



Machine learning based decision support for many-objective optimization problems



João A. Duro^a, Dhish Kumar Saxena^{b,*}, Kalyanmoy Deb^c, Qingfu Zhang^{d,e}

^a Department of Computer Science, Bath University, Bath BA2 7AY, UK

^b Department of Mechanical & Industrial Engineering, Indian Institute of Technology, Roorkee 247667, India

^c Department of Electrical & Computer Engineering, Michigan State University, East Lansing, MI 48824, USA

^d Department of Computer Science, City University of Hong Kong, Hong Kong

^e School of Computer Science and Electronic Engineering, University of Essex, Colchester CO4 3SQ, UK

ARTICLE INFO

Article history:

Received 1 October 2013

Received in revised form

30 May 2014

Accepted 3 June 2014

Available online 24 July 2014

Keywords:

Evolutionary many-objective optimization

Principal component analysis

Maximum variance unfolding

Kernels and multiple criteria

decision-making

ABSTRACT

Multiple Criteria Decision-Making (MCDM) based Multi-objective Evolutionary Algorithms (MOEAs) are increasingly becoming popular for dealing with optimization problems with more than three objectives, commonly termed as many-objective optimization problems (MaOPs). These algorithms elicit preferences from a single or multiple Decision Makers (DMs), *a priori* or interactively, to guide the search towards the solutions most preferred by the DM(s), as against the whole Pareto-optimal Front (POF). Despite its promise for dealing with MaOPs, the utility of this approach is impaired by the lack of—*objectivity; repeatability; consistency; and coherence* in DM's preferences. This paper proposes a machine learning based framework to counter the above limitations. Towards it, the preference-structure of the different objectives embedded in the problem model is *learned* in terms of: a smallest set of conflicting objectives which can generate the same POF as the original problem; the smallest objective sets corresponding to pre-specified errors; and the objective sets of pre-specified sizes that correspond to minimum error. While the focus is on demonstrating how the proposed framework could serve as a decision support for the DM, its performance is also studied vis-à-vis an alternative approach (based on dominance relation preservation), for a wide range of test problems and a real-world problem. The results mark a new direction for MCDM based MOEAs for MaOPs.

© 2014 Elsevier B.V. All rights reserved.

1. Introduction

The need for handling many objectives (four or more) is being increasingly realized in industry. While the existing MOEAs have been well utilized for solving two- or three-objective real-world problems, it is alarming that the performance of most existing MOEAs severely deteriorates with an increase in the number of objectives (M) beyond three [1–3].

The plausible reasons for the performance deterioration of MOEAs, for $M \geq 4$, are the following:

- For a good POF-approximation (complete convergence and full coverage), the requirement of the population size used by the MOEAs grows exponentially with M . Hence, while working within practically reasonable computational resources and

time, a good POF-approximation for MaOPs by the MOEAs is a very difficult task.

- Most existing MOEAs are based on the Pareto-dominance based primary selection, which in the case of MaOPs becomes ineffective, given that almost the entire population becomes non-dominated from the early generations itself [4,5]. In fact, experiments in [6] have revealed that for $M > 12$, 100% population members become non-dominated in very few generations. This results in poor selection pressure for convergence to the POF.

While the difficulty associated with primary selection based on Pareto-dominance is in principle countered by the indicator based MOEAs [7,8], their running time increases exponentially with M [9–11], impeding their wider utility. The decomposition based MOEA/D [12] is also found to cope better with MaOPs [13]. However, owing to its use of the simplex-lattice design [14] for generation of coefficient vectors for aggregation of the original objectives, the population size ought to increase nonlinearly with M , and cannot be fixed at will. While the use of uniform design method [15] promises to counter the above limitations, its performance remains to be tested on a wider range of problems.

* Corresponding author.

E-mail addresses: j.a.duro@bath.ac.uk (J.A. Duro), dhishfme@iitr.ac.in (D. Kumar Saxena), kdeb@egr.msu.edu (K. Deb), qingfu.zhang@cityu.edu.hk, [qzhang@essex.ac.uk](mailto;qzhang@essex.ac.uk) (Q. Zhang).

Given the above difficulties, the merit in guiding an MOEA to only a few DM-preferred solutions, as against the whole POF, is being increasingly realized. This is affirmed by the growing emphasis on the MCDM based MOEAs [16,17]. However, their applicability may be impaired owing to the *cognitive limitations* of the DM [18–20]; and the lack of *objectivity* [21–23] and *coherence* [24] in DM's preferences.

This paper aims to counter the above limitations through a rationally derived decision support that could *inform* the DM about the preference-structure of the objective functions, inherent in the problem model itself. The rationale for this decision support lies in the fact that developing the models for an optimization problem at the first place, requires a lot of domain expertise, time and physical resources. Hence, the preference-structure embedded in these models cannot be treated as trivial, instead it could be used to guide the preference articulation by the DM. To meet its aim, this paper proposes a machine learning based framework that operates on the objective vectors of the non-dominated solutions obtained from an MOEA for the given problem-model; *learns* the preference-structure of the objective functions by preserving the correlation-structure of the solutions; and provides the decision support in terms of:

- I Revelation of an *essential* objective set: Here, for a given M -objective problem denoted by $\mathcal{F}_0 = \{f_1, f_2, \dots, f_M\}$, the framework reveals an *essential* objective set—the smallest set of conflicting objectives¹ ($\mathcal{F}_T, |\mathcal{F}_T| = m(m \leq M)$) that can generate the same POF as the original objective set.
- II Preference-ranking of all the objective functions: The framework facilitates preference-ranking of all the objectives by determining their preference-weights (w_i s), such that $w_i \geq 0$ and $\sum_{i=1}^M w_i = 1$. Each preference-weight essentially represents the normalized error (variance lost) that would be incurred, if the associated objective were to be eliminated.
- III δ -MOSS (δ -minimum objective subset) analysis: Here, the DM may specify the allowable degree of error in terms of δ , where $0 \leq \delta \leq 1$, and may be interested to know the smallest subset of objectives ($\mathcal{F}_{\{\delta\}s}$) which ensures that the error associated with omission of the remaining objectives (in the set $\mathcal{F}_0 \setminus \mathcal{F}_{\{\delta\}s}$) does not exceed δ . $\mathcal{F}_{\{\delta\}s}$ is referred as the δ -minimum² set, and $|\mathcal{F}_{\{\delta\}s}|$ denotes its size.
- IV k -EMOSS (minimum objective subset of size k with minimum error) analysis: Here, the DM may specify an allowable set size k (by specifying $p, 0 \leq p \leq 1$, where $k = \lceil pM \rceil$), and may be interested to know the subset of k objectives ($\mathcal{F}_{\{k\}s}$), such that the error associated with omission of the remaining objectives (in the set $\mathcal{F}_0 \setminus \mathcal{F}_{\{k\}s}$) is the minimum, compared to that corresponding to any other possible combination of k objectives. $\mathcal{F}_{\{k\}s}$ is referred as the k -minimum set, and the associated error, referred as k -minimum error, is denoted by $\mathcal{E}_k^n$.
- V Visual representation: The practical considerations demand that the amount of time sought from the DM should be as short as possible, and the information exchange (what is shown to-, and asked from-) with the DM should be as simple and uncomplicated as possible. Here, a simple yet meaningful visual representation of the above analysis results is presented that could serve as a snap-shot guide for the DM to base his or her preferences on.

While the proposed machine learning based framework is based on the generalization of [25], the distinctive and significant contributions of this paper relate to the following:

1. To the best of the authors' knowledge, this paper is the first in its assertion that in the case of MaOPs *in particular*, a rationally derived decision support is *essential* for the MCDM based MOEAs to be practically relevant. In doing so, this paper identifies:
 - (a) the different features, such as *objectivity*; *repeatability*; *consistency*; and *coherence* (discussed later), that the DM-preferences should be characterized with.
 - (b) that the natural basis for a rationally derived decision support lies in the process of *modeling* an optimization problem at the first place, as it requires multiple domain experts, time and physical resources.
 - (c) that the quest for a rationally derived decision support could be met by employing machine learning techniques on the solution sets obtained from MOEAs for the given problem model.
2. While the focus in [25] is purely limited to the Item I (above), the scope of this paper is extended to include the Items II–V (above). This becomes possible owing to the ability of the proposed framework (unlike in [25]) to determine the preference-weights for *all* the objectives—including the *essential* and *redundant* ones.
3. This paper puts a special emphasis on discussing the accuracy of the decision support. In this context, it may be noted that an objective reduction approach based on dominance relation preservation [26,27] has been generalized to offer the decision support in terms of δ -MOSS and k -EMOSS analysis [28,29]. However, despite its wide publication, no effort has been made by the author(s) to interpret the accuracy/inaccuracy of the reported results in the wake of the approach's ability/inability to handle *nonlinearity* and *noise*³ in the MOEA solutions that serve as the input data. This paper marks the first attempt to investigate the above issue, and offers valuable insights as to why the proposed framework may be more generic and accurate than that based on dominance relation preservation. To support its arguments, this paper considers:
 - (a) a nine-objective, radar waveform optimization problem, and highlights that the decision support offered by the proposed framework is entirely consistent with the physics of the problem, unlike the case of dominance relation preservation approach.
 - (b) a range of test problems that are investigated with regard to solutions directly sampled on to the true POF (*noise-free* solutions, allowing for isolated investigation of the issue of *nonlinearity*), and also those obtained from an MOEA.
4. The results presented in this paper are new and are based on over 1000 simulations⁴ performed on 12 different versions of test problems and a real-world problem.

The remaining paper is organized as follows. Sections 2 and 3 present the rationale for a decision support for MaOPs, and the related research, respectively. The proposed machine learning based framework is presented in Section 4, and demonstrated on

¹ Two objectives are said to be in conflict if there is no single solution that simultaneously optimizes each objective.

² For a given $0 \leq \delta \leq 1$, there may be multiple subsets of objectives, which ensure that the error associated with omission of the remaining objectives does not exceed δ . Each such subset is referred as δ -minimal objective subset. However, the δ -minimal objective subset having the smallest size is referred as δ -minimum objective subset.

³ In the current context, *noise* is defined as the difference in the dominance relations or the correlation-structure, between the Pareto-optimal solutions and those obtained from an MOEA (more details, in Section 4).

⁴ These simulations relate to (i) 20 different seeds for the underlying MOEA (ϵ -MOEA [30]), for each test problem, (ii) one set of solutions directly sampled on the true POF, for each test problem, (iii) the solution sets corresponding to 30 runs conducted by MSOPS-II [31], for the real-world, and (iv) the δ -MOSS and k -EMOSS analysis based on the proposed framework and also a dominance relation preservation based algorithm, for each data set referred above.

Download English Version:

<https://daneshyari.com/en/article/406568>

Download Persian Version:

<https://daneshyari.com/article/406568>

[Daneshyari.com](https://daneshyari.com)