

Learning Joint Models of Prediction and Optimization

James Kotary^{a,1}, Vincenzo Di Vito^{a,1}, Jacob Christopher^a, Pascal Van Hentenryck^b and Ferdinando Fioretto^a

^aUniversity of Virginia

^bGeorgia Institute of Technology

Abstract. The Predict-Then-Optimize framework uses machine learning models to predict unknown parameters of an optimization problem from exogenous features before solving. This setting is common to many real-world decision processes, and recently it has been shown that decision quality can be substantially improved by solving and differentiating the optimization problem within an end-to-end training loop. However, this approach requires significant computational effort in addition to handcrafted, problem-specific rules for backpropagation through the optimization step, challenging its applicability to a broad class of optimization problems. This paper proposes an alternative method, in which optimal solutions are learned directly from the observable features by joint predictive models. The approach is generic, and based on an adaptation of the Learning-to-Optimize paradigm, from which a rich variety of existing techniques can be employed. Experimental evaluations show the ability of several Learning-to-Optimize methods to provide efficient and accurate solutions to an array of challenging Predict-Then-Optimize problems.

1 Introduction

The *Predict-Then-Optimize* (PtO) framework models decision-making processes as optimization problems whose parameters are only partially known while the remaining, unknown, parameters must be estimated by a machine learning (ML) model. The predicted parameters complete the specification of an optimization problem which is then solved to produce a final decision. The problem is posed as estimating the solution $\mathbf{x}^*(\zeta) \in \mathcal{X} \subseteq \mathbb{R}^n$ of a *parametric* optimization problem:

$$\mathbf{x}^*(\zeta) = \arg \min_{\mathbf{x}} f(\mathbf{x}, \zeta) \quad (1a)$$

$$\text{such that: } \mathbf{g}(\mathbf{x}) \leq 0, \quad \mathbf{h}(\mathbf{x}) = 0, \quad (1b)$$

given that parameters $\zeta \in \mathcal{C} \subseteq \mathbb{R}^p$ are unknown, but that a correlated set of observable values $\mathbf{z} \in \mathcal{Z}$ are available. Here f is an objective function, and $\mathbf{g}$ and $\mathbf{h}$ define the set of the problem's inequality and equality constraints. The combined prediction and optimization model is evaluated on the basis of the optimality of its downstream decisions, with respect to f under its ground-truth problem parameters [12]. This setting is ubiquitous to many real-world applications confronting the task of decision-making under uncertainty, such as planning the shortest route in a city, determining optimal power generation schedules, or managing investment portfolios. For example, a vehicle routing system may aim to minimize a rider's total commute time by solving a shortest-path optimization model (1) given knowledge of the transit

times ζ over each individual city block. In absence of that knowledge, it may be estimated by prediction models based on exogenous data $\mathbf{z}$, such as weather and traffic conditions. In this context, more accurately predicted transit times $\hat{\zeta}$ tend to produce routing plans $\mathbf{x}^*(\hat{\zeta})$ with shorter commutes, with respect to the true city-block transit times ζ .

However, direct training of predictions from observable features to problem parameters can generalize poorly with respect to the ground-truth optimality achieved by a subsequent decision model [21, 15]. To address this challenge, *End-to-end Predict-Then-Optimize* (EPO) [12] has emerged as a transformative paradigm in data-driven decision making, where predictive models are trained to directly minimize loss functions defined on the downstream optimal solutions $\mathbf{x}^*(\hat{\zeta})$.

On the other hand, EPO implementations present two key challenges: (i) They require backpropagation through the solution of the optimization problem (1) as a function of its parameters for end-to-end training. The required backpropagation rules are highly dependent on the form of the optimization model and are typically derived by hand analytically for limited classes of models [3, 1]. (ii) Furthermore, difficult decision models involving nonconvex or discrete optimization may not admit well-defined backpropagation rules.

To address these challenges, this paper outlines a framework for training Predict-Then-Optimize models by techniques adapted from a separate but related area of work that combines constrained optimization end-to-end with machine learning. This paradigm, called *Learn-to-Optimize* (LtO), learns a mapping between the parameters of an optimization problem and its corresponding optimal solutions using a deep neural network (DNN), as illustrated in Figure 1(c).

The resulting DNN mapping is then treated as an *optimization proxy* whose role is to repeatedly solve difficult, but related optimization problems in real time [29, 13]. Several LtO methods specialize in training proxies to solve difficult problem forms, especially those involving nonconvex optimization.

The methodology of this paper recognizes that existing LtO methods can provide an array of implementations for producing learned optimization proxies, which can handle hard optimization problem forms, have fast execution speeds, and are differentiable by construction. As such, they can be adapted to the Predict-Then-Optimize setting, offering an alternative to hard optimization solvers with handcrafted backpropagation rules. However, direct transfer of a pretrained optimization proxy into an EPO training loop leads to degradation of accuracy in the proxy solver, for which an end-to-end learning solution is proposed. The resulting *Learning to Optimize from Features* (LtOF) framework extends the Learn-to-Optimize problem setting to encompass that of Predict-Then-Optimize, by learning to construct optimal solutions directly from features, as illustrated in Figure 2(d).

Contributions. This paper makes the following novel contributions:

¹ Corresponding authors with equal contribution. Emails: jk4pn, eda8pc@virginia.edu

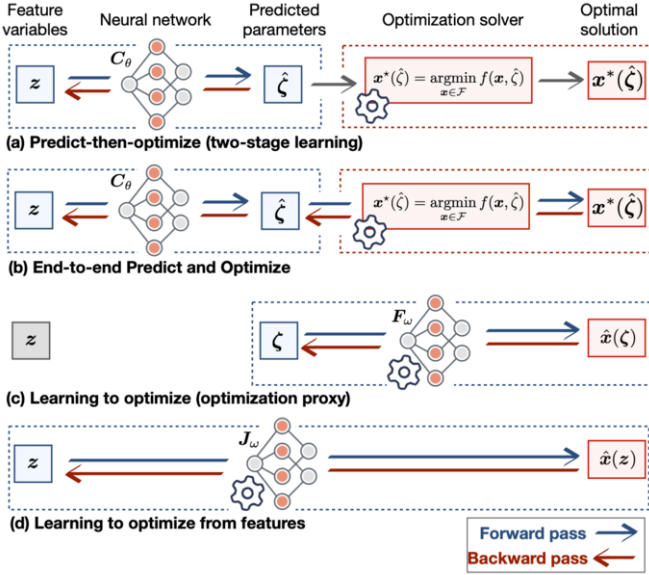


Figure 1: Illustration of Learning to Optimize from Features, in relation to other learning paradigms.

