

Maintaining Diversity in The Bounded Pareto-Set: A Case of Opposition Based Solution Generation Scheme

AKM Khaled Ahsan
Talukder
Department of Computer
Science and Engineering
Michigan State University
East Lansing, MI 48824, USA
talukde1@msu.edu

Kalyanmoy Deb
Department of Electrical and
Computer Engineering
Michigan State University
East Lansing, MI 48824, USA
kdeb@egr.msu.edu

Shahryar Rahnamayan
Department of Electrical,
Computer and Software
Engineering
University of Ontario Institute
of Technology
Oshawa, Canada
shahryar.rahnamayan@uoit.ca

ABSTRACT

For more than two decades, stand-alone evolutionary multi-objective optimization (EMO) methods have been adequately demonstrated to find a set of trade-off solutions near Pareto-front for various multi-objective optimization problems. Despite a long-standing existence of classical generative single-objective based methods, a very few EMO studies have combined the two approaches for a better gain. In this paper, we investigate the effect of seeding the initial population of an EMO algorithm with extreme solutions obtained using a single-objective method. Our proposed approach is further aided with an opposition based offspring creation mechanism which strategically places new solutions on the current Pareto frontier by a simple, yet a novel arbitration policy that utilizes the relative distances from the extreme solutions in the current population members. We conduct an extensive simulation of our proposed approach on a wide variety of two and three-objective benchmark MOP test problems. Results are shown to be remarkably better than the original EMO approach in terms of hyper-volume metric. The results are interesting and should motivate EMO researchers to integrate single-objective focused optimization and an opposition-based concept with diversity-preserving EMO procedures for an overall better performance.

Keywords

Single objective Search, Multiobjective Optimization, Opposition Based Learning (OBL)

1. INTRODUCTION

In the recent years, the idea of Opposition Based Learning (OBL) has been enjoying a noticeable attention among the AI and OR practitioners. The idea of OBL is to accelerate the learning rate (or convergence rate) by imposing an opposite estimate of the current solution, and deliberately

introducing them to influence the search trajectory. This idea was first introduced in [10] and has been demonstrated its effectiveness in different scenarios – from Reinforcement Learning (RL) [11], Differential Evolution (DE) [8] to robotics [5].

In all of the cases, the OBL comes first into play during the initialization phase of the learning algorithm (or optimization algorithm), the argument behind this strategy is that a completely arbitrary (i.e. blind) initialization is no better than an informed boot-strap. There are also several formal proofs that show that, in general, the opposite estimations are more likely to be closer to the optimal solution than the completely arbitrary ones, a formal probabilistic analysis of opposition based learning can be found in [4].

Moreover, especially for the case of optimization algorithms, it has been demonstrated that keeping a small ratio of candidate solutions with the opposite estimate help to converge faster [8] – and most of the recent studies are inspired by this approach. For example, similar strategy has been used in [6]. Another relevant study can be found in [12], where the opposition based initialization has been tested with the Particle Swarm Optimization (PSO) scenario to induce a better convergence. Given all of these studies, still we could not find a good example of this idea of “Opposition-based Learning” applied to the Evolutionary Multi-objective Optimization (EMO) scenarios. In this study we will try to fill this gap in an interesting way.

In our approach, we will address this idea of *opposition* in a different perspective, we will classify¹ the solutions with respect to some *desired traits* that we will like to have. Then we will use this information to deterministically² generate new solutions in a most viable locations in the search space – and this operation will be dictated by our special notion of opposition. In fact, our approach is somewhat similar in spirit to the previous studies done in [9] and [15]. Our approach is extremely simple and does not assume any special property on the underlying search space, moreover, our approach does not build a computationally expensive models.

This paper is organized as follows – first, we will discuss why this idea of OBL needs to be reinterpreted for the Multi-

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

GECCO'16 Companion, July 20-24, 2016, Denver, CO, USA

© 2016 ACM. ISBN 978-1-4503-4323-7/16/07...\$15.00

DOI: <http://dx.doi.org/10.1145/2908961.2931652>

¹not to be confused with the term “classification” as in the machine learning domain

²Here, our method is not “deterministic” in a general sense, the stochasticity incurred by the candidate generation is not equivalent to the generic mutation operation.

objective Optimization (MOP) setting. Next we discuss one interesting limitation that most EMO algorithms suffer from – the *search trajectory bias*, and also give some argument on why such hindrance becomes inevitable. Then we will discuss how the idea of OBL can come into rescue. In our model, we generate the opposite points in a strictly deterministic way, by carrying out the arbitration of *opposition* in a different (and a more MOP relevant manner). To do this, we utilize the extreme solutions on the true Pareto-front (PF), and the next section discusses how to find them efficiently. After that we will describe our main algorithm in details. Then we will show our experiments with different benchmark MOP problem sets.

2. AN ALTERNATIVE INTERPRETATION OF OPPOSITION

As we have already discussed in the previous section, in most of the studies, the idea of *opposition* is employed as an incorporation of new candidate solutions with a certain kind of *opposite traits* into the existing population. Such traits could be interpreted in terms of different problem domain perspectives. For example, the *opposite* solutions could be – i) the ones with complete opposite representation from the current best individual, ii) the ones with the opposite estimates from the other spectrum of the variable bounds (i.e. in the case of real valued optimization). However, the injection of the opposite solutions could cause a re-route from the continuing search trajectory and thus could be a misleading operation – in a sense that the opposite candidate solutions might be useful given that the search space follows a desired pattern.

