. The parallels between the concepts of evolution and approximation have led to the creation of *evolutionary approximation algorithms*—algorithms that attempt to optimize a particular function through a simulated process of evolution. This section will briefly describe the particular field of optimization studied in this thesis (combinatorial optimization), and then will discuss how various paradigms of evolution (Darwinian, Lamarckian, and eugenic) relate to approximation.

### 2.1 Combinatorial Optimization and Evolution

Combinatorial problems (defined below) are often NP-complete. Approaches to solving NP-complete problems can be roughly divided into two categories (Garey and Johnson, 1979). The first category includes more “traditional” techniques, while the second consists of more modern (and radical) algorithms. The traditional techniques focus mainly on reducing the size of the search by eliminating as many potential solutions as possible. For example, “branch-and-bound” and “implicit-enumeration” (Garfinkel and Nemhauser, 1972) utilize a tree-structured search that generates “partial” solutions, subsequently identifying partial solutions that could not possibly be part of a true solution and then eliminating those solutions containing these partial solutions. Other traditional techniques include “dynamic programming” and “cutting-plane” techniques (Garfinkel and Nemhauser,

1972). The second major category of algorithms used for attacking NP-complete problems pertain solely to combinatorial optimization. These algorithms do not attempt to guarantee the finding of an optimal solution, but instead just try to find good solutions as quickly as possible. They are *approximation* algorithms. This thesis focuses on this class of algorithms.

(Garey and Johnson, 1979) provides a concise discussion of combinatorial optimization. Combinatorial optimization problems consist of three components: (1) a set  $D_\pi$  of problem instances, (2) for each problem instance  $I \in D_\pi$  a set  $S_\pi(I)$  of possible solutions for  $I$ , and (3) an *objective function*  $f_\pi$  that assigns a positive rational number  $f_\pi(I, x)$  (a *solution value*) to each problem instance  $I$  and solution  $x \in S_\pi(I)$ . Combinatorial optimization problems can either be *maximization* or *minimization* problems. For a maximization problem instance  $I$ , an *optimal* solution  $x^*$  is one such that  $f_\pi(I, x^*) \geq f_\pi(I, x)$  for all  $x \neq x^*$ . For a minimization problem, the solution value of  $x^*$  must be less than or equal to the solution value of any other solution. A *combinatorial approximation algorithm* is an algorithm that finds a solution  $x$  for any given problem instance  $I \in D_\pi$ . A *combinatorial optimization algorithm* is an algorithm that finds  $x^*$  for *all*  $I \in D_\pi$ . The performance of these algorithms is judged upon how quickly they find good or optimal solutions, respectively.

Combinatorial approximation is similar to the search for “better” organisms that is performed during biological evolution. This similarity is so great that highly successful combinatorial approximation algorithms have been designed that mimic the natural forces that occur during evolution. In this thesis, these algorithms shall be referred to as *evolutionary combinatorial approximation algorithms*, or simply *evolutionary algorithms*. Just as combinatorial approximation algorithms search for better solutions, biological evolution drives systems of populations to generate individuals of increasingly higher “fitness”, where fitness is implicitly *defined* as successful survival and propagation<sup>1</sup>. In evolutionary combinatorial approximation algorithms, the blurred mixing of the concepts of reproductive success and fitness that occurs in biological evolution is replaced by explicit measurements of fitness and explicit laws of reproduction based upon fitness measurements. This allows evolutionary algorithms to produce solutions that are “better” in any sense a practitioner defines.

The biological terminology of natural evolution transfers directly to the mathematical terminology of combinatorial approximation. Solutions are “individuals”, and an individual’s solution value—its “goodness” measurement—is referred to as its “fitness”. In evolutionary combinatorial approximation, solution components get a new name: “genes”, and each different solution component value is termed an “allele”. The aggregate of a solution’s components is called its “genotype”, while a solution’s specific form is called its

---

<sup>1</sup>Propagation is more important than simple reproduction, since production of offspring that do not reproduce themselves would not be “fit”.

“phenotype”. It is commonly necessary to “decode” individuals’ genotypes in order to determine or create their phenotypes. Often, individuals are encoded as binary vectors, although other representations (such as real vectors) are also commonly used. When individuals are represented as vectors, each vector element is considered to be a gene and each vector element value an allele. This thesis will focus on problems that have been encoded in binary vectors, since most GA research deals with such problems.

Biological evolution itself does not necessarily drive the population to generate individuals of higher ability, unless this ability is strictly defined as propagative success. In biological evolution, better individuals are simply those that produce the most reproducing offspring under a variety of environmental conditions. No decisions are made that *directly* select and promote individuals on the basis of their ability to run faster, jump higher, or go farther with less food. In contrast, for an evolutionary algorithm to efficiently approximate or optimize a function, fitness must be clearly defined, and higher fitness individuals must be explicitly promoted. As a result, if any other ability besides propagative success is desired, an evolutionary algorithm must *directly* encourage the formation of individuals with the desired ability.

In addition to these philosophical similarities between biological evolution and combinatorial approximation, there is a very distinct similarity in the *mechanism* of improvement that is utilized by both processes. A heuristic common among many combinatorial approximation algorithms is “neighborhood search”, in which the search proceeds by making small modifications to existing solutions (Garey and Johnson, 1979). This closely parallels the random genotypic mutations that occur during biological reproduction. Some algorithms even go so far as to mimic sexual reproduction in order to share information among solutions. One such algorithm, the genetic algorithm (GA), is probably the most widely recognized evolutionary combinatorial approximation algorithm. Genetic algorithms rely upon fitness-biased simulated sexual reproduction and random mutation of offspring to generate variation. The fitness-biased reproduction ensures that lower fitness solutions have a smaller chance of contributing their genetic material to the next generation (Goldberg, 1989). This process is very similar to that used of Darwinian evolution. Darwinian evolution relies upon random variation to expand its search into unexplored regions of the genotypic search space, and natural selection to discourage exploration into lower fitness regions. Darwinian evolution is not the only type of evolution possible for evolutionary algorithms. In particular, the Eugenic Algorithm (EuA) developed in this paper does not rely on random mutations to create variation; nor does it rely on simple sexual reproduction to combine information from differing previous solutions. Instead, new individuals are purposefully constructed one part at a time, using any genetic material previously encountered.

To better understand the macro-processes that evolutionary approximation algo-

gorithms implicitly utilize, and to understand the process that the EuA will employ, the traditional paradigms of Darwinian and Lamarckian evolution must be understood. The next section will describe these two paradigms, their differences, and how the EuA creates a new paradigm of optimization—eugenic evolution.

## 2.2 Evolution

Evolution is the adaptation of an entire population of replicating individuals to their environment. Two major forms of evolution—Darwinian and Lamarckian—dominate our understanding of adaptation in natural environments and our research in artificial environments. A third form of evolution—eugenic evolution—will be introduced in this thesis. Eugenic evolution differs from Darwinian and Lamarckian evolution by how the genotypes of new individuals are constructed. While both Darwinian and Lamarckian evolution rely on random genetic inheritances, new individuals in a eugenically evolving population are purposefully created using genetic parts that have been carefully selected based on their historical performance. When contrasting these forms of evolution, the specific mechanism of replication is not important, nor are the particular genotypes, phenotypes, or environments present in the evolutionary systems. What is important is how feedback from individuals' performance is used to guide the construction of future genotypes.

### 2.2.1 Darwinian Evolution

In Darwinian evolution (Darwin, 1859), the genetic operators of random mutation and crossover cause genotypic variation in offspring, resulting in variations in phenotype. The phenotypic variations that produce advantages in survival and reproduction, relative to other variations, become increasingly predominant in the population from generation to generation. There is no *direct* feedback from the environment to an individual's genotype; each individual's genotype remains constant, unaffected by the environment. New genomes arise only through the random actions of mutation and crossover. However, the *frequencies* of genotypes adapt in response to the environment's effect on their corresponding phenotypes, and, as a result, higher "fitness" genotypes come to dominate the evolving population. Thus, in Darwinian evolution, entire *populations* adapt to their environment through biased sampling of higher "fitness" genotypes, even though the genotypes themselves do not directly respond or adapt to the environment. Darwinian evolution is the form of evolution prevalent in such well-known evolutionary approximation algorithms as genetic algorithms and simulated annealing, which both rely on random modifications of existing individuals to create new individuals.

### 2.2.2 Lamarckian Evolution

In Lamarckian evolution (Lamarck, 1914), adaptation occurs in not only the genotypic frequencies of the population, but also in the genotypes themselves. In Lamarckian evolution, it is possible for the phenotypic state of an individual to directly modify its current genotypic state. As environmental interaction shapes and molds an individual's phenotypic state, the environmentally specific “adaptation information” encoded in the resultant state (*learned* information) can directly modify and be absorbed by the individual's genotype (Luria et al., 1981; Whitley et al., 1994). These modifications to offsprings' genotypes are a form of “smart” mutation, guided by the parent's environmental encounters, and this smart mutation can supplant or augment the random genetic operations that cause individual variation in Darwinian evolution. Just as with Darwinian evolution, variations that confer higher reproductive success upon individuals increase in frequency in Lamarckian populations, and therefore populations (in addition to genotypes) will adapt to the environment.

### 2.2.3 Eugenic Evolution

A third evolutionary method is proposed in this paper—eugenic evolution. In eugenic evolution, as with both Darwinian and Lamarckian evolution, genotypic frequencies adapt in response to feedback from environmental trials. However, unlike the more “standard” versions of Darwinian and Lamarckian evolution, where an unsupervised (random), trial and error process of reproduction is biased by natural selection, and where the genotypes of new individuals are limited to slightly modified parental inheritances, eugenic evolution is a *supervised* (directed) adaptive process, where the genotypes of newly created individuals are virtually unlimited.

In eugenic evolution, the correlation<sup>2</sup> between occurrences of genotypic “parts” and individual fitness is explicitly analyzed, and new individuals are carefully constructed part by part, by allowing these correlations to bias the probability of choosing each part for inclusion in these new individuals. Environmental feedback is not limited to only *complete* genotypes and phenotypes—environmental feedback can bypass genotypes and be transmitted directly to genotypic parts. In addition, the genotype of a new individual is not constrained to be similar to the genotypes of a handful of parents—the new genotype can be constructed from any alleles found in the entire population. By piecing together the best “parts” of the most elite individuals, the entire population becomes the parent of the new individual.

It should be noted that eugenic and Darwinian evolution do not necessarily require

---

<sup>2</sup>In this paper, “correlation” refers to the qualitative correspondence between two variables, and not necessarily simple linear correlation.

the same level of technological sophistication as Lamarckian evolution, and therefore it is easier to implement eugenic or Darwinian evolutionary combinatorial approximation algorithms. Darwinian evolution requires only random variation and selection, and eugenic evolution requires only fairly simple correlations between alleles or allele combinations and fitness. For example, calculating the average fitness of individuals containing particular alleles or allele combinations can be sufficient to identify patterns of high average fitness. In contrast, Lamarckian evolution requires much more technological sophistication. In order to ensure that offspring share the adaptations of the parent, it must be possible for transients in the phenotypic state of the parent to be encoded in the genotype of the offspring. Therefore, some mechanism for performing specialized modifications to the genotype must exist, and the mapping from phenotypic state to genotype must somehow be available. Such a mechanism is not necessary for eugenic or Darwinian evolution. Furthermore, even if it were possible for a Lamarckian algorithm to *encode* phenotypic adaptations into a genotype, some mechanism for *identifying* the necessary modifications to the genotype to cause the desired pre-adaptations must be available. In addition, the genotype-phenotype mapping may not be rich enough to encode these transients. For example, it may be practically impossible to extend a genotype in order to encode an individual’s acquired memories or learned concepts. Because of these additional complications, Lamarckian evolutionary algorithms are rare and, for many problems, not even possible, and therefore eugenic evolutionary algorithms are much more promising for practical implementations.

## 2.3 Current Methods of Evolutionary Combinatorial Approximation

This section will discuss current methods of evolutionary combinatorial approximation. It will introduce a new way of categorizing evolutionary approximation algorithms into two groups—“mutation-based” and “pattern-based” algorithms, and discuss some of the best known representatives from each group. Afterward, two major characteristics of evolutionary algorithms—the use of a population for search focus and learning, and the processing of schemata, will be discussed.

### 2.3.1 Mutation-based and Pattern-based Evolutionary Algorithms

Modern evolutionary approximation algorithms rely primarily on Darwinian evolution. Depending on how many individuals contribute to the creation of a new individual, they can be roughly divided into two categories: mutation-based methods and pattern-based meth-

ods<sup>3</sup>. Mutation-based methods are those that modify a single existing individual to produce a single new individual. Simulated annealing is probably the best known mutation-based method. In contrast to mutation-based methods, pattern-based methods create new individuals by combining parts of *multiple* existing individuals to produce a new individual. Genetic algorithms are pattern-based algorithms. Current approaches to both mutation- and pattern-based evolutionary combinatorial approximation will now be discussed.

## Mutation-Based Methods

Mutation-based algorithms repeatedly modify a single individual. These modifications (“mutations”) produce new solutions. In all mutation-based methods, tournaments between the original and mutated solution are performed. The loser of each competition is discarded, and a new solution is then created by modifying the winner. The winner and the new solution then compete in a new tournament, and the process is repeated. In simulated annealing (Kirkpatrick and Sherrington, 1988), the tournaments are stochastic—i.e. the winner of each tournament is not necessarily the highest fitness individual. This allows for “stochastic backtracking”, where algorithm can extract itself from dead-ends. The probability of accepting a lower fitness individual over a higher fitness individual (the “acceptance probability”) can be made to decrease over time. Other mutation-based algorithms, including the “two-membered evolution strategy” (1+1)-ES (Bäck et al., 1993) and the roughly equivalent MBSH algorithm (Baluja, 1995), utilize deterministic competitions that are always won by the competitor of higher fitness.

A secondary characteristic of mutation-based methods are their mutation “temperatures”—the average size of modifications made to competition winners to create new individuals. Lower temperatures correspond to fewer mutations, while higher temperatures correspond to more mutations. This temperature can vary over the course of an algorithm’s search. “Cooling” or “heating” schedules control the rate at which temperature decreases or increases. The cooling schedule of simulated annealing can be set in such a way that the entire search space is eventually explored, and can therefore *guarantee* the discovery of an optimal solution. When a search is guaranteed to visit all solutions in a search space, it is called “ergodic search”. However, any algorithm can be modified to simply enumerate the entire search space and therefore be ergodic; therefore ergodicity by itself is a rather uninteresting property. What is more important is *efficiency*—how quickly an algorithm can find good or even optimal solutions. A simulated annealing process can be made more efficient at finding good (but not necessarily optimal) solutions by increasing its cooling rate. When this is done, the process is called “simulated quenching” (SQ). Simulated quenching

---

<sup>3</sup>The methods could also be termed “asexual” and “sexual” methods, respectively.

does not have the same ergodic search property as simulated annealing, and therefore can not guarantee an optimal solution. However, it has been demonstrated previously (Ingber, 1993) and will be shown in this thesis that simulated quenching can compare favorably in efficiency to genetic algorithms (and even the EuA) on some practical combinatorial optimization problems, but suffer from the same premature convergence problems that GAs encounter as a consequence of too rapid a cooling schedule.

One weakness of mutation-based methods is the paucity of information preservation from one iteration to the next. Pattern-based methods attack this weakness by employing a “memory” more sophisticated than just a single stored individual. Pattern based methods will be discussed next.

## Pattern-Based Methods

Pattern-based methods can be considered to be more sophisticated than mutation-based methods, because they utilize *groups* of solutions to generate new individuals. Information is garnered from individuals in the “parent” group and used in the creation of offspring. The parent group, just like mutation-based methods’ “current individual”, serves as a storehouse of knowledge about the search space.

The genetic algorithm (Goldberg, 1989; Holland, 1975) is probably the best known pattern-based evolutionary algorithm. They begin with a (typically random) population of  $n$  solutions. Two parent solutions are chosen and the “recombination” operator is used to combine parts of the parents to form offspring individuals. This process is repeated  $n$  times. The  $n$  generated can either replace individuals in the current population or be added to a completely new parent population. If offspring replace individuals in the current population, the algorithm is called a “steady state” genetic algorithm (SSGA) (Syswerda, 1991). There are a great many variations of the crossover operator, but two forms of recombination—“M-point crossover” and “uniform crossover” are probably the most common (Goldberg, 1989). In M-point crossover, each parent is divided at  $M$  locations into  $M+1$  contiguous sections, numbered 1 through  $M+1$ . Each parent is divided at the same  $M$  locations. Two offspring are created by exchanging every odd section between the two parents. One-point crossover is probably the most commonly used type of M-point crossover. Uniform crossover is very similar to M-point crossover. Uniform crossover can be thought of M-point crossover, where  $M+1$  is the number of genes in each parent. Therefore each gene is a section, and every section is *probabilistically* interchanged between the two parents. Other recombination operators take many shapes and forms, including crossover operators that use more than two parents to generate a single offspring (Eiben et al., 1994). However, these operators typically give no consideration to the *functional interdependence of alleles* (“epistasis”) when deciding how to section the parents. Alleles are functionally interdependent when

their “best” settings (resulting in individuals of highest fitness) depend on the settings of each other—when found in specific combinations, they are “more than the sum of their parts”. Epistasis is considered to be the most significant factor in a problem’s difficulty. Chapter 3 will discuss epistasis in depth.

