

Meta-learning based Alternating Minimization Algorithm for Non-convex Optimization

Jing-Yuan Xia, Shengxi Li, Jun-Jie Huang, Zhixiong Yang, Imad Jaimoukha, and Deniz Gunduz *Fellow, IEEE*

Abstract—In this paper, we propose a novel solution for non-convex problems of multiple variables, especially for those typically solved by an alternating minimization (AM) strategy that splits the original optimization problem into a set of sub-problems corresponding to each variable, and then iteratively optimizes each sub-problem using a fixed updating rule. However, due to the intrinsic non-convexity of the original optimization problem, the optimization can be trapped into spurious local minimum even when each sub-problem can be optimally solved at each iteration. Meanwhile, learning-based approaches, such as deep unfolding algorithms, have gained popularity for non-convex optimization; however, they are highly limited by the availability of labelled data and insufficient explainability. To tackle these issues, we propose a meta-learning based alternating minimization (MLAM) method, which aims to minimize a partial of the global losses over iterations instead of carrying minimization on each sub-problem, and it tends to learn an adaptive strategy to replace the handcrafted counterpart resulting in advance on superior performance. The proposed MLAM maintains the original algorithmic principle, providing certain interpretability. We evaluate the proposed method on two representative problems, namely, bi-linear inverse problem: matrix completion, and non-linear problem: Gaussian mixture models. The experimental results validate the proposed approach outperforms AM-based methods.

Index Terms—Alternating Minimization, Meta-learning, Deep Unfolding, Matrix Completion, Gaussian Mixture Model.

I. INTRODUCTION

ITERATIVE minimization is one of the most widely used approaches in signal processing, machine learning and computer science. Typically, when dealing with multiple variables, these methods follow an alternating minimization (AM) based strategy that converts the original problem of multiple variables into an iterative minimization of a sequence of sub-problems corresponding to each variable while the rest of the variables are held fixed. However, due to the non-convexity

of the problem, the obtained solutions do not necessarily converge to a global optimum, even when all the sub-problems are solved optimally at each iteration, see [1] for examples. The major issue underlying the failure of AM when facing non-convexity is that a greedy and non-adaptive optimization rule is carried out when solving each sub-problem throughout the iterations. Therefore, it lacks sufficient adaptiveness and effectiveness in terms of handling local optimums.

Recent advances in deep learning have highlighted its success in obtaining promising results for non-convex optimization problems [2]–[5]. However, generic deep learning methods have limited generalization ability especially when test data is significantly different from the training data. This problem becomes more crucial in the ill-posed non-convex optimization tasks. The weak explainability of deep neural network behavior also questions its applicability in certain scenarios.

Deep unfolding [6] as an alternative learning-based approach has achieved significant success and popularity in solving various optimization problems. It improves model explainability by mapping a model-based iterative algorithm to a specific neural network architecture with learnable parameters. In this way, the mathematical principles of the original algorithm are maintained, thus leading to better generalization behavior [2], [7]–[12]. We note that most of the deep unfolding algorithms are designed for solving linear inverse problems and require supervised learning. There are also deep unfolding algorithms for solving non-convex optimization problems. Zhang *et al.* [9] propose to unfold alternating optimization for blind image super-resolution, which is typically ill-posed and non-convex. In [10], an unfolded WMMSE algorithm is proposed to estimate the parameter of the gradient descent step size for solving MISO beamforming problem, which is highly non-linear and non-convex. We can see from [9], [10], when solving non-convex problems, deep unfolding algorithms either only retain the iterative framework and replace all components by deep networks [9], or learn a minimum number of parameters but result in less effective performance [10]. There exists a trade off between achieving better performance with highly over-parameterized deep networks and retaining model explainability and generalization ability with minimum learnable parameters. This trade-off is also highly related to the amount of required labelled data and the request of prior knowledge. Therefore, it is essential to design a new approach that enables us to carry learning-based neural network model with interpretable optimization-inspired behavior in an unsupervised learning way.

