

Automated Configuration for Agile Software Environments

Negar Mohammadi Koushki, Sanjeev Sondur, and Krishna Kant

Computer and Information Sciences

Temple University

Philadelphia, USA

{koushki|sanjeev.sondur|kkant}@temple.edu

Abstract—The increasing use of the DevOps paradigm in software systems has substantially increased the frequency of configuration parameter setting changes. Ensuring the correctness of such settings is generally a very challenging problem due to the complex interdependencies, and calls for an automated mechanism that can both run quickly and provide accurate settings. In this paper, we propose an efficient discrete combinatorial optimization technique that makes two unique contributions: (a) an improved and extended metaheuristic that exploits the application domain knowledge for fast convergence, and (b) the development and quantification of a discrete version of the classical tunneling mechanism to improve the accuracy of the solution. Our extensive evaluation using available workload traces that do include configuration information shows that the proposed technique can provide a lower-cost solution (by $\sim 60\%$) with faster convergence (by $\sim 48\%$) as compared to the traditional metaheuristic algorithms. Also, our solution succeeds in finding a feasible solution in approximately 30% more cases than the baseline algorithm.

Index Terms—Configuration Modeling, Resource Allocation, Resource Provisioning, Machine Learning, metaheuristics, Simulated Annealing

I. INTRODUCTION

The DevOps transformation of IT services is fueling a radical change in how cloud services are conceptualized, designed, and implemented [2]. The service-oriented architecture (SOA) and its manifestation in the emerging microservices paradigm [3] attempt to reduce dependencies between services to make their development and deployment more agile and allow for easier maintenance and scalability. The motto of the DevOps phenomenon is *Continuous Integration and Development (CI/CD)* [4]; that is, the *deployed services* are constantly being enhanced and tuned both with respect to their code and run-time settings. Furthermore, the services are increasingly being deployed in their own lightweight containers, allowing even the hardware resources assigned to the service to be fungible at run-time.

The run-time settings, popularly known as *Configuration Parameters* (CPs), play a key role in the functioning of most software systems and are easily misconfigured due to a variety of reasons. This includes a lack of clear documentation/understanding of what they do, interdependencies across them, and a lack of robust automated mechanisms to set them. In traditional systems, incorrect setting of configuration

parameters (or “misconfiguration”) are responsible for up to 80% of down-times [5] and up to 85% of the incidents [6], with *one week on average* for root cause identification [5, 6].

The CI/CD process substantially increases this complexity and the resultant potential for ill effects due to the need to constantly tune the CPs as the service modules undergo changes. In the current agile software systems, configuration changes happen all the time, e.g., reported thousands of times a day in Facebook [7]. Most of these are related to “tuning” the CPs rather than drastic changes. Thus, an automated mechanism that can quickly define the configuration settings (or verify proposed configuration settings) becomes a crucial ingredient of the emerging DevOps driven software development processes that continue to get adopted across a large variety of application domains. This paper aims to develop such a mechanism focused on the performance and cost aspects.

Our Contributions: Owing to the complexity of performance as a function of the configuration parameters, we adopt a combinatorial optimization approach to determine the optimal configuration. Our key contributions in this regard are: (a) extending the metaheuristic for efficiently solving the constrained optimization problem that exploits the application domain knowledge in grouping and choosing the parameters for perturbation and (b) developing the classical notion of “tunneling” in the context of implicit, discrete performance function, and (c) exploring how and when the tunneling helps in an efficient solution. Our extensive evaluation using available workloads shows that the proposed technique can provide a lower-cost solution (by $\sim 60\%$) with faster convergence (by $\sim 48\%$) as compared to the traditional metaheuristic algorithms. Also, our solution succeeds in finding an average of 30% more solutions than the baseline.

The rest of the paper is organized as follows. Section II reviews the configuration selection techniques and challenges. Section III provides an overview of our proposed method. In section IV, we describe the approach for selecting configuration. Section V addresses the issues in configuration modeling and then in section VI, we present the experimental evaluation. We conclude the paper in section VIII.

II. CURRENT ART ON CONFIGURATION SELECTION

Configuration issues are related to resource provisioning and resource management [8, 9] techniques to optimize latency, task completion time, data replication, and impact on cache capacity, delay, and energy consumption. Our work addresses *recommending* a suitable resource allocation (e.g., storage, compute, bandwidth) to achieve the desired goal (e.g., workload performance, energy, cost, size, etc.). Ref. [10] addresses the configuration problem by using a metaheuristics approach to provision Cloud resources for satisfying Quality of Service (QoS). Several studies have used Classification Regression Trees (CART)-based model [11] and ML techniques to design a Performance Influencing Model (PIM) [12]. In our study, PIM is only the first step to build a surrogate function to solve the combinatorial optimization problem. Our work focuses on *choosing a set of CPs* that satisfy user workload/performance demands under given constraints. Satyanarayanan in [13] use an example of Cloudlet infrastructure to show that configurations of Cloudlets are very challenging because of the many unknowns pertaining to the software mechanisms and controls.

Challenges in Configuration and Resource Allocation: In reviewing the importance of resource allocation mechanisms within Cloud/Edge infrastructures, Soumplis in Ref. [8] emphasizes the algorithmic challenges that must be overcome to utilize Cloud resources such as computing, storage, and networking infrastructures efficiently to serve the workload and data. Their study supports our problem statement in that many Cloud resources make it difficult for users to allocate resources to meet the QoS/Service Level Agreement (SLA) requirements, e.g., performance, latency, and execution time, while effectively keeping costs to a minimum. Also, supporting our problem complexity, Sfakianakis [14] states that given the variety and dynamics of applications, it is difficult to achieve optimal resource allocation.

We address the *dynamic resource allocation problem* capable of allocating multiple resources (such as CPU, memory, network bandwidth, and storage (I/O) bandwidth, etc.) to achieve the required QoS/SLA (e.g., throughput, latency, execution time) with minimal cost (e.g., deployment/maintenance cost, energy, power, etc.). Instead of relying on simulation tools, we demonstrate the effectiveness of our solution using publicly available data-sets (see Table. III) from real-world environments like Cloud & Edge applications, HPC workloads, application services, etc.

III. OVERVIEW OF PROPOSED METHOD

In this paper, we focus on the problem of determining configuration parameters that minimize a cost function subject to some minimum performance requirements. Essentially the same methods apply if, instead, we try to maximize performance for a given cost constraint. In a virtualized environment, the cost may refer to either the actual cost charged by the provider (e.g., AWS) or costs that we assign to various resources, including CPU cores, memory bandwidth and size, I/O rate, storage space allocated, etc. The appropriate configuration settings must often be determined quickly and

automatically to cater to the CI/CD needs. Our work addresses some specific questions raised by the DevOps team and further, supported in Ref. [9], viz: (i) How to design the resource allocation mechanism to provide dynamic scalability at CPU, network, application-level, etc., & (ii) How to minimize the cost and optimize the resource allocation simultaneously?

