
MetaBox-v2: A Unified Benchmark Platform for Meta-Black-Box Optimization

Zeyuan Ma¹, Yue-Jiao Gong^{1,*}, Hongshu Guo¹, Wenjie Qiu¹, Sijie Ma¹,
Hongqiao Lian¹, Jiajun Zhan¹, Kaixu Chen¹, Chen Wang¹, Zhiyang Huang¹,
Zechuan Huang¹, Guojun Peng¹, Ran Cheng², Yining Ma³

¹South China University of Technology ²Hong Kong Polytechnic University

³Massachusetts Institute of Technology

Abstract

Meta-Black-Box Optimization (MetaBBO) streamlines the automation of optimization algorithm design through meta-learning. It typically employs a bi-level structure: the meta-level policy undergoes meta-training to reduce the manual effort required in developing algorithms for low-level optimization tasks. The original MetaBox (2023) provided the first open-source framework for reinforcement learning-based single-objective MetaBBO. However, its relatively narrow scope no longer keep pace with the swift advancement in this field. In this paper, we introduce MetaBox-v2 (<https://github.com/MetaEvo/MetaBox>) as a milestone upgrade with four novel features: 1) a unified architecture supporting RL, evolutionary, and gradient-based approaches, by which we reproduce 23 up-to-date baselines; 2) efficient parallelization schemes, which reduce the training/testing time by 10–40x; 3) a comprehensive benchmark suite of 18 synthetic/realistic tasks (1900+ instances) spanning single-objective, multi-objective, multi-model, and multi-task optimization scenarios; 4) plentiful and extensible interfaces for custom analysis/visualization and integrating to external optimization tools/benchmarks. To show the utility of MetaBox-v2, we carry out a systematic case study that evaluates the built-in baselines in terms of the optimization performance, generalization ability and learning efficiency. Valuable insights are concluded from thorough and detailed analysis for practitioners and those new to the field.

1 Introduction

Black-Box-Optimization (BBO) represents challenging optimization tasks in practice. For decades, many BBO optimizers [1–4] are developed and widely discussed. A key limitation of traditional BBO optimizers is that they require human experts to design effective algorithms, which might result in design bias to adapt for novel optimization scenarios [5]. To address this, recent Meta-Black-Box Optimization (MetaBBO) [6, 7] researches propose meta-learning algorithm design policy by a bi-level framework: the meta-level policy is trained on a problem distribution to maximize the performance of low-level BBO optimizer. The trained policy is expected to generalize on unseen problems. Considering MetaBBO lacks decent benchmark, MetaBox [8] in 2023 served as the first benchmark platform for developing and evaluating MetaBBO approaches. In particular, this original version focused on a specific optimization scenario: single-objective optimization, and a specific learning paradigm: MetaBBO with reinforcement learning (MetaBBO-RL). The reason behind is that before 2023, a primary research focus in MetaBBO was exploring how to incorporate RL [9] with BBO optimizers. With 8 popular MetaBBO-RL baselines and 3 basic testsuites, MetaBox provides

*Yue-Jiao Gong is the corresponding author (gongyuejiao@gmail.com).

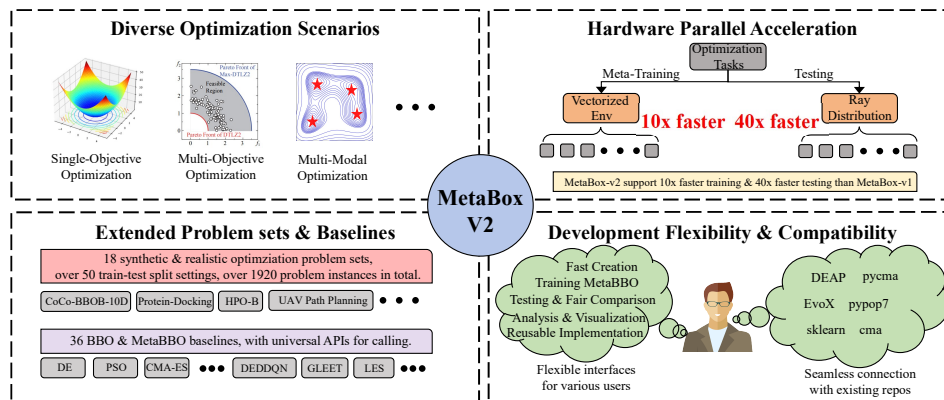


Figure 1: The four novel and user-friendly features of MetaBox-v2.

fully automatic train-test-log workflow with minimal development requirement. It has received considerable recognition in the field and gathered a MATLAB extension recently [10].

Promising as it is, MetaBBO researches grow up rapidly within the last two years. On the one hand, more and more novel ideas came out in terms of flexible learning paradigms. According to the latest survey [6], four major learning paradigms have been discussed: MetaBBO-RL, MetaBBO-SL [11–13] where supervised learning is used to train the meta-level policy, MetaBBO-NE [14, 15] where neuroevolution [16] is adopted, and MetaBBO-ICL [17–19] where LLMs with in-context learning [20] serve as meta-level policies for algorithm design. On the other hand, MetaBBO’s potential has been widely explored in diverse optimization fields such as multi-objective optimization [21], multi-modal optimization [22, 23], large-scale global optimization [24, 25], multi-task optimization [26], etc. The original MetaBox no longer keeps pace with the swift advancement in the field.

We therefore propose MetaBox-v2 through fundamental improvements that systematically address the above limitations while inheriting the benefits of the original Metabox. To summarize, this study presents the following key contributions to advancing the MetaBBO benchmark research:

1. Milestone Framework Upgrade (MetaBox-v2): The upgraded MetaBox-v2 introduces four synergistic enhancements through framework innovations, as illustrated in Figure 1.

1) Unified Integration of All Four MetaBBO Paradigms: Through redefine the algorithmic interfaces, we now propose the *MetaBBO Template* for algorithm development. It serves as the first framework capable of supporting all the four distinct MetaBBO paradigms: MetaBBO-RL, MetaBBO-SL, MetaBBO-NE, and MetaBBO-ICL. Based on the unified framework, we also extend our baseline library from 8 to 23 algorithms.

2) Efficient Parallelization: MetaBBO approaches are typically time-consuming due to the nested bi-level structure. In MetaBox-v2, we introduce two parallel schemes: vectorized optimization environment and instance-level distributed evaluation, to accelerate MetaBBO by 10 – 40x.

3) Rich Benchmarks: To include diverse optimization problem types, we rewrite the *Problem* class as an inheritable class. By inheriting from *Problem*, complex optimization problems such as multi-objective and multi-task ones can be seamlessly integrated, allowing polymorphism in different problem-specific behavior. MetaBox-v2 extend the testsuites from 3 to 18 synthetic/realistic tasks.

4) Plentiful and Extensible Interfaces: We thoroughly upgrade MetaBox’s developer flexibility. Every detailed process data are systematically recorded as metadata, which could be used for customized analysis by users. Furthermore, to match the open-source ecosystem, MetaBox-v2 provides plentiful interfaces to external resources. We prepare a systematic [online documentation](#) to guide the users.

2. Comprehensive Benchmarking Study: A comprehensive benchmarking study is conducted to showcase the practical value of MetaBox-v2, where up-to-date MetaBBO baselines are fairly trained and evaluated in terms of their performance, efficiency, generalization ability, etc. Our analysis yields valuable insights, particularly the significant variance in baseline generalization across testsuites and the critical trade-offs between learning efficiency and performance robustness. These findings establish empirical guidelines for practitioners while identifying promising research directions, accelerating future advancements in MetaBBO algorithm development.

Table 1: Comparison to related benchmarks. *#Optimization Scopes*: supported optimization problem types; *Learning Support*: supported MetaBBO learning paradigms; *Parallel*: hardware-level parallelism support; *#Problem*: the number of problem instances (*#synthetic* + *#realistic*); *#MetaBBO Baseline*: the number of MetaBBO baselines; *Template*: Template coding support; *Auto*: automated train/test workflow; *Custom*: configurable settings; *Visual*: visualization tools support; *Compatibility*: compatibility with open-source resources.