(1) It investigates the use of pretrained LtO proxy models as a means to approximate the decision-making component of an EPO pipeline, and demonstrates a distributional shift effect between prediction and optimization models that leads to loss of accuracy in EPO training. (2) It proposes Learning to Optimize from Features (LtOF), in which existing LtO methods are adapted to learn solutions to optimization problems directly from observable features, circumventing the distributional shift effect over the problem parameters. (3) The generic LtOF framework is evaluated by adapting several well-known LtO methods to solve Predict-then-Optimize problems with difficult optimization components. Besides outperforming conventional two-stage approaches, *the results show that difficult nonconvex optimization components can be incorporated into PtO pipelines naturally*, and illustrate how high decision quality can be reached in such cases even when gradient-based EPO alternatives fail.

2 Problem Setting and Background

In the Predict-then-Optimize (PtO) setting, three distributions of data are assumed. Observable features $z \sim \mathcal{Z}$ are correlated with the *unknown* coefficients $\zeta \sim \mathcal{C}$ of a parametric optimization problem (1), which in turn defines corresponding optimal solutions $x^*(\zeta) \sim \mathcal{X}$. We aim to learn optimal solutions $x^*(\zeta)$ to problem (1) without knowing the objective coefficients ζ , but knowing instead the correlated features z . More precisely, the goal is to learn a mapping $\hat{x}_\theta : \mathcal{Z} \rightarrow \mathcal{X}$ from observable features z to feasible solutions of (1b), while aiming to optimize their objective value (1a) under the ground truth coefficients ζ . Assuming a joint distribution $(z, \zeta) \sim \Omega$, this can be expressed as

$$\text{Minimize } \mathbb{E}_{(z, \zeta) \sim \Omega} [f(\hat{x}_\theta(z), \zeta)]. \quad (2)$$

A deterministic mapping $x^* : \mathcal{C} \rightarrow \mathcal{X}$ from problem coefficients to optimal solutions is defined by optimization problem (1), and can be implemented using any solution method which solves (1). As such, approaches to solving (2) are commonly based on framing $\hat{x}_\theta$ as a composite prediction-and-optimization model $x^* \circ C_\theta : \mathcal{Z} \rightarrow \mathcal{X}$, in which $C_\theta : \mathcal{Z} \rightarrow \mathcal{C}$ is a neural network trained to estimate problem coefficients ζ from the observable features z . The prediction and

optimization components C_θ and x^* are called, respectively, the *first* and *second* stage models.

Two main approaches are typically used to train the predictive component $\hat{\zeta} = C_\theta(z)$, in order to realize the training goal (2):

■ **Two-stage Method.** A conventional approach to training the parameter prediction model $\hat{\zeta} = C_\theta(z)$ is the *two-stage* method. It trains to predict the problem coefficients by MSE regression from their ground-truth values; i.e. with loss function $\ell(\hat{\zeta}, \zeta) = \|\hat{\zeta} - \zeta\|_2^2$, without accounting for the downstream optimization during training. This direct minimization of prediction errors is consistent with the goal (2) of optimizing the objective value $f(x^*(\hat{\zeta}), \zeta)$. However, by ignoring the effect of error propagation from predicted coefficients $\hat{\zeta}$ to downstream optimal solutions $x^*(\hat{\zeta})$ in the second stage, this naive approach is known to result in suboptimal objective values $f(x^*(\hat{\zeta}), \zeta)$ [12], yielding poor performance on the PtO goal (2).

■ **End-to-End Predict-Then-Optimize.** Improving on the two-stage method, the End-to-end Predict-Then-Optimize (EPO) approach trains C_θ directly to optimize the objective $f(x^*(\hat{\zeta}), \zeta)$ as a loss function by gradient descent. This is enabled by finding or approximating the derivatives through $x^*(\hat{\zeta})$. This corresponds to end-to-end training of the PtO goal (2), where $\hat{x}_\theta = x^* \circ C_\theta$:

$$\text{Minimize } \mathbb{E}_{(z, \zeta) \sim \Omega} [f(x^*(C_\theta(z)), \zeta)]. \quad (3)$$

EPO training consistently outperforms two-stage methods with respect to the goal (2), especially when the mapping $z \rightarrow \zeta$ is complex, due to the aforementioned error propagation effect. See Figure 1 (a) and (b) for an illustrative comparison, where the constraint set is denoted with $\mathcal{F}$. An overview of related work on EPO is reported in Section 6.

Challenges in End-to-End Predict-Then-Optimize

Despite their advantages over the two-stage, EPO methods are known to face two key challenges: (1) **Differentiability**: the need for hand-crafted backpropagation rules through $x^*(\zeta)$, which are highly dependent on the form of problem (1), and rely on the assumption of derivatives $\frac{\partial x^*}{\partial \zeta}$ which may not exist, and require that the mapping (1) is unique, producing a well-defined function; (2) **Efficiency**: the need to solve the optimization (1) to produce $x^*(\zeta)$ for each sample, in deployment and at each iteration of training. The results of Section 5 demonstrate a *further* potential pitfall of EPO training: even when differentiable, nonconvexity of (1) can cause its gradients to provide unhelpful descent directions for EPO training.

This paper is motivated by a need to address these disadvantages. To do so, it recognizes a body of work on training DNNs as *learned optimization proxies* which have fast execution, are automatically differentiable by design, and specialize in learning mappings $\zeta \rightarrow x^*(\zeta)$ of hard optimization problems. While the next section discusses why the direct application of learned proxies as differentiable optimization solvers in an EPO approach tends to fail, Section 4 presents a successful adaptation of the approach, in which predictive models are trained to solve the PtO training goal (2) *directly*.

3 EPO with Optimization Proxies

The Learning-to-Optimize problem setting encompasses a variety of distinct methodologies with the common goal of learning to solve optimization problems. This section characterizes the LtO setting, before proceeding to describe an adaptation of LtO methods to the Predict-Then-Optimize setting.

■ **Learning to Optimize.** The idea of training DNN models to emulate optimization solvers is referred to as *Learn-to-Optimize (LtO)*

[15]. Here the goal is to learn a mapping $\mathbf{F}_\omega : \mathcal{C} \rightarrow \mathcal{X}$ from the parameters ζ of an optimization problem (1) to its corresponding optimal solution $\mathbf{x}^*(\zeta)$ (see Figure 1 (c)). The resulting *proxy* optimization model has as its learnable component a DNN denoted $\hat{\mathbf{F}}_\omega$, which may be augmented with further operations $\mathcal{S}$ such as constraint corrections or unrolled solver steps, so that $\mathbf{F}_\omega = \mathcal{S} \circ \hat{\mathbf{F}}_\omega$. While training such a lightweight model to emulate optimization solvers is in general difficult, it is made tractable by restricting the task over a *limited distribution* of problem parameters $\zeta \sim \mathcal{C}$. A variety of LtO methods have been proposed, many specializing in learning to solve problems of a specific form. Some are based on supervised learning, where precomputed solutions $\mathbf{x}^*(\zeta)$ are required as target data in addition to parameters ζ for each sample. Others are *self-supervised*, requiring only knowledge of the problem form (1) along with instances of the parameters ζ for supervision in training. LtO methods employ special learning objectives to train the proxy model $\mathbf{F}_\omega$:

$$\text{Minimize}_{\omega} \mathbb{E}_{\zeta \sim \mathcal{C}} \left[\ell^{\text{LtO}}(\mathbf{F}_\omega(\zeta), \zeta) \right], \quad (4)$$

where ℓ^{LtO} represents a loss that is specific to the LtO method employed. A primary challenge in LtO is ensuring the satisfaction of constraints $\mathbf{g}(\hat{\mathbf{x}}) \leq 0$ and $\mathbf{h}(\hat{\mathbf{x}}) = 0$ by the solutions $\hat{\mathbf{x}}$ of the proxy model $\mathbf{F}_\omega$. This can be achieved, exactly or approximately, by a variety of methods, for example iteratively retraining Equation (4) while applying dual optimization steps to a Lagrangian loss function [13, 24], or designing $\mathcal{S}$ to restore feasibility [11], as reviewed in Section 5.4. In cases where small constraint violations remain in the solutions $\hat{\mathbf{x}}$ at inference time, they can be removed by post-processing with efficient projection or correction methods as deemed suitable for the particular application [15].

EPO with Pretrained Optimization Proxies

Viewed from the Predict-then-Optimize lens, learned optimization proxies have two beneficial features by design: (1) they enable very fast solving times compared to conventional solvers, and (2) are differentiable by virtue of being trained end-to-end. Thus, a natural question is whether it is possible to use a pre-trained optimization proxy to substitute the differentiable optimization component of an EPO pipeline. Such an approach modifies the EPO objective (2) as:

$$\text{Minimize}_{\theta} \mathbb{E}_{(z, \zeta) \sim \Omega} \left[f\left(\underbrace{\mathbf{F}_\omega(\mathbf{C}_\theta(z))}_{\hat{\mathbf{x}}}, \zeta\right) \right], \quad (5)$$

where the solver output $\mathbf{x}^*(\zeta)$ of problem (2) is replaced by the prediction $\hat{\mathbf{x}}$ obtained by LtO model $\mathbf{F}_\omega$ on input $\hat{\zeta}$ (gray highlights a pretrained model, with frozen weights ω).

However, a critical challenge in LtO lies in the inherent limitation that ML models act as reliable optimization proxies *only within the distribution of inputs they are trained on*. This challenges the implementation of the idea of using pretrained LtOs as components of an end-to-end Predict-Then-Optimize model as the weights θ update during training, leading to continuously evolving inputs $\mathbf{C}_\theta(z)$ to the pretrained optimizer $\mathbf{F}_\omega$. Thus, to ensure robust performance, $\mathbf{F}_\omega$ must generalize well across virtually any input during training. However, due to the dynamic nature of θ , there is an inevitable *distribution shift* in the inputs to $\mathbf{F}_\omega$, destabilizing the EPO training.

Figures 2 and 3 illustrate this issue. The former highlights how the input distribution to a pretrained proxy drifts during EPO training, impacting both output and backpropagation. The latter quantifies

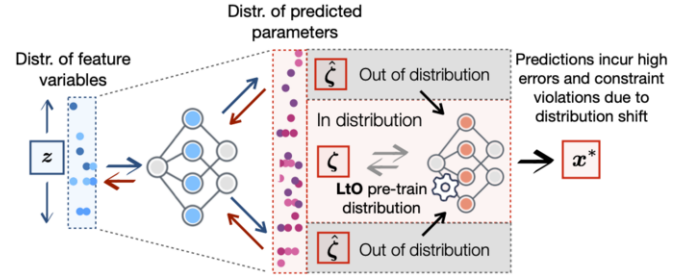


Figure 2: Effect of shifting inputs received by the LtO proxy: a mismatch from its initial training distribution leads to inaccurate solutions when employed in PtO training.

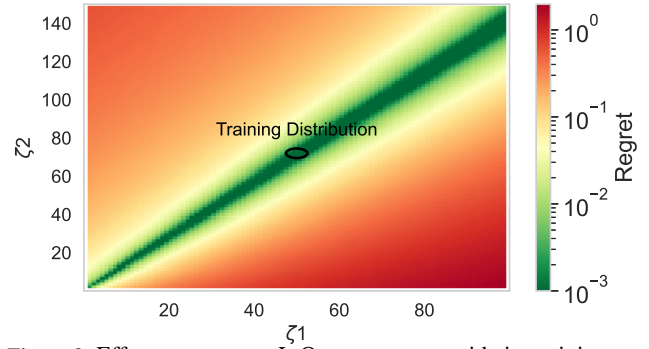


Figure 3: Effect on regret as LtO proxy acts outside its training set.

this behavior, exemplified on a simple two-dimensional problem (described in Appendix A of [17]), showing rapid increase in proxy regret as $\hat{\zeta}$ diverges from the initial training distribution $\zeta \sim \mathcal{C}$ (in black).²

The experimental results of Table 2 reinforce these observations. While each proxy solver performs well within its training distribution, their effectiveness deteriorates sharply when utilized as in equation 5. A step toward resolving this distribution shift issue allows the weights of $\mathbf{F}_\omega$ to adapt to its changing inputs, by *jointly* training the prediction and optimization models:

$$\text{Minimize}_{\theta, \omega} \mathbb{E}_{(z, \zeta) \sim \Omega} \left[f\left(\underbrace{\mathbf{F}_\omega(\mathbf{C}_\theta(z))}_{\hat{\mathbf{x}}}, \zeta\right) \right]. \quad (6)$$

The predictive model $\mathbf{C}_\theta$ is then effectively absorbed into the predictive component of $\mathbf{F}_\omega$, resulting in a *joint* prediction and optimization proxy model $\mathbf{J}_\phi = \mathbf{F}_\omega \circ \mathbf{C}_\theta$, where $\phi = (\omega, \theta)$. Given the requirement for feasible solutions, the training objective (6) must be replaced with an LtO procedure that enforces the constraints on its outputs. This leads to the joint training framework presented next.