Moreover, in the case of black-box optimization scenarios, we hardly have a room to make any assumption about the search space. Keeping this fact in mind, in this paper, we have revised the notion of opposition in terms of the preference criteria imposed on a solution. For example, most EMO algorithms aim to maximize two principal properties – i) the convergence and ii) the diversity, since the quality of a MOP solution depends on these two factors. Let us re-consider the opposite point generation/injection from a different perspective –

- **Opposite Convergence:** A solution *far* from the true Pareto-front is *opposite* to any solution that is *closer* to the true Pareto-front.
- **Opposite Diversity:** An *isolated* solution on the true Pareto-front is *opposite* to a *crowded* solution.

By taking the above two principles into account, we will deterministically generate opposite candidate solutions during the search. Obviously, the deterministic candidate generation scheme will only consider an opposite trait that is *good*. In the next section we will see, how the existing EMO algorithms show the limitations in maintaining this two opposite traits during the search (i.e., solution generation) process.

3. LIMITATIONS OF CANONICAL EMOs: THE SEARCH TRAJECTORY BIAS

Most of the standard EMO algorithms (e.g. NSGA-II, SPEA-II [17] etc.), are elitist by design. They are also “opportunistic” in a sense that the population always tries to

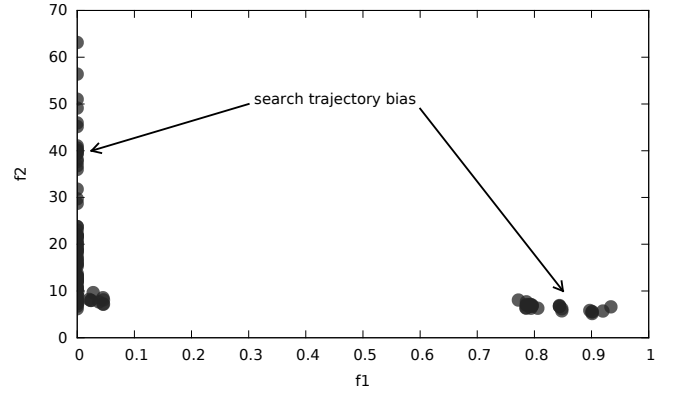


Figure 1: The effect of the *search trajectory bias* could be seen when we try to solve ZDT4 problem with NSGA-II (at generation 46, population size 100). Here we can see a long streak of crowded solution near the objective f_2 , where the distribution of solutions near the objective f_1 is extremely sparse.

converge to a particular portion of the Pareto-front (PF) which seems to be easier to solve at that moment. Therefore, they show preferences over a certain objective function which needs less exploration than the other. We can see a loss in diversity during the search when we try to solve the ZDT4 problem using NSGA-II. In this case, the objective f_2 is easier to optimize than f_1 , the readers can verify this fact from the Figure 1. The search trajectory deliberately accumulates more points over the objective f_2 to explore one particular portion of the Pareto-front. Moreover, while putting more solutions to the vicinity of one particular objective axis, the search trajectory loses the uniformity by forming a crowded streak of points along that axis (i.e. f_2); on the other hand we can see that there is almost no solution on the other spectrum of the objective space (i.e. f_1). This kind of non-symmetric search behaviour is, we think, causes a hindrance to keep a uniformly diverse solutions during the search process. So it would be helpful if we could selectively inject points during the search where the solution distribution is more sparse. In addition, we also believe that this biased nature of the search trajectory could degrade with the addition of more objective functions. Furthermore, this could also lead to a stagnation on the local optima, given that the search space has many of them.

Given this specific scenario, now the main problem is to devise a way to deterministically generate points where the distribution of the solution is sparse. Here we assume that the lesser the number of candidate solutions in a vicinity of objective f_i , the harder it is to solve. This would be easy if we know the exact mapping of the design variable to objective values – however such mapping is always unavailable and above all it is very expensive to infer. Another way could be to mutate the points where the solution is more sparse, but we think as the original algorithm already goes through such step, it is not going to be very effective.

In this paper, we are going to demonstrate a very effective approach to address the above discussed issue, we will show how we can maintain a balanced lateral distribution of solutions along the entire search operation. The proposed approach is also extremely effective in enhancing convergence speed.

4. THE DETERMINISTIC OPPOSITE SOLUTION GENERATION SCHEME

As we have discussed in the Section 2, we will utilize the so-called notion of *opposition* to deterministically generate points into strategically useful places on the search space. In principle, we do not assume any exact mapping over the design variables to objective values, and we apply a linear translation to achieve our goal effectively. The justification is trivial –

if a solution $\mathbf{x}_i$ is better on objective f_i , and if we want it to be better on f_j , then we can translate the vector closer to $\mathbf{x}_j$, which is already better on f_j .

So, the initial step is to approximate the bound of true PF before starting the actual optimization run. To do this, we depend on k number (k = number of objectives) of extreme (or near/approximated extreme) points on the true PF since we can safely assume that the population will eventually reach to the vicinity of those extreme points in the end, because the extreme points belong to the feasible space. These extreme points will be used as a pivot to arbitrate between the opposite traits over the existing solutions. Also note that we are not going to explicitly define which portion of the search space is less easy (or hard) and so forth, rather we will try to devise a technique that will implicitly address and solve such issues on the fly.