Regardless of whether the performance is used as an objective function or constraint, it is likely to be a complex function of various configuration parameters; thus, an accurate analytic or simulation model is usually difficult to construct and time-consuming to run. Therefore, we turn to a machine learning (ML) model that can generally be queried very fast. The main drawback of an ML model is the need for substantial amounts of data for training that mostly covers the parameter ranges that are likely to be used in practice. A beneficial side-effect of the configuration dynamism in DevOps environments is the availability of data with many different configuration settings. At the same time, the dynamism is likely to be limited to sensible ranges for an operational system. Thus the ML model for performance as a function of configuration parameters that are routinely retrained can fulfill our needs well (commonly referred to as PIM [12] and depicted as the “forward problem” in Fig. 3a). Although the cost model could also use the same approach, it is likely to be much simpler, and thus simple analytic expressions for it are generally adequate. We assume this to be the case in the rest of the paper.

A. Metaheuristics Based Modeling

We have demonstrated in [15] that it is difficult to build an ML model directly to solve the “backward problem” in Fig. 3a). Therefore, we devise a method that directly uses the performance and cost models to estimate the optimal configuration parameter values. Because of the lack of convexity of the performance function in general, a suitable approach is to use a discrete combinatorial optimization using metaheuristics [16]. There are numerous such algorithms, the best known perhaps being the Genetic Algorithm (GA) and Simulated Annealing (SA). It turns out that GA is almost universally the slowest [35, 33, 23], whereas SA and a somewhat similar algorithm called the Dynamically Dimensioned Search (DDS) are among the best. In this paper, we develop an enhanced version of SA and also show how the domain knowledge can be used to search the state space more efficiently.

The goal of any metaheuristics is to explore the state space efficiently while avoiding being trapped in local minima, of which there could be many. Fig. 1 illustrates the search process pictorially with the x-axis representing the iterations and the y-axis the solution obtained in each iteration. The algorithm will keep track of the minima obtained so far and may or may not discover the global minima until the maximum iteration count (a hyper-parameter of the algorithm) is reached. Using public domain data from several Cloud environments, we demonstrate the effectiveness of our solution with two distinct yet equally important metrics: (i) the speed in searching the huge configuration space and “quickly” selecting a suitable solution, & (ii) ensure that the selected solution is a minimum

cost solution. We explain this further in the evaluation section (Section VI-A).

Since our problem involves constrained optimization, we also need to ensure that any accepted solution is feasible. In such a setting, it is always helpful to avoid generating infeasible solutions in the first place, but this is not always possible. Regardless of the metaheuristics used, the stochastic optimization techniques move from the current best solution to the next solution that is both feasible (i.e., satisfies the constraints) and are better. Ideally, we would like to choose the next solution that is likely to have these characteristics without considering solutions that are unlikely to be useful. This is where domain knowledge is crucial. Often, domain knowledge consists of an abstract relationship between CPs or rules of thumb that can be evaluated quickly. However, since they are fuzzy and not strictly required, they cannot be used as formal constraints. Another kind of domain knowledge concerns the varying influence of parameters that can guide which parameters need to be perturbed and by how much to get to the next proposed solution. Yet another aspect concerns an estimate of the amount by which one needs to move to get out of the region of local optimality to land in another region that can possibly provide a lower local optimum.

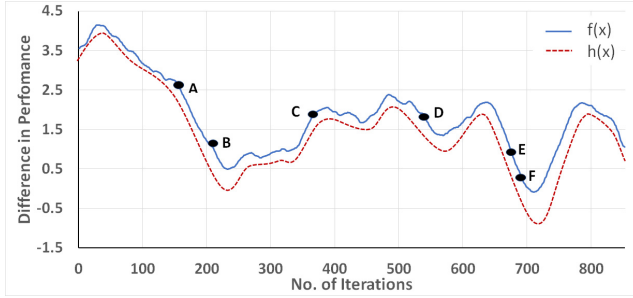


Fig. 1: Stochastic Tunneling.

B. Tunneling in Metaheuristics

A useful notion in stochastic optimization is *tunneling* [17], illustrated in Fig. 1, where we want to “tunnel” from one local minima to a deeper local minima directly [18]. Traditional tunneling works in the continuous parameter state space with explicit objective function $f(x)$ where x is the input vector (i.e., configuration vector in our case). It also assumes that $f(x)$ has the first two derivatives. The method works in two steps: (a) find the local minima, say x^* , using the steepest descent algorithm from the current point, (b) minimize the modified function $h(x) = [f(x) - f(x^*)]/\|x - x^*\|^\alpha$ with $\alpha \geq 1$ to determine the next solution, say x' . It can be seen that with larger α , nearby points are penalized. In general, the choice of α can be problematic, and approaches for its choice have been investigated [19]. Essentially, the use of $h(x)$ instead of $f(x)$ forces us to choose the next point away from x^* and with a lower function value. It thus avoids being caught in the local minima but with some risk that a deeper nearby minima may be missed. Note that $h(x)$ does not need to use $f(x)$ directly, instead, a tight envelope for $f(x)$ from below

would suffice if it is simpler to characterize. Fig. 1 shows such a $h(x)$.

This mechanism cannot be applied directly to our problem due to implicit $f(x)$ and discrete parameter space; we show that the concept can considerably improve the solution performance.

IV. COMBINATORIAL OPTIMIZATION BASED CONFIGURATION SELECTION

A general formulation of the configuration selection problem is as follows. Let CP denote the configuration defined as the vector of user-settable parameters $\vec{x}$ and vendor-selected (usually hidden) parameters $\vec{y}$. These, along with the workload parameters $\vec{w}$, determine the desired objective function ϕ subject to some constraints. That is,

$$CP = \{\vec{x}, \vec{y}\} \quad (1)$$

$$\phi = f(\vec{w}, CP) \quad (2)$$

$$g_i(\vec{w}, CP) \geq 0 \quad i = 1, 2, \dots, K \quad (3)$$

where $f()$ could represent performance or cost, and $g_i()$ is the i th constraint involving the workload and configuration parameters. The functions $f()$ and $g_i()$ ’s are usually quite complex and may not be expressible explicitly. Ref [20] makes a similar point. It is also worth noticing that the configuration space Ω is often discrete, with intermediate values being practically infeasible, even if they are conceptually meaningful. For example, if the memory modules for the systems at hand have a minimum granularity of 16GB, an installed memory of 24GB is infeasible. Thus, defining continuous or differentiable extensions of the functions $f()$ and $g_i()$ is neither straightforward, nor meaningful. This rules out the direct use of (continuous space) tunneling structures. Also, while one could estimate the local gradient by evaluating the functions at nearby feasible points, the value of the local search is less clear.

A. Basic Approach

Fig. 3 shows the overall scheme explored in this paper. Given a set of CPs, the first step is to define an “oracle,” or a model for the **forward problem** of performance prediction based on the settings of CPs. As discussed in [21], statistical ML techniques work quite well for this. Our earlier study [15, 21] shows that typical machine learning methods are not suitable for the **backward problem** of recommending configurations that meet specific criteria (e.g., performance). We explain this deficiency with the results obtained while trying to predict the cache size of ES.

Fig. 2 shows the cache prediction efficiency for ES data-set (explained in section V-A). In the figure, the primary (left side) y-axis labels show the different sample sizes of *train* vs. *test*, the X-axis indicates the different test cases, and their prediction accuracy is shown on the secondary y-axis (right side). The graph shows that the cache prediction accuracy is *barely* 75% for various test cases. The reason for the higher error is self-explanatory by the nature of complex multi-label multi-class parameter prediction with limited training data set [21].