The genetic algorithm also employs mutation; however, the probability of mutation is usually very low. Just as with mutation-based methods, standard GA mutation operators “perturb” offspring by changing a small number of alleles. Some GAs use only the mutation operator—they perform no recombination. These GAs are roughly equivalent to running many simulated annealing algorithms in parallel, and are therefore mutation-based methods. Genetic algorithm use mutation operators for two main purposes: to maintain allele diversity and to “tune” individuals by allowing genetic search to explore those areas of the sample space that have a very small Hamming distance (usually 1) from these individuals (neighborhood search). No attempt is made to determine which genes of selected individuals would be most likely to yield overall fitness improvement through mutation. The use of standard mutation operators for local tuning assumes that the difference in fitness is small among individuals separated by small Hamming distances—that *genotypically* similar individuals are *phenotypically* similar.

As mentioned above, the offspring created by the recombination and mutation operators replace lower fitness individuals in the population or are used to create entirely new populations. By repeatedly creating new offspring more often from the parts of highly fit solutions than from low fitness solutions, a simulated process of natural selection and evolution occurs. The process is usually continued until the chance of generating an offspring that is more fit than the best individual in the current population becomes very small. At this point, the search is terminated, and the best individual in the final population is chosen as the solution to the problem at hand.

### 2.3.2 The Population and Search Focus

Evolutionary algorithms are “weak” search methods; they require no domain knowledge to perform their tasks. When very good heuristics are available to guide the optimization of a given objective function (“knowledge-rich domains”), weak search methods will very often not perform as well or as quickly as traditional algorithms (like branch-and-bound) employing these known heuristics, although there are counterexamples in which GAs performed better than the best known conventional algorithms (H-L. Fang, 1993). However, there exist a great number of practical problems for which few heuristics exist (“knowledge-poor domains”). For such domains, trial and error is the only possibility. But it might be hoped that, at the least, weak search methods learn something from their experiences that may help them guide future search. The populations (whether composed of single or multiple

individuals) utilized by evolutionary algorithms allow these algorithms to accomplish this task.

All evolutionary search algorithms, whether mutation-based or pattern-based, maintain a search focus through the use of populations of sample points. Mutation-based methods (simulated annealing, for instance) typically use populations with only one member, while pattern-based methods typically have populations consisting of many individuals. The population can be viewed as a very basic form of learning; the algorithms are storing explicit representations of “good” parts of the target space and unrepresented regions are implicitly avoided and therefore considered “bad”. The population plays a fundamental role in the search strategy, by biasing the selection probabilities of new samples. As discussed by Syswerda (Syswerda, 1993), it is not easily possible to eliminate the population completely and substitute a simply-updated set of allele statistics, as this strategy would lead to repeatedly generating the same alleles with the same probabilities. A *dynamic* population is needed, as it will adaptively change the search focus and force exploration of the search space, instead of repeated exploitation of a static region.

Although the population serves to focus the search in promising regions of the search space, these algorithms make no explicit attempts to categorize interesting or beneficial properties of particular genotypes; they notice nothing and target nothing. However, it is hypothesized that *implicitly*, important genotypic properties are exploited. This is accomplished through *schemata processing*.

### 2.3.3 Schemata Processing

Pattern-based methods’ primary operator is the recombination operator; it chooses genotypic parts or patterns from parent individuals and combines them to form new individuals. These parts are called “schemata”, as they are patterned constituents of a greater whole. Some of these schemata may be more important than others in forming higher fitness solutions, and so, intuitively, recombination operators that more often choose higher “fitness” schemata are more likely to create higher fitness solutions. This idea is called the “building block hypothesis”. This section will discuss how both mutation-based and pattern-based evolutionary search use schemata, and how schemata should be exploited most effectively.

#### Schemata

The concepts of “schemata” and “schemata processing” have been used to quantify the structure of binary combinatorial problems, explain how evolutionary algorithms process the information contained in their populations, and how the evolutionary algorithms should exploit this information (Holland, 1975). Schemata are generalized descriptors of properties

found in an individual’s genotype; several individuals may share the same property, and it is hoped that an evolutionary algorithm will promote the spread of beneficial properties and curtail the propagation of detrimental properties. A single schema describes a set of individuals by enumerating a group of gene values (alleles); all the individuals that contain these alleles belong to the schema. A standard way to represent schemata is to use vectors whose  $n$ th element either specifies a particular allele for the  $n$ th gene or allows the  $n$ th gene to remain unspecified. In this way, schemata are simple, fixed-length patterns containing symbols specific to certain alleles and possibly wildcard characters. For example, a schema describing binary genotypes could be expressed as a sequence from the set  $\{0,1,\#\}$ —the schema ‘0#0’ would “match” the genotypes ‘000’ and ‘010’. Occurrences of either of these genotypes would provide fitness information on this schema, as well as any other schemata that matched the genotypes. The “order”  $K$  of a schema is the number of fixed genes in the genotype; our example schema ‘0#0’ is order 2. We denote “order  $K$ ” as  $o(K)$ . A schema’s order can range from 0 to  $l$ , where  $l$  is the number of genes in (fixed length) genotype.

*Schemata processing* occurs in both mutation- and pattern-based algorithms; schemata processing occurs when schemata are compared and combined in order to form new solutions. Mutation-based evolutionary algorithms take a single genotype and randomly change parts of it; they never deal with parts of multiple genotypes simultaneously. As a result, these algorithms perform the most basic type of schemata processing since they evaluate and compare only  $o(l)$  schemata—the highest order schema that matches the current individual. Since pattern-based methods combine genotypic parts of various size from multiple individuals, they process all schemata from  $o(1)$  to  $o(l)$ . This allows pattern-based algorithms to be more “conservatively explorative” than mutation-based algorithms, since recombination of *previously tested* genotypic parts is a more cautious strategy than randomly mutating existing parts into *completely new* forms (Syswerda, 1993). But how do pattern-based algorithms choose the “best” parts to recombine? Somehow, the “fitness” of schemata must be gauged. The next section will describe how this can be accomplished.

## Schemata and Fitness Information

Each sample point  $x$  yields info on all schemata  $S_x$  for which  $x \in S_x$ , i.e. all those schemata whose defined alleles match those found in  $x$ . As a result, each  $x$  is a member of  $2^l$  schemata, and therefore very few genotypes are needed to sample large numbers of schemata. However, there exists a trade-off between the number of schemata sampled and the amount of information per schema. The amount of information gathered by one genotype about one schema depends on the schema’s order. A genotype is an instance of  $2^l$  schemata ranging from  $o(1)$  to  $o(l)$ , but yields less information about lower-order schema. For example, if a schema is  $o(l)$ , then every position in that schema is fixed, and so only one genotype will be

a member of the schema. Therefore the matching genotype will yield complete information on such a schema. Conversely, an  $o(0)$  schema will contain  $l$  undefined genes, and so  $2^l$  genotypes will be members of the schema. Each member genotype will sample only  $1/2^l$ th of the schema, and so yield vastly less information about the schema.

Evolutionary algorithms that utilize schemata must rely on averages *observed* in the sample population. It is possible to define the “fitness” of a schema as the average fitness of all individuals belonging to the schema, whether or not these individuals are in the current population. However, except for the highest order schemata, it is impractical to calculate schema fitnesses, due to the large number of individuals that match each schema. As a result, algorithms must rely on limited samples of individuals to collect information on the true average fitness of each schema. The statistics collected in such a way are “observed” statistics, instead of the true statistics of the schemata. “Efficient” processing of schemata through the use of observed statistics remains an open research question, as observed average schemata fitnesses could easily be misleading—low-fitness schemata may have high observed fitnesses, and high-fitness schemata may not even be represented in the population. Because of this “sampling error”, a search algorithm must strike a balance between the exploitation of observed high-fitness schemata and the exploration for more information about observed and unobserved schemata. In the remaining parts of this thesis, the “fitness” of a schema will refer its *observed* average fitness.

The greater the amount of knowledge an evolutionary algorithm has about the search space, the greater its chances of success. (Goldberg, 1989) estimated that a random population of size  $\lambda$  will sample on the order of  $\lambda^3$  schemata. Consequently, the computational burden of explicitly recording statistics about each schema encountered quickly becomes prohibitive as population size is increased. Still, larger populations may be more desirable than smaller populations, as they implicitly provide information on many more schemata than smaller populations, and judicious use of this extra information can increase the probability of discovering highly beneficial genotypic properties. Genetic algorithms deal with this problem of “information overload” through the use of fitness-proportionate schemata representation. Instead of using explicit calculations to keep track of schemata fitness, observed schemata fitnesses are *implicitly* recorded, by including higher fitness schemata in the population more often than lower fitness schemata. A fitness-proportionate selection operator chooses above-average fitness individuals for reproduction more often than below-average fitness individuals. This action engenders a positive feedback phenomena, in which GAs allocate trials to schemata with above-average observed fitness more often than to those of below-average observed fitness, and the population becomes increasingly composed of higher (observed) fitness schemata, which in turn increases the probability of the selection of these schemata. This fitness-proportionate representation of schemata

has the additional benefit of focusing the search by probabilistically limiting the possible combinations of schemata achievable through the application of the crossover operator; the search collects more information on higher fitness schemata, as opposed to collecting equal amounts of information about all schemata. However, only the selection operator performs schemata analysis, and it examines only  $o(l)$  schemata. The mutation and recombination operators of genetic algorithms perform no analysis on schemata; they rely on the selection operator to choose highly fit individuals for manipulation, with the assumptions that these individuals will contain schemata important for the construction of even higher fitness individuals and that enough random mixing of schemata will eventually result in the creation of the “best” combinations of schemata. As a result, the population of the GA allows it to avoid the infeasible task of explicitly recording average schemata performance, while at the same time maintains and adapts the GA’s search focus.

### The “Golden” Heuristic?

A blanket assumption that fitness proportionate schemata selection is the “best” selection method to use is simply not true. The heuristics an algorithm uses to guide its search may have to change, depending on the specific task to be accomplished. Sometimes the goal of evolutionary search is not approximation, but the maximization of the cumulative fitness of all function evaluations. (Holland, 1975) proved that allocating exponentially increasing trials to schemata of above-average observed fitness is a minimum-loss strategy. It must be recognized that this strategy is only a *heuristic* for guiding evolutionary search. Whether high fitness schemata should be propagated or not depends on the type of problem attacked, the operators employed by the algorithm, and the goal of the algorithm itself. In optimization problems, the goal is to find very high fitness individuals in as few function evaluations as possible. While allocating trials on the basis of observed schema fitness probably is a good heuristic for locating this individual, it is quite possible that the highest fitness schemata do *not* contain the highest fitness individuals. This phenomena is called “deception” (Goldberg, 1989). When deception is present, the blanket assumption that above-average fitness schemata should be utilized more often than below-average schemata is incorrect. In such situations, an algorithm must take action to avoid converging permanently to the deceptive local optima. These actions could include restarting the search with a new individual that is maximally different from the local optimum (in a binary problem, for instance, the maximally different individual would be the bitwise complement of the local optimum). Therefore, the “golden” heuristic of fitness proportionate schemata selection must be supplemented by other heuristics in order to effectively deal with deceptive and other difficult situations.

## 2.4 Eugenic Combinatorial Approximation

All of the Darwinian evolutionary search algorithms perform minimal analyses of the information contained in their populations; the only information used in the construction of new individuals is a list of one or more previously encountered individuals and their corresponding fitnesses. These algorithms rely on a large number of trials and repeated applications of random operators to overcome their lack of domain knowledge and “blindly” strike about in the search space (albeit in regions biased by the list of historical individuals). The mutations made by these algorithms may be small or large, increasing or decreasing in size, but all the mutations are *random*, in the sense that the choices of what parts to mutate and by how much are not dependent on any analysis of what changes might be best. The genetic recombinations that take place in pattern-based methods are not guided by any criteria of which parts would work most successfully together; parts are just randomly extracted from higher fitness individuals and patched together. Eugenic evolution hopefully reduces some of the need for many trials by performing reasonable amounts of analysis of the information gathered in previous trials and pro-actively exploiting this information to guide the choices made for choosing and combining genetic material.

At least two algorithms, “Binary Simulated Crossover” (BSC) (Syswerda, 1993) and “Population-Based Iterative Learning” (PBIL) (Baluja, 1994), have attempted to pro-actively and explicitly identify and exploit correlations between fitness and genetic parts or complete genotypes. Both are pattern-based methods, but PBIL also can be considered to be a mutation-based method, even though it makes use of a population of samples. BSC proceeds as a SSGA, but does not construct new individuals through random crossover and mutation. Instead, the alleles of a new individual are generated independently, using modified versions of GA selection techniques such fitness-proportionate selection or rank selection. Instead of using these selection techniques to choose *individuals* to participate in reproduction, these techniques select *alleles* for inclusion in a completely new individual. The “fitness” in “fitness-proportionate selection” would refer to an allele’s fitness, instead of an individual’s fitness. For example, if fitness proportionate selection was being used, and a specific allele of a gene was found in several individuals whose mean fitness was 50, while another allele of the same gene was found in individuals whose mean fitness was 25, then the probability of use of the first allele in new individuals would be twice that of the second—the first allele would be two times more likely to appear in a new individual. The “fitness” of an allele is dependent upon a statistic computed over the individuals in which it is found, and in this case the statistic is the sample mean of the individuals’ fitnesses. BSC relies solely on measurements of *marginal* allele-fitness distributions; it makes no effort to determine the dependence of an allele’s fitness conditional upon the presence or absence

of other alleles. BSC makes no attempt to determine which allele *combinations* go best together; it completely ignore schemata of order greater than one. As a result, BSC only analyzes  $o(1)$  schemata and therefore completely ignores epistasis.

PBIL also forsakes random crossover and mutation for an alternate method of offspring creation. A “probability vector” is used to generate each population. This probability vector is simply a vector whose  $n$ th component represents the probability that the  $n$ th gene will take the ‘1’ form in the next population. An entire population is randomly generated with this single vector. After each population is generated, the single best and worst individuals in the population modify the probability vector so as to increase the likelihood that the alleles found in the best individual but not in the worst individual will be produced in the next generation. In a somewhat similar way to simulated annealing’s cooling schedule, the amount that the probability vector is modified by decreases with time. PBIL is pattern-based in the sense that the probability vector represents a pattern of “good” individuals, while it is also mutation-based in that a single pattern that is the “starting point” to generate new solutions. PBIL’s use of a single probability vector is basically equivalent to a BSC algorithm that combines the allele-fitness averages of the highest and lowest fitness individuals from many generations. Since PBIL only processes entire genotypes and probability vectors, and never attempts to analyze parts of these genotypes or vector, it processes only  $o(l)$  schemata.

In both BSC and PBIL, all alleles participate equally in updating the fitness averages and selection probabilities, and no attempt is made to sort out the possibly conflicting information coming from very diverse individuals. BSC processes only  $o(1)$  schemata, and PBIL processes only  $o(l)$  schemata—neither algorithm processes intermediate-order schemata. It will be demonstrated in later chapter that the eugenic algorithm, developed independently of both BSC and PBIL, processes schemata of all orders, through the use of measurements of both marginal and conditional allele-fitness distributions. Therefore the EuA additionally has the extremely important ability (lacking in both BSC and PBIL) to track and exploit epistasis by explicitly examining allele behavior that is dependent upon other alleles, even when these dependencies only occur in schemata of less than order  $l$ .

## Chapter 3

# Problem Structure and Difficulty

A clear understanding of the optimization problems faced by evolutionary algorithms is necessary to fully appreciate the mechanism of the EuA. In this chapter, five major qualitative characteristics of discrete combinatorial optimization problems (attractors, hill-climbability, deception, epistasis, and fitness distribution) are defined and discussed. Although none of these concepts are new, they are usually not defined very concretely. This section will attempt to define and relate them into a unified view of problem structure. Attractors are optima (local or global) that have basins of attraction. Hill-climbable problems are ones that can be solved easily using exploitation—when no suboptimal attractors exist. When suboptimal attractors do exist, both deception and epistasis can occur. Deception is encountered when strong suboptimal attractors exist that are very genotypically different from the global optima. Problems with high epistasis have a great number of small attractors. Greater deception and higher epistasis lead to greater problem difficulty. Unlike deception and epistasis, the relationship between fitness distribution and problem difficulty is not so concrete. Extremely small or large fitness variance or ranges can either adversely affect the success of evolutionary optimization algorithms, decrease problem difficulty, or have no effect at all. It can be very difficult to determine by observation the extent to that these problem characteristics manifest themselves in a particular problem, but quantitative methods for measuring a problem’s epistasis and fitness distribution will be proposed in this thesis. However, before a full discussion of any of these characteristics is possible, the difference between the concepts of *genotypic distance* and *algorithmic distance* must be explored and clarified.