4 Learning to Optimize from Features

The distribution shift effect described above arises due to the disconnect in training between the first-stage prediction network $\mathbf{C}_\theta : \mathcal{Z} \rightarrow \mathcal{C}$ and the second-stage optimization proxy $\mathbf{F}_\omega : \mathcal{C} \rightarrow \mathcal{X}$. However, the Predict-Then-Optimize setting (see Section 2) ultimately only requires the combined model to produce a candidate optimal solution $\hat{\mathbf{x}} \in \mathcal{X}$ given an observation of features $z \in \mathcal{Z}$. Thus, the intermediate prediction $\hat{\zeta} = \mathbf{C}_\theta(z)$ in Equation (6) is, in principle, not needed. This motivates the choice to learn direct mappings from features to optimal solutions of the second-stage decision problem. The joint

² The Appendix [17] can be found online at <https://arxiv.org/pdf/2311.13087>.

model $\mathbf{J}_\phi : \mathcal{Z} \rightarrow \mathcal{X}$ is trained by LtO procedures, employing

$$\underset{\phi}{\text{Minimize}} \quad \mathbb{E}_{(\mathbf{z}, \zeta) \sim \Omega} \left[\ell^{\text{LtO}}(\mathbf{J}_\phi(\mathbf{z}), \zeta) \right]. \quad (7)$$

This method can be seen as a generalization of the Learn-to-Optimize framework, to the Predict-then-Optimize setting. The key difference from the typical LtO setting is that problem parameters $\zeta \in \mathcal{C}$ are not known as inputs to the model, but the correlated features $\mathbf{z} \in \mathcal{Z}$ are known instead. Therefore, estimated optimal solutions now take the form $\hat{\mathbf{x}} = \mathbf{J}_\phi(\mathbf{z})$ rather than $\hat{\mathbf{x}} = \mathbf{F}_\omega(\zeta)$. Notably, this causes the self-supervised LtO methods to become *supervised*, since the ground-truth parameters $\zeta \in \mathcal{C}$ now act only as target data while the feature variable $\mathbf{z}$ takes the role of input data.

We refer to this approach as *Learning to Optimize from Features (LtOF)*. Figure 1 illustrates the key distinctions of LtOF relative to the other learning paradigms studied in this work. Figures (1c) and (1d) distinguish LtO from LtOF by a change in model’s input space, from $\zeta \in \mathcal{C}$ to $\mathbf{z} \in \mathcal{Z}$. This brings the framework into the same problem setting as that of the two-stage and end-to-end PtO approaches, shown in Figures (1a) and (1b). The key difference from the PtO approaches is that they produce an estimated optimal solution $\mathbf{x}^*(\hat{\zeta})$ by using a true optimization solver, but applied to an imperfect parametric prediction $\hat{\zeta} = \mathbf{C}_\theta(\mathbf{z})$. In contrast, LtOF directly estimates optimal solution $\hat{\mathbf{x}}(\mathbf{z}) = \mathbf{J}_\phi(\mathbf{z})$ from features $\mathbf{z}$, circumventing the need to represent an estimate of ζ .

Advantages of LtOF Section 5 demonstrates various LtOF implementations which can greatly outperform two-stage methods in terms of the optimality of their learned solutions, and are competitive with EPO training based on exact differentiation through $\mathbf{x}^*(\zeta)$. In contrast to EPO, this is achieved *without* access to exact optimization solvers, nor models of their derivatives. Two advantages of LtOF over EPO training, demonstrated in Section 5, are emphasized next.

4.1 Efficiency Benefits

Because the primary goal of the Learn-to-Optimize methodology is to achieve *faster solving times* than are possible with conventional optimization solvers, the LtOF approach broadly inherits this advantage. As learned neural network mappings, LtOF models of joint prediction and optimization can have order-of-magnitude lower runtimes than other PtO methods which require to solve the full optimization problem (1) at inference time (see Table 1). This enables the design of *real-time* PtO models within the LtOF framework.

4.2 Modeling Benefits

While EPO approaches require the use of problem-specific backpropagation rules, the LtOF framework instead requires an existing LtO implementation which can learn to solve the PtO problem’s second-stage optimization component. Section 5.2 shows how several LtOF implementations can succeed where EPO training fails, on a problem with a nonconvex oscillating objective term (see Figure 5). This result is significant but intuitive, since derivatives through a nonconvex function often do not correspond to useful descent directions for minimization. By contrast, the LtOF approach learns to construct solutions to the nonconvex problem directly from features, without relying on their derivatives, by drawing from existing LtO methods [13, 24, 11] that reliably learn to solve both convex and nonconvex optimization.

5 Experimental Results

The LtOF approach is evaluated against *two-stage* and *EPO* baselines, on three Predict-Then-Optimize tasks, each with a distinct second stage optimization component $\mathbf{x}^* : \mathcal{C} \rightarrow \mathcal{X}$, as in equation 1. These include a convex quadratic program (QP), a nonconvex QP variant, and a nonconvex program with sinusoidal constraints, to showcase the flexibility of LtOF over various problem forms.

Performance Criteria. Each PtO method considered in this section is evaluated on the basis of its downstream decisions $\hat{\mathbf{x}}$, which are required to be *feasible* to the problem constraints (1b). Subject to feasibility, the object is to minimize the expected ground-truth objective $f(\hat{\mathbf{x}}, \zeta)$ as per (2). This is equivalent to minimizing expected *regret*, defined as the magnitude of suboptimality of a solution $\hat{\mathbf{x}}$ to problem (1) with respect to the ground-truth parameters:

$$\text{regret}(\hat{\mathbf{x}}, \zeta) = f(\hat{\mathbf{x}}, \zeta) - f(\mathbf{x}^*(\zeta), \zeta). \quad (8)$$

LtOF methods. Three different LtOF implementations are evaluated on each PtO task, based on distinct Learn-to-Optimize methods, reviewed in detail in Section 5.4:

- **Lagrangian Dual Learning (LD)** [13], which augments a regression loss with penalty terms, updated to mimic a Lagrangian Dual ascent method to encourage the satisfaction of the problem’s constraints.
- **Self-supervised Primal Dual Learning (PDL)** [24], which uses an augmented Lagrangian function to perform joint self-supervised training of primal and dual networks for solution estimation.
- **Deep Constraint Completion and Correction (DC3)** [11], which relies on a completion technique to enforce constraint satisfaction, while maximizing the empirical objective function in self-supervised training.

While several other Learn-to-Optimize methods have been proposed in the literature, the above-described collection represents diverse subset which is used to demonstrate the potential of adapting the LtO methodology as a whole to the PtO setting.

Feature generation. End-to-End Predict-Then-Optimize methods integrate learning and optimization to minimize the propagation of prediction errors—specifically, from feature mappings $\mathbf{z} \rightarrow \zeta$ to the resulting decisions $\mathbf{x}^*(\zeta)$ (regret). It’s crucial to recognize that *even methods with high error propagation* can yield low regret *if the prediction errors are low*. To account for this, EPO studies often employ synthetically generated feature mappings to control prediction task difficulty [21]. Accordingly, for each experiment, we generate feature datasets $(\mathbf{z}_1, \dots, \mathbf{z}_N) \in \mathcal{Z}$ from ground-truth parameter sets $(\zeta_1, \dots, \zeta_N) \in \mathcal{C}$ using random mappings of increasing complexity. A feedforward neural network, $\mathbf{G}^k$, initialized uniformly at random with k layers, serves as the feature generator $\mathbf{z} = \mathbf{G}^k(\zeta)$. Evaluation is then carried out for each PtO task on feature datasets generated with $k \in \{1, 2, 4, 8\}$, keeping target parameters ζ constant.