Meta-learning has witnessed increasing importance in terms

Jing-Yuan Xia, Zhixiong Yang and Jun-Jie Huang are with College of Electronic Engineering and College of Computer, National University of Defense Technology, Changsha, 410073, China.

Shengxi Li is with College of Electronic Engineering, Beihang University, Beijing, 100191, China.

Imad Jaimoukha is with the Department of Electrical and Electronic Engineering, Imperial College London, SW7 2AZ, UK.

Deniz Gunduz is with the Department of Electrical and Electronic Engineering, Imperial College London, SW7 2AZ, UK and the Department of Engineering EnzoFerrari, University of Modena and Reggio Emilia (UNIMORE), Italy.

E-mail: {j.xia16, shengxi.li17, j.huang15, i.jaimoukha, d.gunduz}@imperial.ac.uk.

*Jun-Jie Huang is the corresponding author.

This work is supported by National Natural Science Foundation of China, projects 61921001 and 62022091

of strong adaptation ability in solving new tasks [13]–[19]. Unlike the standard supervised learning solutions, meta-learning does not focus on solving a specific task at hand but aims to learn domain-general knowledge in order to generate an adaptive solution for a series of new tasks. Typically, a meta-network collects domain-specific knowledge when solving each specific task, and then extracts domain-general knowledge across solving different tasks. Most of the popular meta-learning algorithms, such as model-agnostic meta-learning (MAML) [20], metric-based meta-learning [21] and learning-to-learn [22], share a hierarchical optimization structure that is composed of *inner* and *outer* procedures: the meta-network performs as an optimizer to solve specific tasks at inner procedure and the parameters of meta-network are then updated through the outer procedure. We note that the AM-based iterative method can also be regarded as sharing this bi-level optimization framework. However, the optimization strategy in an AM-based algorithms is commonly frozen and the inner optimization behavior is independent to the outer alternating procedure. Therefore, we are inspired to propose a meta-learning based alternating minimization (MLAM) algorithm for solving non-convex problems, which will enable inner optimization to be continuously updated with respect to the mutual knowledge extracted over outer steps.

The proposed MLAM algorithm is composed of two-level meta-learning, namely, the *upper-level* and *ground-level* meta learning. The upper-level meta-learning learns on a set of non-convex problems and aims to enhance the adaptability towards new problems. In contrast, the ground-level meta-learning learns on a sequence of sub-problems within each non-convex problem from the upper-level; and therefore, aims to find an adaptive and versatile algorithm for the sequential sub-problems. The overview structure of the proposed MLAM method is shown in Fig.1.

Specifically, the upper-level meta-learning learns to leverage the optimization experiences on a series of problems, while the learned algorithm in the ground-level meta-learning maintains the original inner-and-outer iterative structure as well as the algorithmic principles, but replaces the frozen and handcrafted algorithmic rule by a dynamic and adaptive meta-learned rule. In other words, it aims to learn an optimization strategy that is able to provide a "bird's-eye view" of the mutual knowledge extracted across outer loops for those sub-problems being optimized in the inner loops. Therefore, the learned strategy does not optimize each sub-problem locally and exhaustively through minimization; instead, it optimizes them by incorporating the global loss information with superior adaptability. Moreover, the proposed MLAM algorithm is able to solve optimization problems in an unsupervised manner. As a result, the proposed MLAM algorithm achieves better performance in non-convex problems, while requiring less (and even no) labelled data for training.

The main contributions of this paper are mainly three-fold:

- The core contribution is the proposed meta-learning based alternating minimization (MLAM) approach for non-convex optimization problems. In an unsupervised manner, MLAM achieves a less-greedy and adaptive

optimization strategy to learn a non-monotonic algorithm for solving non-convex optimization problems.