	<i>#Optimization Scopes</i>	<i>Learning Support</i>	<i>Parallel</i>	<i>#Problems</i>	<i>#MetaBBO Baselines</i>	<i>Template</i>	<i>Auto</i>	<i>Custom</i>	<i>Visual</i>	<i>Compatibility</i>
COCO [42]	4	×	×	481+0	×	✓	✓	×	✓	few
CEC [43]	1	×	×	30+0	×	×	×	×	×	none
IOHprofiler [44]	3	×	×	55+0	×	✓	×	✓	✓	few
Bayesmark [45]	1	×	×	0+228	×	✓	✓	×	×	few
Zigzag [46]	1	×	×	4+0	×	×	×	✓	×	none
Engineering [47]	1	×	×	+57	×	×	×	×	×	none
MA-BBOB [48]	1	×	×	1000+0	×	×	×	×	×	none
BBOPlace [49]	1	×	×	0+14	×	×	✓	×	×	none
PyPop7 [50]	2	×	×	92+11	×	×	×	×	×	few
EvoX [51]	2	×	✓	44+50	×	✓	×	✓	✓	rich
MetaBox [8]	1	RL	×	54+280	8	✓	✓	✓	✓	few
MetaBox-v2	5	RL,SL,NE,ICL	✓	541+1393	23	✓	✓	✓	✓	rich

2 Related Works

MetaBBO. We first illustrate the bi-level paradigm of Meta-Black-Box-Optimization (MetaBBO) [6] in Figure 2. In low-level optimization environment, a BBO optimizer $\mathcal{A}$ is maintained to optimize a problem p sampled from distribution $\mathcal{P}$. At each optimization step t , optimization status features are extracted from the current optimization process (such as population and objective values information). Then in meta level, an algorithm design policy π_θ (with learnable parameters θ) outputs a desired design ω_i^t by $\omega_i^t = \pi_\theta(s_i^t)$. $\mathcal{A}$ optimizes p by ω_i^t for one step. A performance measurement function r_t is used to evaluate the performance gain obtained by this algorithm design decision. Suppose T optimization steps are allowed for the low-level optimization process, then π_θ is meta-trained to maximize a meta-objective formulated as: $J(\theta) = \mathbb{E}_{p \in \mathcal{P}}[\sum_{t=1}^T r_t]$, which is expectation of accumulated single step performance gain over all problem instances in $\mathcal{P}$. In practice, a training problem set serves as the distribution $\mathcal{P}$. It is important to note that in this brief introduction we use ω as an abstract notation for algorithm design, considering that MetaBBO provides a universal concept framework for diverse algorithm design tasks such as algorithm selection, algorithm configuration etc. ω is hence a flexible and abstract concept to represent specific design choice under the given context. For a comprehensive summary of different algorithm design categories, we refer to MetaBBO survey [6] for further reading.

Following this paradigm, a wide array of MetaBBO approaches have been proposed, which further diverge into four same-end branches according to the learning techniques they adopt [6, 7]. The four branches are: 1) MetaBBO-RL: those first model the algorithm design process as a Markov Decision Process (MDP), then employ effective RL techniques to learn well-performing policies. Initial works such as DEDDQN [27], LDE [28], DEDQN [29] and RLEPSO [30] focus on designing dynamic configuration strategy for BBO optimizer. Following them, in-depth exploration include high-capacity neural policy [31, 32], complete optimizer generation [33, 34], optimization feature learning [35, 36] and efficient offline learning [37, 38]; 2) MetaBBO-SL: this branch originates from RNN-opt [11], where given a solution as input, a RNN is used to auto-regressively iterate it for better solution. The RNN is trained by minimizing the differentiable objective function. Although this paradigm requires white-box (differentiable) problems for training, recent works such as GLHF [13] and B2Opt [12] demonstrate that only training on synthetic problems is sufficient for generalization towards unseen problems; 3) MetaBBO-NE: where a the meta-level policy is learned by evolutionary optimization on its net parameters [14, 15, 39]; 4) MetaBBO-ICL: where a general LLM serves as either the low-level optimizer [17–19], or a meta-level configuration policy [40, 41]. The textual optimization process information is regarded as the context and learned by the LLM to propose algorithm designs. The fast development in MetaBBO field anticipates corresponding benchmarks.

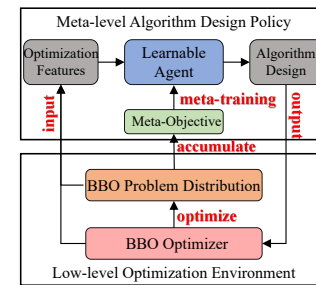


Figure 2: Bi-level Paradigm of existing MetaBBOs.

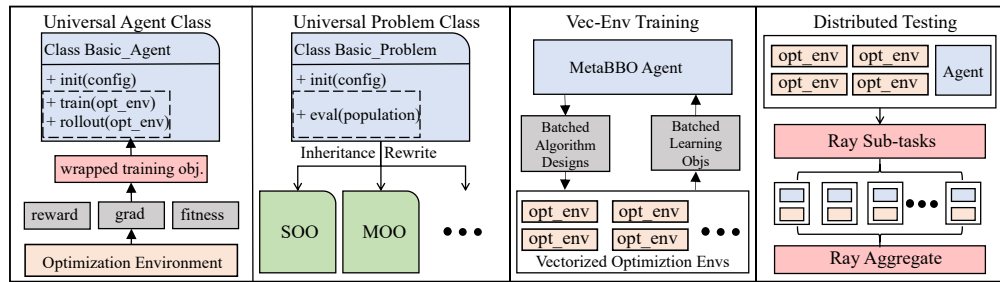


Figure 3: Major architecture adjustments in MetaBox-v2.

Related Benchmarks. A comparison of MetaBox-v2 to representative and up-to-date BBO benchmarks is presented in Table 1 to show the novelty of our work. Apart from those have been reviewed and compared in MetaBox [8], latest efforts on developing BBO benchmark include: 1) Engineering [47]: a collection of real world engineering optimization problems such as heat exchanger network design, industrial chemical process optimization, etc; 2) MA-BBOB [48]: a many-affine problem sets constructed from COCO [42] by interpolation operations on COCO’s problem instances, resulting in diversified synthetic instances; 3) BBOPlace [49]: chip placement tasks which represents complex and challenging optimization scenarios; 4) PyPop7 [50]: a comprehensive benchmark platform featured by massive number of BBO optimizers, which include decades of advanced optimizers with different types and specialized scenarios. 5) EvoX [51]: a high-efficiency benchmark platform featured by its distributed GPU-accelerated evaluation framework, with over 100x speedups than traditional BBO benchmarks such as CoCo and PyPop7. Notably, MetaBox, across its two versions, remains the sole platform that supports MetaBBO’s bi-level framework to streamline the development and benchmarking processes in this research domain.

3 MetaBox-v2

3.1 Unified MetaBBO Interface

Compatibility with Diverse MetaBBO. As reviewed in Section 2, the rapid development in MetaBBO has witnessed the exploration of various learning paradigms. A fundamental challenge is that, while sharing the bi-level paradigm, their underlying learning forms differ with each other. MetaBBO-RL is built on MDP, necessitating a reward signal from the low-level optimization environment to train the meta-level RL agent through trial-and-error. MetaBBO-SL and MetaBBO-NE require gradient information and fitness-like feedback, respectively. To achieve this, in MetaBox-v2, we replace the original RL-specific agent class with a unified *Basic_Agent* class featuring universal train and rollout interfaces. This is achieved through a wrapper function to transform different learning objective forms into a universal data object (shown in the left of Figure 3).

By such a novel design, we reproduce 15 more representative MetaBBO baselines upon the original MetaBox, including **1) MetaBBO-RL**: NRLPSO [52], RLDAS [32], SYMBOL [33], GLEET [31], RLDEAFL [36], Surr_RLDE [37], MADAC [21], PSORLNS [53], RLEMMO [22], L2T [26]; **2) MetaBBO-SL**: GLHF [13], B2OPT [12]; **3) MetaBBO-NE**: LES [15], LGA [14]; **4) MetaBBO-ICL**: OPRO [17]. In summary, MetaBox-v2 now supports 36 baselines including 23 MetaBBO baselines and 13 traditional BBO baselines. It is capable of providing not only comprehensive comparisons and analysis usages, but also a formal tutorial for those new to this field.

Scalable Testsuites Library. While the original MetaBox supported three synthetic/realistic testsuites for single-objective optimization (SOO), the up-to-date MetaBBO approaches have been initiated to multi-objective optimization (MOO) [21, 54], large-scale global optimization (LSGO) [24, 25], multi-modal optimization (MMO) [22] and multi-task optimization (MTO) [26, 55]. To keep pace with MetaBBO’s advancement so as to embrace users from diverse optimization sub-domains, we make a key adjustment to generalizing the SOO-specific problem class in MetaBox into a polymorphic *Basic_Problem* parent class (shown in the second column of Figure 3). This abstract base class with its core `eval()` interface enables problem specialization through inheritance, namely, users implement domain-specific evaluation logic by overriding this method.

Table 2: Diverse BBO testsuites in MetaBox-v2.