### 3.1 “Distance”: Genetic and Algorithmic

In GA literature, the term “distance” has been associated with both genotypic and algorithmic distance. Although both measures of distance are often directly proportional, this is not always the case. A well known measure of genotypic distance is Hamming distance. Hamming distance is a suitable metric for computing genotypic distance for binary genotypes, as it simply counts the number of gene positions that differ between two individuals. As Hamming distance increases, genotypic similarity decreases. Algorithmic distance, on the other hand, is the amount of search an algorithm must perform to generate one point from another. It is inversely related to transitional probability—the probability that a particular offspring will be generated from a given parent. The algorithmic distance between point A and point B increases as the likelihood decreases that an algorithm’s search will progress directly from point A to point B. If the algorithmic distance between point A and point B is infinite, then the transitional probability from point A to point B is zero. Although algorithmic and genotypic distance are often related, they are independent of each other, since algorithmic distance varies depending on the algorithm in use, while genotypic distance remains constant.

There are cases in which algorithmic distance is directly proportional to genotypic distance. This is true for some of the most popular evolutionary algorithms. In these cases, the use of the term “distance” to simultaneously refer to both distance measures has been both convenient and justified. For mutation-based algorithms that use small mutations, Hamming distance can easily be used to estimate the algorithmic distance of two points. These methods rely on small perturbations of the current solutions to generate new solutions, and so points separated by small Hamming distances are highly likely to be generated from each other, while points separated by large Hamming distances very likely will never be generated from each other—they are truly “far away” from each other. Therefore, for these algorithms, the qualitative relationship between genotypic and algorithmic distance is simple and static. The concepts of genotypic proximity and transitional *improbability* are truly interchangeable, and both concepts can be easily unified into a single idea of “distance”.

For pattern-based algorithms, the concept of distance is much more complicated. An intuitive assumed equivalence between genotypic and algorithmic distance is not valid, since extremely large perturbations to the current solution(s) can and do occur from one iteration (generation) to the next. It is often very likely that a new individual will differ significantly from every individual in its parent population, even though the new offspring was created with only a single iteration of the search. All that is likely to remain similar between parents and offspring are the special patterns that the algorithm promotes. When these patterns

include most of an individual’s alleles, as in mutation-based algorithms, Hamming distances remain small from generation to generation—genotypic distance is small when algorithmic distance is small. When the exploited patterns are more complex, such as the random-length schemata of GAs, the Hamming distances between offspring and their immediate parents can often be very large. Therefore, although the genotypic distance between two individuals may be very large, the *algorithmic* distance could still be very small. The simple relationship between genotypic and algorithmic distance that is encountered in many mutation-based methods breaks down as the complexity of the exploited patterns increase. For pattern-based methods, and for mutation-based methods that utilize extremely high mutation rates, genotypic distance is often largely unrelated to algorithmic distance.

## 3.2 Attractors

Because of the breakdown of the relationship between genotypic and algorithmic distance, the ideas of “attractors” and “local optima” must be approached very carefully, since the very idea of “locality” depends on the distance measure being used. Local optima are *suboptimal* points or regions that optimization algorithms frequently converge to, for one reason or another. Accompanying these sub-optima are their “basins of attraction”, which are regions of the search space such the algorithm is likely to converge to the sub-optima if it enters the region. Whether optimal or suboptimal, points (or regions) having basins of attraction are called “attractors”. An attractor for mutation-based algorithms is not necessarily an attractor for pattern-based algorithms. Attractors for mutation-based algorithms (*mutation-based attractors*) are points that are superior to points in the *genotypic* neighborhood. Fitness in this neighborhood monotonically decreases with increasing genotypic distance from the attractor. Points that are more similar to the attractor have higher fitness, while points that are less similar have lower fitness. Therefore, a mutation-based algorithm that encounters a point in the basin of attraction can easily increase the fitness of the solution by mutating the point to become more similar to the attractor. Once inside the basin of attraction, a mutation-based algorithm would converge to the attractor, unless the algorithm performed large mutations (“jumps”) that made the current solution less similar to the attractor. The “strength” of a mutation-based attractor depends on how far out its basin of attraction extends (i.e. at what genotypic distance does fitness stop monotonically decreasing), and on how many points inside the basin of attraction violate the rule of monotonically decreasing fitness.

In contrast, attractors for pattern-based algorithms *pattern-based attractors* are rooted in suboptimal *patterns*, not suboptimal points. When a genotypic pattern is observed to have fitness superior to that of its alternatives, and therefore promoted throughout the

search, then this *pattern* can be considered to be an attractor for a pattern-based algorithm (see “Deception”, below). The fitness of a pattern is generally considered to be the average fitness of all individuals encountered that match the pattern, but the choice of the statistic used (whether mean, mean-squared, etc.) is arbitrary. Patterns arise through observations of many different individuals, but the patterns do not necessarily lead to any particular highly-fit point. The observations are incomplete, because they often do not sample *all* the points matching a particular pattern, and therefore any statistics calculated using the observations are prone to sampling error. Therefore, the patterns are “noisy” attractors that help guide the algorithm towards *seemingly* good groups of points, not necessarily directly toward specific superior individuals. The strength of a pattern-based attractor increases as the mean fitness of the pattern increases and its variation decreases. Strength decreases with increasing pattern length, because the pattern becomes more specific and therefore matches fewer individuals. As a result, it is less likely that the pattern will be encountered and, when encountered, sampling error will be higher since fewer matching individuals will be available. In effect, by increasing pattern length, the basin of attraction of the pattern decreases in size.

### 3.3 Hill-Climbability

For many evolutionary algorithms, there are some types of discrete optimization problems that are easier to solve than others. In particular, the characteristic of “hill-climbability” can be found in many of the simplest problems. In hill-climbable problems, the exploitation of observed superior properties leads easily to the creation of an optimal individual. In such problems, there is little or no need for exploration and, specifically, there is no need to discard an observed superior pattern or individual in favor of a inferior one. Varying degrees of hill-climbability exist. Problems that have more suboptimal attractors are less hill-climbable, as are problems whose suboptimal attractors have larger basins of attraction, since exploitative algorithms are more likely to converge to the sub-optima in these cases. Therefore hill-climbability decreases with increasing numbers and strength of suboptimal attractors.

There are two types of hill-climbability—mutation-based and pattern-based. In a mutation-based hill-climbable (MBHC) problem, all the mutation-based attractors are global optima. This means that from every non-optimal point in the search space, there are many single-bit mutations that lead to points of higher fitness. Therefore, there is an *exploitative*, single-bit “mutation-trajectory” from every point in the space to a global optimum. A MBHC problem can be solved very easily by a mutation-based “hill-climbing” algorithm, which performs single-bit mutations and replaces the current individual only

when the new individual has higher fitness. Later in this thesis, a specific mutation-based hill-climbing algorithm will be developed to determine the MBHC of any problem.

Pattern-based hill-climbability (PBHC), better known as “building block constructivity”, has been discussed extensively in GA literature (Goldberg, 1989). In theory, a PBHC problem could easily be solved by exploiting alleles and allele sets that have superior observed fitness averages. Such alleles or allele sets would surely be present in a global optimum, and therefore it would be a simple matter for an algorithm to combine superior patterns into an optimum individual. However, in practice, even toy problems that have been designed specifically to be building block constructive are *not* that easily optimized by pattern-based algorithms. However, it has been demonstrated that GAs can perform relatively poorly on problems that are specifically designed to be building-block constructive (the “Royal-Road” functions described in (Forrest and Mitchell, 1993)). In particular, far simpler approaches such as random bit-climbing algorithms have been shown to easily outperform GAs on several of these problems. It is now clear that there are many other factors involved in GA performance, and that building-block constructivity cannot be so easily measured. It is usually the case that only “toy” problems can be clearly seen to be optimizable through the recombination of short, low-order schemata. Similar to MBHC, PBHC decreases when the number and strength of pattern-based suboptimal attractors increase.

In order to more fully understand both mutation- and pattern-based hill-climbability, it is useful to study a problem that is completely mutation- and pattern-based hill climbable. The BIT-COUNT problem is such a problem. It is probably the simplest toy problem used to test and illustrate the actions of evolutionary algorithms. In BIT-COUNT, the fitness of each point ( $\mathbf{x}$ ) is its Hamming distance from the zero-vector ( $\mathbf{x} = 0...0$ ). The optimum is therefore reached when all the bits are set to '1' ( $\mathbf{x} = 1...1$ ), since this point has the highest Hamming distance from the zero-vector. It is mutation-based hill-climbable because any mutation from a '0' allele to a '1' allele leads to an individual of higher fitness that is also closer to the optimum, while mutations from '1' alleles to '0' alleles lead to individuals of lower fitness that are farther away from the optimum. Therefore a completely exploitative mutation-based algorithm will easily and quickly locate the optimum, and the problem is therefore mutation-based hill-climbable. In addition, all schemata that contain more '1' alleles than their rivals will have higher average fitnesses than their rivals, if the number of observations over which these averages are computed is high enough. As a result, the complete exploitation of these higher fitness schemata will lead to optimization. Therefore BIT-COUNT is also pattern-based hill-climbable. It should, however, be noted that pattern-based algorithms do not usually have access to complete schemata information—they must rely on only small samplings of schema members to compute schema fitness, and

are therefore vulnerable to the errors associated with incomplete sampling. The fitness averages observed in smaller samples may incorrectly reflect relative fitnesses, and therefore an exploitative pattern-based algorithm may erroneously promote patterns not matching the optimum. As a result, such algorithms will not always find the optimum of BIT-COUNT as quickly as a completely exploitative mutation-based algorithm.

### 3.4 Deception

Even though a problem may display a high degree of *observed* hill-climbability—where exploitation leads to significant and reliable increases in fitness—the problem may still have local optima. When local optima have very large hill-climbable basins of attraction and differ greatly from the global optima, the problem is “deceptive”<sup>1</sup>. In this case, the easy “progress” made by hill-climbing algorithms is deceptive, as it leads only to local optima that are very genotypical different than the global optima. For example, a particular problem may have one schema that is observed to have a greater average fitness than its competitors. However this “better” schema might not be part of the optimum genotype, even though the promotion of this schema above its competitors results in individuals of increasingly higher fitness. In fact, the schema may match individuals that differ significantly from the optimum individuals. In this case, the problem would be pattern-based deceptive. A similar example can be constructed for mutation-based deception, where mutations towards a local attractor would lead farther and farther away from the true global optimum.

Depending on a problem’s attractors, both pattern-based deception as well as mutation-based deception are possible, but do not necessarily always coexist. “DECEPT” is standard GA test problem that is both pattern- and mutation-based deceptive. DECEPT is very similar in structure to BIT-COUNT, but DECEPT is a much harder problem for an evolutionary algorithm to solve. DECEPT (equation 3.4) is basically identical to “BIT-COUNT”, except that the zero-vector ( $\mathbf{x} = 0...0$ ) itself is the optimum, with slightly higher fitness than the one-vector ( $\mathbf{x} = 1...1$ ). In both problems, when all non-optimal individuals are examined<sup>2</sup>, and then the average fitness of both the ‘0’ and ‘1’ alleles are calculated, it will be found that the average fitness of the ‘1’ allele is higher than the average fitness of the ‘0’ allele. For the BIT-COUNT problem, this intuitively confirms the presence of the optimum at  $\mathbf{x} = 1...1$ . However, for the DECEPT problem, it contradicts the intuition that the observed “average best” allele (allele ‘1’) will be the one present in the optimum, since

---

<sup>1</sup>Deception is typically said to exist in problems with “misleading” (or locally optimal) *schemata* (Goldberg, 1989), and therefore this definition does not distinguish between mutation- and pattern-based deception.

<sup>2</sup>It is reasonable to exclude the optimum from the calculation of allele or schema fitness, since we are assuming that the optimum is not known *a priori* (that’s why we need an optimization algorithm!).

$$l + 1 \text{ when } \mathbf{x} = \mathbf{0} \quad (3.1)$$

$$Hamm(\mathbf{x}, \mathbf{0}) \text{ otherwise} \quad (3.2)$$

Figure 3.1: The DECEPT problem

the optimum contains only '0' alleles. As a result, DECEPT is significantly harder for both mutation-based and pattern-based algorithms. It is hard for mutation-based algorithms to optimize since fitness will increase each time a '0' is mutated to a '1' (except when starting from the optimum). Each seemingly “good” mutation is actually an “incorrect” mutation, in the sense that it will lead to a sub-optimal solution and away from the true global optimum. DECEPT is hard for pattern-based algorithms because, in populations that do not contain the optimum, every “schemata competition”<sup>3</sup> is won by the schema with the greater number of '1' alleles, and so more '1' alleles are propagated through the population. A pattern-based algorithm will allocate increasing amounts of trials to “misleading” schemata that can never be recombined to form the optimum.

### 3.5 Epistasis

In general, the degree to which a problem exhibits epistasis is the degree to which the “best” setting of an allele is dependent on the settings of other alleles—it is the amount of interdependence among a genotype’s alleles (Goldberg, 1989). However, little attention has been given to the interaction between attractors and epistasis. Epistasis occurs when a problem has an extremely high number of very small mutation-based attractors, to such an extent that there are almost no pattern-based attractors, since the basins of attraction of the sub-optima contain so few individuals over which patterns can be discerned. If the problem has only one (global) attractor, then the best setting of every allele is always independent of the settings of the other alleles, and the problem is completely unepistatic—it exhibits 0% epistasis. There is no need to coordinate the selection of alleles; each gene has an elite allele that is consistently found in superior individuals, and so it is a simple matter to construct high fitness individuals using these alleles. BIT-COUNT is an example of a 0% epistatic problem. It has no suboptimal attractors. For any given individual, the best setting of every allele is completely independent of the other alleles.

---

<sup>3</sup>I define a “schemata competition” as a comparison between two schemata with the same order and fixed positions

Conversely, if the best setting of an allele can never be specified without the examination of all the other allele settings, then the problem is 100% epistatic. The best allele choice for a gene can rely heavily on a suitable and “harmonious” setting of other genes. 100% epistatic problems are hard for both mutation-based and pattern-based optimization algorithms, since the settings of almost all of the genes must be coordinated simultaneously in order to create a superior individual. For example, in a mutation-based algorithm, such as SA, the alleles found in the current sample roughly indicate “good” settings. As the search progresses, fewer alleles are changed and correspondingly more alleles are left untouched, as it is hoped that higher and higher fitness combinations have been found. But in a 100% epistatic problem, the “best” setting of each allele changes any time another allele changes values. Therefore SA’s heuristic of mutating fewer alleles as better solutions are found is much less effective. The heuristic of promoting good sets of alleles (schemata), used by pattern-based algorithms, is also less effective in these circumstances, as the fitness of each schemata becomes too dependent on the presence or absence of other “companion” schemata. As a result, random recombination of schemata must rely on the random chance that two entire schemata will work well with each other. Because of the combinatorial explosion of the number of possible combinations of schemata, this random chance is very small. Consequentially, 100% epistatic problems are the hardest problems for evolutionary algorithms.