- The proposed MLAM takes a step further towards enhanced interpretability. The algorithmic principles of the original model-based iterative algorithm is fully maintained without the need to replace iterative operations with black-box deep neural networks.
- With extensive simulations, we have validated that the proposed MLAM algorithm achieves promising performances on the challenging problems of matrix completion and Gaussian mixture model (GMM). It is able to effectively solve these extremely difficult non-convex problems even when the traditional approaches fail.

The rest of this paper is organized as follows. Section II gives the background and a brief review of previous approaches. Section III introduces our proposed MLAM approach and presents an LSTM-based MLAM method. Section IV illustrates two representative applications of our MLAM method. Section V provides simulation results and Section VI concludes the paper.

II. RELEVANT PRIOR WORK

In this section, we will first briefly introduce the general problem formulation for multi-variable non-convex optimizations and the solution approaches. Then, we will demonstrate our motivation of proposing MLAM and review the relevant meta-learning approaches.

Non-convex optimization problems that involve more than one variable are of great practical importance, but are often difficult to be well accommodated. The underlying relationship between variables can be linear (e.g., product, convolution) or non-linear (e.g., logarithmic operation, exponential kernel). For illustration convenience, we consider a general optimization formulation over an intersection of two variables in the matrix form, which can be expressed in the form of:

$$(\hat{\mathbf{W}}, \hat{\mathbf{X}}) = \arg \min_{(\mathbf{W}, \mathbf{X}) \in \mathcal{W} \times \mathcal{X}} F(\mathbf{W}, \mathbf{X}), \quad (1)$$

where $F : \mathcal{W} \times \mathcal{X} \rightarrow \mathbb{R}$ is a non-convex function that describes the mapping between the observations $\mathbf{Y} = F(\mathbf{W}, \mathbf{X})$ and two variables $\mathbf{W} \in \mathcal{W}$ and $\mathbf{X} \in \mathcal{X}$.

A. Model-based Solutions

The model-based iterative algorithms [23]–[31] typically solve (1) by adopting an AM-based strategy. The basic idea is to sequentially optimize a sub-problem corresponding to each variable whilst keeping the other variable fixed. That is, starting from an arbitrary initialization $\mathbf{W}_0 \in \mathcal{W}$, the AM-based algorithm sequentially solves two sub-problems at the t -th iteration via:

$$\begin{aligned} \mathbf{X}_t &= \arg \min_{\mathbf{X} \in \mathcal{X}} f_{\mathbf{W}_{t-1}}(\mathbf{X}), \\ \mathbf{W}_t &= \arg \min_{\mathbf{W} \in \mathcal{W}} f_{\mathbf{X}_t}(\mathbf{W}), \end{aligned} \quad (2)$$

where $f_{\mathbf{W}_{t-1}}(\mathbf{X})$ and $f_{\mathbf{X}_t}(\mathbf{W})$ are the functions corresponding to $\mathbf{X}$ and $\mathbf{W}$, respectively, while fixing the other one to the value obtained in the previous iterations, i.e., $f_{\mathbf{W}_{t-1}}(\mathbf{X}) =$

$f(\mathbf{W}_{t-1}, \mathbf{X})$ and $f_{\mathbf{X}_t}(\mathbf{W}) = f(\mathbf{W}, \mathbf{X}_t)$. The solution of each sub-problem in (2) could be attained by a gradient-descent based iterative process,

$$\mathbf{X}_i = \mathbf{X}_{i-1} + \phi(\{\mathbf{X}_k\}_{k=0}^{i-1}, \nabla f_{\mathbf{W}_{t-1}}(\mathbf{X}_{i-1})), \quad (3)$$

where $\{\mathbf{X}_k\}_{k=0}^{i-1}$ represent the historical values of the parameters for i optimization steps, $\nabla f_{\mathbf{W}_{t-1}}(\mathbf{X}_{i-1})$ is the gradient of objective function on $\mathbf{X}_{i-1}$, and $\phi(\cdot)$ defines the variable updating rule of different algorithms. Algorithms such as ISTA [7] and WMMSE [10] formulate a closed-form solution when iteratively solving each sub-problem in (2), which is essentially equal to a first order stationary point obtained by gradient descent based methods as well.