Another reason to fixate over the k extreme points is that we also wanted to keep the algorithm simple so that it can only utilize the “minimal information” to approximate the bound of true PF. We also think it’s valid to assume that any PF could be bounded by at least k -extreme points for any k -objective problem. Although, if we could supply other intermediary points within this bound, we will be able to see a better performance gain with the existing model. But a supply of 1 extra solution (within the bound) comes with an added cost of running a full blown single objective optimization run. As the extreme (or near extreme) points are the pivot to define the notion of *opposite* in our case, we will discuss how to find them efficiently in the following section.

4.1 Finding the Extreme Solutions

Extreme points on the Pareto-front could be found using global search as well [3], however our goal was to save the extra computational cost as much as possible³. Therefore, we resort to the classical single-objective optimization methods to solve this problem. Our choice of such algorithms were limited to, namely, the Interior Point (IP) method and the Mesh Adaptive Direct Search (MADS)⁴. As we also did not want to spend the valuable solution evaluations for this purpose, we have conducted this extreme solution search as a fixed budget operation. Depending on the difficulty of the problems, appropriate routine parameters were empirically found out which are problem dependent. These settings can be summarized as follows:

³The readers might be aware that efficiently finding the extreme points on the true Pareto-front is itself a separate research problem [3].

⁴We have used the `fmincon()` and the `patternsearch()` routine in MATLAB (v. R2014a) for IP method and MADS respectively. A detailed discussion on these algorithms can be found elsewhere [7] [1].

Algorithm 1 Find Extreme Points

```

1:  $k \leftarrow$  no. of objectives
2:  $N_p \leftarrow$  population size
3:  $N_{\text{gen}} \leftarrow$  maximum generation
4:  $T \leftarrow \frac{1}{k}(\frac{1}{4}N_pN_{\text{gen}})$ 
5:  $E^* \leftarrow \emptyset$ , an empty solution set
6: for  $i$  from 1 to  $k$  do
7:    $f_i \leftarrow$   $i$ -th objective function
8:    $\mathbf{z} \leftarrow$  random initial vector
9:   repeat
10:     $\mathbf{z} \leftarrow$  solve  $f_i$ 
11:    until  $\frac{T}{2}$  solution evaluation reached
12:    $f_{\text{aasf}} \leftarrow$  construct AASF function from  $f_i$ 
13:   repeat
14:     $\mathbf{x} \leftarrow$  solve  $f_{\text{aasf}}$  with  $\mathbf{z}$ 
15:    until  $\frac{T}{2}$  solution evaluation reached
16:    $E^* \leftarrow \{E^* \cup \mathbf{x}\}$ 
17: end for
18: return  $E^*$ 

```

- If the problem has no local optima then we use IP method, it has been found (empirically) to be comparatively less expensive even if the variable size is large.
- If the problem has local optima, MADS is faster for finding more accurate extreme points. Also, this observation is confirmed with empirical experiments.

The actual extreme point computation algorithm was conducted in two steps – given a particular objective function f_i , first we try to solve it directly using either IPM or MADS (depending on the problem type); then after some $\frac{T}{2}$ iterations once we find a reference solution $\mathbf{z}$ that is hopefully close to f_i ’s optima, then we construct a so-called Augmented Achievement Scalarizing Function (AASF) [14] from f_i as $f_{\text{aasf}} = \max_{j=1}^k w_j(f_j(\mathbf{x}) - f_j(\mathbf{z})) + \rho \sum_{j=1}^k w_j(f_j(\mathbf{x}) - f_j(\mathbf{z}))$ and solve it again for $\frac{T}{2}$ iterations. Here, we set $w_i = 0.9$, $w_{j \neq i} = \frac{1}{10(k-1)}$ and $\rho = 0.0001$. The choice of w_i values are trivial, and ρ is followed from [3].

To limit the solution evaluations, we kept T to a constant value (as a budget). For all problems, we have fixed this maximum iteration count to the $\frac{1}{4}$ -th of the total generation specified. More precisely, $T = \frac{1}{k}(\frac{1}{4}N_pN_{\text{gen}})$, where k = no. of objectives, N_p = population size and N_{gen} = maximum generation. A basic listing for this routine is presented in Algorithm 1. The set of the extreme points E^* generated from this algorithm may not contain all the unique solution, and also they might not be the true extreme always, they can be weakly dominated solutions by the true PF extremes. However, our approach can utilize them efficiently to converge to the true PF extremes.

4.2 Generating the Opposite Solutions

Once the extreme points are discovered, now we utilize them to generate the so-called *opposite* points during the main evolutionary runs. On each generation, we select 25% of the best individuals (front-wise) from the current population and deterministically change them to generate opposite solutions – in such a way that we can address the strategically preferable places. And to conduct this variation, we will utilize the points in the set E^* as pivot points. We call these points as “pivot” since we will selectively try to gener-