Name	Type	Dimension	maxFEs	Size	Scenario	Description	License
<i>bbob-10D</i> [42]	SOO	10D	2E4	24	synthetic	Single-objective instances in CoCo	BSD-3-Clause
<i>bbob-30D</i> [42]	SOO	30D	5E4	24	synthetic	Single-objective instances in CoCo	BSD-3-Clause
<i>bbob-noisy-10D</i> [42]	SOO	10D	2E4	24	synthetic	<i>bbob-10D</i> with gaussian noise	BSD-3-Clause
<i>bbob-noisy-30D</i> [42]	SOO	30D	5E4	24	synthetic	<i>bbob-30D</i> with gaussian noise	BSD-3-Clause
<i>hpo-b</i> [63]	SOO	2-16D	2E3	935	realistic	Hyper-parameter optimization	MIT License
<i>uav</i> [65]	SOO	30D	2.5E3	56	realistic	UAV path planning tasks	Attribution 4.0
<i>protein</i> [64]	SOO	12D	2E3	280	realistic	Simplified protein-docking instances	Attribution 4.0
<i>lsgo</i> [60]	LSGO	$\geq 905D$	3E6	20	synthetic	Large-scale problem instances	GPL-3.0
<i>ne</i> [51]	LSGO	$\geq 1000D$	2.5E3	66	realistic	Neuroevolution for control tasks	GPL-3.0
<i>zdt</i> [56]	MOO	10-30D	5E3	5	synthetic	A group of bi-objective problems	Apache-2.0
<i>uf</i> [58]	MOO	30D	5E3	10	synthetic	Multi-objective problem instances	Apache-2.0
<i>dtlz</i> [57]	MOO	6-29D	5E3	46	synthetic	Scalable multi-objective problems	Apache-2.0
<i>wfg</i> [59]	MOO	12-38D	5E3	117	synthetic	Complex multi-objective problems	Apache-2.0
<i>moo-uav</i> [56]	MOO	30D	2.5E3	56	realistic	Multi-objective form of <i>uav</i>	Apache-2.0
<i>mmo</i> [61]	MMO	1-20D	5E4-4E5	20	synthetic	Standard multi-modal problems	Simplified BSD
<i>cec2017mto</i> [62]	MTO	25-50D	2.5E4	9	synthetic	Multi-task problems in CEC2017	-
<i>wcci2020</i> [62]	MTO	50D	6.25E5	10	synthetic	Multi-task problems in WCCI2020	-
<i>wcci2020-aug</i>	MTO	50D	1.25E5	127	synthetic	Flexible combinations of <i>wcci2020</i>	-

Specifically, we have integrated 18 testsuites with over 1900 problem instances from diverse problem types into MetaBox-v2. These problems include not only representative synthetic benchmark functions for SOO [42], MOO [56–59], LSGO [60], MMO [61] and MTO [62], but also realistic testsuites with challenging optimization characteristics widely collected from AutoML [63], protein science [64], UAV system [65] and robotics [51]. We present basic information of them in Table 2. More details are accessible at the [online documentation](#).

Open-Source Ecosystem. MetaBox-v2 exemplifies exceptional extensibility through strategic integration with established optimization frameworks. For example, some built-in BBO baselines are implemented by calling powerful platforms such as DEAP [66], PyCMA [67], PyPop7 [50] etc. Some testsuites are borrowed from emerging benchmark platforms such as EvoX [51] to further acquire in-testing acceleration. We provide point-to-point [tutorial documentation](#) to connect users with these flexible usages.

3.2 Efficiency Optimization

Parallel Training. The time-consuming training caused by MetaBBO’s bi-level nested paradigm is a critical but understudied bottleneck in current literature. Our preliminary experiments reveal that serialized environment evaluations in the original MetaBox lead to prohibitive training times when handling modern testsuite scales, posing barriers to the rapid development of MetaBBO field. In MetaBox-v2, to address this efficiency issue, we propose a novel parallel scheme termed as *vectorized optimization environment* to accelerate MetaBBO’s training. In specific, as illustrated in the third column of Figure 3, during the training, we simultaneously construct a batch of low-level optimization environments and wrap them into a vectorized environment based on Tianshou [68]. Then the meta-level agent could perform batched algorithm designs via multi-processing parallelization in the vectorized environment. This allows parallel collection of learning signals (rewards/gradients/fitness measures) across multiple environments and problem instances, which are aggregated into mini-batches for efficient meta-policy updates. To the best of our knowledge, MetaBox-v2 is the first development exemplar to make MetaBBO’s training parallel.

Parallel Testing. MetaBox-v2 provides Ray-based parallel scheme [69] for distributed testing of MetaBBO/BBO baselines. As illustrated in the right of Figure 3, given a MetaBBO’s meta-level agent and a testsuite, we first copy the agent for each testing run and use Ray to construct the corresponding sub-tasks. Then all sub-tasks are distributed into independent CPU/GPU cores for parallel testing. The testing results in these sub-tasks are aggregated automatically by Ray’s handler. To show the detail, we first assume a testing scenario where N problem instances compose the testing set, B baselines to be tested, and R independent runs to ensure the statistical robustness. Naive testing procedure requires nested 3-layer loop to complete all these $N \times B \times R$ testing runs. In contrast, in MetaBox-v2, we provide four parallel testing modes to achieve fine-grained testing efficiency optimization: 1) mode-1: N cpu cores are used to distribute the N testing problem instances, on each core, line2-line3 are executed as two-layer loop; 2) mode-2: R cpu cores are used to distribute the R independent runs, while on each core, line1-line2 are executed as two-layer loop; 3) mode-3: $N \times B$ cpu cores are used to distribute the N problem instances and B baselines, while on each core, line3 are executed as one-layer loop; 4) mode-4: $N \times B \times R$ cpu cores are used to distribute all evaluation tasks, hence there is no loop anymore. The hardware requirement from mode-1 to mode-4 is incremental.

By decomposing parallelism into two orthogonal dimensions as above: (a) distributed solving across *problem instances*, and (b) parallel execution of *independent test runs*, we provide sufficient flexibility for users to accelerate their programs based on their specific hardware conditions.

3.3 Novel Evaluation Metrics

Metadata. MetaBox-v2 inherits the automatic train-test-log workflow of original MetaBox. However, the original MetaBox does not provide interfaces for users to operate on the process data observed from both the training and testing. Instead, it provides users post-processed data such as comparison tables and optimization progress figures. As an emerging topic, it is still an open question how to measure different MetaBBO approaches with fairness and objectivity. To this end, we open a data acquirement interface `get_metadata()` for users who would like to custom their own metrics in analysis. For example, consider evaluating a pre-trained MetaBBO approach $\mathcal{A}$ on a testsuite $\mathbb{D}$ containing N problem instances $\{p_1, \dots, p_N\}$. For each problem instance p_i , we execute K independent runs. Then the metadata md_i for p_i is structured as a json object:

$$\begin{aligned} \{ & \text{"problem_id": } p_i, \text{ "data": } \{ \text{"run_1": } \{ \text{"X": List, "Y": List, "T": 5.23} \}, \\ & \dots, \\ & \text{"run_K": } \{ \text{"X": List, "Y": List, "T": 4.96} \} \} \end{aligned} \quad (1)$$

where "X" is a list of each generation's solution population, "Y" is a list of the objective values, "T" is the wall-clock time consumed for optimizing p_i . Then the overall metadata $md(\mathcal{A}, \mathbb{D})$ is aggregated from each p_i : $\{ \text{"problem_type": } S00, \text{ "all_data": } [md_1, \dots, md_N] \}$, where "problem_type" is the optimization types of $\mathbb{D}$. Notably, $md(\mathcal{A}, \mathbb{D})$ provide significant convenience for computing various metrics. For example, a normalized performance metric widely used in existing MetaBBO approaches [13, 15, 31, 70] can be easily computed as $\text{Perf}(\mathcal{A}, \mathbb{D}) = \frac{1}{N \times K} [\sum_{i=1}^N \sum_{j=1}^K \frac{Y_{i,j}^* - p_i^*}{Y_{i,j}^0 - p_i^*}]$, where $Y_{i,j}^0$ and $Y_{i,j}^*$ are the initial and final objective values, p_i^* is the optimal. We next showcase two customized metrics based on the metadata.

Learning Efficiency Indicator. We provide a novel built-in metric in MetaBox-v2: *learning efficiency*, to measure how efficiently a MetaBBO approach learns an effective meta-level policy. Specifically, during the training of a given algorithm $\mathcal{A}$, we save a series of its model snapshots $\{\mathcal{A}^{(g)}\}_{g=0}^G$ where G is the number of training epochs. Evaluating all snapshots on the testsuite $\mathbb{D}$, we obtain corresponding metadata: $\{md(\mathcal{A}^{(g)}, \mathbb{D})\}_{g=0}^G$. Suppose training $\mathcal{A}^{(g)}$ consumes $T^{(g)}$ hours, then the *learning efficiency* of $\mathcal{A}^{(g)}$ is computed as: $\frac{\text{Perf}(\mathcal{A}^{(g)}, \mathbb{D})}{T^{(g)}}$. This metric could fairly reflect the training efficiency of $\mathcal{A}$ at different time slots g .