Epistasis can be partially studied by examining single-allele competitions between otherwise identical genotypes. Such a competition is performed when a single allele of an individual is mutated into an alternative allele. The fitnesses of both the original and the mutated individual are determined, and the allele associated with the higher fitness individual wins the competition. These competitions can be performed on any individual in the search space. In a 0% epistatic problem, each gene will have a special allele that it *always* wins its competitions with its alternate allele, no matter how the rest of the genotype is configured. In a 100% epistatic problem, however, the winners of these competitions will be equally balanced among all the competing alleles. For example, if the genotypes are binary, then the value of ‘0’ for a particular gene will yield superior fitness for half the genotypes, while ‘1’ will be better for the other half. In effect, the search space is partitioned into two sets by a particular gene—the first set containing all those genotypes for which ‘0’ yields the highest fitness, and the other set that contains those genotypes for which ‘1’ is best. In a 100% epistatic problem, these two sets are equal in size—each allele is best only half the time. However, as one set increases in size and the other shrinks, a “pattern” emerges in which one allele more frequently is a better choice, resulting in a corresponding decrease in epistasis. Therefore, by keeping track of the frequency with which alleles yield higher fitness than their complements (i.e. win a competition), epistasis can be estimated. More

“balanced” win frequencies imply more epistasis, as they imply that it is very difficult to decide which allele is better, without examining the rest of the genotype. However, this epistasis estimate only deals with  $o(1)$ -schemata; it ignores competitions among higher-order schemata and how unbalanced the win frequencies of these schemata are. It is unclear as to whether or not this “ $o(1)$ -epistasis” is an under-estimator or over-estimator of overall epistasis; figure 3.5 shows an example problem that exhibits almost 100%  $o(1)$ -epistasis, but in which higher order competitions are unbalanced, and therefore epistasis is lower for higher-order schemata. Equation 3.3 defines the relationship between genotypes and fitness; it was designed so that just about every  $o(1)$  competition was perfectly balanced.

$$f(\mathbf{x}) = \begin{cases} 2\text{INT}(\mathbf{x}) & \text{when } \text{INT}(\mathbf{x}) < 2^{l-1} \\ 2(2^l - \text{INT}(\mathbf{x})) & \text{otherwise} \end{cases} \quad (3.3)$$

| Individual       |                 | Gene 1 ( $x_1$ ) |           | Gene 2 ( $x_2$ ) |           | Gene 3 ( $x_3$ ) |           |
|------------------|-----------------|------------------|-----------|------------------|-----------|------------------|-----------|
| $\mathbf{x}$     | $f(\mathbf{x})$ | $x_1='0'$        | $x_1='1'$ | $x_2='0'$        | $x_2='1'$ | $x_3='0'$        | $x_3='1'$ |
| 000              | 1               | 1                | -         | 1                | -         | 1                | -         |
| 001              | 3               | 3                | -         | 3                | -         | -                | 3         |
| 010              | 5               | 5                | -         | -                | 5         | 5                | -         |
| 011              | 7               | 7                | -         | -                | 7         | -                | 7         |
| 100              | 8               | -                | 8         | 8                | -         | 8                | -         |
| 101              | 6               | -                | 6         | 6                | -         | -                | 6         |
| 110              | 4               | -                | 4         | -                | 4         | 4                | -         |
| 111              | 2               | -                | 2         | -                | 2         | -                | 2         |
| $f_\mu(x_i = a)$ |                 | 4.0              | 5.0       | 4.5              | 4.5       | 4.5              | 4.5       |
| $p_\mu(x_i = a)$ |                 | 0.44             | 0.56      | 0.5              | 0.5       | 0.5              | 0.5       |

Figure 3.2: A

3-bit problem with extremely high  $o(1)$ -epistasis; the mean fitness of almost every allele is equal to that of its competitor.  $f_\mu(x_i = a)$  is the mean fitness of allele  $a$  for gene  $x_i$ , and  $p_\mu(x_i = a)$  is ratio of allele of this mean fitness over the sum of all the mean fitness of all alleles for gene  $x_i$ . When fitness-proportionate selection is being used,  $p_\mu(x_i = a)$  would be the probability of selecting allele  $a$  for gene  $x_i$ . Note that for genes 2 and 3, the probability of selecting each allele is identical, and for gene 1, the allele selection probabilities are fairly equal. The allele selection probabilities for gene 1 would become increasingly similar as the number of genes is increased. Since the allele selection probabilities are almost all balanced, it is very hard to decide which alleles are the most responsible for increased fitness.

A measure is proposed here that can be used to directly estimate  $o(1)$ -epistasis. It measures how hard it is to immediately solve a problem by randomly sampling a small number of genotypes and then create a very high fitness individual by setting each allele the the average winner of each  $o(1)$  competition. In a 100%  $o(1)$ -epistatic problem, allele win frequencies would be balanced, and so the decision to use one allele or the other would be difficult. Since there may be hundreds of genes in genotype, it would not be adequate to simply count the win frequencies of the alleles of a single gene; the win frequencies of all the alleles of all genes must be recorded, and how much this distribution of frequencies differs from a completely balanced distribution (expected in a 100% epistatic problem) must somehow be measured. The chi-squared ( $\chi^2$ ) test (equation 3.4) can be used to measure this difference between the observed win frequencies and expected win frequencies. The chi-squared test is used to determine how well theoretical distributions fit empirical distributions. In this case, the theoretical distribution ( $e_i$ ) of 100% epistasis expects each allele to win 50% of its competitions. Therefore, if there are  $N$  competitions per gene, each allele should win  $N/2$  competitions (if a binary genotype is being examined). The empirical distribution ( $o_i$ ) is the actual number of times each allele wins an  $o(1)$  competition over the course of the  $N$  competitions. For binary genotypes, there are two alleles per gene, and therefore alleles can be numbered from 1 to  $(2l)$ .  $\chi^2$  is calculated by summing the frequency “errors”  $\frac{(o_i - e_i)^2}{e_i}$  for all  $(2l)$  alleles. The minimum possible value (0) for  $\chi^2$  is achieved when the observed distribution  $o_i$  of every allele  $i$  exactly matches its expected distribution  $e_i$ —when there is 100% epistasis. The maximum possible value ( $l \times N$ ) occurs when the allele competitions are completely unbalanced—either an allele wins all  $N$  competitions, or it wins zero competitions and its complementary allele wins all  $N$  competitions. When the allele competitions are completed unbalanced in this way,  $\chi^2$  is a sum of  $2l$  integers, of which  $l$  have the value  $N$  and  $l$  have the value zero. Therefore the maximum possible value of  $\chi^2$  is  $(l \times N)$ . Equation 3.5 is used to calculate a value for  $o(1)$   $\chi^2$ -epistasis  $E_{\chi^2}$  that ranges from zero to one. A value of zero would indicate 0%  $o(1)$   $\chi^2$ -epistasis, while a value of one would indicate 100%  $o(1)$   $\chi^2$ -epistasis. The extension of this method, to measure the dependence of higher order schemata on the settings of alleles outside those schemata, is straightforward but would require exponentially increasing amounts of computation with each increase in schemata order.

To calculate  $o(1)$   $\chi^2$ -epistasis, single allele competitions must occur between genotypes that differ by only one allele. An algorithm for calculating  $o(1)$   $\chi^2$ -epistasis is specified in figure 3.4.

Looking back at the DECEPT example, we see that, even though the problem is fairly difficult, there is actually very little epistasis in the problem since there is only one suboptimal attractor, and since most allele competitions are unbalanced (the ‘1’ allele most

$$\begin{aligned}
N &= \text{number of competitions performed} \\
o_i &= \text{observed number of wins for allele } i \\
e_i &= \text{expected number of wins for allele } i = N/2 \\
i &\in 1..(2l) \\
E_{\chi^2} &= o(1) \chi^2\text{-epistasis} \\
\chi^2 &= \sum_{i=1}^{2l} \frac{(o_i - e_i)^2}{e_i} \\
max_{\chi^2} &= l \times N \\
E_{\chi^2} &= 1 - \frac{\chi^2}{max_{\chi^2}}
\end{aligned} \tag{3.4}$$

$$E_{\chi^2} = 1 - \frac{\chi^2}{max_{\chi^2}} \tag{3.5}$$

Figure 3.3: Calculating  $o(1) \chi^2$ -Epistasis

```

DO WHILE (N < max - N)
1. Choose random genotype x
2. Evaluate f(x)
3. Choose random gene g in x and change its allele from i to i', yielding x'
4. Evaluate f(x')
5. If f(x) < f(x'), increment oi', else increment oi
6. N++
OD

```

Figure 3.4: Pseudo-Code for Calculating  $o(1) \chi^2$ -Epistasis

often beats the '0' allele). DECEPT is an example of a difficult problem with practically 0% epistasis, but high deception. It is therefore clear that a problem need not be highly epistatic in order to be difficult to optimize.

In summary, if a problem has no suboptimal attractors (whether mutation- or pattern-based), it is hill-climbable, because exploitation of observed higher-fitness schemata leads easily to global optima. If a problem has an extremely large number of suboptimal attractors, it is very difficult to determine when a schema with higher observed fitness is local to a suboptimal attractor, or if the feature is actually present in a global optimum.

Such a problem is highly epistatic. If a suboptimal attractor leads an algorithm to a solution that is significantly (genotypically) different from the global optima, then the attractor is deceptive. If this attractor is strong enough—if it has a large enough basin of attraction—then the problem is deceptive. Hill-climbability, deception and epistasis are all problem characteristics that arise from the genotypic topology of the search space. If this topology changes, then all three characteristics can radically change. All three of these characteristics are relatively independent of raw fitness values, but they are very sensitive to changes in individuals' fitness *ranking*. The next section will discuss “fitness distribution”, which is a problem characteristic that is completely dependent on raw fitness values.

### 3.6 Fitness Distribution

The distribution of a problem's fitness function can be characterized by many different measures: the range of fitness values, their mean and variation, their ratios, their concentration around specific values, etc.. Since most evolutionary algorithms directly use fitness values to calculate important search probabilities such as selection probabilities and mutation rates, a mismatch between a problem's fitness distribution and the probability-calculating functions used by the particular algorithm can cause dismal performance. For example, if an algorithm relies too much on absolute differences in individual's fitnesses, it can have difficulty solving problems that have a very small range of fitnesses. For instance, a GA using fitness proportionate selection would perform miserably on a problem in which fitness ranged from  $10^6$  to  $(10^6 + 1)$ , as the algorithm would allocate trials to the best and worst individuals with virtually the same frequency. The same algorithm can run into difficulty when the range of fitness is too wide. If fitnesses ranged from 0 to  $10^6$ , but most individuals had fitness less than 10, the first individual encountered with fitness  $10^2$  or  $10^4$  would “swamp” the population, as fitness proportionate selection would be likely to allocate all trials to this individual. This “bad match” between the algorithm's use of fitness values and the fitness range in the problem being attacked would quickly cause premature convergence.

Bad matches can also exist between a problem's fitness distribution and algorithms that do not use fitness-proportionate selection. For example, simulated annealing has an “acceptance probability” that determines whether lower fitness individuals will replace the current individual. This probability changes in relation to the difference in fitness between a new sample and the old sample. When the difference in fitness between two samples is very low, the probability of acceptance can be very high. If a search space is dominated by a dense concentration of low fitness points, the search may spend much of its time wandering among these points, without ever exploiting small increases in fitnesses necessary to climb out of the low fitness region. Almost every mutation will be accepted, no matter what its

effect on fitness (since the changes in fitness will be “small”). However, it might be these tiny changes in fitness are what is necessary to guide the search to the best regions of the space. An example of this phenomena will be demonstrated later in this thesis (see the discussion of F31 in the “Experiments” chapter).

Various methods (such as “fitness scaling” (Goldberg, 1989)) have been proposed that alleviate some of the problems associated with mismatches in fitness distributions and specific algorithms. These methods usually try to convert algorithms that rely on fitness-proportionate selection into hill-climbing algorithms that allocate trials on the basis of the fitness-rank-order of the different individuals or schemata in the population.

With respect to evolutionary optimization, the major characteristic of fitness distributions is their variance—in what ranges can the majority of fitnesses be found, and how concentrated are these fitnesses around particular values? Most importantly, from the standpoint of an optimization algorithm, is the fitness variation caused by a particular genotypic feature? How much effect does a particular allele, gene or schema have on fitness—how “significant” are these genotypic features? In the following, a possible method for the analysis of variance (ANOVA) of allele significance will be proposed.

### 3.6.1 ANOVA of Allele-Fitness

The more predictable an individual’s fitness is, given that it contains a particular allele, the easier it should be to construct individuals with desired levels of fitness. An individual’s fitness is “predictable” from the presence of a particular allele if the variation of fitnesses of individuals containing that allele (the allele’s fitness variation) is low. When the allele’s fitness variation increases, it becomes harder to decide whether or not that that allele has a large effect on fitness, as the fitness of individuals containing that allele become less consistent. When *all* the alleles have high fitness variances—when no single allele is consistently associated with a particular small range of fitness values, it becomes very hard to predict fitness from the entire genotype. One way to measure how predictable fitness is from complete genotypes is to sum all the allele fitness variances. I shall call this sum “within-allele fitness variance” ( $S_w^2$ ). A low  $S_w^2$  implies that many alleles’ fitness variances are low, and therefore predictability of fitness from genotypes is high. High  $S_w^2$  values imply that there are many alleles that appear in individuals of wildly different fitness; allele fitness variance is high overall. High  $S_w^2$  values therefore indicate that the genotype-fitness mapping is hard to predict from individual alleles.

How are “high” and “low” values of  $S_w^2$  defined? An allele’s fitness variation should be measured relative to the overall fitness variation of the problem. One measure of the overall variation in a problem is “between-allele fitness variance”  $S_b^2$ . It measures how much each gene’s allele mean fitnesses differ from each other. It is the sum of the square of

the differences between all pairs of allele mean fitnesses. This quantity is mathematically equivalent to the sum of squared differences between a population's mean fitness and all allele fitness means.

$$\begin{aligned}
i, j &\in 1..2l \\
f_\mu(i) &= \text{mean fitness of allele } i \\
f_\mu(.) &= \text{mean of all allele mean fitnesses}
\end{aligned}$$

$$\begin{aligned}
S_b^2 &= \sum_{i,j} (f_\mu(i) - f_\mu(j))^2 \\
&= \sum_i (f_\mu(i) - f_\mu(.))^2
\end{aligned}$$

When allele means are very different from the population mean, the alleles most highly correlated with high or low fitness become more evident—the *significance* of the alleles becomes more pronounced. For example, consider an  $l$ -bit genotype that simply encodes the binary integers from 0 to  $2^l - 1$ . The mean phenotypic value (fitness) of this problem would be  $f_\mu = \frac{2^l - 1}{2}$ . The most significant bit (gene) could have one of two allele values—'0' or '1'. Allele '0' would be found in individuals having mean fitness half that of the population at large, while allele '1' would have mean fitness one and half times that of the mean population fitness. Therefore both alleles for this gene would be highly significant, as they both would have a significant influence on the fitness of the individuals they were found in.  $S_b^2$  is the sum of differences in allele pairs, and so when many genes have significant alleles,  $S_b^2$  increases. Higher  $S_b^2$  implies that it is easier for an algorithm to identify alleles that are commonly associated with high fitness individuals.

By combining both between- and within-allele variance of a problem into a single ratio,  $\frac{S_b^2}{S_w^2}$ , we can measure the overall ease of predicting an individual's fitness, given that specific alleles are present in that individual, and therefore we have taken a large step in categorizing the difficulty of the optimization problem. The numerator,  $S_b^2$ , increases when variations in gene significance increase, resulting more predictability. The denominator,  $S_w^2$ , indicates how unpredictable fitness is from single alleles, so as predictability decreases,  $S_w^2$  will increase, and thus the ratio  $\frac{S_b^2}{S_w^2}$  will decrease. The implications of these observations will be discussed next.

### 3.6.2 Gene Significance and Fitness ANOVA

In most mutation-based algorithms, each gene has the same probability of mutation. However, in many problems, some genes have greater significance than others. Changes to these genes result in larger changes in fitness than changes to other genes. Most mutation-based algorithms do not exploit differences in significance. Pattern-based algorithms (and particularly GAs) indirectly exploit differences in significance because genes that are higher in significance lose diversity much faster than lower significance genes, and therefore become less and less likely to change settings as the search progresses. Problems with a wide variety of gene significances will have large  $S_b^2$  values, as many alleles will have mean fitnesses differing significantly from the population mean. Problems with high variation in gene significances might be considered to be easier than problems for which all gene significances are equal, since, for the first type of problems, an algorithm could quickly determine what the best settings were for the most significant genes, and thus make large progress initially in the search. Once this was accomplished, the algorithm could then concentrate on determining the best settings of the lower-significance genes.

Although phenotypic fitness may be highly predictable from single alleles, this does not mean that the problem is easily optimizable. Choices must be made among competing alleles. If complementary alleles have the same mean, at least one of their variances must be high, since the entire range of fitness values must be included in the mean calculations. It could be the case that one of the alleles had very low fitness variation, while the other allele was frequently found in both the best and worst individuals and therefore had very high fitness variation, but still had the same mean as the first individual. What would be the best allele to use when constructing a new individual? If it could somehow be determined *why* the second allele was present sometimes in high fitness individuals, and other times in low fitness individuals, then it might be possible to carefully construct a new individual that met all the conditions found in the high fitness individuals, and therefore the second allele could be used with confidence. An optimization algorithm—the eugenic algorithm—will be presented in the next chapter that accomplishes this goal through the use of an operator (the “restriction” operator) that causes *conditional* allele-fitness distributions to be measured and therefore determines what *combinations* of alleles lead to extremes in fitness.

## Chapter 4

# The Eugenic Algorithm

The Binary Simulated Crossover (BSC) (Syswerda, 1993) algorithm explicitly analyzes and recombines  $o(1)$  schemata to form new individuals. Population-Based Iterative Learning (PBIL) (Baluja, 1994) uses a probability vector as an explicit “fuzzy” representation of a “good”  $o(l)$  schema from which similar  $o(l)$  schemata are constructed. What is missing from both these algorithms is the explicit analysis and processing of intermediate order schemata— $o(2)$  to  $o(l - 1)$ . There are a great many of these schemata, and an explicit analysis of every one of them is infeasible. If a search algorithm is to perform explicit analyses of intermediate order schemata, it must limit its computational burden to only a small subset of these possibilities. The chosen subset should include only those schemata that are most relevant to reducing uncertainty and furthering the search.