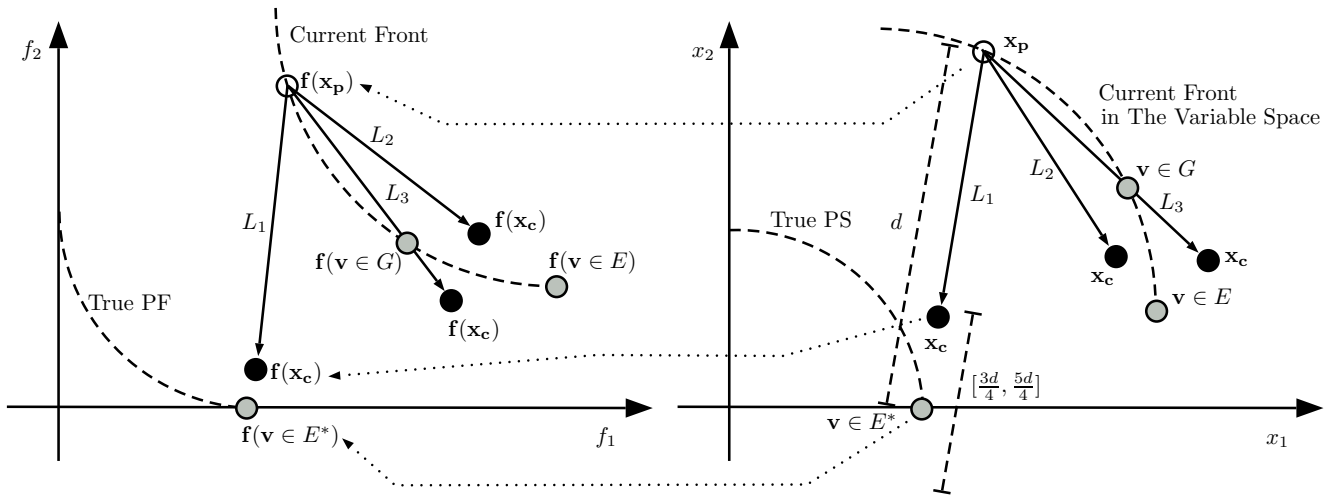


Figure 2: The illustration of lines 9–12 in Algorithm 3. The right axes are the variable space and the left axes are the corresponding objective space. The point $\mathbf{x}_c$ is the child (black circle) and $\mathbf{x}_p$ is the parent (white circle). The point $\mathbf{v}$ are the pivot points (grey circles). The operation will choose one of the directions denoted by L_1 , L_2 or L_3 arbitrarily. If $\mathbf{x}_c$ violates the variable bound then it is reverted back to the vicinity of the corresponding pivot point $\mathbf{v}$.

ate points around these pivots. However, before doing this, we will *refine* our pivot points E^* in a certain way.

The *refinement* starts by finding the current population extreme points E_c and merging them with the set E^* such that $E = \{E_c \cup E^*\}$. Next we apply the non-dominated sort on E to find the Pareto-front within this set. We apply this sorting because we are not still sure if E^* contains true PF extremes. This sorting will keep the true extreme points if ones are found in the later generations. After this step, we select the points from E that are on the best front and with ∞ crowding distances. Let us denote these selected points as E' . Now at this point, two situations are possible:

- The set E' contains only the solutions E^* when we are not converged to PF yet, or
- The set E' contains the solutions E_c , if the whole population reach to PF, then E_c are the PF extreme.

However, during the intermediate generations, it is possible that we may include some solutions into E that weakly dominate a subset of points already in E , this inclusion will reduce the expected spread of the pivot points – that may diminish the effect of maintaining the diversity. For example, if the actual true PF is a broken Pareto-front, and if E contains the extreme points from one broken edge, then we need to expand the current edges so that the refinement procedure can include points from the further extreme ends. Therefore, if there exist a point in $E - E'$ that is on the best front and also weakly dominated by any point in E' , then we replace the weakly dominating point from E' with the one from the set $E - E'$. The readers might have already noticed that $|E'| \leq |E|$.

Now, at this point, we can say that the set E' contains either approximated PF extremes (or points near them). Now if we can generate new points near E' , they will induce both better convergence and diversity. In section 3, we have discussed a scenario where we can see how the bias in the search trajectory is introduced. However, the difference in the relative difficulty of the objective functions may not

Algorithm 2 Generate Pivot Points

Require: true PF extreme points E^* from Algorithm 1

- 1: $E_c \leftarrow$ the extreme points from the current PF
 - 2: $E \leftarrow \{E^* \cup E_c\}$
 - 3: rank points in E , $E \rightarrow \{\mathcal{F}_1, \mathcal{F}_2, \dots, \mathcal{F}_n\}$
 - 4: take the best front in E' , $E' \leftarrow \mathcal{F}_1$
 - 5: **for** all points p_i in $E - E'$ **do**
 - 6: **if** p_i weakly dominates any $p_j \in E'$ **then**
 - 7: replace p_j by p_i
 - 8: **end if**
 - 9: **end for**
 - 10: update E^* , $E^* \leftarrow E'$
 - 11: $G \leftarrow$ find k intermediary gaps from the current PF
 - 12: $E' \leftarrow \{E' \cup G\}$
 - 13: **return** E'
-

be the only reason for a bias. The imbalance in the solution distribution could happen for other reasons as well. For example, a disconnected Pareto-front, a local optimal front or a specific portion of the Pareto-front being more difficult to solve than the rest. In such cases, we loose the diversity in the trade-off solutions and *gaps* start to form over the Pareto-front during the search, we can see such pathologies in many problems (especially in ZDT4 problem). To address such diversity loss, we also find the solutions with k -highest ($k = \text{no. of objectives}$) crowding distance from the best front that are not ∞ , and call them as set G . Clearly, the G solutions are those that reside on the edge of the broken front. Now we add the G to the set E' , thus we make E' ($|E'| > k$) as the final “pivot” solutions to generate the opposite points (see Algorithm 2).