Baselines. In our experiments, LtOF models use feedforward networks with k hidden layers. For comparison, we also evaluate two-stage and, where applicable, EPO models, using architectures with m hidden layers for each $m \in \{1, 2, 4, 8\}$. Additionally, for each task and each $m \in \{1, 2, 4, 8\}$, we evaluate EPO model with pretrained optimization proxies, which predictive model $\mathbf{C}_\theta$ is, or is not, pretrained by MSE regression, as the *two-stage* predictive model. Further training specifics are provided in Appendix B of [17].

Comparison to LtO setting. It is natural to ask how solution quality varies when transitioning from LtO to LtOF in a PtO setting, where solutions are learned directly from features. To address this question, each PtO experiment includes results from its analogous Learning to Optimize setting, where a DNN $\mathbf{F}_\omega : \mathcal{C} \rightarrow \mathcal{X}$ learns a mapping from

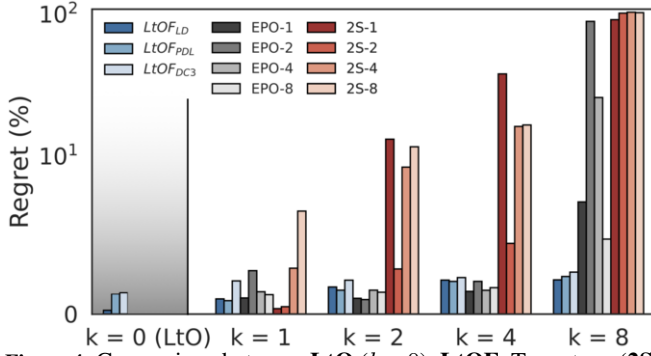


Figure 4: Comparison between **LtO** ($k=0$), **LtOF**, Two-stage (**2S**) and **EPO** ($k > 1$) on the portfolio optimization. **2S(EPO)- m** indicates that the prediction model of the respective PtO method is an m layer ReLU neural network. Plot y-axis is in semi log-scale.

	Method	Portfolio	N/conv. QP	AC-OPF
LtOF	LD it	0.0003	0.0000	0.0004
	LD fct	0.0000	0.0032	0.0516
	PDL it	0.0003	0.0000	0.0006
	PDL fct	0.0000	0.0032	0.0207
	DC3 it	0.0011	0.0001	-
	DC3 fct	0.0003	0.0000	-
PtO	PtO-1 et	0.0054	0.0122	2.5922
	PtO-2 et	0.0059	0.0104	2.5841
	PtO-4 et	0.0062	0.0123	2.5835
	PtO-8 et	0.0067	0.0133	2.5907

Table 1: Execution (*et*), inference (*it*), and feasibility correction (*fct*) times for **LtOF** and **PtO** (in seconds) for each sample. **Two-stage** methods execution times are comparable to PtO's ones.

the parameters ζ of an optimization problem to its corresponding solution $x^*(\zeta)$. This is denoted $k=0$ (LtO), indicating the absence of any feature mapping. All figures report the regret obtained by LtO methods for reference, although they are not directly comparable to the Predict-then-Optimize setting.

Finally, all reported results are averages across 20 random seeds and we refer the reader to Appendix B of [17] for extensive additional details regarding experimental settings, architectures, data generation and hyperparameters adopted.

5.1 Convex Quadratic Optimization

A well-known problem combining prediction and optimization is the Markowitz Portfolio Optimization [26]. This task has as its optimization component a convex Quadratic Program (QP):

$$x^*(\zeta) = \arg \max_{x \geq 0} \zeta^T x - \lambda x^T \Sigma x, \quad \text{s.t. } \mathbf{1}^T x = 1 \quad (9)$$

in which parameters $\zeta \in \mathbb{R}^D$ represent future asset prices, and decisions $x \in \mathbb{R}^D$ represent their fractional allocations within a portfolio. The objective is to maximize a balance of risk, as measured by the quadratic form covariance matrix Σ , and total return $\zeta^T x$. Historical prices of $D = 50$ assets are obtained from the Nasdaq online database [23] and used to form price vectors ζ_i , $1 \leq i \leq N$, with $N = 12,000$ individual samples collected from 2015-2019. In the outputs $\hat{x}$ of each LtOF method, any feasibility violations are restored, at *low computational cost*, by first clipping $[\hat{x}]_+$ to satisfy $x \geq 0$, then dividing by its sum to satisfy $\mathbf{1}^T x = 1$. The convex solver *cvxpy* [10] is used as the optimization component in each baseline method.

Results. Figure 4 shows the percentage regret due to LtOF variants based on *LD*, *PDL* and *DC3*. Two-stage and EPO models are evaluated for comparison, with predictive components given various numbers of layers. For feature complexity $k > 1$, each LtOF model outperforms the best two-stage model, increasingly with k and up to nearly *two orders of magnitude* when $k = 8$. The EPO model, trained using exact derivatives through (9) using the differentiable solver in *cvxpylayers* [1] is competitive with LtOF until $k = 4$, beyond which its best variant is outperformed by each LtOF variant. This shows the ability of LtOF models to reach high accuracy under complex feature mappings *without* access to optimization solvers *or* their derivatives, in training or inference, in contrast to EPO methods.

Table 1 presents LtOF inference times (*it*) and feasibility correction times (*fct*), compared with the per-sample execution times (*et*) for PtO methods. Run times for two-stage methods are closely aligned with those of EPO, and thus omitted. Notice how LtOF methods are at least an order of magnitude faster than the PtO baselines.

5.2 Nonconvex QP Variant

As a step in difficulty beyond convex QPs, this experiment considers a generic QP problem augmented with an oscillating objective term, resulting in a *nonconvex* optimization problem:

$$x^*(\zeta) = \arg \min_x \frac{1}{2} x^T Q x + \zeta^T \sin(x) \quad (10a)$$

$$\text{s.t. } Ax = b, Gx \leq h. \quad (10b)$$

A variant of this formulation was used to evaluate the LtO methods proposed both in [11] and in [24]. Following those works, $0 \preceq Q \in \mathbb{R}^{n \times n}$, $A \in \mathbb{R}^{n_{\text{eq}} \times n}$, $b \in \mathbb{R}^{n_{\text{eq}}}$, $G \in \mathbb{R}^{n_{\text{ineq}} \times n}$, $h \in \mathbb{R}^{n_{\text{ineq}}}$ and each ζ_i has elements drawn from a normal distribution. The EPO baseline for comparison differentiates (10) via the fixed-point conditions of a locally convergent Projected Gradient Descent method, implemented using the *fold-opt* library [16]; details on this EPO model are in Appendix B of [17]. Feasibility is restored by a projection onto the feasible set, efficiently calculated by as a *convex* QP. The problem dimensions are $n = 50$, $n_{\text{eq}} = 25$, and $n_{\text{ineq}} = 25$.