In addition, it should be noted that the above results do not include simultaneously predicting multiple inter-dependent CPs or use of additional conditions such as optimizing the cost of cache or performance constraints.

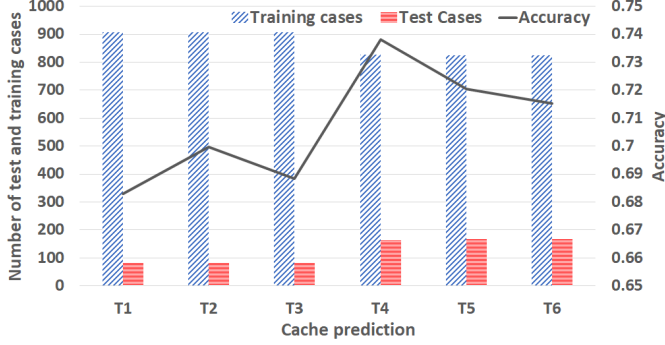


Fig. 2: Efficiency Metrics for Configuration Prediction [21]

We overcome the aforementioned problem by utilizing domain knowledge, which is vital to making reasonable conclusions from limited data. Despite the risks associated with using domain knowledge, we reduce these risks by asking domain experts only for the nature of the behavior rather than its numerical parameters, which are still determined by the limited available data. Principal Component Analysis (PCA) or similar analysis techniques can be used as shown in Fig. 3b to utilize the corresponding training data in order to determine the relative importance of various parameters. It assists both in *confirming* and applying the domain knowledge, such as (i) various relationships, e.g., higher CPU speeds require lower memory latencies, (ii) generic rules of thumb, e.g., an additional 64MB of memory per additional VDI client, (iii) system-specific ones that have been observed or (iv) can be deduced from the training data. It is important to underscore the importance of domain knowledge here since a blind application of these techniques is likely to yield incorrect or misleading results.

The problem to be solved is now to select a configuration $\vec{CP}$, i.e., $\vec{x}$ a few user-settable configurations that provide a performance p above the user desired lower bound performance p_u while minimizing the cost of the solution. The choice of parameter values in $\vec{x}$ can directly or indirectly affect the objective function, i.e., ω cost of the configuration. This is shown as $z()$ depicting the “Backward Problem” in Fig. 3a.

We now define the constraint as the desired performance p_u :

$$p = \phi \geq p_u \quad (4)$$

where p is the expected performance from configuration $\vec{x}$. The objective is to find such a configuration $\vec{x}$ at a minimum cost.

$$\min(g(\vec{x})) \quad (5)$$

The “cost” of a configuration can represent a user’s desired metric, such as the deployment cost, power consumption, cooling requirements, etc. Data for cost function can be derived from vendor specification for hardware server and allocated resources (e.g., disk capacity) or other suitable

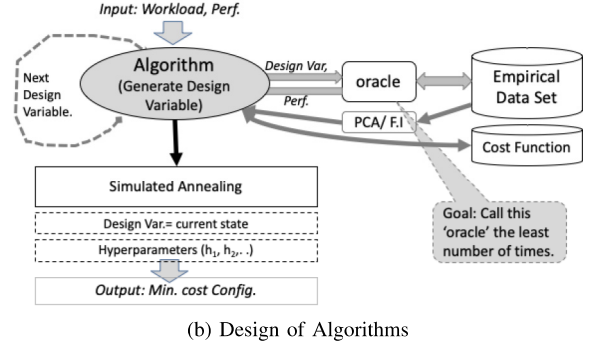
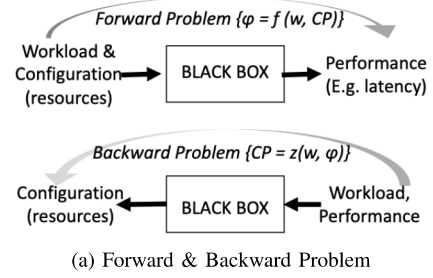


Fig. 3: Basic Approach: Design of Solution.

function $g()$. For example, k^{th} configuration $\vec{x}_k$ for some choice of parameters such as number of CPU cores, core speed, memory bandwidth, IO bandwidth, storage capacity, etc. has a cost $g(\vec{x}_k)$.

Given the non-convex and non-monotonic influences of various parameters, the use of combinatorial optimization is natural for solving the *backward problem* (shown as $z()$ in Fig. 3a). In general, this optimization could be either *deterministic* or *stochastic*, where the latter allows for uncertainty in the objective function. Although our interest is in the deterministic case, uncertainties arise naturally in real-world problems (e.g., the cost of the solution better described by a distribution rather than a single value). In the stochastic case, the objective is generally to use a statistical measure (e.g., expected value) so that essentially the same methods apply in both cases.

All stochastic methods explore a sequence of next states to find a better solution with different techniques to make a trade-off between the expense of exploration (i.e., number of iterations, cost of evaluation) and the quality of the solution. For the latter, the algorithm must necessarily consider states where the objective function is worse than the optimum found so far, which means that a monotonic convergence is generally impossible.

Because of their stochastic nature, these algorithms cannot guarantee any measure of the quality of the solution, although the solution quality should generally improve as the number of iterations of the algorithm is increased. However, note that again, as a result of the stochastic search, the entire state space may not be searched even with an infinite number of iterations. Furthermore, with the same number of iterations, each run could potentially produce a different solution.

Many comparative studies have been put forward to demon-

strate the efficiency of a particular method [22]. Therefore, the focus of our paper is *not* to compare algorithms; rather, we show that embedding domain knowledge into stochastic methods can help the algorithm to converge into a solution faster than an uninformed algorithm. Our work focused on studying algorithms for their convergence speed and their capacity to find good objective function minima. On this front, we adapted our solution to Simulated Annealing as shown in Table. II, which we will explain next. We first discuss the notion of tunneling as referred to in section III-B and then explain the modifications to the SA algorithm.

B. Emulating Tunneling

To apply the Tunneling approach to our context, we first seek the local minima (step 1), followed by an intelligent perturbation based on the sensitivity to take us further away from the current minima (step 2). To prepare for this, we first determine the relative ranking of each CP in terms of its influence on the objective function and the constraints. A pure data-driven method for doing this is the standard PCA and should work well, assuming a sufficiently large data-set. Starting with the most dominant parameter, we form a group by pulling in other parameters known to be related to it based on domain knowledge. We then start with the next most dominant parameter and form the next group until we have covered all the parameters. Then the “smart” tunneling can be described as follows:

- 1) We approach the local minima by considering the variation concerning the first group leader. The gradient needs to be determined by taking a few samples.
- 2) The tunneling phase then perturbs each group leader in proportion to the PCA metric and adjusts all others in the group based on the known relationships.

A stochastic optimization process works by randomly jumping from the current state x_i to a new state x_{i+1} based on some probability factor ρ , with an aim to find local minima x^* that minimizes the objective function $f(x)$. We apply the above approach with tunneling by enabling the stochastic algorithm to circumvent the local minima points and rapidly move from an area of shallow minima (point (a) in Fig. 1) to a region of deeper minima (point (d) in the same figure), thereby allowing for faster exploration of solution space and faster convergence to a good solution. With a configuration problem at hand, as our objective function cannot be characterized by a direct analytical function, we use the performance prediction oracle (a black box) function as an objective function.