Anti-NFL Indicator. Recall that the core motivation of MetaBBO is to learn generalizable policy towards unseen problems. This is somewhat against the well-known *no-free-lunch* (NFL) theorem [5]. We hence propose a novel indicator named as Anti-NFL, which could reflect the performance variance of a given algorithm $\mathcal{A}$ on unseen testsuites apart from the one it was trained. Suppose $\mathcal{A}$ is trained on $\mathbb{D}_{\text{train}}$, and tested on other B testsuites: $\{\mathbb{D}_{\text{test}}^{(b)}\}_{b=1}^B$. After obtaining all of the metadata by testing the final model $\mathcal{A}^{(G)}$ on these testsuites, the Anti-NFL is computed as:

$$\text{Anti-NFL} = \exp \left(\frac{1}{B} \sum_{b=1}^B \frac{\text{Perf}(\mathcal{A}^{(G)}, \mathbb{D}_{\text{test}}^{(b)}) - \text{Perf}(\mathcal{A}^{(G)}, \mathbb{D}_{\text{train}})}{\text{Perf}(\mathcal{A}^{(G)}, \mathbb{D}_{\text{train}})} \right) \quad (2)$$

A larger Anti-NFL indicator indicates that $\mathcal{A}$ performs robustly under problem-shifts, and vice versa.

4 Benchmarking Study

We present a comprehensive case study on up-to-date MetaBBO approaches through MetaBox-v2, addressing four critical research questions: **RQ1:** Does MetaBox-v2 significantly enhance train/test efficiency compared to the conventional version? **RQ2:** How effectively do up-to-date MetaBBO methods generalize under standardized training protocols and diverse test scenarios? **RQ3:** How to objectively compare the dominance relationship of learning efficiency and generalization ability of baselines? **RQ4:** How does MetaBBO perform under extreme problem shift in practice?

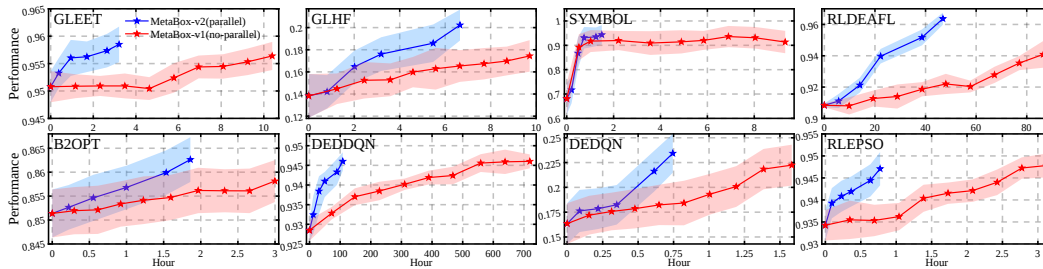


Figure 4: Training improvement curves of baselines on Metabox-v2 and MetaBox.

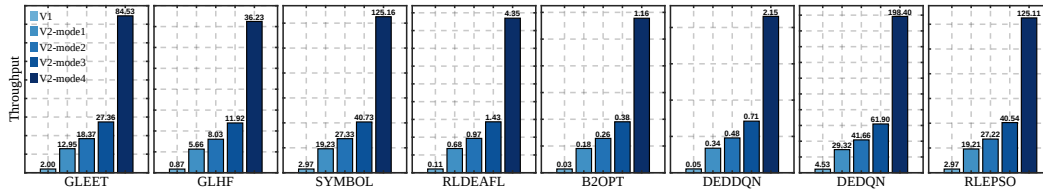


Figure 5: Testing efficiency comparison of MetaBox-v2 (4 parallel modes) and MetaBox.

4.1 Experimental Setup

Baselines. We select 20 baselines from the library of MetaBox-v2 for case study, including 5 traditional BBO optimizers: PSO [2], DE [3], SHADE [71], JDE21 [72], MadDE [73], and 15 up-to-date MetaBBO baselines from all of four learning paradigms: RNNOPT [11], DEDDQN [27], DEDQN [29], LDE [28], RLPSO [74], RLEPSO [30], NRLPSO [52], LES [15], GLEET [31], GLHF [13], RLDAS [32], SYMBOL [33], OPRO [17], B2OPT [12], RLDEAFI [36]. The settings follows their original papers.

Testsuites. First, we employ the *bbob-10D* testsuite and split its 24 problem instances into 8 training instances and 16 testing instances, with the latter serving as in-distribution evaluation. Then, the out-of-distribution evaluation involves four other testsuites: *bbob-noisy-30D*, *protein*, *uav*, and *hpo-b*. To ensure the fairness of training, all algorithms undergo 1600 episodes for each training instance. The testing phases employ 51 independent runs with seed-controlled reproducibility. All experiments are conducted with 2 AMD EPYC 7H12 CPU, a RTX 3080 GPU and 512GB RAM.

4.2 Platform's Acceleration Performance (RQ1)

We use vectorized environment with *batch_size* as 16 to accelerate the training of all involved MetaBBO baselines. Due to the space limitation, we selectively illustrate in Figure 4 the performance improvement curves of 8 baselines, where y-axis denotes normalized performance on testing set of *bbob-10D*. Compared to original MetaBox, MetaBox-v2 consistently accelerates MetaBBO baselines by at most 10x. We can observe that the concrete acceleration may varies on different baselines, this is because the differences of the internal logic and communication cost among the baselines. An important note is that the irregular record points in Figure 4 is due to the unstable multi-processing. For parallel cases, we draw 5 points (every 20 training epochs), which might present irregular patterns since the x-axis denotes the training time.

We also illustrate the acceleration performance of MetaBox-v2 compared to original MetaBox in Figure 5 in terms of testing efficiency, where y-axis denotes the throughput of evaluation process measured by the number of instance test runs per second. We compare the throughput of the 4 Ray modes in MetaBox-v2 to original MetaBox, and the results show that even the simplest distribution *Mode-1* could significantly accelerate the testing workflow. If users have advanced hardware, the distribution *Mode-4* could introduce no less than 40x acceleration.

4.3 Generalization Performance Comparisons among Baselines (RQ2)

In-distribution Test. Table 3 shows the average results and error bars on the testing set of *bbob-10D* across 51 independent runs. We additionally summarize the average ranks among the baselines at the bottom of the table. The in-distribution test aims at validating the basic learning effectiveness of

Table 3: In-distribution optimization performances of baselines over *bbob-10D*, with gray box labeling the best. Due to the space limitation, results for 8 of 16 problem instances in *bbob-10D*'s testing set are presented here while the complete results can be accessed at [this online page](#).