Before proceeding further, we first introduce some concepts that will be used throughout this paper. We define the *overall problem* as the optimization problem with objective function $F(\mathbf{W}, \mathbf{X})$, which will be called the global loss function. We refer to the optimization problems over each of the variables as the *sub-problem*, and their objective functions $f_{\mathbf{W}_{t-1}}(\mathbf{X})$ and $f_{\mathbf{X}_t}(\mathbf{W})$ for $t \geq 1$, as local loss functions. We define the alternative iterations over the sub-problems as the outer-loop, and the iterative iterations for solving each sub-problem as the inner-loop.

The AM-based algorithm attempts to solve the overall problem by sequentially minimizing the two sub-problems $f_{\mathbf{W}_{t-1}}(\mathbf{X})$ and $f_{\mathbf{X}_t}(\mathbf{W})$. However, the AM-based algorithm does not necessarily converge to a global optimal solution. This could be due to two main reasons: 1) AM-based methods optimize over the local loss functions without fully utilizing the information from the global loss function, and 2) AM-based methods usually solve the local loss function greedily using the first order information, which may not necessarily lead to the best solution, that is, the global optima, in terms of the global loss function.

Addressing these two issues of the AM is non-trivial. The key difficult is that the variable optimized rules of the model-based solutions are frozen during the iteration, with respect to a certain update function, i.e., the ϕ in (3). Typically, ϕ is designed for optimizing each sub-problem greedily, as a result, is not expected to reach the global optimal solutions for the overall problem.

B. Learning-based Solutions

Different from the model-based approaches, the recent deep learning based methods typically require to train an over-parameterized deep neural networks in an end-to-end learning fashion with a large labelled dataset [7], [32]–[36]. During testing, the trained deep neural network is fed by the observations and directly outputs the estimated variables. The performance of these methods is highly bounded by the training datasets; however, the ground-truth data are neither sufficient nor even exist in realistic non-convex tasks, such as the GMM problems. Another shortcoming that is common to these learning-based methods is the weak explainability of the end-to-end deep neural network behavior.

Different from the generic deep learning approaches, deep unfolding algorithms try to combine model-based and

learning-based approaches. They have three major features: 1) They map the iterative optimization algorithm into a specific unfolded network architecture with trainable parameters; 2) each layer in the deep unfolding network corresponds to one iteration of the original iterative algorithm, while the number of layers, especially the iterations, is frozen; 3) similar to the other deep learning approaches, deep unfolding also requires pair-wise labelled data for training.

Mathematically, deep unfolding approaches follow the iterative framework as in equation (2), but replace the analytical minimization algorithm (or specific operators such as soft-thresholding or singular value thresholding) by neural networks in the form of:

$$\begin{aligned} \mathbf{X}_t &= \text{Layer}_t^{\mathbf{X}}(\mathbf{W}_{t-1}), \\ \mathbf{W}_t &= \text{Layer}_t^{\mathbf{W}}(\mathbf{X}_t), \end{aligned} \quad (4)$$

where the number of iterations is fixed with $t = T$. Hence the whole deep unfolded network is composed by T layers, where each layer is composed of several operators that reflect the mathematical behavior of the original iterative algorithm.

In [36], Zhang et al. propose a deep unfolded network for image super-resolution in which the solver for one variable is a generic deep network while the other one keeps consistent with the model-based solver. In deep alternating network [9], two networks, referring to an estimator and a restorer, work as two solvers for the splitted sub-problems. Therefore, the whole unfolding algorithm alternates between two network operators. In a different way, the unfolded network in [10] almost retains the original model-based iterative algorithm to solve a non-linear problem, but takes a network-based generator to learn the hyper-parameters for gradient descent step size by unsupervised learning. However, its performance does not surpass the counterpart model-based algorithm.