We explore ways to tunnel through the shallow minima (point (B) in the same figure) and avoid the slow dynamics of the complex objective function. Such tunneling mechanisms should invariably use the domain knowledge to intelligently jump the local barriers and avoid uninteresting space (i.e., avoid points (D and E) in the same figure). We group the design variables together as a discrete space tunneling mechanism to aid the algorithm from being trapped at local minima and jumping through or tunneling out of the minimum.

We present the details in the next sections in context to the algorithm we explore.

C. Generic Design Approach Summarized

The approach described above can be generalized independent of the data-set and domain as follows:

- 1) Run experiments to collect the data with relevant configurable parameters and observable outcomes.
- 2) In the absence of a clear analytical function to describe the relationship between the configurable parameters to the outcome use a suitable ML model to design an “oracle” as a prediction engine (a.k.a PIM).
- 3) Use PCA metrics from the data to determine the relative importance of design variables and group attributes based on domain knowledge to avoid exploring undesired spaces.
- 4) Use the problem knowledge to define the objective function (e.g Eq. 5) and the required constraint function (e.g. Eq. 4).
- 5) To explore new design states, use PCA metrics as probability factors and perturb the variables in groups.
- 6) Verify new state satisfies constrain using ML-based oracle as a tool and accept/reject the current design state.

D. Modified Simulated Annealing (mSA)

SA is a general probabilistic local search algorithm generally used to solve difficult optimization problems. The pseudocode [23] for generic SA (gSA) is given in Algorithm 1. In SA, a state refers to a set of design variables, and a neighboring state refers to a set of values relatively closer to current design variables. In SA, entropy is represented as the cost function that has to be minimized. An acceptable state is a solution to the problem at hand.

Algorithm 1: Pseudocode for SA [23]

```

1 initialize(temperature T, random starting point) ;
2 for i in T do
3   p = select_point_from_neighborhood(i) ■ ;
4   currentCost = compute_currentCost_at(p) ;
5    $\delta$  = currentCost - previousCost ;
6   if  $\delta \leq 0$  then
7     accept_neighbor_point(p) ;
8   else
9     accept with probability  $\exp(-\delta/T)$  ♦ ;
10  T =  $\beta * T$  ■ ;
```

The SA method has been widely used since the cost function can be easy to put into practice [23]. Our SA algorithm uses design variables from the configuration ($\vec{x}$) to represent the state. The entropy of the system is defined as the cost of the current state (i.e., cost of the configuration $g(\vec{x})$). The gSA steps in Algorithm 1 can be summarized as follows: (i) we first start with an initial annealing temperature (T_0) and a random design state (line 1), (ii) we search for the next state depending on the annealing temperature T_k and a random

distribution (line 4), (iii) we compute the difference in entropy (δ) between the current state and the previous state (line 5), and (iv) probabilistically accept the current state depending on Boltzmann probability factor (line 6...9). In line 9, if the current solution is accepted, we apply tunneling logic to search for a better local minima. The annealing scheme is defined in line 10. The algorithm stops after reaching a defined cooling temperature (line 2).

Our solution is based on very fast simulated annealing (VFSA) presented by Xu [24], which enhances both the annealing temperature (line 10) and the perturbation model (line 4). Lee [25] and others have discussed VFSA in detail and show the advantages of VFSA over SA. To speed up the convergence rate of SA, VFSA uses the Cauchy distribution function as the perturbation [25]. This can realize a narrower search as the iterative solution approaches an optimum solution, which accelerates the convergence speed [24]. Our enhancements to the basic generic algorithm are illustrated in Table. II.

TABLE I: Notations used in functions of gSA and mSA

Notation	Description
k	Current iteration
n	Number of design variables
T_0	Initial annealing temperature
α	Damping coefficient ($0 < \alpha < 1$)
μ and U	Uniform random variables between 0 and 1
$(B_j - A_j)$	Range of j^{th} design variable ($1 \leq j \leq n$)
$\vec{x}_i$	Configuration (design variables) at i^{th} state
c_i	Configuration cost at i^{th} state
p_i	Predicted performance of configuration $\vec{x}_i$ at i^{th} state
p_u	User given performance

TABLE II: Very Fast Simulated Annealing Functions

Entity	gSA	mSA
Annealing temp T_k	$T_0 * \exp(-\alpha(k-1)^{1/n})$	
Entropy change δ	$\exp(\frac{c_i - c_{i-1}}{T_k})$	$(\frac{c_i - c_{i-1}}{T_k})^3$
Acceptance Probability ρ	1, if $p_i \geq p_u$ 1, if $\delta \leq U(0, 1)$ 0, otherwise	1, if $p_i \geq p_u \& c_i \leq c_{i-1}$ 1, if $\delta \leq U(0, 1)$ 0, otherwise
Perturbation Model ζ_j	$T_k(\mu - 0.5) \left[\left(1 + \frac{1}{T_k}\right)^{ 2\mu-1 } - 1 \right] (B_j - A_j)$	
Selecting neighboring state (s_{i+1})	$\begin{cases} \text{random_new_state}(), & \text{if } \rho = 1 \\ \zeta_j, & \text{otherwise} \end{cases}$	
Design Variables	Individually varied	Varied as a group

We discuss the supporting functions of gSA and mSA in Table II using the notations described in Table I. mSA uses the annealing temperature T_k and Cauchy distribution perturbation model ζ from VFSA (see Table II). For acceptance probability ρ , mSA makes a slight modification to accommodate the case where the next solution has the same performance but lower cost. If the acceptance probability for the current state is 1, a new random state is chosen to avoid getting stuck in local minima; else, a new state in the neighborhood is chosen.

V. WORKLOADS FOR CONFIGURATION MODELING

We applied the above design to several publicly available data-sets as shown in Table. III. As these are published data-sets from various studies, we have no control over the data collected; experiments run, CPs, variability, etc. We explain these data-sets briefly in the following.

A. Cloud/Edge Storage Data-set

Edge computing [8] offers computation and storage close to where data is produced. It has recently emerged as a way to reduce latency and limit the load that is carried to higher layers of the infrastructure hierarchy[8, 21]. The Edge Storage (ES) is deployed at a branch office or remote location and has access to a rather limited local compute/storage, and is connected to a Cloud data center over the Internet. ES essentially uses local storage as a cache for the remote Cloud storage to bridge the gap between the demand for low-latency/high-throughput local access and the connection to the Cloud that may experience unpredictable bandwidth and latency [8, 21]. Resource allocation for the ES is challenging since it has to dynamically adapt to the requirements of the end-users' applications and end-devices (e.g., cars, drones), taking into account the resources' characteristics in terms of processing latency, capacity, security, location, and cost.

We address the configuration selection problem in Edge platforms using data published in Ref. [21]¹. We examine the performance of various configuration settings on the ES server for different data-transfer request sizes, i.e., different workloads such as Health monitors, MRI/CT Scans, Mammography images, etc. The observed performance of the ES, denoted as ϕ , is influenced by its CPs ($\vec{x}$ in Eq. 2) and the given workload ($\vec{w}$). A full description of the ES system, various CPs influencing the behavior, workloads, etc., is given in Ref. [21]. We represent the ES configuration as a combination of required compute and storage resource - number of cores nc , core speed cs , memory capacity mc , memory bandwidth bw , and disk IO rate di . Workload can be defined by the request arrival rate ar , request size rs , and the metadata size ms . Now, in the context of an ES system being studied, we can express Eq. 2 more clearly as: $\vec{x} = \{nc, cs, mc, bw, di\}$ and $\vec{w} = \{ar, rs, ms\}$. The research question would then be to find the suitable values for $\vec{x}$ for a given $\vec{w}$ that satisfies the given constraint (Eq. 4) at a minimum cost (Eq. 5).