	Sharp_Ridge	Different_Powers	Schaffers_HC	Composite_GR	Schwefel	Gallagher_21	Katsuura	Lunacek_BR
PSO(1995) [2]	1.905E+02 ± 2.156E+01	6.802E-01 ± 1.760E-01	5.600E+00 ± 1.368E+00	3.290E+00 ± 5.796E-01	2.560E+00 ± 3.067E-01	6.803E+00 ± 6.472E+00	1.272E+00 ± 2.933E-01	6.139E+01 ± 5.747E+00
DE(1997) [3]	8.588E-01 ± 1.054E+00	8.180E-04 ± 2.537E-04	9.454E-02 ± 6.483E-02	2.577E+00 ± 4.860E-01	9.156E-01 ± 3.039E-01	3.393E+00 ± 4.999E+00	1.467E+00 ± 2.734E-01	4.210E+01 ± 3.043E+00
SHADE(2013) [71]	1.442E+00 ± 4.321E-01	2.721E-04 ± 4.192E-05	2.649E-01 ± 6.818E-02	2.238E+00 ± 3.476E-01	1.338E+00 ± 1.957E-01	1.155E+00 ± 9.320E-01	1.553E+00 ± 3.454E-01	4.248E+01 ± 4.209E+00
JDE21(2021) [72]	3.476E+00 ± 6.350E+00	4.398E-04 ± 8.807E-04	4.496E-01 ± 3.700E-01	2.542E+00 ± 6.355E-01	5.777E-01 ± 2.246E-01	1.604E+00 ± 1.641E+00	1.416E+00 ± 3.359E-01	4.059E+01 ± 7.940E+00
MADDE(2021) [73]	1.736E+00 ± 3.300E-01	5.830E-04 ± 2.318E-04	9.538E-01 ± 2.897E-01	1.077E+00 ± 3.709E-01	8.049E-01 ± 1.997E-01	5.458E-01 ± 7.264E-01	1.350E+00 ± 2.395E-01	4.308E+01 ± 4.974E+00
RNNOPT(2017) [11]	1.822E+03 ± 0.000E+00	2.297E+01 ± 0.000E+00	4.645E+01 ± 0.000E+00	3.609E+00 ± 0.000E+00	9.297E+03 ± 1.819E-12	8.431E+01 ± 0.000E+00	2.186E+00 ± 0.000E+00	1.142E+02 ± 0.000E+00
DEDDQN(2019) [27]	1.841E-03 ± 1.841E-03	4.224E-09 ± 4.069E-09	1.080E-02 ± 7.097E-03	2.480E+00 ± 5.250E-01	1.720E+00 ± 4.164E-01	1.574E+00 ± 9.236E-01	1.344E+00 ± 2.839E-01	4.039E+01 ± 4.264E+00
DEDQN(2021) [29]	9.538E+02 ± 1.548E+02	1.115E+01 ± 2.837E+00	2.709E+01 ± 5.790E+00	1.268E+01 ± 2.131E+00	4.880E+03 ± 3.385E+03	5.711E+01 ± 1.366E+01	3.286E+00 ± 6.136E-01	1.591E+02 ± 2.132E+01
LDE(2021) [28]	5.955E-01 ± 5.103E-01	5.159E-05 ± 3.700E-05	2.156E-01 ± 1.238E-01	2.024E+00 ± 1.812E-01	1.071E+00 ± 1.603E-01	4.292E-01 ± 7.059E-01	1.306E+00 ± 2.245E-01	3.616E+01 ± 3.494E+00
RLPSO(2021) [74]	2.769E+02 ± 7.000E+01	1.481E+00 ± 9.514E-01	1.429E+01 ± 2.968E+00	3.629E+00 ± 1.115E+00	2.722E+00 ± 2.998E-01	1.597E+01 ± 1.719E+01	2.225E+00 ± 3.550E-01	6.525E+01 ± 7.460E+00
RLPSO(2022) [30]	6.388E+00 ± 6.093E+00	2.554E-04 ± 1.396E-04	1.687E+00 ± 7.471E-01	1.387E+00 ± 4.516E-01	1.261E+00 ± 2.497E-01	7.703E+00 ± 1.223E+01	1.017E+00 ± 2.993E-01	2.413E+01 ± 7.015E+00
NRLPSO(2023) [52]	1.968E+02 ± 8.105E+01	6.449E-01 ± 3.607E-01	5.710E+00 ± 2.194E+00	3.367E+00 ± 1.081E+00	2.631E+00 ± 4.837E-01	7.478E+00 ± 5.155E+00	1.599E+00 ± 4.433E-01	7.007E+01 ± 1.466E+01
LES(2023) [15]	1.099E+03 ± 1.516E+02	1.273E+01 ± 2.222E+00	3.812E+01 ± 6.523E+00	1.215E+01 ± 2.095E+00	8.044E+03 ± 4.741E+03	5.777E+01 ± 2.074E+01	4.099E+00 ± 9.875E-01	1.793E+02 ± 2.481E+01
GLEET(2024) [31]	4.464E+00 ± 7.370E+00	1.130E-04 ± 8.072E-05	2.137E+00 ± 1.618E+00	8.624E-01 ± 3.202E-01	1.481E+00 ± 1.765E-01	8.632E+00 ± 1.209E+01	4.839E-01 ± 2.181E-01	2.717E+01 ± 8.473E+00
GLHF(2024) [13]	9.652E+02 ± 1.286E+02	1.074E+01 ± 1.796E+00	3.163E+01 ± 5.541E+00	1.027E+01 ± 1.857E+00	6.827E+03 ± 4.036E+03	4.923E+01 ± 1.678E+01	3.527E+00 ± 9.244E-01	1.582E+02 ± 2.103E+01
RLDAS(2024) [32]	1.627E+00 ± 1.073E+00	3.740E-04 ± 2.542E-04	9.798E-01 ± 5.450E-01	1.650E+00 ± 4.859E-01	5.505E-01 ± 3.074E-01	4.698E-01 ± 5.763E-01	1.296E+00 ± 2.623E-01	3.630E+01 ± 1.035E+01
SYMBOL(2024) [33]	1.344E+01 ± 9.453E+00	5.332E-03 ± 2.537E-03	4.256E+00 ± 2.238E+00	1.383E+00 ± 4.880E-01	1.732E+00 ± 2.436E-01	5.611E+00 ± 4.981E+00	6.371E-01 ± 3.070E-01	3.188E+01 ± 1.164E+01
OPRO(2024) [17]	2.003E+03 ± 1.562E+02	3.007E+01 ± 3.326E+00	5.099E+01 ± 9.258E+00	1.434E+01 ± 6.317E+00	9.299E+03 ± 4.804E+03	9.031E+01 ± 2.184E+01	6.113E+00 ± 2.771E+00	1.812E+02 ± 2.562E+01
B2OPT(2025) [12]	2.581E+02 ± 3.724E+01	2.510E+00 ± 3.641E-01	6.057E+00 ± 1.696E+00	8.728E-01 ± 2.494E-01	2.543E+00 ± 1.547E-01	1.130E+01 ± 7.585E+00	1.575E+00 ± 2.701E-01	5.814E+01 ± 6.917E+00
RLDEAF(2025) [36]	1.136E+01 ± 1.345E+01	1.487E-04 ± 9.165E-05	4.166E+00 ± 2.214E+00	2.535E+00 ± 1.036E+00	1.397E+00 ± 3.326E-01	5.452E+00 ± 6.106E+00	1.199E+00 ± 6.034E-01	3.231E+01 ± 7.111E+00
Rank	1:LDE, 2:DEDDQN, 3:RLDAS, 4:SHADE, 5:MADDE, 6:GLEET, 7:RLPSO, 8:RLDEAF, 9:JDE21, 10:DE, 11:SYMBOL, 12:PSO, 13:B2OPT, 14:NRLPSO, 15:RLPSO, 16:GLHF, 16:DEDQN, 18:RNNOPT, 19:LES, 20:OPRO							

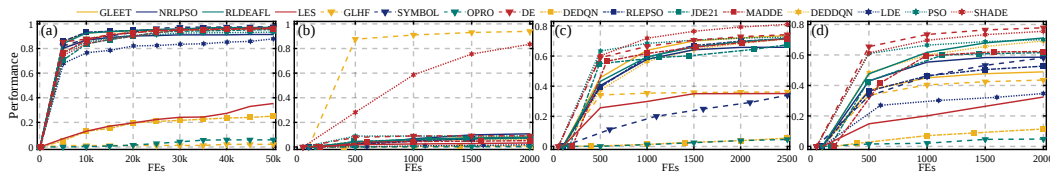


Figure 6: Out-of-distribution generalization performance of baselines on: (a) *bbob-noisy-30D*; (b) *protein*; (c) *uav*; and (d) *hpo-b*.

MetaBBO since the synthetic problem instances within *bbob-10D* show certain similarity in landscape features [75]. Several key observations can be obtained: 1) Overall, on 14 of all 16 testing instances, MetaBBO baselines attain the best optimization results and advance traditional BBO baselines by orders of magnitudes. 2) So far, the MetaBBO-RL baselines generally outperform MetaBBO-SL, MetaBBO-NE and MetaBBO-ICL baselines. 3) We notice that a 2019 method DEDQN [27] outperforms other baselines in 6 of 18 testing instances and ranks the first place on average.

Out-of-distribution Test. Figure 6 presents comparative evaluation results across MetaBBO and traditional BBO baselines using four out-of-distribution testsuites (*bbob-noisy-30D*, *protein*, *uav*, and *hpo-b*), with the y-axis representing instance-normalized performance averages. Combining the results of in-distribution test, several valuable insights are obtained: 1) Generally speaking, traditional BBO algorithms such as DE, SHADE demonstrate empirical robustness across testsuites, contrasting with MetaBBO baselines that appear to overestimate their generalization potential; 2) The out-of-distribution performance of DEDQN on *protein* is reversed from its good performance in the in-distribution test, which might indicates that the overfitting issue should be addressed for future MetaBBO researches. 3) Almost all of MetaBBO baselines show certain level of performance oscillation in diverse testsuites, which further underscores that how to define and measure similarity across diverse BBO problems is crucial to ensure MetaBBO's generalization.

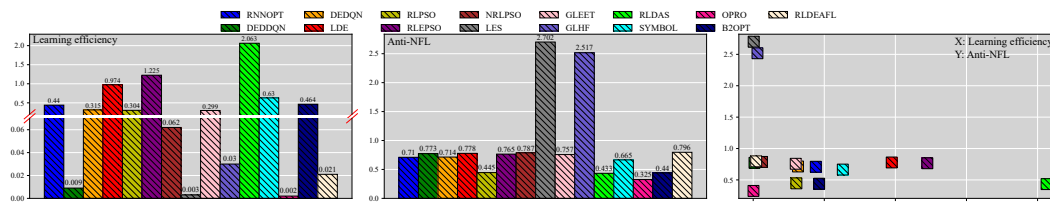


Figure 7: Left: Learning efficiency comparison of MetaBBO baselines, larger is better. Middle: Anti-NFL indicator of MetaBBO baselines, larger is better. Right: Domination relationship among MetaBBO baselines considering learning efficiency and Anti-NFL indicator.