Results. Figure 5 (top) shows regret due to LtOF models based on *LD*, *PDL* and *DC3*, along with two baseline PtO methods. The best two-stage models perform poorly for all values of k , implying that the regret is particularly sensitive to prediction errors in the oscillating term. As alluded earlier, EPO training based on differentiation of (10) performs even worse. This is intuitive, since gradients of the objective in (10a) do not correspond to descent directions, due to its nonconvex $\sin$ term [7]. Thus, backpropagation through (10) guides the EPO model (3) to poor local minima in gradient descent training. The LtOF models achieve over $4\times$ times lower regret than the best baselines, suggesting strong potential for LtOF when predicting parameters of nonconvex objective functions. Additionally, the LtOF methods execute several times *faster* than both baselines, after restoring feasibility.

5.3 Nonconvex AC-Optimal Power Flow

Given a vector of marginal costs ζ for each generator in an electrical grid, the AC-Optimal Power Flow problem optimizes the generation and dispatch of electrical power from generators to nodes with predefined demands. The objective is to minimize cost, while meeting demand exactly. The full optimization problem is specified in Appendix A of [17], where a quadratic cost objective is minimized subject to

		Portfolio			Nonconvex QP			AC-OPF		
Method		$k = 2$	$k = 4$	$k = 8$	$k = 2$	$k = 4$	$k = 8$	$k = 2$	$k = 4$	$k = 8$
LtOF	LD Regret	1.7170	2.1540	2.1700	9.9279	9.7879	9.5473	0.0748	0.3762	0.7231
	LD Regret (*)	1.5739	2.0903	2.1386	9.9250	9.8211	9.5556	0.0013	0.0071	0.0195
	LD Violation (*)	0.0010	0.0091	0.0044	0.0148	0.0162	0.0195	0.0020	0.0037	0.0042
	PDL Regret	1.5150	2.0720	2.3830	7.2699	10.747	7.6399	0.8714	0.8012	0.8373
	PDL Regret (*)	1.4123	1.9372	2.0435	7.2735	10.749	7.6394	0.0260	0.0243	0.0242
	PDL Violation (*)	0.0001	0.0003	0.0003	0.0028	0.0013	0.0015	0.0000	0.0002	0.0002
	DC3 Regret	2.1490	2.3140	2.6600	14.271	11.028	10.666	-	-	-
	DC3 Regret (*)	2.1490	2.3140	2.6600	13.779	11.755	10.849	-	-	-
	DC3 Violation (*)	0.0000	0.0000	0.0000	0.5158	0.5113	0.5192	-	-	-
EPO	EPO Regret (Best)	0.9220	1.4393	4.7495	103.978	117.651	130.379	-	-	-
	EPO w/ Proxy Regret (Best)	154.40	119.31	114.69	812.75	804.26	789.50	389.04	413.89	404.74
	EPO w/ Proxy and pretrained prediction Regret (Best)	14.783	23.052	22.158	214.53	198.02	241.93	52.344	55.195	60.052
	Two-Stage Regret (Best)	2.8590	4.4790	91.326	36.168	37.399	38.297	1.4090	1.5280	2.4740

Table 2: Percentage Regret and Constraint Violations for all experiments. (*) denotes “Before Restoration”. Missing entries denote experiments not achievable with the associated method.

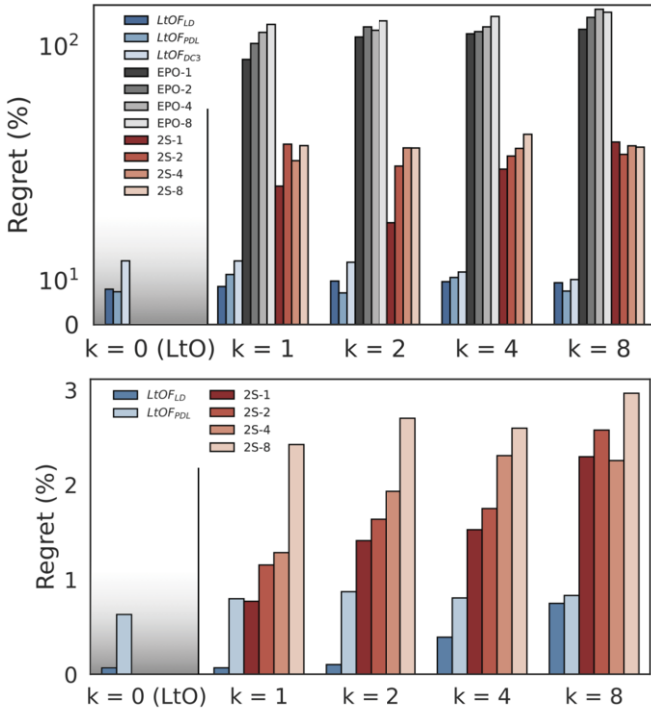


Figure 5: Comparison between LtO ($k = 0$), LtOF, and Two Stage Method (2S) on the nonconvex QP (top) and AC-OPF case (bottom). Top plot y-axis is in semi log-scale.

nonconvex physical and engineering power systems constraints. This experiment simulates an energy market situation in which generation costs are as-yet unknown to the power system planners, and must be estimated based on correlated data. The goal is to predict costs so as to minimize cost-regret over an example network with 54 generators, 99 demand loads, and 118 buses taken from the NESTA energy system test case archive [9]. Cost data are generated as perturbations around the test case’s cost values as described in Appendix B of [17]. Feasibility is restored after LtOF by Newton’s method on the constraint violations as specified in Appendix A of [17].

Results. Figure 5 (bottom) presents regret percentages, comparing LtOF to a two-stage baseline. Note that no general EPO exists for handling such nonconvex decision components. It is found empirically that the Lagrangian Hessian is not positive-semidefinite, preventing

even a differentiable convex QP approximation at its optima, as leveraged in [30]. Further, the solving times reported in Table 1, due to Casadi solver [4], are prohibitively slow for EPO training. *DC3* is also omitted, following [24], due to its incompatibility with the LtO variant of this experiment. Notice how the *best* two-stage model is outperformed by the *LD* variant of LtOF for $k > 1$, and also by the *PDL* variant for $k > 2$. Notably, PDL appears robust to increases in the feature mapping complexity k . On the other hand, it is outperformed in each case by the LD variant. Notice how, in the complex feature mapping regime $k > 1$, the best LtOF variant achieves up to an order of magnitude improvement in regret relative to the most competitive two-stage method. Additionally, LtOF methods report orders-of-magnitude speed advantages in Table 1.

Table 2 collects an abridged set of accuracy results due to each LtOF implementation and PtO baseline, across all experimental tasks. In particular, average constraint violations and objective regret are shown in addition to regret after restoring feasibility. Infeasible results are shown in grey for purposes of comparing the regret loss due to restoration. Best results on each task are shown in bold. Additionally, regret achieved by the EPO framework with pretrained proxies (discussed in Section 3) are included. Average regrets between 10 and 1000 times higher than LtOF illustrate the effect of their distributional shifts on accuracy. Notice how, in the context of complex feature mappings $k \geq 2$, LtOF is competitive with EPO, while bringing substantial computational advantages, and consistently outperforms two-stage methods, often, beyond an order of magnitude in regret.