B. Configuration Modeling for BitBrains Data Center

The other publicly available data-set used in this research is BitBrains² workload trace [26] containing the performance logs of 1,750 VMs from a distributed data center from BitBrains, which are collected over 5000 cores and 5 million CPU hours accumulated over four months. This data-set (see Table III) provides specialized interactive services and batch processing workloads in a Cloud environment for managed

¹[ES] https://www.kkant.net/config_traces/CHIproject

²[BB] <http://gwa.ewi.tudelft.nl/datasets/gwa-t-12-bitbrains> (RND500)

TABLE III: CPs, Workload and Output of the Data-Sets.

Data-Set	Domain	CPs $\vec{x}$	Workload Characteristics $\vec{w}$	Output ϕ
ES [21]	Cloud Storage	No. of Cores, Core Speed, Memory Capacity, Memory Bandwidth, Disk IO Rate	Request Arrival Rate, Request Size, Metadata Size	Performance
BB [26]	Virtual Machines	No. of Cores, Core Speed, Memory Capacity, Network Data Rcvd., Network Data Transmit	Disk Read Throughput, Disk Write Throughput	CPU Usage (%)
Apache [12]	Web Server	Base, KeepAlive, Handle, HostnameLookups, EnableSendfile, FollowSymLinks, AccessLog, ExtendedStatus, InMemorI	N/A	Performance
SQLite server [27]	SQL Server	SetCacheSize, StandardCacheSize, LowerCacheSize, HigherCacheSize, LockingMode, ExclusiveLock, NormalLockingMode, PageSize, StandardPageSize, LowerPageSize, HigherPageSize, HighestPageSize.	N/A	Performance
Berkeley DB C [28]	Embedded database	havecrypto, havehash, havereplication, haveverifl, havevsequence, havestatistics, diagnostic, pagesize, ps1k, ps4k, ps8k, ps16k, ps32k, cachesize, cs32mb, cs16mb, cs64mb, cs512mb	N/A	Performance
MIT [29]	HPC Env.	CPU Frequency, Resident Memory Size, Virtual Memory Size	Amount of Data ReadWrite (MB)	CPU Util. (%)

hosting and business computation, including leading banks, insurance companies, credit card operators, etc.

Workloads for Evaluating BB: Iosup and Et al. [26] conduct a comprehensive characterization of both requested and actually used resources, using data corresponding to CPU, memory, disk, and network resources. The initial configuration *CP* of each VM present in these traces is characterized by the attributes shown in Table III. With limited knowledge of the details of the data-set, we formulate the BB VM configuration as a combination of required compute, storage, and network resource - number of cores nc , memory capacity mc , network receive bandwidth nwr , and network transmit bandwidth nwt . We characterize the workload as the load on the storage disks as read request rate $dskrd$ and write request rate $dskwr$ and observed behavior (ϕ) as the CPU usage (%).

We can now represent the configuration problem as selecting the right combination of configuration values (i.e. resources $\vec{x} = \{nc, mc, nwr, nwt\}$) for a given workload ($\vec{w} = \{dskrd, dskwr\}$) to satisfy the user defined conditions (Eq. 4 and Eq. 5).

C. Configuration Modeling for Enterprise Data-set

We evaluate our work using three Cloud applications (Apache, SQLite, and Berkeley DB) from Ref. [12, 27, 28, 30], by commonly grouping them as Enterprise Data-sets^{3,4} (EE). Apache HTTP Server is a highly popular web server. Xu et al. [31] report that the Apache server has more than 550 parameters, and many of these parameters have dependencies and correlations, which further complicates the configuration problem we address here. Reference [12, 30] narrows the CPs down to only nine CPs as given in Table. III. Berkeley DB (C) [28] is an embedded key-value-based database library that provides scalable high performance database management services to applications. SQLite [27] is the most popular lightweight relational database management system used by

several browsers and operating systems as an embedded database. In producing the data-set, Nair [30] stresses the application to maximum workload and observes performance data for various configurations. We use their data-set with 18 CPs for BDBC and 29 CPs for SQLite data-sets to evaluate our proposed solution. We demonstrate the efficacy of our solution and its application in data-sets with large configuration spaces Ω .

D. Configuration Modeling for MIT Cloud Data Set

MIT published a rich data-set⁵ from their Supercloud petascale cluster [29] running various HPC workloads. This huge (2TB) data set contains time-series data of the scheduler, file system, compute nodes, CPU, GPU, and sensor data from physical monitoring of the facility housing the cluster itself. We used the 2nd partition data-set with 480 CPU nodes (2x24-core Intel Xeon Platinum 8260 processor), each with 192GB of RAM and a Lustre high-performance parallel file system running on a 3-petabyte Cray L300 parallel storage array (See Table.IV Slurm Time Series Data). The data attributes in this work comprise CPU frequency, residual memory, virtual memory size, CPU utilization, disk IO, etc. We now propose the resource allocation (configuration) question as: "finding the design variables (node number, VM-Size, CPU Frequency, RSS Memory Size) for HPC workload (given as ReadMB and WriteMB) for a required CPU Utilization (Constraint)." We demonstrate the efficacy of the proposed solution for large data-sets.

We refer readers to the detailed literature at Ref. [12, 27, 28, 29, 30] for a full description of the data-set(s). With our problem at hand, the problem (Eq. 2 and Eq. 3) reduces to finding the best configuration ($\vec{x}$) for a given workload ($\vec{w}$) and a user given performance (p_u) at minimum possible cost.

³[EE] https://github.com/ai-se/Reimplement/tree/cleaned_version/Data

⁴[EE] <https://goo.gl/689Dve> (RawData/PopulationArchives)

⁵[MIT] <https://dcc.mit.edu>

VI. EVALUATION

A. Metrics for Evaluation

Using the above data-sets, we evaluate the efficacy of our solution in finding a satisfying solution with two key metrics: **(M1)** the number of calls to the performance function (a.k.a oracle), and **(M2)** the minimum cost of the selected configuration (i.e., Eq. 5). Metric M1 is important as it relates to how fast the algorithm can *find an optimal set of parameters* from the vast configuration space Ω . Metric M2 may refer to the monetary cost (\$\$) of the selected physical configuration, resource consumption of the selected configuration in a virtualized environment, or some other attribute (e.g., energy consumption, provisioning difficulty, etc.). Naturally, metric M2 is generally much more important than M1, but two situations make M1 very important: (a) frequent changes in configuration, which is quite common in current Clouds; for example, Facebook [7] reports thousands of configuration changes per day and (b) models (oracles) with long running times.

We executed 100s of test cases across all the data-sets, each test case T_i refers to a unique combination of $\vec{w}_j$ and ϕ_k in the data-set. We discuss the evaluation results using M1 and M2 metrics w.r.t the three approaches discussed above, i.e. (a) Generic Simulated Annealing (gSA), (b) Modified Simulated Annealing (mSA), and (c) Modified Simulated Annealing with Tunneling (mSA(T)).