4.3 Integrating into an Elitist EMO

Algorithm: NSGA-II

As we have mentioned at the beginning that we select front-wise best 25% of the current population for opposite point generation. We go through each of them randomly

and every time we pick k number of random points from E' and pick the pivot point that is the furthest from it, and find the opposite vector using a linear translation. A straightforward way – given a pivot vector $\mathbf{v}$ and a parent vector $\mathbf{x}_p$, we generate an opposite child $\mathbf{x}_c$ as $\mathbf{x}_c = \mathbf{x}_p + \mathbf{U}[(\frac{3d}{4}, \frac{5d}{4})] \circ (\frac{1}{d}(\mathbf{x}_p - \mathbf{v}))$. Here, $d = \|\mathbf{v} - \mathbf{x}_p\|$, $\mathbf{U}[(d, u)]$ is a uniform random vector where each element is within the range $[d, u]$ and $\circ$ denotes *Hadamard* product. The overall procedure is presented in Algorithm 3 in line 9–12 and illustrated in the Figure 2. The lines 9–12 in Algorithm 3 can be recapped as follows: $\mathbf{v}$ is on the true PF extreme and $\mathbf{x}_i$ is far from $\mathbf{v}$, therefore, move $\mathbf{x}_i$ closer to $\mathbf{v}$ – *opposite of far is close*. Similar interpretation can be made when the vector $\mathbf{v}$ is an intermediary *gap*.

Moreover, upon generating the vector $\mathbf{x}_c$, one of the variable values might go beyond the variable bounds (i.e. $x_j > x_{j_H}$ or $x_j < x_{j_L}$), in that case, we replace the overshoot value from the corresponding variable value v_j from the pivot point $\mathbf{v}$. Therefore, if a certain vector $\mathbf{x}_c$ can't make a successful translation, then $\mathbf{x}_c$ is reverted back to the vicinity of $\mathbf{v}$. Thus, we assure a local best estimated translation of the parent vector $\mathbf{x}_p$. This process is done on the line 13 of the Algorithm 3.

When we apply this algorithm to NSGA-II, we follow the obvious way, the generated opposite population will be inserted into the child population Q_t , the Algorithm 3 also shows how to integrate everything in NSGA-II. Moreover, this algorithm is “pluggable” in a sense that we can integrate it to any other elitist EMO algorithm. In the following section, we are going to see in details, how our opposite generation algorithm drastically improves the convergence rate. The readers should be aware that the 25% allocation was found by empirical experiments, and also to be better for most of the test problems.

5. EXPERIMENTS AND RESULTS

First we have tested the performance of Algorithm 3 on five 2-objective problems [16], namely ZDT1, ZDT2, ZDT3, ZDT4 and ZDT6, and we have set NSGA-II as the control. To maintain a fair comparison, we have compensated the extra solution evaluations by the Algorithm 1 for the NSGA-II runs, and compared NSGA-II and Algorithm 3 side by side. The performance measure for our test was Hypervolume (HV) [13], and we are interested to see which algorithm can reach to a desired HV within less solution evaluations (SE). For all problems, we have seen our algorithm can demonstrate a very steep convergence to the true PF, even when the extra SE from Algorithm 1 are compensated for NSGA-II.

Moreover, as we already know, the standard test problems are designed from some of the well known single-objective optimization functions. For example, in the case of ZDT4, $g(\vec{x})$ is the well known *Rastrigin's function*, and the two objectives are related as $f_2 = 1 - \sqrt{f_1}$. Therefore, each solution on the true PF follows a well-defined pattern – the variables $x_2, x_3, \dots, x_n$ are all the same (solution to the *Rastrigin's problem*) and the first variable x_1 defines the trade-off. As a result, if our algorithm can find any solution on the true-PF at a certain time instance, the overshoot-correction on the line 13 of Algorithm 3 have a good chance of creating a valid neighbouring solution. To mute such an un-warranted exploration, we have changed the problem definition in such a way that all the solutions on the true-PF will have $x_1 = 0.5$,