In summary, deep unfolding has made a step further towards better explainability; however, these approaches still perform as an end-to-end network behavior and mostly only enable interpretable alternating structure while replacing the original optimization-based algorithm with deep neural networks. Consequently, its interpretability is limited when applied to ill-posed non-convex problems, and it is based on a data-driven optimization strategy. Besides, most deep unfolding algorithms require a large number of labelled data for supervised learning. Hence, the performance of learning-based approaches on solving ill-posed non-convex problems, especially those without sufficient labelled data, is still limited.

III. META-LEARNING BASED ALTERNATING MINIMIZATION (MLAM) APPROACH

In this section, we introduce the proposed meta-learning based alternating minimization (MLMA) approach for solving optimization problems with multiple variables which are highly non-convex and ill-posed. As aforementioned, towards solving these problems, the existing model-based methods typically struggle, while on the other hand the lack of sufficient labelled training samples also restricts the performance of learning-based methods. Therefore, the key idea of our MLAM approach is to design a novel way that is not only capable of benefitting from both the optimization-based algorithmic

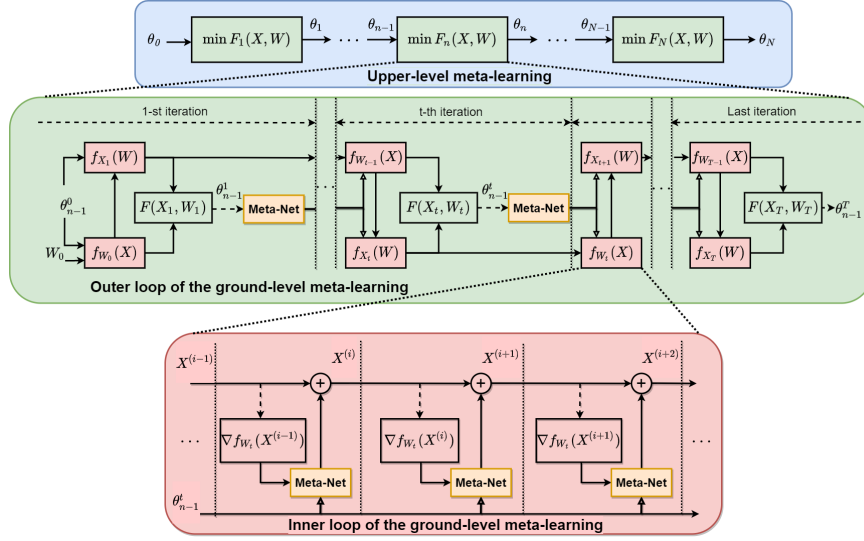


Fig. 1. The overall structure of MLAM algorithm. The parameters of the applied meta network are denoted by θ . The upper-level meta-learning is implemented on a set of non-convex problems $\{F_n(\mathbf{X}, \mathbf{W})\}_{n=1}^N$ as shown in the upper row. θ is continuously updated across different problems. The ground-level meta-learning is applied to solve each non-convex problem $F_n(\mathbf{X}, \mathbf{W})$ through an alternative outer loop between two inner loops corresponding to two variables, respectively, presented in the middle row. θ is dynamically updated with respect to the global loss $F(\mathbf{X}, \mathbf{W})$, where the variable is iteratively updated by the meta network. θ is frozen at inner loop iterations. The dashed lines denote gradient operator occurs and the solid line represents information flow along the edge.

principle and the superior performances by learning, but also surpass the AM-based strategy through meta-learning.

A. Overall Structure of the MLAM Approach

The proposed MLAM approach consists of two levels of meta-learning. The upper-level meta-learning operates on a set of overall problems and the ground-level meta-learning optimizes the sequential sub-problems within an overall problem. Both upper-level and ground-level operations contain a bi-level optimization structure that is composed of outer and inner procedures. The outer loop of the upper-level meta-learning continuously updates the parameters of the MetaNet θ across different overall problems, and each inner loop equals to one ground-level meta-learning on solving an overall problem. For both the ground-level and the upper-level meta-learning processes, we denote the inner loop index at superscripts and the outer loops index at subscripts. We will then introduce the details of MLAM in a bottom-up manner.