4.4 Other In-depth Analysis

Although there is continuous discussion on the fairness and objectivity of BBO benchmarks [76–78], it is still an open challenge for optimization community to agree on a certain golden standard, let alone for the MetaBBO field. We hence provide the following discussions on the objective profiling of efficiency and generalization of MetaBBO, and the impact analysis of extreme problem shift.

Learning Efficiency (RQ3). Following the computation detail introduced in Section 3.3, we compute the *learning efficiency* indicators of all baselines on all testsuites. Their average efficiency values is shown in the left of Figure 7. Combining the results with the average ranks of baselines in Table 3, we could conclude that RLDAS [32] is a remarkable MetaBBO baseline since it uses less computational resource to achieve better optimization performance. In contrast, DEDDQN [27] achieves the best performance while consuming hundreds of hours for training, which might not be favorable when the computational resource is limited. It is also worthy to note that MetaBBO-NE approach such as LES [15] and MetaBBO-ICL approach such as OPRO [17] hold the lowest efficiency due to the nested evolutionary optimization and the expensive LLM calling, respectively.

Anti-NFL Performance (RQ3). The Anti-NFL indicator computed in Eq. 2 reflects the robustness of a MetaBBO approach when being generalized to diverse BBO problems. The middle of Figure 7 reports the Anti-NFL indicators of MetaBBO baselines on all testsuites. The conclusions could be obtained here seems to be different with the aforementioned performance metrics. Two MetaBBO baselines GLHF [13] and LES [15] have much higher Anti-NFL values than others, while they hold relatively low absolute performance (Table 3). This point deserves in-depth analysis in future works. While RLDAS has favorable performance and efficiency, its Anti-NFL is among the lowest due to its meta-level policy’s architecture, which is not generalizable across different problem dimensions.

Combining the results of *learning efficiency* and Anti-NFL indicator, as illustrated in the right of Figure 3, we analyze the domination relationship of MetaBBO baselines. It can be observed that, so far, no baseline participating in this case study dominates all the others. This reflects that there is certain design tradeoff in existing MetaBBO between the efficiency and effectiveness.

Extreme Problem Shift Analysis (RQ4). We evaluate MetaBBO’s adaptability through an extreme domain shift scenario: algorithms trained on low-dimensional synthetic problems (*bbob-10D*) are directly deployed on high-dimensional neuroevolution tasks (*ne*), a robotic control benchmark integrated from EvoX [51] where optimizers must tune neural networks (thousands of parameters) to maximize robotic returns. The results from three Ant tasks (Ant-3, Ant-4, Ant-5 with 4328, 5384, 6440 parameters) are reported in Figure 8, which reveal that: 1) Certain MetaBBO methods (e.g., GLEET [31] with Transformer meta-policy) achieve performance parity or superiority over advanced BBO baselines (CMAES [4] and GLPSO [79]) despite the exclusive training on simple problems. 2) MetaBBO performance correlates with policy architecture complexity—Transformer-based GLEET maintains robustness across scaling dimensions, while MLP (RLEPSO [30]) and LSTM (LDE [28]) agents degrade sharply in higher-dimensional tasks.

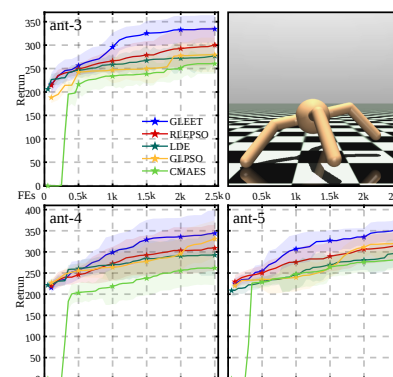


Figure 8: Comparison on neuroevolution task by return curves.

5 Discussion

Takeaways. This work proposes MetaBox-v2 as a milestone upgrade for its predecessor, introducing several key architecture adjustments that establish a comprehensive benchmark platform for the MetaBBO research. The fundamental advancements not only enable unified development and evaluation for various MetaBBO paradigms and diverse optimization problem types; but also support streamlined parallel acceleration for training/evaluation by 10x-40x. With the expanded baseline library (8 $\rightarrow$ 23) and testsuite library (3 $\rightarrow$ 18, 1900+ problem instances), we conduct a rigorous case study, which discloses insights including but not limited to: 1) Current literature overestimates MetaBBO capabilities through narrow evaluation practices, with MetaBox-v2 exposing substantial performance gaps in cross-domain settings. 2) Out-of-distribution generalization demands special attention, since our analysis reveals that overfitting persists across baseline algorithms. 3) Effective MetaBBO assessment requires multidimensional analysis (optimization efficacy, learning efficiency, generalization robustness) beyond traditional convergence metrics.

Future Work. We outline the following directions to continuously improve MetaBox-v2: 1) Maintain cutting-edge benchmarks through continuous integration of emerging algorithms and problem suites, with an open-source ecosystem welcoming community contributions; 2) Optimize the parallel computing framework to achieve higher resource utilization; 3) Lower adoption barriers through enhanced tutorials and beginner-friendly interfaces. We aim to establish MetaBox-v2 as both a powerful research platform and an accessible educational resource for the MetaBBO field.

Acknowledgments and Disclosure of Funding

This work was supported in part by National Natural Science Foundation of China (Grant No. 62276100), in part by Guangzhou Science and Technology Elite Talent Leading Program for Basic and Applied Basic Research (Grant No. SL2024A04J01361), in part by the Guangdong Provincial Natural Science Foundation for Outstanding Youth Team Project (Grant No. 2024B1515040010), and in part by the Fundamental Research Funds for the Central Universities (Grant No. 2025ZYGXZR027).

References

- [1] John H Holland. Genetic algorithms. *Scientific American*, 1992.
- [2] James Kennedy and Russell Eberhart. Particle swarm optimization. In *Proceedings of ICNN'95-international conference on neural networks*, 1995.
- [3] Rainer Storn and Kenneth Price. Differential evolution—a simple and efficient heuristic for global optimization over continuous spaces. *Journal of Global Optimization*, 1997.
- [4] Hans-Georg Beyer and Hans-Paul Schwefel. Evolution strategies—a comprehensive introduction. *Natural Computing*, 2002.
- [5] David H Wolpert and William G Macready. No free lunch theorems for optimization. *IEEE Transactions on Evolutionary Computation*, 1997.
- [6] Zeyuan Ma, Hongshu Guo, Yue-Jiao Gong, Jun Zhang, and Kay Chen Tan. Toward automated algorithm design: A survey and practical guide to meta-black-box-optimization. *arXiv preprint arXiv:2411.00625*, 2024.
- [7] Xu Yang, Rui Wang, Kaiwen Li, and Hisao Ishibuchi. Meta-black-box optimization for evolutionary algorithms: Review and perspective. *Swarm and Evolutionary Computation*, 2025.
- [8] Zeyuan Ma, Hongshu Guo, Jiacheng Chen, Zhenrui Li, Guojun Peng, Yue-Jiao Gong, Yining Ma, and Zhiguang Cao. Metabox: A benchmark platform for meta-black-box optimization with reinforcement learning. *Advances in Neural Information Processing Systems*, 2023.
- [9] Richard S Sutton, Andrew G Barto, et al. *Reinforcement learning: An introduction*. 1998.
- [10] Xu Yang, Rui Wang, Kaiwen Li, Wenhua Li, Tao Zhang, and Fujun He. Platmetax: An integrated matlab platform for meta-black-box optimization. *arXiv preprint arXiv:2503.22722*, 2025.