A new algorithm that accomplishes these objectives is proposed in this chapter. It achieves *smart recombination*, in addition to *smart mutation*. This “Eugenic Algorithm” (EuA) consists of three major elements: (1) selection of promising alleles, based on an explicit analysis of allele fitness distributions, (2) restriction of the population to relevant individuals and (3) steady-state replacement. This chapter will begin with an illustrative description of example execution of the EuA, followed by discussions and justifications of how good alleles are selected, relevant individuals are identified, and why steady-state replacement is used.

### 4.1 Example of Eugenic Algorithm Execution

Since many of the concepts introduced in the EuA have never previously been integrated into the same algorithm, an illustrative example will be used to help explain the workings of the EuA. I shall list EuA pseudo-code before beginning the example.

```

L01: pop = random-pop() /* Randomly initialize the population: */
    /* N counts the number of individuals created by the EuA: */
    N = 0
    /* Keep on creating new individuals until N exceeds  $N_{max}$  */
L02: DO WHILE ( $N < N_{max}$ )
    u = {  $x_1, x_2, x_3, \dots, x_l$  }
    rpop = pop
L03:     DO WHILE nonempty(u)
        /* Select and set a gene: */
L04:          $x_g$  = most-significant-gene(rpop, u)
L05:          $x_{g,a}$  = select-allele( $x_g$ )
L06:         Remove  $x_g$  from u
        /* Restrict population when epistasis of unset genes is high: */
L07:          $E$  = epistasis-of-population(rpop, u)
L08:          $p$  = probability-of-restriction( $E$ )
L09:         if (uniform-random[0,1]  $\leq p$ )
L10:             rpop = restrict-population(rpop,  $x_{g,a}$ )
        OD
        /* Replace lowest-fitness individual with the new individual */
L11:     w = lowest-fitness-individual(pop)
L12:     popw = x
    N++
OD

```

Figure 4.1: Pseudo-Code for the Eugenic Algorithm

#### 4.1.1 Pseudo-Code for the Eugenic Algorithm

In figure 4.1, pseudo-code is shown that outlines the eugenic algorithm. The subroutines it relies upon will be described in the following sections.

The EuA begins by generating a random population of size  $\lambda$  in step L01. In the simplest case, the population size  $\lambda$  will remain constant during the entire execution of the algorithm. Once the random population is created, the algorithm repeatedly creates a new individual (steps L04 through 10) and replaces the lowest-fitness individual in the population with the new individual (steps L11 and L12) until the maximum number of individuals have been created. There are obviously many other possible termination criteria, such as stopping when an individual with particular level of fitness has been encountered.

The EuA constructs new individuals one gene at a time. Initially, all the genes have

no set value—they are “unset” with  $u$  the set of unset genes (see steps L02, L04, and L10). The EuA orders genes—from most significant to least significant. Once the most significant gene remaining in  $u$  is determined (step L04), an allele value is selected for that gene (step L05). Once a gene’s value is set, that gene is removed from  $u$  (step L06).

After setting the value of a gene, the EuA estimates the amount of epistasis  $E$  that the population exhibits among the remaining unset genes (step L07). The EuA uses  $E$  to determine  $p$ , the probability of restricting the population to only “relevant” individuals (step L08). If epistasis is high, then the necessity of restriction will be high, and therefore  $p$  will be high. In step 09, the algorithm probabilistically decides whether or not to restrict the population. If a random sample from a uniform variable in the range  $[0,1]$  has a lower value than  $p$ , the algorithm then temporarily removes from the population all those individuals that do not carry allele  $x_{g,a}$  (step L10). In effect, the population is “restricted” to only those genotypes that are relevant to the allele choices for the remaining unset genes, i.e. those individuals that contain the allele that was just chosen. These individuals are relevant since they contain information about how other alleles interact with allele  $x_{g,a}$ ; individuals that do not contain allele  $x_{g,a}$  have no information on allele  $x_{g,a}$ ’s epistatic interactions. At the start of the creation of each individual, all individuals are included in **rpop**, the restricted population. After each restriction, more individuals are removed from **rpop**. If a restriction would cause **rpop** to have fewer than **minpop** individuals, **rpop** is not restricted.

Once a population is restricted, all the statistics used in steps L03-L10 are calculated using only the restricted population. In the next iteration, step L04 determines the gene  $x_g$  that is most significant in the restricted population. Step L05 calculates allele selection probabilities for  $x_g$ , based on the average fitnesses of alleles found in the restricted population. Step L07 estimates the epistasis among the unset genes of the restricted population, and, by limiting the scope of these calculations of significance, allele selection probability, and epistasis to only those individuals in the restricted population, the EuA increases their relevancy to determining the best alleles to combine with the alleles already chosen.

A description of an example execution of the EuA will be given below.

### 4.1.2 Example Execution

In this section, I will give an example of the creation of a single individual from the population shown in figure 4.1.2. The EuA creates a single new individual in steps L03-L10.

Initially, the population contains 5 individuals, and the set of unset genes ( $u$ ) contains all the genes ( $u = \{x_1, x_2, x_3, x_4\}$ ).

The first step in creating a new individual is to determine the most significant unset gene (step L04). The most significant gene is the gene whose allele values seem to have the greatest effect on individuals’ fitness. A simple way to estimate gene significance is to use

| Individual       |                 | Gene 1 ( $x_1$ ) |           | Gene 2 ( $x_2$ ) |           | Gene 3 ( $x_3$ ) |           | Gene 4 ( $x_4$ ) |           |
|------------------|-----------------|------------------|-----------|------------------|-----------|------------------|-----------|------------------|-----------|
| $\mathbf{x}$     | $f(\mathbf{x})$ | $x_1='0'$        | $x_1='1'$ | $x_2='0'$        | $x_2='1'$ | $x_3='0'$        | $x_3='1'$ | $x_4='0'$        | $x_4='1'$ |
| 1100             | 25              | -                | 25        | -                | 25        | 25               | -         | 25               | -         |
| 1110             | 23              | -                | 23        | -                | 23        | -                | 23        | 23               | -         |
| 1101             | 22              | -                | 22        | -                | 22        | 22               | -         | -                | 22        |
| 0100             | 20              | 20               | -         | -                | 20        | 20               | -         | 20               | -         |
| 1111             | 2               | -                | 2         | -                | 2         | -                | 2         | -                | 2         |
| $f_\mu(x_{i,a})$ |                 | 20.0             | 18.0      | -                | 15.0      | 22.3             | 12.5      | 22.7             | 12.0      |
| $p_\mu(x_{i,a})$ |                 | 0.53             | 0.47      | 0.0              | 1.0       | 0.64             | 0.36      | 0.65             | 0.35      |

Figure 4.2: Complete Population

$\mathbf{x}$  : an individual's genotype

$x_i, i \in 1...l$  : the  $i$  th gene of the genotype

$a \in \{'0', '1'\}$  : the possible allele values

$x_{i,0}, x_{i,1}$  : the two possible alleles ('0' or '1') for gene  $i$

$f_\mu(x_{i,a})$  : mean fitness of all individuals containing allele  $a$  for gene  $i$

$p_\mu(x_{i,a})$  : probability of selecting allele  $a$  for gene  $i$

the absolute difference between its alleles' mean fitnesses. This quantity is  $|20.0 - 18.0| = 2.0$  for  $x_1$  (gene 1),  $|22.3 - 12.5| = 9.8$  for  $x_3$ , and  $|22.7 - 12.0| = 10.7$  for  $x_4$ . Gene 2 ( $x_2$ ) has no significance since it had no allele diversity and therefore none of its allele could have any observable effect on the fitness of individuals in the populations. Genes with no allele diversity are always assigned the lowest significance possible. As a result, the most significant gene in this population is  $x_4$ , followed by  $x_3$ ,  $x_1$ , and  $x_2$ .

Since alleles are assigned to the most significant genes first, an allele is chosen for  $x_4$  in step L05. An allele's selection probability is equal to its mean fitness divided by the sum of allele mean fitnesses for all the alleles associated with this particular gene. For gene  $x_4$ , the probability of choosing allele '0' ( $x_{4,0}$ ) is 0.65, since  $\frac{f_\mu(x_{4,0})}{f_\mu(x_{4,0}) + f_\mu(x_{4,1})} = \frac{22.7}{22.7 + 12.0} = 0.65$ . The corresponding probability of choosing allele  $x_{4,1}$  is 0.35. Say that allele  $x_{4,0}$  is chosen. In the individual being created, gene  $x_4$  is therefore assigned the allele value of '0'. This occurs in step L05.

Once a gene's allele value is set, the gene is removed from the set of unset genes ( $u$ ) in step L06. Therefore,  $x_4$  is removed from  $u$ , resulting in  $u = \{x_1, x_2, x_3\}$ .

When epistasis is high, the "correct" choice of alleles depends heavily on the alleles already present in an individual. Therefore the EuA should not always use allele selection statistics gathered over the entire population; it should use statistics gathered over individ-

uals that contain the same alleles as those that have been chosen for the new individual. In steps 07 through 10, the EuA decides whether or not to restrict the population to such individuals. Step 07 determines the probability of restriction by measuring the level of epistasis ( $E$ ) in the *unset* genes remaining in the current (possibly restricted) population. The epistasis in the unset genes is used, since it reflects the uncertainty in the remaining allele decisions to be made, depending upon the allele choices that have already been made. A simple way to calculate  $E$  is to use the maximum difference in allele selection probabilities of the remaining unset genes ( $D_{max}$ ). When  $D_{max}$  is very low, the allele selection probabilities are all about equal, indicating high uncertainty about what allele choices would be best—all allele choices seem to be equally good. When  $D_{max}$  is high, there are unset genes whose best allele choice is very evident. High uncertainty among allele choices implies heavy interdependence among alleles—high epistasis. The largest difference in allele selection probabilities in unset genes occurs in  $x_3$ ; this value is  $|0.64 - 0.36| = 0.28 = D_{max}$ . Since this difference ranges in the interval  $[0.0, 1.0]$ , an easy way to calculate  $E$  is to use  $E = 1.0 - D_{max} = 0.72$ . The calculation of  $D_{max}$  and  $E$  occurs in the call to the subroutine `epistasis-of-population()` in step L07.

A simple way of determining the probability of restriction ( $p$ ) is to directly use  $p = E$ , since  $E$  ranges in the interval  $[0.0, 1.0]$ , with  $E = 1.0$  indicating the highest epistasis possible and therefore the highest justification for restriction. Therefore the probability of restriction is 0.72. Say restriction is not performed in steps L09 and L10. The algorithm then loops back to step L03. Since  $x_3$  is the most significant unset gene, an allele is chosen for  $x_3$ —say allele '0'.  $x_3$  is therefore assigned the value of '0' in the new individual in step L05.

Again, the probability of restriction is determined in steps L07 and L08. This time,  $D_{max} = |0.53 - 0.47| = 0.06$  and  $E = (1.0 - 0.06) = 0.94 = p$ . Say restriction is performed this time in steps L09 and L10. Therefore, individuals that do not contain allele '0' for  $x_3$  are temporarily removed from the population, leaving individuals 1100, 1101, and 0100. These individuals have respective fitnesses 25, 23, and 20. Figure 4.1.2 shows the restricted population.

Note that statistics are not calculated for genes 3 and 4 at this point, as genes 3 and 4 have already been assigned values. Alleles for genes 1 and 2 of the new individual must still be chosen. The most significant gene is now  $x_1$ . Previously, the mean fitness of the alleles for gene 1 were 20.0 for  $x_{1,0}$  and 18.0 for  $x_{1,1}$ . The mean fitness of  $x_{1,0}$  does not change from 20.0 (since, by coincidence, all the individuals containing allele '0' for gene 1 remain in the population), but the mean fitness of  $x_{1,1}$  increases to 23.5, and now exceeds that of allele  $x_{1,0}$ . By restricting the population to only those individuals that contain  $x_{3,0}$ , the statistics calculated are now “relevant” to the remaining allele choices. Given that

| Individual       |                 | Gene 1 ( $x_1$ ) |           | Gene 2 ( $x_2$ ) |           | Gene 3 ( $x_3$ ) |           | Gene 4 ( $x_4$ ) |           |
|------------------|-----------------|------------------|-----------|------------------|-----------|------------------|-----------|------------------|-----------|
| $\mathbf{x}$     | $f(\mathbf{x})$ | $x_1='0'$        | $x_1='1'$ | $x_2='0'$        | $x_2='1'$ | $x_3='0'$        | $x_3='1'$ | $x_4='0'$        | $x_4='1'$ |
| 1100             | 25              | -                | 25        | -                | 25        | 25               | -         | 25               | -         |
| 1101             | 22              | -                | 22        | -                | 22        | 22               | -         | -                | 22        |
| 0100             | 20              | 20               | -         | -                | 20        | 20               | -         | 20               | -         |
| $f_\mu(x_{i,a})$ |                 | 20.0             | 23.5      | 0.0              | 15.0      | n/a              | n/a       | n/a              | n/a       |
| $p_\mu(x_{i,a})$ |                 | 0.46             | 0.54      | 0.0              | 1.0       | n/a              | n/a       | n/a              | n/a       |

Figure 4.3: Restricted Population,  $x_3 = '0'$

$x_3$  is assigned the value '0', the conditional fitness of  $x_{1,1}$  is greater than its complement,  $x_{1,0}$ , while globally  $x_{1,1}$ 's fitness is lower than that of  $x_{1,0}$ . Had the population not been restricted, this change in relative fitnesses would not have been noticed. By restricting the population to only relevant individuals, the EuA can track epistatic dependencies of various alleles.

Through restriction, higher-order schemata statistics are being collected. The global fitness of allele  $x_{0,1}$  is identical to the average observed fitness of schema '1###'. This average of this schema is 18.0, which is lower than that of its competitor '0###'. However, when higher-order schemata are examined, we find that the average observed fitness of schema '1#0#' is 23.5, higher than that of '0#0#'. Therefore the EuA's restriction operator allows the algorithm to calculate schema fitness averages "on demand", when these averages are necessary to making correct allele decisions.

The algorithm still must choose alleles for genes 1 and 2. Say  $x_{1,1}$  is chosen for  $x_1$ . Since there is no allele diversity in the remaining unset gene,  $x_2$ , there is no point to restricting the population, since restriction will not allow us to compare the conditional fitnesses of the alleles for  $x_2$ . The restricted population only contains information for allele '1'. At this point, the `select-allele` subroutine can either choose allele '1' with the hope that it is better than the '0' allele, since the '0' allele is not present, or the algorithm can experiment a little and choose the '0' allele. The EuA uses a parameter called "creation rate", which is the probability that an allele that is not present in the restricted population will be chosen ("created"). Without a "creation rate",  $x_2$  would always be set to  $x_{2,1}$ . With a creation rate, there is a chance that  $x_{2,0}$  would be chosen. Say that  $x_{2,0}$  is chosen, completing the construction of the new individual. The new individual has genotype '1000', and therefore is a completely new individual.

Once a new individual is created, the population is unrestricted—all individuals are added back into the population. Then the EuA replaces the lowest fitness individual in the population with the new individual. In this example, individual '1111' has the lowest fitness, and is therefore replaced by the individual '1000' (steps L11 and L12).

This example illustrated the behavior of a "standard" EuA. Many options of the

algorithm can be changed in order to increase performance. For example, the policy of always replacing the lowest-fitness individual in the current population ( $pop_w$ ) with the new individual may be slightly modified so that the new individual replaces  $pop_w$  only when the new individual has equal or greater fitness than  $pop_w$ . In the sections to follow, the motivation and mechanism behind each step of this algorithm are explained. Allele selection will be discussed first, then population restriction, and then finally replacement policy.

## 4.2 Allele Selection

In this section, I will discuss how the EuA performs step L05, `select-allele`. In the EuA, the selection of single alleles to use in the construction of new individuals is based upon the statistical analysis of the correlations between alleles and fitness<sup>1</sup>. New individuals are constructed one gene at a time, by choosing a single allele for each gene from several possible allelic forms. As with BSC, there are many choices for the selection technique used to choose new alleles, and certainly the performance of each technique is heavily, if not completely, problem dependent. Some examples of allele selection techniques are fitness-proportionate selection (equation 4.2) and rank selection. In fitness-proportionate selection, the probability of selecting an allele varies directly with the proportion of that allele's fitness to the sum of all of the allele fitnesses of that gene. In rank selection, an allele's selection probability for a given gene varies with its rank among the other alleles corresponding to that gene—the highest ranking allele has the greatest probability of being selected, followed by the second rank allele, etc.. Since most GA research has dealt with fitness-proportionate selection, this thesis will be limited to EuAs using fitness-proportionate selection.

$$\begin{aligned}
 f_{\mu}(x_{i,a}) &= \text{Mean fitness of all individuals having allele value } a \text{ for gene } i \\
 &= \text{The "fitness" of allele } x_{i,a} \\
 p_{\mu}(x_{i,a}) &= \text{Probability of selecting allele value } a \text{ for gene } i \\
 &= \frac{f_{\mu}(x_{i,a})}{\sum_{b \in \{0',1'\}} f_{\mu}(x_{i,b})}
 \end{aligned}$$

---

<sup>1</sup>Although developed independently, this strategy is basically identical to that of BSC.