Algorithm 3 NSGA-II with Opposition

Require: true PF extreme points E^* from Algorithm 1

```

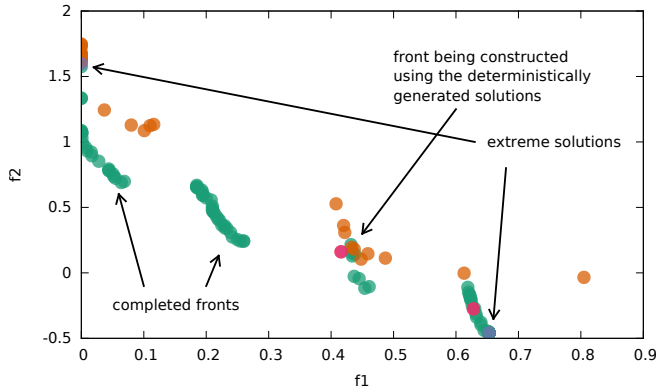
1:  $N \leftarrow$  population size  $|P_t|$ 
2:  $N_{\text{gen}} \leftarrow$  maximum generation
3:  $t \leftarrow 1$ 
4: while  $t \leq N_{\text{gen}}$  do
5:    $P'_t \leftarrow$  front-wise best 25% from  $P_t$  and shuffle
6:    $E'_t \leftarrow$  construct pivot set  $E'$  using algorithm 2
7:    $O_t \leftarrow \emptyset$ 
8:   for each solution  $\mathbf{x}_i \in P'_t$  do
9:      $S \leftarrow$  pick  $k$  random solutions from  $E'_t$ 
10:     $\mathbf{v} \in S$  such that  $\mathbf{v}$  is the furthest point from  $\mathbf{x}_i$ 
11:     $d \leftarrow \|\mathbf{v} - \mathbf{x}_i\|$ 
12:     $\mathbf{x}_c \leftarrow \mathbf{x}_i + \mathbf{U}[(\frac{3d}{4}, \frac{5d}{4})] \circ (\frac{1}{d}(\mathbf{x}_i - \mathbf{v}))$ 
13:     $x_j \leftarrow v_j \in \mathbf{v}$  if  $x_j \in \mathbf{x}_c > x_{j_H}$  or  $x_j \in \mathbf{x}_c < x_{j_L}$ 
14:     $O_t \leftarrow \{O_t \cup \mathbf{x}_c\}$ 
15:   end for
16:    $P_t \leftarrow \{P_t \cup E^*\}$ 
17:    $R_t \leftarrow \{P_t \cup Q_t\}$ 
18:   rank  $R_t$  into fronts,  $R_t \rightarrow \{\mathcal{F}_1, \mathcal{F}_2, \dots, \mathcal{F}_n\}$ 
19:    $P_{t+1} \leftarrow \emptyset$ 
20:    $i \leftarrow 1$ 
21:   while  $|P_{t+1}| + |\mathcal{F}_i| \leq N$  do
22:     assign crowding distances on the front  $\mathcal{F}_i$ 
23:      $P_{t+1} \leftarrow \{P_t \cup \mathcal{F}_i\}$ 
24:      $i \leftarrow i + 1$ 
25:   end while
26:   sort  $\mathcal{F}_i$  in descending order using  $\prec_n$ 
27:    $P_{t+1} \leftarrow$  the first  $N - |P_{t+1}|$  solutions from  $\mathcal{F}_i$ 
28:    $Q_{t+1} \leftarrow$  select, crossover and mutate  $P_{t+1}$ 
29:   randomly insert all  $x_i \in O_t$  into  $Q_{t+1}$ 
30:    $t \leftarrow t + 1$ 
31: end while
```

so that the Pareto-set is consisted of a straight line located at $x_1 = 0.5$ instead of $x_1 = 0.0$. Such modifications are done on the problems where appropriate.

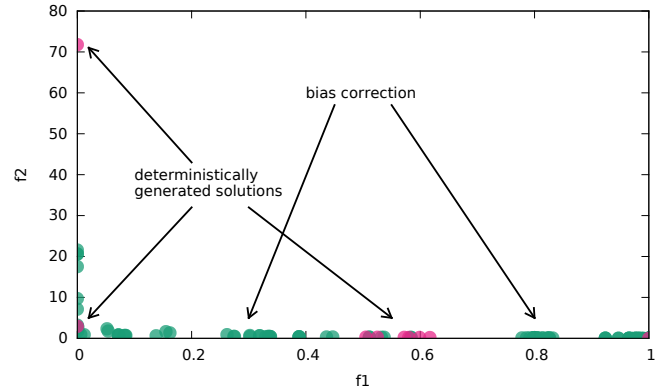
All the results are collected from 31 independent runs started with non-identical random seeds. In all plots, *onsga2r* stands for Algorithm 3. The extra cost to find the extreme points are indicated with a “T” arrow on the x-axis. During computation of the HV measure, we have set the reference point to $\{2.0, 2.0\}$ for all problems except ZDT6, where it has been set to $\{4.0, 4.0\}$. The Table 1 summarizes the performance measuring parameters, the second column shows the *ideal* hypervolume of each of the PF of the corresponding problems. The second and the third column presents the number of SE have been spent to reach 90% of the HV by each of the algorithms. To ensure a completely fair comparison, for all instances, we have excluded the solutions in E^* during the HV computation.

The experiment with ZDT1 is illustrated in Figure 5, here we can see that the Algorithm 1 takes up to around 2K of solution evaluations. Given that, NSGA-II still lags behind with a multiple factors to reach the desired PF. We have seen similar effect on all the rest of the problems ZDT2, ZDT3, ZDT4 and ZDT6. Except for ZDT3, we have observed some fluctuations due the disconnected nature of the true PF.

Moreover, our approach can also efficiently solve the issue of *search trajectory bias*, if we look at the Figure 3b, we can see that the new solutions are deterministically generated where the explorations are not done thoroughly yet. Due to the space constraints, we present a subset of the results in the Figure 6. In this paper, we only present the results for those problems that are comparatively harder to solve,