5.4 Learning to Optimize Methods

Finally, this section describes in more detail those LtO methods which were adapted to solve PtO problems by LtOF, in the above experiments. Each description below assumes a DNN model $\mathcal{F}$ and its weights ω , which acts on problem parameters ζ specifying an instance of problem (1), to produce an estimate of the optimal solution $\hat{x} := F_{\omega}(\zeta)$, so that $\hat{x} \approx x^*(\zeta)$.

Lagrangian Dual Learning (LD). Fioretto et al. [13] constructs the following modified Lagrangian as a loss function for training the predictions $\hat{x} = F_{\omega}(\zeta)$:

$$\mathcal{L}_{LD}(\hat{x}, \zeta) = \|\hat{x} - x^*(\zeta)\|_2^2 + \lambda^T [g(\hat{x}, \zeta)]_+ + \mu^T h(\hat{x}, \zeta). \quad (11)$$

At each iteration of LD training, the model F_ω is trained to minimize the loss $\mathcal{L}_{LD}$. Then, updates to the multiplier vectors λ and μ are calculated based on the average constraint violations incurred by the predictions $\hat{x}$, mimicking a dual ascent method [8]. In this way, the method minimizes a balance of constraint violations and proximity to the precomputed target optima $x^*(\zeta)$.

Self-Supervised Primal-Dual Learning (PDL). Park and Van Hentenryck [24] use an augmented Lagrangian loss function

$$\mathcal{L}_{PDL}(\hat{x}, \zeta) = f(\hat{x}, \zeta) + \hat{\lambda}^T g(\hat{x}, \zeta) + \hat{\mu}^T h(\hat{x}, \zeta) + \frac{\rho}{2} \left(\sum_j \nu(g_j(\hat{x})) + \sum_j \nu(h_j(\hat{x})) \right), \quad (12)$$

where ν measures the constraint violation. At each iteration of PDL training, a separate estimate of the Lagrange multipliers is stored for each problem instance, and updated by an augmented Lagrangian method [8] after training $\hat{x} = F_\omega(\zeta)$ to minimize equation 5.4. In addition to the primal network F_ω , a dual network $\mathcal{D}_\zeta$ learns to store updates of the multipliers for each instance, and predict them as $(\hat{\lambda}, \hat{\mu}) = \mathcal{D}_\zeta(\zeta)$ to the next iteration. The method is self-supervised, requiring no precomputation of target solutions for training.

Deep Constraint Completion and Correction (DC3). Donti et al. [11] use the loss function

$$\mathcal{L}_{DC3}(\hat{x}, \zeta) = f(\hat{x}, \zeta) + \lambda \| [g(\hat{x}, \zeta)]_+ \|^2 + \mu \| h(\hat{x}, \zeta) \|^2 \quad (13)$$

which combines a problem's objective value with two additional terms which aggregate the total violations of its equality and inequality constraints. The scalar multipliers λ and μ are not adjusted during training. However, feasibility of predicted solutions is enforced by treating $\hat{x} = \hat{F}_\omega(\zeta)$ as an estimate for only a subset of optimization variables. The remaining variables are completed by solving the under-determined equality constraints $h(x) = 0$ as a system of equations. Inequality violations are corrected by gradient descent on the their aggregated values $\| [g(\hat{x}, \zeta)]_+ \|^2$. These completion and correction steps form the function S , where $F_\omega(\zeta) = S \circ \hat{F}_\omega(\zeta)$.

6 Related Work

This section gives an overview of related work in the Predict-Then-Optimize setting. While the idea is general and has broader applications, differentiation through the optimization of (1) is central to EPO approaches. Backpropagation of parametric quadratic programming problems was introduced by [3], which implicitly differentiates the solution via its KKT conditions of optimality [7]. Agrawal et al. [2] followed by proposing a differentiable cone programming solver, which uses implicit differentiation of problem-specific optimality conditions. That framework is leveraged by [1] to solve and differentiate general convex programs, by pairing it with a symbolic system for conversion of convex programs to canonical convex cone programs.

For many practical problems with discrete structure, such as linear programs, the mapping defined by (1) does not have well-defined derivatives, necessitating a suitable approximation in EPO training. [12, 19] propose a surrogate loss function for (2) in cases where f is linear, which admits useful subgradients for stochastic gradient descent training. [31] proposes backpropagation through linear programs by adding a smooth quadratic term to the objective and differentiating the resulting QP problem via [3]. [6] also propose backpropagation through linear programs but by smoothing the mapping (1) through

random noise perturbations to the objective function. [25] form approximate derivatives through linear optimization of discrete variables, by using a finite difference approximation between a pair of solutions with perturbed input parameters. Finally, [18] use the barrier function of an interior point method with early stopping to provide a smoothed surrogate model for differentiable linear programming.

A few recent works have addressed the topic of PtO learning without computing or approximating derivatives through optimization. In [27], [28], [32] neural networks are trained to function as surrogate models of solution regret, for use in EPO training. In [20], PtO is recast as a learning to rank (LTR) problem and solved with various LTR methods, in favor of EPO training.

7 Discussion and Conclusions

While the typical role of Learning to Optimize is to accelerate the solution of optimization problems, this paper demonstrates a novel use case: solving problems in the Predict-Then-Optimize scope. The adaptations of LtO described in this paper bring distinct advantages in the PtO setting, including real-time inference and enhanced ability to handle some PtO problems with nonconvex optimization.

Another advantage of the Learning to Optimize from Features approach to PtO settings is its generic framework, which enables it to leverage a variety of existing techniques from the LtO literature. On the other hand, as such, a particular implementation of LtOF may inherit any limitations of the specific LtO method that it adopts. Future work should focus on understanding to what extent a broader variety of LtO methods can be applied to PtO settings; given the large variety of existing works in the area, such a task is beyond the scope of this paper. In particular, this paper does not investigate the use of *combinatorial* optimization proxies in learning to optimize from features. Such methods tend to use a distinct set of approaches from those studied in this paper, often relying on training by reinforcement learning [5, 14, 22], and are not suited for capturing broad classes of optimization problems. As such, this direction is left to future work.

One *disadvantage* inherent to LtOF, compared to EPO, is the inability to recover parameter estimations from the predictive model, since optimal solutions are predicted directly from features. Although it is not required for the PtO problem setting, this may pose a challenge if transferring the learned parameters to external solvers is desirable. Furthermore, LtOF cannot be applied to PtO problems whose optimization component does not have an effective LtO solution.