### 4.2.1 Defining Allele Fitness: Average Fitness versus Maximum Fitness

Fitness-proportionate selection methods can rely on different definitions of allele fitness. Two ways of defining allele fitness will be discussed in this section: mean allele fitness ( $\mu$ -fitness) and maximum allele fitness ( $X_n$ -fitness). In  $\mu$ -fitness, an allele's fitness is defined as the *average* fitness of individuals that contain that allele. When using fitness-proportionate selection and  $\mu$ -fitness, the EuA's allele selection mechanism is identical to the technique described previously in relation to BSC. In  $X_n$ -fitness, the *maximum* fitness of all the individuals containing a particular allele defines that allele's fitness. Figure 4.4 illustrates how the allele selection probabilities of  $\mu$ -fitness and  $X_n$ -fitness differ for an example 4-bit genotype. Note that, for this example, the  $X_n$ -fitness selection probabilities for genes 3 and 4 are much closer to 0.5 than those of  $\mu$ -fitness, and therefore  $X_n$ -fitness would select alleles for these genes more randomly than  $\mu$ -fitness. A different example could easily be generated in which these effects were reversed.

One of the main reasons  $\mu$ -fitness has been so widely used in the past is because it has been shown to be the best strategy for maximizing the cumulative “payoff” of repeated trials (Holland, 1975)— $\mu$ -fitness is the best choice, on average, to select “good” alleles. Since the goal in this thesis is to design an algorithm to optimize a fitness function, instead of finding “good” solutions on average,  $X_n$ -fitness may be a better choice than  $\mu$ -fitness.  $X_n$ -fitness focuses the search only on the best performance of each allele, whereas  $\mu$ -fitness incorporates information about all aspects of a single allele's performance, including its worst performance, its typical performance, and its best performance.  $X_n$ -fitness may lead to more “risky” choices, as these choices are not tempered by all the information available. In addition,  $X_n$ -fitness may be more likely to cause the EuA to prematurely converge than  $\mu$ -fitness, because the alleles belonging to a single super-individual could completely dominate every allele selection choice, and therefore every new individual created would be identical to the super-individual. More theoretical and experiment exploration of “what allele fitness definition is best” is necessary, but for the meantime, the standard EuA will use  $\mu$ -fitness.

After a gene in the new individual is fixed to a specific allele value, the next most significant gene is examined and each of its allele's probabilities of selection are calculated. Since the domain of the current implementation of the EuA is limited to binary genotypes, it is only required to calculate the selection probability  $p$  for one of the alleles from the most significant gene, as the probability of selecting the complementary allele would simply be  $(1 - p)$ .

$\mathbf{x}$  : an individual's genotype  
 $f(\mathbf{x})$  : the fitness of individual with genotype  $\mathbf{x}$   
 $x_i, i \in 1 \dots l$  : the  $i$ th gene of the genotype  
 $a \in \{ '0', '1' \}$  : the possible allele values  
 $f_\mu(x_{i,a})$  : mean fitness of all individuals containing allele  $a$  for gene  $i$   
 $p_\mu(x_{i,a})$  : using  $\mu$ -selection, probability of selecting allele  $a$  for gene  $i$   
 $f_{X_n}(x_{i,a})$  : maximum fitness of all individuals containing allele  $a$  for gene  $i$   
 $p_{X_n}(x_{i,a})$  : using  $X_n$ -selection, probability of selecting allele  $a$  for gene  $i$

$$p_\mu(x_{i,a}) = \frac{f_\mu(x_{i,a})}{\sum_{b \in \{ '0', '1' \}} f_\mu(x_{i,b})}$$

$$p_{X_n}(x_{i,a}) = \frac{f_{X_n}(x_{i,a})}{\sum_{b \in \{ '0', '1' \}} f_{X_n}(x_{i,b})}$$

| Individual         |                 | Gene 1 ( $x_1$ ) |           | Gene 2 ( $x_2$ ) |           | Gene 3 ( $x_3$ ) |           | Gene 4 ( $x_4$ ) |           |
|--------------------|-----------------|------------------|-----------|------------------|-----------|------------------|-----------|------------------|-----------|
| $\mathbf{x}$       | $f(\mathbf{x})$ | $x_1='0'$        | $x_1='1'$ | $x_2='0'$        | $x_2='1'$ | $x_3='0'$        | $x_3='1'$ | $x_4='0'$        | $x_4='1'$ |
| 1100               | 25              | -                | 25        | -                | 25        | 25               | -         | 25               | -         |
| 1110               | 23              | -                | 23        | -                | 23        | -                | 23        | 23               | -         |
| 1101               | 22              | -                | 22        | -                | 22        | 22               | -         | -                | 22        |
| 0100               | 20              | 20               | -         | -                | 20        | 20               | -         | 20               | -         |
| 1111               | 2               | -                | 2         | -                | 2         | -                | 2         | -                | 2         |
| $f_\mu(x_{i,a})$   |                 | 20.0             | 18.0      | 0.0              | 15.0      | 22.3             | 12.5      | 22.7             | 12.0      |
| $p_\mu(x_{i,a})$   |                 | 0.53             | 0.47      | 0.0              | 1.0       | 0.64             | 0.36      | 0.65             | 0.35      |
| $f_{X_n}(x_{i,a})$ |                 | 20.0             | 25.0      | 0.0              | 25.0      | 25.0             | 23.0      | 25.0             | 22.0      |
| $p_{X_n}(x_{i,a})$ |                 | 0.44             | 0.56      | 0.0              | 1.0       | 0.52             | 0.48      | 0.53             | 0.47      |

Figure 4.4: How the allele selection probabilities of  $\mu$ -selection and  $X_n$ -selection differ for an example 4-bit genotype.

### 4.2.2 Adaptive Allele Selection

To balance exploration and exploitation, it would be desirable for an evolutionary search algorithm to perform a “smart mutation” when a clear winner among a gene’s alleles is apparent, but random mutation/search when the allele fitnesses are roughly equal. When one of a gene’s allele has a very high fitness, and the alternate alleles have very low fitness, the high-fitness allele should be exploited. When the fitnesses of the gene’s alleles are roughly equal, the choice between them is (for the current population) probably not that important, and so the choice should be more random. The EuA accomplishes these things automatically, when using fitness-proportionate selection. With fitness-proportionate selection, the amount of allele exploitation varies proportionately with the difference between a gene’s alleles’ fitnesses. Alleles with equal fitnesses have equal probabilities of being selected; for binary alleles, each has a 50% selection probability. In such a case, each allele is equally exploited and therefore the algorithm will equally explore the alternatives. However, an allele

with  $x$  times the fitness of its alternates has  $x$  times the selection probability, and therefore experiences  $x$  times more exploitation. The EuA explores when uncertainty among allele choices is high, and exploits when clearly superior alleles exist. In this way, adaptive allele selection is achieved, as the EuA adapts its selection probabilities to the differing situations. This is unlike most mutation-based algorithms, which always randomly choose a gene to mutate, no matter what information has been gathered in the past.

Note that, by relying on an explicit calculation of allele fitness, the EuA differs significantly from the GA. The GA implicitly records allele fitness by relying on an allele's frequency of occurrence in its population. For example, if all the individuals in a GA population had the same fitness, and if allele A occurred in 99 out of 100 individuals, while allele B occurred in only a single individual, allele A would be *99 times more likely* to be used than allele B, even though the average fitness and maximum fitness of both alleles were equal! As a result, by relying on explicit calculations of allele fitness, the EuA more reasonably allocates trails to and pursues promising alleles.

### 4.2.3 Modifications to Fitness-Proportionate Selection

As a result of fitness-proportionate selection, alleles with relatively low fitness will infrequently be chosen for inclusion in a new individual. Alleles that repeatedly cause low fitnesses in individuals containing them will eventually be eliminated from the population. However, it is not always desirable to have a particular allele completely eliminated from consideration. Consider two alleles, A and B. Allele A consistently has higher fitness than allele B, which consistently has very low fitness. It may be the case that allele B is needed to construct an optimal individual, even though it is detrimental to suboptimal individuals. In this case, allele A would be “deceptive”, as it would lead the algorithm away from the allele that was found in the optimum. In order to avoid convergence to deceptive alleles, the EuA must occasionally attempt to use alleles that have been eliminated from the population. The EuA uses a parameter called “creation rate” in order to reintroduce such alleles. This parameter ( $C$ ) specifies the probability that an allele lost from the population will be reintroduced by including it in a new individual being created. In addition, creation rate allows alleles absent from *restricted* populations to be included in a new individual. If, during the course of creating a new individual, the population is restricted to such an extent that extremely little allele diversity exists—all the individuals in the restricted population are extremely similar—the creation rate parameter allows the EuA to occasionally try out an allele that is not found in these individuals. By doing so, the EuA is afforded a small measure of protection against deceptive and/or premature convergence.

In cases where the chosen selection mechanism is too exploitative, another parameter, “selection noise” ( $N$ ), can be used by the EuA to encourage exploration. Selection noise is

the probability that the allele choice made by the allele selection operator is overridden, and a different allele is chosen. This parameter just adds noise to the allele choices made by the particular allele selection mechanism in place. It can be increased to encourage exploration. We shall see that this is rarely necessary, and otherwise almost always detrimental, when using fitness-proportionate selection.

Even with adaptive allele selection probabilities and a correctly set creation rate, the EuA still faces a possibly large obstacle to performance when dealing with large genotypes. For many problems, even one mutation can cause large fluctuations in fitness. Therefore the *number* of low-fitness alleles present in an individual may be much more important than the *percentage* of low-fitness alleles. As the number of genes of a genotype increases, the number of low-fitness alleles that are included in new individuals increases (assuming that selection probabilities remain constant with respect to genotype size). For example, if, for each gene, there was a 10% chance of choosing a low-fitness allele instead of a high-fitness allele, the number of low-fitness alleles appearing in a new 100-gene individual would be ten times on average the number that would appear in a 10-gene individual, even though the percentage of low-fitness allele would be equal between the two genotypes. Although it may seem reasonable for exploration of low-fitness alleles to remain proportionally constant with respect to problem size, it could easily be the case that this effect is detrimental to the EuA's performance on a large number of problems. Therefore, as genotypes increase in size, the balance of exploration and exploitation may have to be tipped in favor of exploitation. In the standard EuA, no effort is made to counter the problems that may be associated with larger genotypes. However, if desired, allele selection probabilities could easily be modified (using fitness scaling, perhaps) from problem to problem, so as to increase exploration versus exploitation of higher fitness alleles. Selection probabilities could be based upon one-tailed statistical tests of significance, or could be adjusted based upon the observed variance in a particular allele's fitness (H-L. Fang, 1993). Such modifications will not be attempted in this thesis, as no experiments have been conducted to determine their necessity.

#### 4.2.4 Clustering and Allele Significance

The EuA needs to determine gene significance for two reasons. First, the EuA sets genes in order of their significance, so that more important allele decisions are made before less important choices. Second, the EuA uses gene significance as a measure of epistasis. It then uses this measure of epistasis to determine when to restrict the population (subsection 4.3.2). Several different measures of gene significance are available to the EuA, and the choice of this measure will heavily affect the EuA's behavior.

One way to quantify gene significance would be to observe how well a gene's alleles "cluster" individuals into separate groups. An individual containing a particular allele would

belong to that allele’s group. As the fitness variance of each group decreased and their fitness means became increasingly different, the clustering would improve. As the means diverged, one group would become the “high-fitness group” and the other the “low-fitness group” for that gene. When the difference between the means was high, then allele settings for that gene would obviously have a large effect on fitness, and therefore the gene would be very significant. The standard EuA defines gene significance as the distance between a gene’s alleles’ means. Future versions could easily incorporate the variance of each allele group into the definition of gene significance.

### 4.3 Population Restriction

Smart recombination requires not only the identification and comparison of the best genotypic parts, but also the determination of how to best combine these parts to form very high-fitness individuals. Independently setting each allele value is not good enough, as this strategy ignores the epistasis so frequently encountered in nontrivial problems. Such a strategy is equivalent to randomly recombining  $o(1)$  schemata, without respect for the high levels of dependence that may exist among the best settings of alleles—such a strategy ignores higher-order interactions. Simple allele fitness averages do not provide enough information to make the right decisions. What is needed is an analysis of *conditional* allele fitness distributions. In one of these distributions, the fitness distribution of a single allele is analyzed, given that other alleles have fixed values. Since the number of possible conditional allele fitness distributions is *more* than even the number of possible schemata, not all conditional distributions can be examined. Therefore only the most important distributions should be examined. The EuA uses the *population restriction* operator to accomplish this goal and to separate global allele fitness distributions into conditional distributions.

#### 4.3.1 The Restriction Operator

The EuA intelligently selects sets of alleles from individuals in the population, but, unlike the GA-crossover operator, these genotypic parts can come from the entire population, instead of from just two parents. Due to epistasis, it is not reasonable to simply mate the entire population, as BSC does, because doing so would ignore allele dependencies found in each individual. Therefore, any population-wide recombination operator must determine which genotypic parts work well together, and selectively choose individuals that contain the most information about these parts.

As the EuA builds a new individual, the choice of what genotypic parts will be most appropriate depends on both gene significance and the already constructed parts of the new individual. The eugenic algorithm builds new individuals from scratch, by proceeding

through the genes, from the most significant to the least significant, and probabilistically selecting an allele value for each. At some point, gene significance decreases to such a point that the best choice between two complementary alleles becomes highly uncertain, because the alleles have very similar fitnesses. The EuA can reduce the uncertainty in the choice of the next allele by focusing on those individuals most relevant to this choice. It accomplishes this by restricting the population to those individuals that contain information about the relationship between the unset genes and the alleles already fixed in the new individual. These individuals are those that contain the same alleles as those fixed in the new individual. They contain information on how to best fine tune the remaining alleles, since they belong to the same region of the search space as the new individual. The *restrict-schema*, composed of the allele fixed in the new individual, is used to identify these individuals. Individuals that contain the restrict-schema make up the restricted population.

Once the most relevant individuals are identified, allele fitness statistics are recomputed using only the individuals from this restricted subset. These statistics are no longer global measures, they are *conditional fitness statistics*. They contain information on the performance of the unset alleles in the region of space the restrict-schema defines; they reflect the conditional allele fitness distributions of the “don’t care” genes in the restrict-schema.

The use of all of the restrict-schema’s fixed positions to limit the population poses a major problem, however; the restricted population can quickly become too small. For diverse populations, in which alleles ‘0’ and ‘1’ occur with about the same frequency for each gene, each application of the restriction operator will reduce the population size by approximately half, since only about half the individuals will have the allele restricted on. And so a restrict schema of order- $n$  will limit the population to only  $\lambda/2^n$  individuals. For even a single individual in the population to match the restrict-schema, either  $\lambda$  must be very large, or the restrict-schema must not contain too many fixed positions. Since increasing population size increases the computational demands of the EuA, the latter choice is preferable. One way to limit the number of fixed positions in the restrict-schema is to include only those alleles that were set immediately preceding a restriction, instead of all the alleles fixed in the new individual. Genes are fixed in order of significance, from most significant to least significant. Therefore the last gene to be fixed before a restriction occurs is the least significant gene of the restrict-schema. The lower the significance of a gene, the more likely good values for it depend on other alleles. Consequently, it is more likely that less significant genes take part in heavily epistatic allele interactions. More significant genes are more independent from other genes; they need less help from other allele settings. As a result, if all the genes of the restrict-schema cannot be used to limit the population, the least significant genes should be the ones used, since the values of the remaining unset genes should be specifically tuned for the alleles that need the most help. This is the strategy the

current implementation of the EuA employs.

### 4.3.2 The Probability of Restriction

The restriction operator cannot be applied after every gene that is fixed, since each restriction reduces the population size by roughly half. In addition, while highly significant genes remain, restriction is not necessary. Restriction is necessary when gene significance is low. When gene significance is low, epistasis is high. In order to determine the necessity of restriction, the EuA needs to estimate the level of epistasis among the unset genes in the current restricted population. When epistasis is high, it becomes very important to calculate conditional fitness distributions, since high epistasis implies that global distributions are not that useful. In order to calculate these conditional distributions in situations of high epistasis, the population must be restricted.