- [11] Vishnu TV, Pankaj Malhotra, Jyoti Narwariya, Lovekesh Vig, and Gautam Shroff. Meta-learning for black-box optimization. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*. Springer, 2019.
- [12] Xiaobin Li, Kai Wu, Xiaoyu Zhang, and Handing Wang. B2opt: Learning to optimize black-box optimization with little budget. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2025.
- [13] Xiaobin Li, Kai Wu, Xiaoyu Zhang, Handing Wang, Jing Liu, et al. Pretrained optimization model for zero-shot black box optimization. *Advances in Neural Information Processing Systems*, 2024.
- [14] Robert Lange, Tom Schaul, Yutian Chen, Chris Lu, Tom Zahavy, Valentin Dalibard, and Sebastian Flennerhag. Discovering attention-based genetic algorithms via meta-black-box optimization. In *Proceedings of the Genetic and Evolutionary Computation Conference*, 2023.
- [15] Robert Lange, Tom Schaul, Yutian Chen, Tom Zahavy, Valentin Dalibard, Chris Lu, Satinder Singh, and Sebastian Flennerhag. Discovering evolution strategies via meta-black-box optimization. In *Proceedings of the Companion Conference on Genetic and Evolutionary Computation*, 2023.
- [16] Kenneth O Stanley, Jeff Clune, Joel Lehman, and Risto Miikkulainen. Designing neural networks through neuroevolution. *Nature Machine Intelligence*, 2019.
- [17] Chengrun Yang, Xuezhi Wang, Yifeng Lu, Hanxiao Liu, Quoc V Le, Denny Zhou, and Xinyun Chen. Large language models as optimizers. *arXiv preprint arXiv:2309.03409*, 2023.
- [18] Zeyuan Ma, Hongshu Guo, Jiacheng Chen, Guojun Peng, Zhiguang Cao, Yining Ma, and Yue-Jiao Gong. LLaMoCo: Instruction tuning of large language models for optimization code generation. *arXiv preprint arXiv:2403.01131*, 2024.
- [19] Caigao Jiang, Xiang Shu, Hong Qian, Xingyu Lu, Jun Zhou, Aimin Zhou, and Yang Yu. Llmopt: Learning to define and solve general optimization problems from scratch. *arXiv preprint arXiv:2410.13213*, 2024.
- [20] Qingxiu Dong, Lei Li, Damai Dai, Ce Zheng, Jingyuan Ma, Rui Li, Heming Xia, Jingjing Xu, Zhiyong Wu, Tianyu Liu, et al. A survey on in-context learning. *arXiv preprint arXiv:2301.00234*, 2022.
- [21] Ke Xue, Jiacheng Xu, Lei Yuan, Miqing Li, Chao Qian, Zongzhang Zhang, and Yang Yu. Multi-agent dynamic algorithm configuration. *Advances in Neural Information Processing Systems*, 2022.
- [22] Hongqiao Lian, Zeyuan Ma, Hongshu Guo, Ting Huang, and Yue-Jiao Gong. Rlemmo: Evolutionary multimodal optimization assisted by deep reinforcement learning. In *Proceedings of the Genetic and Evolutionary Computation Conference*, 2024.
- [23] Zeyuan Ma, Hongqiao Lian, Wenjie Qiu, and Yue-Jiao Gong. Accurate peak detection in multimodal optimization via approximated landscape learning. In *Proceedings of the Genetic and Evolutionary Computation Conference*, 2025.
- [24] Hongshu Guo, Wenjie Qiu, Zeyuan Ma, Xinglin Zhang, Jun Zhang, and Yue-Jiao Gong. Advancing cma-es with learning-based cooperative coevolution for scalable optimization. *arXiv preprint arXiv:2504.17578*, 2025.
- [25] Shuai Shao, Ye Tian, Yajie Zhang, and Xingyi Zhang. Knowledge learning-based dimensionality reduction for solving large-scale sparse multiobjective optimization problems. *IEEE Transactions on Cybernetics*, 2025.
- [26] Sheng-Hao Wu, Yuxiao Huang, Xingyu Wu, Liang Feng, Zhi-Hui Zhan, and Kay Chen Tan. Learning to transfer for evolutionary multitasking. *IEEE Transactions on Cybernetics*, 2025.
- [27] Mudita Sharma, Alexandros Komninos, Manuel López-Ibáñez, and Dimitar Kazakov. Deep reinforcement learning based parameter control in differential evolution. In *Proceedings of the Genetic and Evolutionary Computation Conference*, 2019.

- [28] Jianyong Sun, Xin Liu, Thomas Bäck, and Zongben Xu. Learning adaptive differential evolution algorithm from optimization experiences by policy gradient. *IEEE Transactions on Evolutionary Computation*, 2021.
- [29] Zhiping Tan and Kangshun Li. Differential evolution with mixed mutation strategy based on deep reinforcement learning. *Applied Soft Computing*, 2021.
- [30] Shiyuan Yin, Yi Liu, GuoLiang Gong, Huaxiang Lu, and Wenchang Li. RLEPSO: Reinforcement learning based ensemble particle swarm optimizer. In *Proceedings of the 2021 4th International Conference on Algorithms, Computing and Artificial Intelligence*, 2021.
- [31] Zeyuan Ma, Jiacheng Chen, Hongshu Guo, Yining Ma, and Yue-Jiao Gong. Auto-configuring exploration-exploitation tradeoff in evolutionary computation via deep reinforcement learning. In *Proceedings of the Genetic and Evolutionary Computation Conference*, 2024.
- [32] Hongshu Guo, Yining Ma, Zeyuan Ma, Jiacheng Chen, Xinglin Zhang, Zhiguang Cao, Jun Zhang, and Yue-Jiao Gong. Deep reinforcement learning for dynamic algorithm selection: A proof-of-principle study on differential evolution. *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, 2024.
- [33] Jiacheng Chen, Zeyuan Ma, Hongshu Guo, Yining Ma, Jie Zhang, and Yue-Jiao Gong. Symbol: Generating flexible black-box optimizers through symbolic equation learning. In *The Twelfth International Conference on Learning Representations*, 2024.
- [34] Qi Zhao, Tengfei Liu, Bai Yan, Qiqi Duan, Jian Yang, and Yuhui Shi. Automated metaheuristic algorithm design with autoregressive learning. *IEEE Transactions on Evolutionary Computation*, 2024.
- [35] Zeyuan Ma, Jiacheng Chen, Hongshu Guo, and Yue-Jiao Gong. Neural exploratory landscape analysis for meta-black-box-optimization. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [36] Hongshu Guo, Sijie Ma, Zechuan Huang, Yuzhi Hu, Zeyuan Ma, Xinglin Zhang, and Yue-Jiao Gong. Reinforcement learning-based self-adaptive differential evolution through automated landscape feature learning. *arXiv preprint arXiv:2503.18061*, 2025.
- [37] Zeyuan Ma, Zhiyang Huang, Jiacheng Chen, Zhiguang Cao, and Yue-Jiao Gong. Surrogate learning in meta-black-box optimization: A preliminary study. *arXiv preprint arXiv:2503.18060*, 2025.
- [38] Zeyuan Ma, Zhiguang Cao, Zhou Jiang, Hongshu Guo, and Yue-Jiao Gong. Meta-black-box-optimization through offline q-function learning. In *Forty-second International Conference on Machine Learning*, 2025.
- [39] Robert Lange, Yingtao Tian, and Yujin Tang. Evolution transformer: In-context evolutionary optimization. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion*, 2024.
- [40] Xingyu Wu, Yan Zhong, Jibin Wu, Bingbing Jiang, and Kay Chen Tan. Large language model-enhanced algorithm selection: Towards comprehensive algorithm representation. In *The 33rd International Joint Conference on Artificial Intelligence*, 2024.
- [41] Niki van Stein and Thomas Bäck. Llamea: A large language model evolutionary algorithm for automatically generating metaheuristics. *IEEE Transactions on Evolutionary Computation*, 2024.
- [42] Nikolaus Hansen, Anne Auger, Raymond Ros, Olaf Mersmann, Tea Tušar, and Dima Brockhoff. Coco: A platform for comparing continuous optimizers in a black-box setting. *Optimization Methods and Software*, 2021.
- [43] JJ Liang, BY Qu, PN Suganthan, and Q Chen. Problem definitions and evaluation criteria for the cec 2015 competition on learning-based real-parameter single objective optimization. In *Proceedings of the 2015 IEEE Congress on Evolutionary Computation*, 2015.