B. Performance Oracle

The efficiency of ML algorithms depends on a variety of factors, including the input attributes and hyper-parameters (e.g., regularization parameters, learning rate, etc.), and it is generally not possible to characterize which algorithm works the best in a given situation [32]. Therefore, we tried several models and ultimately settled on Logistic Regression, as it consistently performed well and beat others in most workloads. An extensive analysis (e.g., k-fold validation) of the model ensured that it did not suffer from under-fit or over-fit.

C. Using Domain Knowledge to Group Attributes

We incorporate domain knowledge in the algorithm by dividing the design variables into groups based on their level of interdependencies. That is, the design variables within a group show strong interdependence and thus should be set collectively, but the settings across groups can be done independently. In theory, such grouping can be done purely in a data-driven manner (e.g., by using clustering techniques), but this is likely to result in spurious groups unless we have a large amount of data covering full ranges of various CPs and the clustering algorithm does not result in anomalies. The value of domain knowledge is to do a suitable grouping either entirely manually or by coercing the clustering algorithm to prefer certain groupings over others.

In any configuration context, we are likely to have several generic and usage-specific insights into the system. For example, in computing infrastructure, a faster CPU must be paired with a faster DRAM; else, the CPU will simply stall, waiting

TABLE IV: Grouping Design Variables for Various Data-Sets

Data-Set	Group	Design Attribute Pairs
ES	Group G1	Number of Cores, Memory capacity
	Group G2	Core speed, Memory bandwidth
	Independently varied	Data cache, Disk IO rate
BB	Group G1	Number of Cores, Memory capacity
	Independently varied	CPU capacity, Network data transmit, Network data received
Apache BDBC SQLite	Independently varied	All CPs
MIT	Group G1	Virtual memory used by process, Resident memory footprint set size
	Independently varied	CPU clock frequency

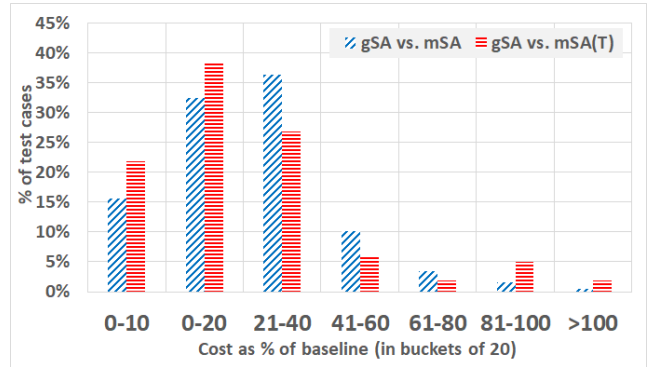


Fig. 4: Histogram of Cost for ES

for the memory. A faster disk is also important, but much less so, since the IOs involve a context switch, whereas the memory access does not. Similarly, more CPU cores doing independent work will likely need more memory, and for workloads involving remote IO, both network and IO speeds must increase in tandem. *Grouping* of CPs based on such insights avoids exploration of states that are unlikely to be useful and thus is expected to both speed up the convergence and lead to better solutions within a given number of iterations. As shown in Table IV, ES, BB, and MIT data-sets are grouped according to the interdependence between design variables. Since the trace description doesn't give much information

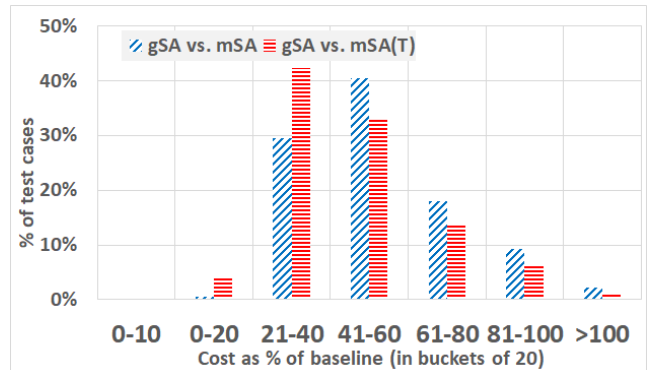


Fig. 5: Histogram of Cost for MIT

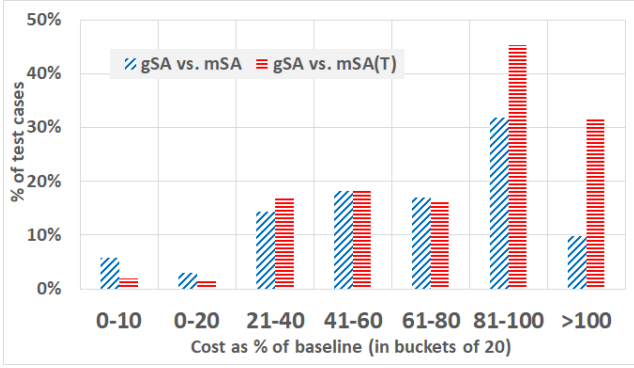


Fig. 6: Histogram of Cost for BB

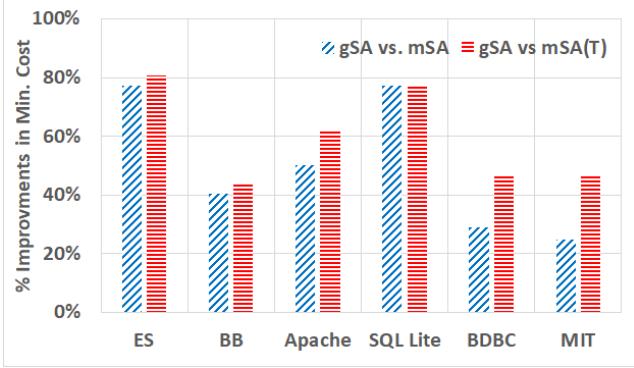


Fig. 7: Percent Improvement in Solution Cost (M2)

about the configuration or workload, we have grouped the other data-sets independently (Apache, BDBC, and SQLLite).

D. Efficacy of the algorithms (gSA, mSA and mSA(T))

We use the gSA algorithm as the baseline, as it presents the naive (or *uninformed*) stochastic method of searching a wide configuration space for a set of suitable CPs.

In our evaluation for metric M2, the minimum cost of the solutions with mSA(T) was much less (hence better/desired) than the solution found by mSA or gSA (shown in Fig 7).