By showing that effective Predict-Then-Optimize models can consist solely of Learning-to-Optimize methods, this paper has aimed to provide a unifying perspective on these as-yet distinct problem settings. The flexibility of its approach has been demonstrated by showing superior performance over PtO baselines with diverse problem forms. As the advantages of LtO are often best realized in combination with application-specific techniques, it is hoped that future work can build on these findings to maximize the practical benefits offered by LtO in settings that require data-driven decision-making.

Acknowledgements

This research is partially supported by NSF grants RI-2007164, RI-2232054, EPCN-2242931, AI Institute-2112533, and CAREER-2143706. Its views and conclusions are those of the authors only.

References

- [1] A. Agrawal, B. Amos, S. Barratt, S. Boyd, S. Diamond, and J. Z. Kolter. Differentiable convex optimization layers. *Advances in neural information processing systems*, 32, 2019.
- [2] A. Agrawal, S. Barratt, S. Boyd, E. Busseti, and W. M. Moursi. Differentiating through a cone program. *arXiv preprint arXiv:1904.09043*, 2019.
- [3] B. Amos and J. Z. Kolter. Optnet: Differentiable optimization as a layer in neural networks. In *International Conference on Machine Learning*, pages 136–145. PMLR, 2017.
- [4] J. A. Andersson, J. Gillis, G. Horn, J. B. Rawlings, and M. Diehl. Casadi: a software framework for nonlinear optimization and optimal control. *Mathematical Programming Computation*, 11:1–36, 2019.
- [5] I. Bello, H. Pham, Q. V. Le, M. Norouzi, and S. Bengio. Neural combinatorial optimization with reinforcement learning. *arXiv:1611.09940*, 2017.
- [6] Q. Berthet, M. Blondel, O. Teboul, M. Cuturi, J.-P. Vert, and F. Bach. Learning with differentiable perturbed optimizers. *Advances in neural information processing systems*, 33:9508–9519, 2020.
- [7] S. Boyd, S. P. Boyd, and L. Vandenberghe. *Convex optimization*. Cambridge university press, 2004.
- [8] S. Boyd, N. Parikh, E. Chu, B. Peleato, J. Eckstein, et al. Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends® in Machine learning*, 3(1): 1–122, 2011.
- [9] C. Coffrin, D. Gordon, and P. Scott. Nesta, the nicta energy system test case archive. *arXiv preprint arXiv:1411.0359*, 2014.
- [10] S. Diamond and S. Boyd. Cvxpy: A python-embedded modeling language for convex optimization. *The Journal of Machine Learning Research*, 17(1):2909–2913, 2016.
- [11] P. L. Donti, D. Rolnick, and J. Z. Kolter. Dc3: A learning method for optimization with hard constraints. *arXiv preprint arXiv:2104.12225*, 2021.
- [12] A. N. Elmachtoub and P. Grigas. Smart “predict, then optimize”. *Management Science*, 2021.
- [13] F. Fioretto, P. V. Hentenryck, T. W. Mak, C. Tran, F. Baldo, and M. Lombardi. Lagrangian duality for constrained deep learning. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 118–135. Springer, 2020.
- [14] W. Kool, H. van Hoof, and M. Welling. Attention, learn to solve routing problems! In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2019.
- [15] J. Kotary, F. Fioretto, P. Van Hentenryck, and B. Wilder. End-to-end constrained optimization learning: A survey. In *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, IJCAI-21*, pages 4475–4482, 2021. doi: 10.24963/ijcai.2021/610. URL <https://doi.org/10.24963/ijcai.2021/610>.
- [16] J. Kotary, M. H. Dinh, and F. Fioretto. Backpropagation of unrolled solvers with folded optimization. *arXiv preprint arXiv:2301.12047*, 2023.
- [17] J. Kotary, V. D. Vito, J. Christopher, P. V. Hentenryck, and F. Fioretto. Learning joint models of prediction and optimization, 2024.
- [18] J. Mandi and T. Guns. Interior point solving for lp-based prediction+optimisation. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- [19] J. Mandi, P. J. Stuckey, T. Guns, et al. Smart predict-and-optimize for hard combinatorial optimization problems. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pages 1603–1610, 2020.
- [20] J. Mandi, V. Bucarey, M. M. K. Tchomba, and T. Guns. Decision-focused learning: Through the lens of learning to rank. In *International conference on machine learning*, pages 14935–14947. PMLR, 2022.
- [21] J. Mandi, J. Kotary, S. Berden, M. Mulamba, V. Bucarey, T. Guns, and F. Fioretto. Decision-focused learning: Foundations, state of the art, benchmark and future opportunities. *arXiv preprint arXiv:2307.13565*, 2023.
- [22] H. Mao, M. Schwarzkopf, S. B. Venkatakrishnan, Z. Meng, and M. Alizadeh. Learning scheduling algorithms for data processing clusters. In *Proceedings of the ACM special interest group on data communication*, pages 270–288, 2019.
- [23] Nasdaq. Nasdaq end of day us stock prices. <https://data.nasdaq.com/databases/EOD/documentation>, 2022. Accessed: 2023-08-15.
- [24] S. Park and P. Van Hentenryck. Self-supervised primal-dual learning for constrained optimization. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pages 4052–4060, 2023.
- [25] M. V. Pogančić, A. Paulus, V. Musil, G. Martius, and M. Rolinek. Differentiation of blackbox combinatorial solvers. In *International Conference on Learning Representations (ICLR)*, 2020.
- [26] M. Rubinstein. Markowitz’s “portfolio selection”: A fifty-year retrospective. *The Journal of finance*, 57(3):1041–1045, 2002.
- [27] S. Shah, K. Wang, B. Wilder, A. Perrault, and M. Tambe. Decision-focused learning without decision-making: Learning locally optimized decision losses. *Advances in Neural Information Processing Systems*, 35:1320–1332, 2022.
- [28] S. Shah, A. Perrault, B. Wilder, and M. Tambe. Leaving the nest: Going beyond local loss functions for predict-then-optimize. 2024.
- [29] N. Vesselinova, R. Steinert, D. F. Perez-Ramirez, and M. Boman. Learning combinatorial optimization on graphs: A survey with applications to networking. *IEEE Access*, 8:120388–120416, 2020.
- [30] K. Wang, A. Perrault, A. Mate, and M. Tambe. Scalable game-focused learning of adversary models: Data-to-decisions in network security games. In *AAMAS*, pages 1449–1457, 2020.
- [31] B. Wilder, B. Dilkina, and M. Tambe. Melding the data-decisions pipeline: Decision-focused learning for combinatorial optimization. In *Proceedings of the AAAI Conference on Artificial Intelligence (AAAI)*, volume 33, pages 1658–1665, 2019.
- [32] A. Zharmagambetov, B. Amos, A. Ferber, T. Huang, B. Dilkina, and Y. Tian. Landscape surrogate: Learning decision losses for mathematical optimization under partial information. *Advances in Neural Information Processing Systems*, 36, 2024.