Various methods are available to measure epistasis and gene significance. One measure that could be used to determine restriction necessity is  $\chi^2$ -epistasis (see section 3.5). Alternatively, if allele fitnesses are believed to be drawn from approximately normal distributions, the  $t$ -score could be used to test the significance of the differences in mean allele fitnesses. However, the epistasis most relevant to determining restriction necessity is that which prevents a clear choice from being made among competing alternatives. Therefore probably the most appropriate measure of epistasis would be the difference in allele selection probabilities ( $D$ ), since the selection probabilities directly indicate how much uncertainty is present in the remaining allele choices. For example, when  $D$  is very high, it is obvious that there are very clear winners among the remaining unused alleles. This implies that there is a clear dominance of particular genotypic patterns in the remaining unfixed regions of the search space; these patterns have relatively high conditional fitnesses. If, however, the difference in selection probabilities was very low, then the information contained in the population is too contradictory to be useful to the selection method being employed. If it has already been decided that this selection mechanism will be useful for successfully solving the problem, then when this mechanism falters, some action (like restricting the population) must be taken to reduce the interference among competing possibilities.

In the experiments conducted in this thesis, the EuA used this difference in selection probability of a gene's alleles in its measure of epistasis, with epistasis ( $E$ ) equal to  $(1 - D_{max})$ , where  $D_{max}$  was the largest difference in allele selection probabilities for any gene. Since  $E$  varied in the interval  $[0, 1]$ , with  $E = 0$  implying no epistasis and therefore no need to restrict, and  $E = 1$  implying extremely high epistasis and a great need to restrict,  $E$  could possibly be used directly as the probability of restriction. However, it was observed that this measure of  $E$  was often not very close to zero, and so restriction would occur very often. As previously discussed, too much restriction can quickly lead to populations that

are too small. Therefore, modifications were made to decrease the probability of restriction. First, a basic probability of restriction was calculated—this value was just  $E$ . Then, this value was squared, therefore mapping  $E = 0$  to  $p_{restrict} = 0$ ,  $E = 0.5$  to  $p_{restrict} = 0.25$ , and  $E = 1$  to  $p_{restrict} = 1$ . In addition, the probability of restriction was explicitly reduced by half after each restriction to prevent the overuse of restriction on the most significant bits. Therefore, as the population was reduced in size by restriction,  $p$  was also reduced.

### 4.3.3 Consequences of Population Restriction

#### Incremental Schema Evaluation

What is the order of the schemata explicitly analyzed by the EuA? After each restriction, higher and higher order schemata are analyzed by the EuA. Initially, before any restrictions have occurred, only  $o(1)$  schemata are compared. After the first restriction, at least one allele becomes fixed in the entire population. Each allele's fitness distribution is now measured conditionally on the value of the fixed allele. Each schemata analyzed contains the fixed allele. Therefore, after one restriction, schemata of at least order-2 are analyzed. In general, after  $r$  restrictions, the number of genes for which only a single allele is present in the population is *at least*  $r$ , because each time the population is restricted, there can be a corresponding loss of diversity in alleles *other than* the one restricted on. For example, after 6 restrictions, it could easily be possible that 10 or 15 genes exhibit no allele diversity. Therefore the schema being analyzed would have at least 10 to 15 fixed positions.

There is, however, an upper limit on the number of restrictions that occur when generating a new individual. Since, on average, each restriction reduces the population size ( $\lambda$ ) by half, only  $\lambda/2^r$  individuals will be in the population after  $r$  restrictions. This means that, for a population size of  $2^r$ , it is highly likely that at least order- $r$  conditional allele fitnesses will be explicitly analyzed and compared by the EuA, and very likely that conditional fitnesses of order even more than  $r$  will be analyzed as the search progresses and focuses on smaller regions of the search space.

#### Maximization of *Relevant* Information

Although the *relevancy* of the information contained in the population is increased by restriction, the overall *amount* of the information is decreased, since the number individuals over which information is collected decreases with each restriction. How can the tradeoff between information amount and relevancy be managed?

The EuA maintains the amount of relevant information at the highest possible level by only restricting the population when necessary. If it is clear which alleles are associated with higher fitness, then restriction is not necessary. In these cases, it is most likely that

allele *combinations* are not as important as the presence or absence of single alleles. The entire population would provide better, more statistically significant information about these alleles than a small subpopulation. However, in cases of higher epistasis, the conditional allele fitness distributions become increasingly different from the marginal (global) distributions. Larger samples of individuals are not necessarily better samples in these situations. The most relevant information is found in subpopulations, which can be used to more accurately determine conditional distributions than a large diverse population could. Larger populations would contain less relevant information. Since the probability of restriction is based upon the perceived level of epistasis, the EuA adaptively maximizes the amount of relevant information in the population.

### **Diversity, Schemata Competition, and Subpopulations**

Because of the restriction operator, the EuA can and will maintain extremely high levels of diversity in the population, without bogging down the pace of the search with many ineffective schemata recombinations. The restriction operator ensures successful matings among individuals in the population. When epistasis is low, alleles can be drawn from the entire population without ill effect. However, when epistasis is high, if the genotypes of a group of mating individuals are too different, their offspring will most likely have low fitness. The restriction operator attempts to guarantee the compatibility of prospective mates. When high epistasis is detected, the EuA restricts the population to a group of similar individuals, and therefore reduces the probability of incorrect schemata recombination and low-fitness offspring. Because of the randomness of the application of restriction, the EuA causes many different subpopulations to intermate. Unlike many GA-related methods, reproductive compatibility and the subgrouping of individuals is not limited to a single, simple overall measure such as Hamming distance. The randomness of the application of the restriction operator causes many different population subgroups to be matched for reproduction, as compatibility is determined by different genes as each new individual is constructed. As a result, not only does the restriction operator *allow* large amounts of diversity to exist in a single population, it *causes* diversity to be maintained. It is a “multicultural dating service”.

Diversity is not always desirable. If it is fairly easy to generate good individuals with very diverse genotypes, then subpopulations of individuals can and will easily coexist in a single large population, and diversity is helpful to generate more representative statistics. In this case, the EuA can simulate the action of a “parallel GA”. A parallel GA is a set of GAs that occasionally share individuals between their populations and, by doing so, can collectively maintain search focus in many more regions of the space than a single-population GA. However, there are situations in which homogeneity is preferable to diversity. In search

spaces where the number of good individuals is extremely low, it will be difficult to generate even a single good individual and therefore practically impossible to generate a diverse set of good individuals. If the fitness of the single good individual exceeds the population's average fitness by a great enough amount, the EuA will experience a large drop in diversity of bits with very high significance with the occurrence of this first good individual, as the algorithm will have no choice but to exploit the only information it has that can lead to higher fitness. This decrease in diversity is not necessarily detrimental, as the drop in diversity of some genes allows the algorithm to much more thoroughly explore smaller, more promising regions of the search space. This is exactly what is needed when high-fitness individuals are very difficult to locate.

The EuA's smart recombination allows it to maintain much higher levels of diversity in a single population than a GA could maintain. In GAs, schema fitness is implicitly recorded by how often it is found in the population—fitness proportionate representation. Fitness is not explicitly defined using averages, maximum values, or any statistics whatsoever. It is therefore difficult for two or more competing schemata to exist in the population for long, because of the positive feedback between representation and selection. As soon as one schemata achieves higher representation than its competitors, it becomes more likely to be included in reproduction than its competitors, and its representation increases even more. Schemata whose representation fall below that of the others become less likely to be propagated, and their representation falls even more until they vanish from the population. In contrast, competition among schemata in the EuA is not won by higher amounts of representation, but by fitness itself. As a result, a high-fitness schema can exist in just one or two individuals in a large population, but still have a high probability of being used in the construction of new individuals. For example, consider a situation in which 99 individuals of a population contained schema A, and one individual contained a competing schema B. If the mean fitness of the schema A individuals was the same as the fitness of the schema B individual, the probability of the occurrence of each schema in new individuals would be equal, even though only one individual contained schema B. In a fitness-proportionate GA, schema A would be 99 times more likely to occur in offspring than schema B, since the sum of fitnesses of all schema A individuals would be 99 times that of the schema B individual. As a result of the EuA's managed competition of schemata, and the restriction operator's compatibility checking, many diverse groups of individuals can and will be maintained in the population during the search. The experiments to follow will demonstrate this empirically, illustrating that very genotypical different individuals coexist peacefully, yet competitively, in EuA populations.

It should be pointed out that the restriction operator performs exactly the opposite of what has been recommended by previous researchers—"incest prevention" (Eshelman

and Schaffer, 1991). Since the problem of premature convergence is very significant in the performance of GAs, some method such as incest prevention must be used to maintain high levels of diversity, without introducing random mutation. When incest prevention is used to prevent diversity loss, individuals mate only when they are significantly *different*—when the Hamming distance between them is large. It has been demonstrated that this is a very good heuristic for the GA, since its highly effective prevention of premature convergence overshadows the increase in incompatible matings. It allows the GA’s search to increase exploration without the use of higher mutation rates. Since the amount of exploration the EuA performs is adaptive, and diversity is promoted only when necessary, this artificial method of preventing diversity loss is not needed. Furthermore, it would most likely be detrimental to the performance of the EuA, because it would increase the chance of incompatible schemata being paired together.

## 4.4 Steady-State Replacement

Instead of generating an entirely new population from the existing population (“Generational Replacement”), the EuA uses “Steady-State Replacement” (Eshelman, 1991; Syswerda, 1991), in which only one individual from the existing population is replaced at a time. Since feedback can be gained from the evaluation of only one new individual, Steady State Replacement is a superior strategy to Generational Replacement because new information can be immediately exploited for use in new individuals. In addition, the action of the EuA is very genotypically “disruptive”, as it frequently generates new individuals that have very high Hamming distances from those in the current population. Steady State Replacement guarantees that good schemata found in the current population will be preserved as the search proceeds, while at the same time allows the EuA search to be very explorative. If Generational Replacement were used, then valuable information would be much more likely to be lost from generation to generation, and therefore the search would have to be much more cautious and conservative.

The choice of which individual to replace is yet another heuristic to determine for the EuA. The current implementation of the EuA always replaces the lowest-fitness individual in the population with the newly created individual. “Elitism” can be employed to prevent the possibility of a lower-fitness new individual from replacing a higher-fitness individual already in the population. Yet another variation of the replacement strategy is to probabilistically replace individuals based on their fitness. However, such a strategy would increase the chance that high-fitness schemata would be lost, and, given the high amount of exploration performed by the EuA (as will be demonstrated in the experiments), the current strategy of always replacing the worst is likely a better heuristic for most stationary problems.

## 4.5 Summary of the Eugenic Algorithm

The EuA is composed of two main parts: individual creation and individual replacement. Individual creation is composed of smart allele selection and population restriction. In the standard EuA, fitness-proportionate selection is used to choose alleles, where an allele's fitness is defined as the mean fitness of all the individuals that contain that allele. Mean fitness selection is a superior strategy to representation-proportionate fitness selection, since only a single instance of a high-fitness allele is necessary to guarantee its exploitation. In addition, it allows much greater levels of schemata diversity to exist. Population restriction is a major innovation of the EuA that is not present in any previous algorithm. Population restriction allows the EuA to purposefully (not randomly) juxtapose compatible schemata and also maintain many diverse subpopulations inside of a single large population. These subpopulations allow the EuA to focus its search in many differing regions. By using population restriction, the EuA can calculate only those conditional allele fitness distributions that are necessary. These calculations lead to the computation of increasingly higher order schema fitnesses, and only the most important of these schemata are analyzed. Population restriction allows epistasis to be intelligently tracked, analyzed and exploited. In effect, population restriction allows the EuA to perform smart recombination. The EuA uses Steady-State Replacement to prevent the loss of valuable genetic information—information on high-fitness alleles and schemata. In Steady-State Replacement, a newly created individual replaces the lowest-fitness individual in the EuA population.

In the chapter to follow, EuA performance will be compared against that of several other optimization algorithms on six different combinatorial optimization problems. We shall see that the EuA performs extremely well in relation to the other algorithms.

## Chapter 5

# Experiments

In this chapter, six combinatorial optimization problems will be attacked by several evolutionary algorithms, including the EuA. Five of these problems are drawn from the genesys-2.0 test suite<sup>1</sup>, and the sixth has been studied in simulated annealing literature comparing the performance of simulated annealing to that of genetic algorithms (Ingber, 1993). These problems are considered to be very hard, as even customized methods (such as direction set, branch and bound, etc.) have great trouble finding the global optima. At least four are NP-complete. Trivial toy problems, such as BIT-COUNT (described in section 3.3), will not be studied, as they would provide little feedback on how these algorithms would perform in real-world applications. A detailed description of each of the problems will be provided in section 5.3.

This thesis focuses on combinatorial optimization, and so all of the independent variables of the problems were encoded using binary genotypes, even though some of these problems might be more easily solved using a straightforward real-number representation (F02 and F38 in particular). It is extremely important to realize that the entire topological structure of the problem search space changes with such a transformation from  $m$ -dimensional real space to  $n$ -dimensional binary space. For example, there are a great many points that are mutation-based neighbors in the real space, but are separated by great (Hamming) distance in the binary space. This phenomenon is known as “Hamming cliffs”. For instance, consider 4-bit binary genotypes that encoded equidistant real numbers in the range  $[0, 15]$ . The genotypes ‘0111’ and ‘1000’ would encode the real genotypes 7.0 and 8.0, respectively. While their distance in binary space is the greatest possible for a 4-bit

---

<sup>1</sup>The genesys-2.0 software was provided in pre-release form by Joerg Heitkoetter. In the genesys-2.0 test suite, there are thirty-seven problems, designated F01 through F37. To aid researchers in the future, I will use genesys-2.0’s nomenclature. As a result of this nomenclature, the one test problem studied in this thesis that was not included in the genesys-2.0 package shall be designated “F38”.

binary genotype, their genotypic distance in real space is fairly small. Therefore genotypes that are neighbors in real space do not necessarily maintain this relationship when they are encoded into another representation. As a result, the difficulty of an optimization problem can be greatly increased or decreased when the representation of the inputs change. In particular, local optima present in real space do not necessarily retain their same neighborhoods of attraction in binary space, and therefore they may no longer be locally optimal, since they are no longer superior to their immediate neighbors. Even a transformation from one binary representation to another that maintains local neighborhoods can cause local optima to appear or disappear, since the bit patterns that matched specific individuals in the original space may no longer match the individuals in the new space. In problem F02, we shall see that the change from real to binary space greatly increases the number of local optima (and therefore problem difficulty). In contrast, problem F38, a problem with a great many strong local optima in real space, seems to lose this characteristic upon transformation of its inputs into a binary encoding. Almost all the algorithms tested were never trapped by F38’s real-space local optima. Furthermore, to make a fair comparison, all the algorithms searched using the same genotype-phenotype mapping, since changing this mapping could easily alter a problem’s epistatic and local optima properties. By comparing each algorithm’s performance on a number of differing problems, it is hoped that the special bias of each algorithm can be identified and eliminated from the comparison. In order to make this comparison, the algorithms must be tested on the same problems using the same representation.

Since the ideas of epistasis and problem difficulty are so closely linked, one algorithm (IMHC) was specifically designed to be the “silver bullet” for zero epistasis problems (specifically  $0.0 \text{ } o(1) \text{ } \chi^2$ -epistasis). The relative speed at which IMHC solves a problem, compared to the other algorithms used, will be considered a measure of problem triviality (and therefore problem insignificance). As we shall see, IMHC easily optimized one of the problems (F38), and very quickly solved a second problem (F32), relative to the more complicated algorithms used. Hence, we may assume that F32 and F38 have little epistasis, and are therefore “easier” problems than the other four. We shall see that the main hurdle to IMHC’s performance is *not* high epistasis (i.e. a great many local optima), but deception.

Finding the best solutions for these six problems is not the goal of this study. The goal is to discover the problem-independent characteristics of the best performing algorithms, and subsequently determine whether the theorized “important” characteristics of the EuA—smart mutation and recombination—are useful when attacking hard discrete optimization problems, and, in addition, whether or not these characteristics are efficiently performed by the EuA. It will be shown that this is indeed the case. The next section will describe what is meant by the phrase “best performing algorithm”, and will be followed by

descriptions of the algorithms tested and their performance on the test problems.

## 5.1 Quantifying Performance

Optimization algorithm performance is usually measured by two quantities: the quality of solutions found and the speed at which high quality solutions are found. A standard method of representing an optimization algorithm's performance on a particular problem has been to construct "learning curves", which plot the change in quality of the best solutions found versus algorithm resource expenditure. Solution quality is usually measured with an average of the fitnesses of the best individuals discovered over a number of different runs, and resource expenditure is usually measured using raw computation time or the number of function evaluations conducted. While computation time and function evaluations are both very practical and valid measures of resource usage, the use of mean fitness values as a measure of quality must be used with caution. For instance, it may be the case that the mean fitness of the best individuals found by one algorithm was twice the mean fitness of the best individuals found by another algorithm. In spite of this, it must not be assumed that, the first algorithm performed twice as well, or even much better than the second algorithm. For example, consider a problem in which the optimal individual had fitness of 100.0, the second best individual had a fitness of 95.0, and all the other individuals had fitness ranging from 0.0 to 5.0. Assume that algorithm  $A_1$  can locate the optimum in half of the runs, but in the other runs  $A_1$  can only find points of fitness 5.0 or less. Algorithm  $A_2$  can locate the *second* best point 100% of the time, but can never locate the optimum. If we evaluate algorithms on the basis of the mean fitness of the best individuals they found, we would rate algorithm  $A_2$  to perform almost twice as well as algorithm  $A_1$ , even though  $A_2$  could never find the optimum. Thus the use of mean fitness to judge algorithm performance can be misleading, since it does not clearly demonstrate an algorithm's probability of finding optimal or almost-optimal points. This problem could be made even worse if the range of all fitness values was small, relative to their absolute values. For example, if we added  $10^6$  to the fitness of all the individuals in the search space, then the two algorithms would seem to perform identically, since the difference in their mean best fitness values would be almost nonexistent and easily attributed to sampling error.