Detailed results are shown in Fig. 4, 5, 6 for a few data-sets (ES, MIT, and BB, respectively), where we plot the improvement in the solution cost provided by our two

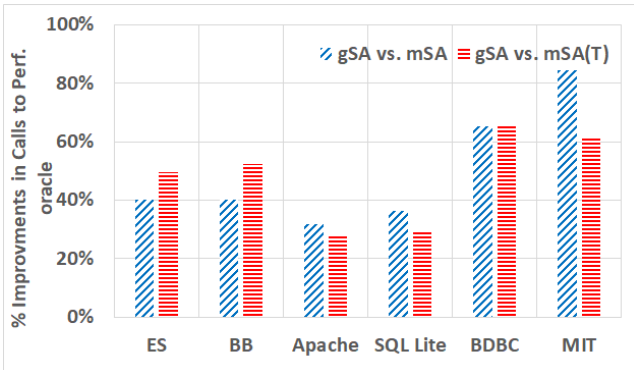


Fig. 8: Percent Improvement in #calls to Oracle (M1)

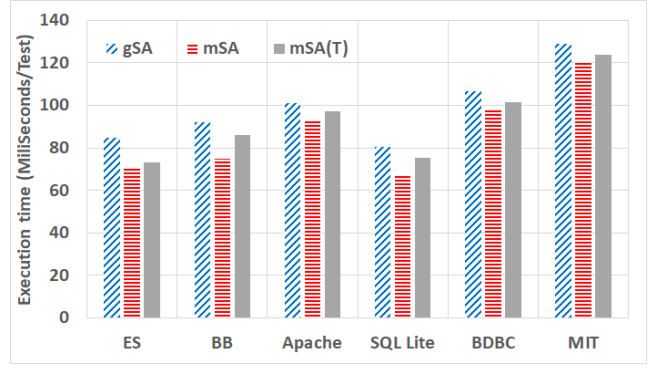


Fig. 9: Execution Time for Different Data-sets

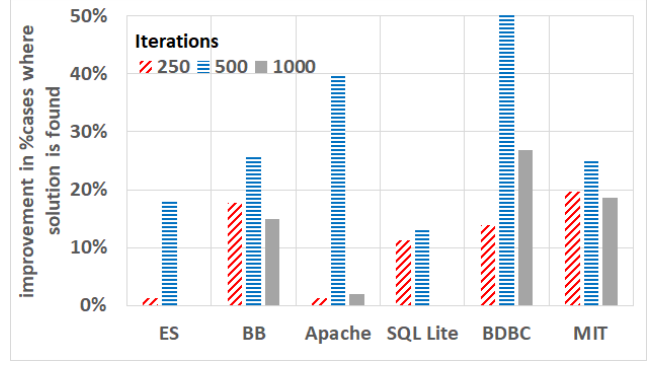


Fig. 10: Percent Improvement in % cases that provide a solution (gSA vs. mSA(T))

algorithms (mSA and mSA(T)) over the baseline uninformed algorithm gSA. Here the X-axis refers to the cost of mSA and mSA(T) divided by the solution cost of gSA, and expressed as a percentage. That is, the buckets for $\leq 100\%$ represent an improvement over gSA, and $> 100\%$ represent a degradation. The y-axis is the count of test cases whose solution cost falls into that bucket – again, normalized so that it all adds up to 100% of the test cases. These charts are produced by considering 100s of test cases and thus represent an extensive exploration of the configuration space.

Fig. 4 in ES results shows that in the 1st group (0-10%), mSA achieved the solution with $\leq 10\%$ of the baseline cost for 15% of test cases; and mSA(T) further improved this to 22% of test cases. The maximum improvement observed was in a few cases where the cost of mSA(T) was only 2% of the cost provided by gSA! Fig. 5 in MIT results also shows that in the 2nd group, mSA achieved the solution with $\leq 20\%$ of the baseline cost for 0% of test cases, but mSA(T) improved on that to 4% of test cases.

Similarly, in Fig. 6 in BB results, in the 6th group (81 to 100%), mSA cost was about 61-80% of the cost of the baseline in about 17% test-cases; and mSA(T) improved this in about 22% of test-cases. The final group (> 100) in all the sub-graphs shows cases where gSA cost was better than mSA or mSA(T); however, these cases were small in the case of an ES and MIT. With BB, the evaluation showed that mSA(T) failed to get minimum cost in about 30% of the cases (compared

to gSA). Note that no stochastic algorithm can provide a universally better result in all cases because of the inherent randomness in the way the states are explored.

Fig. 7 shows that both of our algorithms (mSA and mSA(T)) provide a lower-cost solution in comparison with gSA by $\sim 50\%$ and $\sim 60\%$, respectively. The solution cost for the ES, BB, Apache, SQLite, JDBC, and MIT data-sets is improved in range 25% – 77% by mSA, and in range 47% – 81% by mSA(T).

Fig. 8 shows the improvement in the number of calls to the performance Oracle. It is again seen that mSA/mSA(T) consult the performance Oracle significantly fewer times, which can be significant if running the performance model becomes expensive.

Fig. 9 shows the run-time of the three algorithms (gSA, mSA, and mSA(T)) for various test-cases. The figure shows the execution time in millisecond/test-case for various data-sets for max 500 iterations. The results show that mSA and mSA(T) beat gSA even here, although by rather small amounts of 12% and 6% respectively. However, the more significant observation is that the run-times are fairly small in all cases, which is essential for frequent configuration changes. Thus we expect that even with much more complex situations, the mechanism would be able to determine the suitable configuration rather quickly, thereby satisfying the needs of DevOps-related auto-configuration.

Finally, Fig. 10 compares the ability to find a solution within a certain number of iterations. For this, we choose 250, 500, and 1000 as the limits on #iterations. The key reason to consider three different values is to ascertain that the results are not an artifact for a given iteration count. For 1000 iterations, gSA is successful in only 75% of the cases, but mSA/mSA(T) are successful in 98% and 97% of the cases, respectively. Fewer iterations show an even better improvement of mSA/mSA(T) over gSA, although the absolute success rate will surely decrease with #iterations. All in all, unlike gSA, mSA/mSA(T) succeed in finding the solution in almost all cases, find solutions of significantly lower cost (see Fig 7), and even run somewhat faster (see Fig. 9).

VII. LESSONS LEARNT

During our research on tunneling to find a better local minima, we discovered that computing the gradient from the previous two data points and then sliding along that gradient offered a better objective solution than other alternatives, such as hill climbing and binary search.

One of the persistent issues in our configuration management (CM) research has been the lack of workload traces with carefully documented configuration information. With only a few exceptions, publicly available workload traces fail to provide any configuration information and are thus useless for CM research. We call upon the community to pay serious attention to recording the configuration when generating traces, as this would be crucial in tackling the crucial problem of proper

configuration and detecting misconfigurations in large-scale systems.

Owing to the complexity of performance as a function of the CPs, we have used an ML model for the forward problem. In general, a system may have numerous CPs, and training a comprehensive ML model may need a lot of data that may not be available or may not be of adequate quality. Even more important is the pitfall of blind faith in data. We have demonstrated that simply throwing in all the CPs in the model hurts its accuracy rather than improving it. This further emphasizes the importance of the “domain knowledge” in pre-selecting the most relevant parameters for training the forward model and intelligently searching the state space for the backward problem.

VIII. CONCLUSIONS

In this paper, we presented an efficient methodology to recommend optimal configurations for the emerging DevOps environments with implicit performance function and cost constraints. We proposed an improved metaheuristics-based approach enhanced by both domain knowledge and smart tunneling techniques. We applied the technique to several real-world traces from various publicly available Cloud environments where the configuration information was included in the data-set. The results show that the proposed mechanism outperforms a standard uninformed approach by 44-81% (depending on the domain and data-set) in terms of the cost of the solution, converges faster by 28-65%, and still runs somewhat faster.