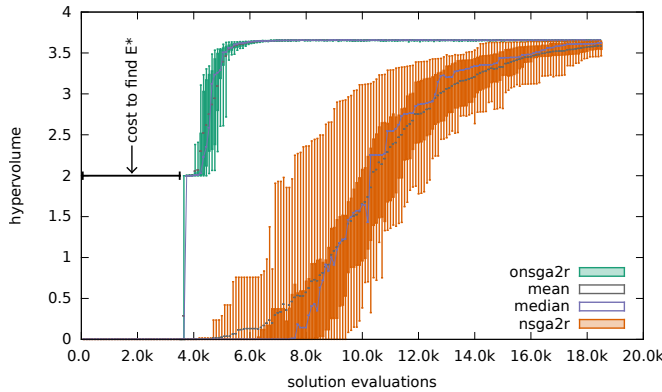


(a) Exploration of ZDT3 at generation 18

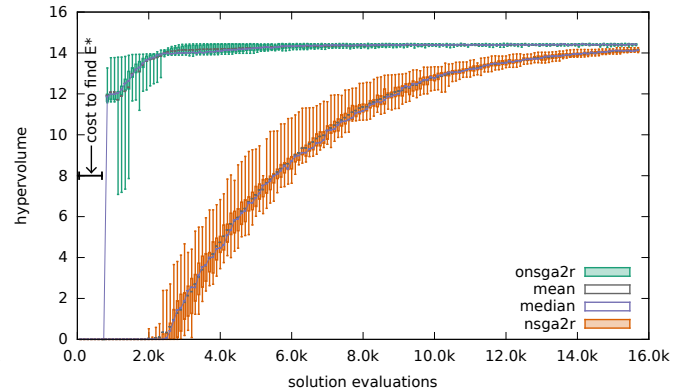


(b) Exploration of ZDT4 at generation 21

Figure 3: This figure illustrates how our algorithm deterministically identifies which front needs to be explored first and gradually discovers the entire PF. The example here demonstrates a case of ZDT3 problem (Figure 3a). The outlier dots (orange) represents the deterministically generated solutions that did not survived because they are weakly dominated. The light green dots are those that are deterministically generated and survived. In the case of ZDT4 (Figure 3b), we can see how the bias and gaps have been corrected by our approach.



(a) Convergence test for ZDT4 problem



(b) Convergence test for ZDT6 problem

Figure 4: These plots illustrates the comparative analysis of the convergence rates for 2-objective problems (ZDT4 and ZDT6), the curves are actually consisted of box-plots. Here onsga2r denotes our algorithm and nsga2r is NSGA-II.

Problem	Ideal HV (IHV)	Reference Points: (f_1, f_2, f_3)	NSGA-II 90%-IHV / SE	Algorithm3 90%-IHV / SE
ZDT1	3.67	(2.0, 2.0)	5100	1909
ZDT2	3.34	(2.0, 2.0)	7700	2164
ZDT3	4.82	(2.0, 2.0)	4900	8497
ZDT4	3.67	(2.0, 2.0)	14500	4950
ZDT6	15.35	(4.0, 4.0)	10500	1539
DTLZ1	999.98	(10.0, 10.0, 10.0)	24800	11602
DTLZ2	7.48	(2.0, 2.0, 2.0)	1200	10602
DTLZ3	3374.98	(15.0, 15.0, 15.0)	30600	11802
DTLZ4	7.48	(2.0, 2.0, 2.0)	1800	2022
DTLZ5	6.1	(2.0, 2.0, 2.0)	1200	1226
DTLZ6	55.6	(4.0, 4.0, 4.0)	29400	24602
DTLZ7	134.20	(10.0, 10.0, 10.0)	9200	3606

Table 1: The ideal HV, corresponding reference points. The last two columns represent the total number of SEs to attain 90% of the ideal HVs for both algorithms over the benchmark problems. For DTLZ7 problem, the ideal HV is approximated using a Monte-Carlo sampling.

i.e. problems with local-optima, disconnected fronts, non-uniform densities in the true PF etc.

6. CONCLUSIONS AND FUTURE WORKS

The main contribution of this paper is the incorporation of opposite point generation scheme in a different perspective. Our approach also shows that a simple and a deterministic scheme can aid the EMO optimization algorithm in a very interesting way. Our technique is also easy to implement and offers less overhead to the host algorithm. This approach can also correct the search bias introduced by the problem difficulty in an automated and predictable manner. The original idea of the *Opposition Based Learning (OBL)* is quite interesting; and we can make more of it if this idea is utilized in a more meaningful way – we think this is the main contribution of our study.

However, in the future we want to address some other interesting issues with our current study, especially in the

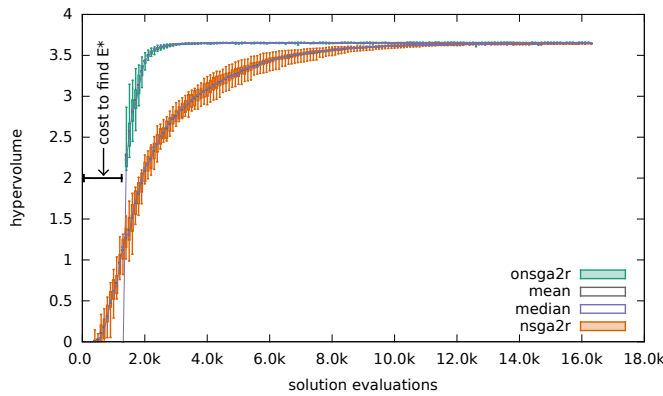


Figure 5: The convergence test of Algorithm 3 (onsga2r) vs. NSGA-II on problem ZDT1. The curves are actually consisted of box-plots, and for the Algorithm 3 (onsga2r), the medians are joined with a line. Here we can see that HV is 0 upto 1.8k SE and then abruptly reaches to 2 (the long vertical line at the beginning of Algorithm 3’s curve).