A better way to quantify an algorithm's performance on a particular problem might be to compute statistics using the overall *ranks* of the best points found, instead of the absolute fitnesses of these points. By using rank instead of raw fitness, relative search capabilities of the algorithms could be compared without being affected by order-preserving distortions of the fitness space. However, determining the rank of each individual encountered is practically infeasible, as it would require enumeration of almost the entire search

space. In addition, reducing performance to a simple statistic (such as an average) will almost inevitably lead to a loss of important performance-related information. The cumulative distribution function of the ranks of points found by a given algorithm after a specific number of function evaluations is probably one of the best practical representations of an algorithm’s performance, as it simultaneously demonstrates both the quality and frequency of discovery of solutions.

In spite of the problems with mean fitness values, the standard practice of plotting mean fitness versus time (function evaluations) will be used in this thesis to represent the behavior of the algorithms, for the sake of simplicity and ease of understanding. In future studies, it might be practical to determine the ranks of created individuals. At this time we will be limited to mean fitness metrics and, in a few cases, a limited study of points close to the optima. Algorithms will be judged by how many function evaluations they require to solve the given problems, and not by their raw computational usage, as it is assumed that, in “real-world” optimization problems, the cost of a function evaluation is much greater than the computational cost of a particular search algorithm.

The next section will include a detailed description for each algorithm, as well as the algorithm parameters used in the experiments. Then, each experiment and its results will be discussed in depth.

## 5.2 The Algorithms

Three mutation-based and two pattern-based algorithms are evaluated in this section. The mutation-based algorithms are “Increasing Mutation Hill Climbing” (IMHC), “Decreasing Mutation Hill Climbing” (DMHC), and simulated annealing (SA) (Kirkpatrick and Sherrington, 1988). Both IMHC and DMHC were derived from “Random Mutation Hill Climbing” (RMHC) (Forrest and Mitchell, 1993) especially for this thesis. The mutation based-algorithms can be broken down into two classes—heating algorithms (IMHC) and cooling algorithms (DMHC and SA). Heating algorithms start with a random individual and slowly increase the “heat” (the size of mutations) using an arbitrary schedule. Heating algorithms’ main emphasis is exploitation. Therefore they can quickly find excellent solutions in some instances, but are prone to become trapped in local optima in others (at least until heat increases to a point where the search can escape the local optimum’s basin of attraction). Cooling algorithms start with high mutation rates and decrease these rates over time. Cooling algorithms are initially very explorative, with the hope of initially avoiding the basins of attraction of local optima. It is assumed that the basins of attraction of true optima will be stronger than those of local optima. However, high initial exploration may lead to very little progress, as promising areas are ignored in favor of further

research. Cooling algorithms also are liable to be trapped in local minima, if exploration decreases too much at the wrong time. If this occurs, the algorithm has little chance of escaping. The pattern-based methods studied can also be divided into two classes—genetic and eugenic. The genetic algorithm studied relies on random mutation and recombination, while the eugenic algorithm uses smart, statistically justified mutation and recombination. The EuA will be the only eugenic algorithm studied, although tests will be conducted with EuAs that do not use the restriction operator, and are therefore effectively equivalent to BSC (Syswerda, 1993).

Different parameter settings for these algorithms were tested in the experiments. For SA and DMHC, two settings of the “cooling rate” parameter are studied. For the GA, two recombination methods are tested. IMHC has no parameters, and therefore it wasn’t possible to study more than one incarnation of it. The parameters varied for the EuA were population size, selection noise, creation rate, and the use of the restriction operator. The goal of the study was not to determine the “absolute best” algorithm for use on a particular test problem, but to identify the prevalent biases, strengths and weakness of the algorithms, and to develop a unified theory of why a particular algorithm *type* would perform well on some problems, but not on others. Therefore no effort was made to tune any of the algorithms parameters used, even when hindsight suggested that some parameters were obviously ill-chosen for a particular problem. Please note that there are *many* variations of SA and GAs; the ones used for the experiments, as well as their parameters, were chosen to represent the “standard” incarnations, for which a high level of performance should be expected.

### 5.2.1 Increasing Mutation Hill Climbing (IMHC)

In IMHC, a random individual is initially chosen. Random locus mutation of varying size (from one-bit to  $l$ -bit) is then applied to this individual. Mutations are accepted when they increase fitness or cause no change to fitness, otherwise the individual remains unchanged. Initially, mutation size ( $i$ ) is one. Mutation size is slightly increased when no increases in fitness are encountered after a number of attempts. In particular, when no increase in fitness is encountered after  $l \cdot 2^i$   $i$ -bit mutations are attempted, mutation size is increased to  $i + 1$  bits. For example, if  $l \cdot 2^1$  one-bit mutations cause no increases in fitness, then  $l \cdot 2^2$  two-bit mutations are attempted, and if still no increases in fitness are encountered,  $l \cdot 2^3$  three-bit mutations are attempted, and so on.. The number  $i$  can increase until  $i = l$ . If this happens,  $i$  is reset to one. In order for IMHC to escape a local optimum, a point outside that local optimum’s basin of attraction must be discovered that has a fitness equal to or greater than that of the local optimum’s.

IMHC is the “silver bullet” against trivial toy problems such as BIT-COUNT. Out

of all the algorithms tested, IMHC will find the optimum the fastest if a problem is 0% epistatic, since its one-bit mutations will quickly hill climb to points of higher and higher fitness. In such a case, the expected number of fitness evaluations it will have to perform is  $O(l)$ . However, if it encounters a very large group of genotypically similar individuals with *equal* fitnesses, it can waste large amounts of function evaluations by mutating back and forth among members of this group. A slight optimization to IMHC (and all the mutation-based algorithms) would be to keep track of previous mutations, and prevent the algorithm from retrying the same mutation on the same individual more than once.

### 5.2.2 Decreasing Mutation Hill Climbing (DMHC)

DMHC is a very simple cooling algorithm. After an initial individual is randomly chosen, random locus mutations of decreasing size are attempted. When mutation causes no change or an increase in fitness to current individual, the mutated individual replaces the current individual. The mutation probability of each bit,  $p_M$ , is a function of the “cooling rate”  $r_c$ , the number of function evaluations  $t$  and the maximum number of function evaluations  $t_{max}$ :

$$p_M = 0.5(1.0 - \frac{t}{t_{max}})^{r_c} \quad (5.1)$$

At  $t = 0$ ,  $p_M$  is 0.5, and so half the bits in a genotype will be mutated on average. As  $r_c$  is increased,  $p_M$  decreases more quickly, with a resulting decrease in exploration, in favor of exploitation. DMHC is very similar to simulated annealing (see below), except that newly generated points with lower fitness than the current individual have absolutely no chance of replacing the current individual. In the discussions to follow, the notation DMHC- $x$  will be used to denote DMHC with a cooling rate of  $r_c = x$ .

### 5.2.3 Simulated Annealing (SA)

The simulated annealing algorithm used in the experiments is only a rough approximation of “true” simulated annealing. It is identical to DMHC, except that points of fitness lower than the current individuals’ fitness have a non-zero probability of being accepted and replacing the current individual. This probability decreases as more function evaluations are performed. This is the “acceptance probability” ( $p_A$ ) is defined as follows:

$$p_A = \frac{1.0}{1.0 + \exp(\frac{\Delta_f}{0.5(1.0 - \frac{t}{t_{max}})^{r_c}})} \quad (5.2)$$

$$(5.3)$$

$$\Delta_f = \frac{|f(\mathbf{x}_1) - f(\mathbf{x}_2)|}{|f(\mathbf{x}_1)|} \quad (5.4)$$

$$(5.5)$$

where  $f(\mathbf{x}_1)$  is the fitness of the current point,  $f(\mathbf{x}_2)$  is the fitness of the mutated point, and  $r_c$  controls the rate at which the acceptance probability and mutation size decreases. The desired benefit of accepting mutations that cause decreases in fitness is to decrease the probability of being trapped in a local minima and increase exploration. For four of the six problems (F31, F32, F33, and F38), this tactic appeared to be slightly detrimental, as SA under-performed DMHC, and, for the three more difficult problems (F31, F32, and F33), the tactic was severely harmful to the SA's performance relative to DMHC's performance. Against a problem (F2) that was specifically designed to have an extremely significant local minimum, this tactic was clearly beneficial. However, it must be noted that F2 is an artificial test problem created to *simulate* hard problems, while F31-F33 are classic combinatorial optimization problems. In the discussions to follow, the notation SA- $x$  will be used to denote SA with a cooling rate of  $r_c = x$ .

#### 5.2.4 The Genetic Algorithm (GA)

The GA used in these tests was “genesys-2.0” (Bäck, 1992). Fitness-proportionate selection and elitism were used. The population size was fixed 50 individuals, while the crossover probability and per-bit mutation probability were set to 0.6 and 0.001, respectively. The fitness values of the population were rescaled every 5 generations. Both simple 2-point crossover (2X) and uniform crossover (UX) were tested independently, with UX seeming to have a slight advantage. In all the experiments, the GA's consistently converged to unimpressive populations, in spite of the fitness rescaling used.

#### 5.2.5 The Eugenic Algorithm (EuA)

A complete description of the eugenic algorithm was given in chapter 4. Many different combinations of parameters were studied:

Population Size: 20,50, and 100

Selection Noise: 0.01 to 0.25 (usually 0.05)

Creation Rate: 0.01 to 0.25 (usually 0.05)

Restriction Operator: off or on

The primary parameters studied were population size and selection noise. Also, the EuA was tested with the restriction operator shut off. Removal of the restriction operator degraded performance for all problems except for F02.

### 5.3 Combinatorial Optimization Experiments

In this section, six experiments are described. Each experiment features a particular combinatorial optimization problem and the performance and behavior of all five optimization algorithms on this problem. All of the problems are *minimization* tasks, and so “higher fitness” individuals produce smaller phenotypic values for these functions. For a given individual, I shall refer to the actual function value as the individual’s phenotype, while I will use fitness to refer to the “goodness” of the individual.

All of the results were averaged over 100 runs of a particular algorithm/parameter combination. The mean phenotypic value of the best individuals found by the 100 runs was computed every 100 function evaluations, and learning curves were constructed that graphed the number of function evaluations versus the average phenotypic value of the best individual found by each run. Both axes of the learning curves are plotted using base ten logarithms of the actual numbers encountered. This is done in order to more clearly display differences in behavior. When the range of phenotypic values is huge, differences in the average phenotypic values can be more easily seen on log-log plots. Also, most of the learning curves of the experiments were highly exponential in nature, with initial extremely large changes in phenotypic value followed by very little change for many function evaluations. Few performance differences could be detected when examining linear plots. If the logarithms of both the number of function evaluations and average phenotypic value were not used, most of the plots would resemble many “L”-shaped curves squashed together on the left side of the graphs. The appearance of the learning curves is more linear using log-log plots, and behavioral differences are clearer.

Each experiment in this chapter will be presented with four accompanying graphs. All the graphs show the performance of three “standard” EuAs, which use a 5.0% selection noise and creation rate. The three standard EuAs use different population sizes—20, 50 and 100 individuals. Each graph compares the performance of these three standard EuAs with either “nonstandard” EuAs or different algorithms altogether. Nonstandard EuAs are those that have extremely high selection noise, creation rates, or do not use restriction. The

notation EuA- $\lambda$  is used to denote eugenic algorithms with population size  $\lambda$ . The variables  $N$  and  $C$  will denote EuA selection noise and creation rates, respectively. The four graphs presented for each experiment are:

1. **Eugenic vs Genetic Search:** Graphs the behavior of the standard EuAs relative to two or more GAs that use different recombination operators (uniform or simple 2-point crossover).
2. **Eugenic vs Mutation-Based Search:** Graphs the behavior of the standard EuAs relative to SA, DMHC, and IMHC. Two cooling rates, 1.0 and 4.0, were studied with both SA and DMHC.
3. **Selection Noise and Creation Rate Tests:** Graphs the behavior of the standard EuAs relative to EuAs that use various differing combinations of population size, selection noise and creation rate.
4. **Restriction Operator Tests:** Graphs the behavior of the standard EuAs relative to EuAs that do not use the restriction operator.

The  $o(1)$ -epistasis ( $E$ ) and fitness variation ( $F$ ) of each problem was estimated for each problem. A random sampling of individuals was generated and then epistasis ( $o(1)\chi^2$ -Epistasis) and fitness variation ( $\frac{S_b}{S_w}$ ) were calculated using the random sampling. For each paired  $o(1)$ -competition, two genotypes that differed by a single bit were compared. Five hundred  $o(1)$ -competitions occurred per gene (bit). Therefore the total number of random samples taken for an  $l$ -bit problem was  $l \cdot 500$ .

### 5.3.1 F02: The Two-Dimensional Rosenbrock Function

#### Problem Description

The two-dimensional Rosenbrock function has one global optimum and one major local optimum, when encoded in 2-dimensional real space (De Jong, 1993). It is an artificial problem that was intentionally created to have a local optimum with a basin of attraction much larger than the global optimum's basin of attraction. It is therefore very deceptive in real space. The two dependent variables can vary in the region  $[-5.12, 5.12]$ . This problem's phenotypic values can range from 0 to approximately  $10^5$ .

$$F02(\mathbf{y}) = (100 \cdot (y_2 - y_1^2)^2 + (y_1 - 1)^2) \quad (5.6)$$

$$y_i \in [-5.12, 5.12] \quad (5.7)$$

The genotype  $[1.0, 1.0]$  yields the global optimal phenotype (0.0), while the local optimum (0.017) occurs at  $[1.13111, 1.28]$ . When  $y_1 = 1.13111$  and  $y_2 = 1.28$ ,  $y_1^2 \approx 1.28 = y_2$ , and therefore both the first term in F02 ( $100 \cdot (y_2 - y_1^2)^2$ ) and the second term ( $(y_1 - 1)^2$ ) are very close to zero.

32,000 samples were generated for the random sampling conducted for estimating epistasis and fitness variance. The measured  $o(1)\chi^2$ -Epistasis) was 0.70 and the fitness variation ( $\frac{S_b}{S_w}$ ) was 1.84. The average phenotypic value was 14418.5; the minimum value encountered in the random sampling was 0.0159001 and the maximum was 97766.4.

When the two real variables are encoded in  $n$ -dimensional binary space, however, the problem becomes rife with many additional mutation-based and pattern-based local optima (due to hamming cliffs). Obviously, the global optimum still occurs when the both groups of 32 bits are decoded to yield the real number 1.0. However, due to limited precision, this number can be coded in two different ways:

$$\begin{aligned} G_1 &= 00000000000000000000000010011001 \quad (= 1.0) \\ G_2 &= 111111111111111111111111100011001 \quad (= 1.0) \end{aligned}$$

Note that the rightmost bits are the most significant in this binary encoding.

In addition, the strong local optimum 0.017 still occurs when the first group of 32 bits decodes to 1.13111 and the second group decodes to 1.28. The coding for this local optimum is:

$$\begin{aligned} y_1 &= 01010010010001001110001000111001 \quad (= 1.13111) \\ y_2 &= 0000000000000000000000000000101 \quad (= 1.28) \end{aligned}$$

The reason why this is a local optimum is because a large mutation would be required to decrease the phenotypic value, while intermediate-sized mutations would cause *large* increases in the phenotypic value. As a result, even though the second group of bits differs by only 4 bits from the  $G_1$  setting (causing 1.0), changing these 4 bits would require a simultaneous change of either 11 bits (to the  $G_1$  setting) or 16 bits (to the  $G_2$  setting) in  $y_1$ 's bits, so that  $y_1$ 's bits would also decode to the value 1.0. Therefore, to decrease the phenotypic value from 0.017 down to 0.0 would require a fairly precise jump of  $(4 + 11) = 15$  or  $(4 + 16) = 20$  bits—roughly 25%-30% of the bits. The jump would also have to be well coordinated, because the large first term in F02 causes large increases in the phenotypic value when  $y_1^2$  differs too much from  $y_2$ . Because of the transformation from a real space genotype to a binary space genotype, many additional local optima are introduced into the problem (including another fairly strong one at phenotypic value 1.0, requiring a perfect 12-bit mutation). Many of the optimization algorithms tested frequently converged to one of these other local optima (see below).