The key reasons for such performance gains include the following. First, we compute entropy as a quadratic function to give us a wider choice of acceptance which is able to avoid better getting stuck at local minima. Second, we intelligently group the attributes, which avoids exploring unnecessary portions of the search space. Third, by adding the tunneling logic, the algorithm avoids *jumping out of the local minima too quickly*; instead, it explores a few additional states *closer to current local minima*. We also show that the proposed approach can determine desired configuration very quickly, which is essential in highly dynamic DevOps and microservices environments.

Our future work will examine how domain knowledge may be extracted semi-automatically from the best practice specifications or by using human-experts as a part of the interactive optimization technique [36]. We will also explore how our technique can be applied to more dynamic environments, such as Kubernetes environments with rapidly changing container placements, and possibly benefit from the additional data collected from performance and systems logs.

REFERENCES

- [1] S. Sondur, K. Kant, “Performance Health Index for Complex Cyber Infrastructures,” in *ACM Trans. Model. Perform. Eval. Comput. Syst.*, ACM, 2022,.
- [2] M. Shahin, “Architecting for devops and continuous deployment,” in *Proc. of ASWE*, ACM, 2015, pp. 147–148.

- [3] S. Newman, *Building microservices – Designing Fine Grained Systems, 2nd Edition*, "O'Reilly Media", 2021.
- [4] J. Humble and D. Farley, *Continuous delivery: reliable software releases through build, test, and deployment automation*, Pearson Education, 2010.
- [5] F. Connolly, "Production operations – the last mile of a devops strategy," *LMC Report*, Mar 2014.
- [6] W. Cappelli, "Causal analysis makes availability and performance data actionable," *Gartner Report*, Oct 2015.
- [7] C. Tang, T. Kooburat, P. Venkatachalam, A. Chander, Z. Wen, A. Narayanan, P. Dowell, and R. Karl, "Holistic Configuration Management at Facebook," in *Proc. of SOSP*, ACM, 2015, pp. 328–343.
- [8] P. Soumplis, P. Kokkinos, A. Kretsis et.al., "Resource Allocation Challenges in the Cloud and Edge Continuum," in *Advances in Computing, Informatics, Networking and Cybersecurity*, Springer, 2022, pp. 443–464.
- [9] S. Singh and et al., "Cloud resource provisioning: survey, status and future research directions," *Knowledge and Information Systems*, vol. 49, no. 3, pp. 1005–1069, 2016.
- [10] S. S. Gill and Et al., "Chopper: an intelligent qos-aware autonomic resource management approach for cloud computing," *Cluster Computing*, vol. 21, 2018.
- [11] M. Wang and Et al., "Storage device performance prediction with CART models," in *MASCOTS*, 2004.
- [12] N. Siegmund, A. Grebhahn, S. Apel, and C. Kastner, "Performance-influence models for highly configurable systems," in *Proc. of ESEC/FSE*, 2015.
- [13] M. Satyanarayanan, "The emergence of edge computing," *Computer*, vol. 50, no. 1, pp. 30–39, 2017.
- [14] Y. Sfakianakis, M. Marazakis, and A. Bilas, "Skynet: Performance-driven Resource Management for Dynamic Workloads," in *IEEE CLOUD*, 2021.
- [15] S. Sondur, K. Kant, S. Vucetic, and B. Byers, "Storage on the edge: Evaluating cloud backed edge storage in cyberphysical systems," in *IEEE Intl Conf. on MASS*, 2019, pp. 362–370.
- [16] K. Hussain, M. N. M. Salleh, S. Cheng, and Y. Shi, "Metaheuristic research: a comprehensive survey," *Artificial Intelligence Review*, 2019.
- [17] A. V. Levy and A. Montalvo, "The tunneling algorithm for the global minimization of functions," *SIAM J. Sci. and Stat. Comput.*, vol. 6, no. 1, p. 15–29, 1985.
- [18] F. Schoen, "Stochastic techniques for global optimization: A survey of recent advances," *J. of Global Optimization*, vol. 1, no. 3, pp. 207–228, 1991.
- [19] Z. Li and Y. Yang, "A modified tunnelling algorithm for global minimization with box constrained," in 5th Intl. conf on Computational Sciences & Optimization, 2012, pp. 423–427.
- [20] O. Alipourfard, H. H. Liu, J. Chen et. al., "Cherrypick: Adaptively unearthing the best cloud configurations for big data analytics," *Proc. of NSDI*, 2017, pp. 469–482.
- [21] S. Sondur and K. Kant, "Towards automated configuration of cloud storage gateways: A data driven approach," in Intl. conf on Cloud Computing, Springer, 2019, pp. 192–207.
- [22] M. Barrette and Et al., "Statistical multi-comparison of evolutionary algorithms," *Bioinspired Optimizaiton Methods and their Applications*, 2008.
- [23] K. P. Ferentinos, K. G. Arvanitis, and N. Sigrimis, "Heuristic optimization methods for motion planning of autonomous agricultural vehicles," *J. Global Optimization*, vol. 23, no. 2, pp. 155–170, 2002.
- [24] Y. Xu, Q. Ye, and G. Meng, "Hybrid phase retrieval algorithm based on modified very fast simulated annealing," *Intl. J. of Microwave & Wireless Technologies*, vol. 10, no. 9, pp. 1072–1080, 2018.
- [25] C.-Y. Lee, "Fast simulated annealing with a multivariate cauchy distribution and the configuration's initial temperature," *J. Korean Physical Society*, 2015.
- [26] A. Iosup and Et al., "The grid workloads archive," *Future Generation Computer Systems*, vol. 24, 2008.
- [27] V. Nair, Z. Yu, T. Menzies, N. Siegmund, and S. Apel, "Finding faster configurations using flash," *IEEE Transactions on Software Engineering*, 2018.
- [28] V. Nair, T. Menzies, N. Siegmund, and S. Apel, "Faster discovery of faster system configurations with spectral learning," *Automated Software Engineering*, 2018.
- [29] S. Samsi, , and Et al., "The MIT Supercloud Dataset," in *IEEE HPEC*, 2021.
- [30] V. Nair, T. Menzies, N. Siegmund, and Apel, "Using bad learners to find good configurations," in *Proc. of Joint Meeting on Foundations of Software Eng.*, 2017.
- [31] T. Xu and Et al., "Hey, you have given me too many knobs!" in *Proc. of Joint Meeting on Foundations of Software Eng.*, 2015.
- [32] M. S. Sorower, "A literature survey on algorithms for multi-label learning," *OSU, Corvallis*, vol. 18, 2010.
- [33] Kerr, Alexander and Mullen, Kieran, "A comparison of genetic algorithms and simulated annealing in maximizing the thermal conductance of harmonic lattices", *Computational Materials Science, Elsevier*, vol. 157, 31–36, 2019.
- [34] Legendre, Pierre and Legendre, Louis, "Numerical ecology", *Elsevier*, 2012.
- [35] Sondur, Sanjeev and Alazzawe, Anis and Kant, Krishna, "Optimal Configuration of High Performance Systems", *The 2020 International Conference on High Performance Computing & Simulation (HPCS)*, 2020.
- [36] Meignan, David and Knust, Sigrid and Frayret, Jean-Marc and Pesant, Gilles and Gaud, Nicolas, "A review and taxonomy of interactive optimization methods in operations research", *ACM Transactions on Interactive Intelligent Systems (TiIS)*, 2015.