- [44] Carola Doerr, Hao Wang, Furong Ye, Sander Van Rijn, and Thomas Bäck. Iohprofiler: A benchmarking and profiling tool for iterative optimization heuristics. *arXiv preprint arXiv:1810.05281*, 2018.
- [45] R.Turner and D.Eriksson. *Bayesmark: Benchmark framework to easily compare bayesian optimization methods on real machine learning tasks*, 2019. URL <https://github.com/uber/bayesmark>.
- [46] Jakub Kudela and Radomil Matousek. New benchmark functions for single-objective optimization based on a zigzag pattern. *IEEE Access*, 2022.
- [47] Abhishek Kumar, Guohua Wu, Mostafa Z Ali, Rammohan Mallipeddi, Ponnuthurai Nagaratnam Suganthan, and Swagatam Das. A test-suite of non-convex constrained optimization problems from the real-world and some baseline results. *Swarm and Evolutionary Computation*, 2020.
- [48] Diederick Vermetten, Furong Ye, Thomas Bäck, and Carola Doerr. Ma-bbob: Many-affine combinations of bbob functions for evaluating automl approaches in noiseless numerical black-box optimization contexts. In *International Conference on Automated Machine Learning*. PMLR, 2023.
- [49] Ke Xue, Ruo-Tong Chen, Rong-Xi Tan, Xi Lin, Yunqi Shi, Siyuan Xu, Mingxuan Yuan, and Chao Qian. Bboplace-bench: Benchmarking black-box optimization for chip placement. 2024.
- [50] Qiqi Duan, Guochen Zhou, Chang Shao, Zhuowei Wang, Mingyang Feng, Yuwei Huang, Yajing Tan, Yijun Yang, Qi Zhao, and Yuhui Shi. Pypop7: A pure-python library for population-based black-box optimization. *Journal of Machine Learning Research*, 2024.
- [51] Beichen Huang, Ran Cheng, Zhuozhao Li, Yaochu Jin, and Kay Chen Tan. Evox: A distributed gpu-accelerated framework for scalable evolutionary computation. *IEEE Transactions on Evolutionary Computation*, 2024.
- [52] Wei Li, Peng Liang, Bo Sun, Yafeng Sun, and Ying Huang. Reinforcement learning-based particle swarm optimization with neighborhood differential mutation strategy. *Swarm and Evolutionary Computation*, 2023.
- [53] Jiale Hong, Bo Shen, and Anqi Pan. A reinforcement learning-based neighborhood search operator for multi-modal optimization and its applications. *Expert Systems with Applications*, 2024.
- [54] Ye Tian, Xiaopeng Li, Haiping Ma, Xingyi Zhang, Kay Chen Tan, and Yaochu Jin. Deep reinforcement learning based adaptive operator selection for evolutionary multi-objective optimization. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 2022.
- [55] Shuijia Li, Wenying Gong, Ling Wang, and Qiong Gu. Evolutionary multitasking via reinforcement learning. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 2023.
- [56] Eckart Zitzler, Kalyanmoy Deb, and Lothar Thiele. Comparison of multiobjective evolutionary algorithms: Empirical results. *Evolutionary Computation*, 2000.
- [57] Kalyanmoy Deb, Lothar Thiele, Marco Laumanns, and Eckart Zitzler. Scalable multi-objective optimization test problems. In *Proceedings of the 2002 Congress on Evolutionary Computation. CEC'02 (Cat. No. 02TH8600)*. IEEE, 2002.
- [58] Qingfu Zhang, Aimin Zhou, Shizheng Zhao, Ponnuthurai Nagaratnam Suganthan, Wudong Liu, Santosh Tiwari, et al. Multiobjective optimization test instances for the cec 2009 special session and competition. 2008.
- [59] Simon Huband, Philip Hingston, Luigi Barone, and Lyndon While. A review of multiobjective test problems and a scalable test problem toolkit. *IEEE Transactions on Evolutionary Computation*, 2006.
- [60] Xiaodong Li, Ke Tang, Mohammad N Omidvar, Zhenyu Yang, Kai Qin, and Hefei China. Benchmark functions for the cec 2013 special session and competition on large-scale global optimization. *gene*, 2013.

- [61] Xiaodong Li, Andries Engelbrecht, and Michael G Epitropakis. Benchmark functions for cec'2013 special session and competition on niching methods for multimodal function optimization. *RMIT University, evolutionary computation and machine learning Group, Australia, Tech. Rep*, 2013.
- [62] Bingshui Da, Yew-Soon Ong, Liang Feng, A Kai Qin, Abhishek Gupta, Zexuan Zhu, Chuan-Kang Ting, Ke Tang, and Xin Yao. Evolutionary multitasking for single-objective continuous optimization: Benchmark problems, performance metric, and baseline results. *arXiv preprint arXiv:1706.03470*, 2017.
- [63] Sebastian Pineda Arango, Hadi S Jomaa, Martin Wistuba, and Josif Grabocka. Hpo-b: A large-scale reproducible benchmark for black-box hpo based on openml. *arXiv preprint arXiv:2106.06257*, 2021.
- [64] Howook Hwang, Thom Vreven, Joël Janin, and Zhiping Weng. Protein-protein docking benchmark version 4.0. *Proteins: Structure, Function, and Bioinformatics*, 2010.
- [65] Mhd Ali Shehadeh and Jakub Kudela. Benchmarking global optimization techniques for unmanned aerial vehicle path planning. *arXiv preprint arXiv:2501.14503*, 2025.
- [66] Félix-Antoine Fortin, François-Michel De Rainville, Marc-André Gardner Gardner, Marc Parizeau, and Christian Gagné. Deap: Evolutionary algorithms made easy. *The Journal of Machine Learning Research*, 2012.
- [67] Nikolaus Hansen, Youhei Akimoto, and Petr Baudis. CMA-ES/pycma on Github. Zenodo, DOI:10.5281/zenodo.2559634, February 2019. URL <https://doi.org/10.5281/zenodo.2559634>.
- [68] Jiayi Weng, Huayu Chen, Dong Yan, Kaichao You, Alexis Duburcq, Minghao Zhang, Yi Su, Hang Su, and Jun Zhu. Tianshou: A highly modularized deep reinforcement learning library. *Journal of Machine Learning Research*, 2022.
- [69] Philipp Moritz, Robert Nishihara, Stephanie Wang, Alexey Tumanov, Richard Liaw, Eric Liang, Melih Elibol, Zongheng Yang, William Paul, Michael I Jordan, et al. Ray: A distributed framework for emerging {AI} applications. In *13th USENIX symposium on operating systems design and implementation (OSDI 18)*, 2018.
- [70] Minyang Chen, Chenchen Feng, and Ran Cheng. Metade: Evolving differential evolution by differential evolution. *IEEE Transactions on Evolutionary Computation*, 2025.
- [71] Ryoji Tanabe and Alex Fukunaga. Success-history based parameter adaptation for differential evolution. In *2013 IEEE Congress on Evolutionary Computation*, 2013.
- [72] Janez Brest, Mirjam Sepesy Maučec, and Borko Bošković. Self-adaptive differential evolution algorithm with population size reduction for single objective bound-constrained optimization: Algorithm j21. In *2021 IEEE Congress on Evolutionary Computation (CEC)*, 2021.
- [73] Subhodip Biswas, Debanjan Saha, Shuvodeep De, Adam D Cobb, Swagatam Das, and Brian A Jalaian. Improving differential evolution through bayesian hyperparameter optimization. In *2021 IEEE Congress on Evolutionary Computation (CEC)*, 2021.
- [74] Di Wu and G Gary Wang. Employing reinforcement learning to enhance particle swarm optimization methods. *Engineering Optimization*, 2022.
- [75] Olaf Mersmann, Bernd Bischl, Heike Trautmann, Mike Preuss, Claus Weihs, and Günter Rudolph. Exploratory landscape analysis. In *Proceedings of the 13th Annual Conference on Genetic and Evolutionary Computation*, 2011.
- [76] Ramses Sala and Ralf Müller. Benchmarking for metaheuristic black-box optimization: perspectives and open challenges. In *2020 IEEE Congress on Evolutionary Computation (CEC)*, 2020.
- [77] Manuel López-Ibáñez, Diederick Vermetten, Johann Dreö, and Carola Doerr. Using the empirical attainment function for analyzing single-objective black-box optimization algorithms. *IEEE Transactions on Evolutionary Computation*, 2024.

- [78] Sietse Schröder, Mitra Baratchi, and Jan N van Rijn. Overfitting in combined algorithm selection and hyperparameter optimization. In *International Symposium on Intelligent Data Analysis*, 2025.
- [79] Yue-Jiao Gong, Jing-Jing Li, Yicong Zhou, Yun Li, Henry Shu-Hung Chung, Yu-Hui Shi, and Jun Zhang. Genetic learning particle swarm optimization. *IEEE transactions on cybernetics*, 2015.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The contributions of MetaBox-v2 includes not only its motivations accumulated from the users, but also the extensive development efforts of all authors. These aspects are all reflected in abstract and introduction.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We have provided a brief yet objective discussion on the limitations of MetaBox-v2 in the end of the main text (Conclusion section).

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: This paper mainly introduce our proposed behcmarking platform MetaBox v2.0, which does not include theoretical results.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: We have provided sufficient details about the experimental settings in Experimental Results section and Appendix. Assisted by these reproduction details and the project of MetaBox v2.0, one can easily reproduce the results and observations within this paper.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: MetaBox has not only opensourced its codes and assets on Github, but also provided well-organized and detailed documentation for users to learn and use MetaBox step by step. We have provided these resources links in the main text.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: These details can be found in the beginning of Experimental Results section and in the Appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: MetaBox is closely related with optimization domain, hence the optimization results presented in this paper all include error bars to reflect the statistical significance of multiple independent runs.

Guidelines:

- The answer NA means that the paper does not include experiments.

- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We have provided this information in the beginning of Experimental Results section.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Yes, we have double-checked our paper and related code resources, which conform with the NeurIPS Code of Ethics, in every respect.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: Yes, we have added an informative discussion at the end of the main text to list such potential issues, where we found that there is no critical negative societal impact.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper poses no such risks

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We have cited every existing assets we used in MetaBox in the main text. Furthermore, we have included a separate section in Appendix to list these assets and their corresponding project links & licenses.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.

- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: We have proposed several problem sets used for benchmarking MetaBBO approaches. For the usage, problem specification and benchmarking procedure details, we provide a comprehensive online documentation.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.