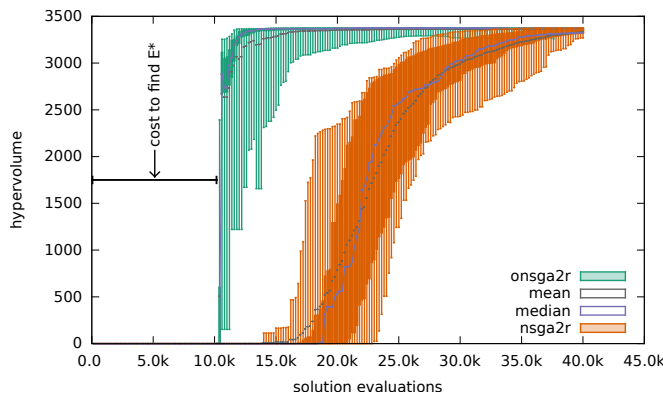


Figure 6: These plots illustrates the comparative analysis of the convergence rates for a 3-objective problem DTLZ3, the curves are actually consisted of box-plots. Here onsga2r denotes our algorithm and nsga2r is NSGA-II.

“many-objective” problems. Moreover, we think there are also a possible scope in changing the idea *opposition* for reference point based many-objective algorithms like MOEA/D and NSGA-III [2]. In the future research, we hope to investigate these ideas.

7. REFERENCES

- [1] C. Audet and J. J. E. Dennis. Analysis of generalized pattern searches. *SIAM Journal on Optimization*, 13(3):889–903, 2002.
- [2] K. Deb and H. Jain. An Evolutionary Many-Objective Optimization Algorithm Using Reference-Point-Based Nondominated Sorting Approach, Part I: Solving Problems With Box Constraints. *Evolutionary Computation, IEEE Transactions on*, 18(4):577–601, Aug 2014.
- [3] K. Deb, K. Miettinen, and S. Chaudhuri. Toward an Estimation of Nadir Objective Vector Using a Hybrid of Evolutionary and Local Search Approaches. *Evolutionary Computation, IEEE Transactions on*, 14(6):821–841, Dec 2010.
- [4] M. Ergezer and D. Simon. Probabilistic properties of fitness-based quasi-reflection in evolutionary algorithms. *Computers & Operations Research*, 63:114 – 124, 2015.
- [5] J. Kulk and J. Welsh. Using redundant fitness functions to improve optimisers for humanoid robot walking. In *Humanoid Robots (Humanoids), 2011 11th IEEE-RAS International Conference on*, pages 312–317, Oct 2011.
- [6] X. Ma, F. Liu, Y. Qi, M. Gong, M. Yin, L. Li, L. Jiao, and J. Wu. MOEA/D with opposition-based learning for multiobjective optimization problem. *Neurocomputing*, 146:48 – 64, 2014.
- [7] M. J. D. Powell. *Numerical Analysis: Proceedings of the Biennial Conference*, chapter A fast algorithm for nonlinearly constrained optimization calculations, pages 144–157. Springer, Jun 1978.
- [8] S. Rahnamayan, H. Tizhoosh, and M. Salama. Opposition-Based Differential Evolution. *Evolutionary Computation, IEEE Transactions on*, 12(1):64–79, Feb 2008.
- [9] V. Tirronen, F. Neri, and T. Rossi. Enhancing Differential Evolution frameworks by scale factor local search - Part I. In *Evolutionary Computation, 2009. CEC '09. IEEE Congress on*, pages 94–101, May 2009.
- [10] H. Tizhoosh. Opposition-Based Learning: A New Scheme for Machine Intelligence. In *Computational Intelligence for Modelling, Control and Automation, 2005 International Conference on*, volume 1, pages 695–701, Nov 2005.
- [11] H. R. Tizhoosh and S. Rahnamayan. Learning Opposites with Evolving Rules. *CoRR*, abs/1504.05619, 2015.
- [12] H. Wang, Z. Wu, S. Rahnamayan, Y. Liu, and M. Ventresca. Enhancing particle swarm optimization using generalized opposition-based learning. *Information Sciences*, 181(20):4699 – 4714, 2011. Special Issue on Interpretable Fuzzy Systems.
- [13] L. While, L. Bradstreet, and L. Barone. A Fast Way of Calculating Exact Hypervolumes. *Evolutionary Computation, IEEE Transactions on*, 16(1):86–95, Feb 2012.
- [14] A. Wierzbicki. The Use of Reference Objectives in Multiobjective Optimization. In G. Fandel and T. Gal, editors, *Multiple Criteria Decision Making Theory and Application*, volume 177 of *Lecture Notes in Economics and Mathematical Systems*, pages 468–486. Springer, 1980.
- [15] X. Zhang and S. Y. Yuen. A directional mutation operator for differential evolution algorithms. *Applied Soft Computing*, 30:529 – 548, 2015.
- [16] E. Zitzler, K. Deb, and L. Thiele. Comparison of Multiobjective Evolutionary Algorithms: Empirical Results. *Evolutionary Computation*, 8(2):173–195, 2000.
- [17] E. Zitzler, M. Laumanns, and L. Thiele. SPEA2: Improving the Strength Pareto Evolutionary Algorithm for Multiobjective Optimization. In *Evolutionary Methods for Design, Optimisation, and Control*, pages 95–100. CIMNE, 2002.