1) *Upper-level Meta-learning*: The upper-level meta-learning is depicted in Fig. 1 (the boxes in blue). It aims to extract a general knowledge of updating rules across different overall problems.

Each overall problem is considered as a single task, and the whole learning process is accommodated on a set of tasks within the same dimension. Therefore, the learnt algorithm is no longer designed for solving a single task, but a set of tasks $\{F_i(\mathbf{W}, \mathbf{X})\}_{i=1}^N$. Although the proposed MLAM model follows the AM structure, it establishes a new bridge between the upper-level and ground-level meta-learning by replacing the frozen update functions with meta-networks. This allows for variable updated at the inner loops being guided by global loss information from the outer loops, thus achieving a global scope optimization.

2) *Ground-level Meta-learning*: The general structure of the ground-level meta-learning is shown in Fig. 1 (the boxes in green and red). A ground-level meta-learning is performed by solving a specific overall problem $F_n(\mathbf{W}, \mathbf{X})$. This process contains an alternating optimization process on a sequence of sub-problems (boxes in green) which is named as the outer loop, and each sub-problem is solved by an iterative gradient descent based operation, which is defined as the inner loop (boxes in red).

In the inner loop, the variable (e.g., $\mathbf{X}$ and $\mathbf{W}$) to be optimized is updated based on the iterative gradient descent process using a meta network and the parameter of the MetaNet is frozen. Taking the optimization on $f_{W_{i-1}}(\mathbf{X})$ as an example, the i -th step update on the variable $\mathbf{X}$ at the inner loop can be expressed in the following form:

$$\mathbf{X}^{(i)} = \mathbf{X}^{(i-1)} + \text{MetaNet} \left(\nabla f_{W_{i-1}}(\mathbf{X}^{(i-1)}) \right), \quad (5)$$

where $\text{MetaNet}(\cdot)$ is a neural network with learnable parameter θ , and performs as an optimizer for variable update.

MetaNet replaces the handcrafted update function $\phi(\cdot)$ in (3) as a learnable and adaptive update function. Its parameter θ is frozen at the inner loop and will be updated at the outer loop with respect to the global loss function. Specifically, different from the traditional AM-based algorithms, MLAM establishes an extra update cue for the parameters of MetaNet at the outer loops, by minimizing the global loss $F(\mathbf{X}, \mathbf{W})$. This enables that the gradients of $f_{\mathbf{W}}$ and $f_{\mathbf{X}}$ on variables $\mathbf{X}$ and $\mathbf{W}$ and the gradient of $F(\mathbf{X}, \mathbf{W})$ on parameter θ are integrated into one circulating system. In this way, the optimization behavior on each sub-problem is no longer independent to the others, all of which interact through the MetaNet. In Fig.1, the dashed arrows at the outer loop of the ground-level meta-learning indicate the back-propagation of global loss to update the

parameters θ based on gradient descent. Taking the $t - th$ iteration as an instance, the parameter update is expressed as

$$\theta_{n-1}^t = \theta_{n-1}^{t-1} + \alpha \cdot \Phi(\nabla_{\theta_{n-1}^{t-1}} F(\mathbf{X}_1, \mathbf{W}_1)), \quad (6)$$

where α is the learning rate, and $\Phi(\cdot)$ denotes the learning algorithms for network update (e.g., the Adam method [37]).

Because the optimization landscapes for different sub-problems can be significantly different, MetaNet tends to extract a general knowledge on generating descent steps for variable updates across different sub-problems. Consequently, MetaNet learns to provide a superior and adaptive estimation for a sequence of sub-problems through this meta-learning process.

In Fig.1 the outer loop shows how the parameters at each iteration are updated. We define the update interval as the number of iterations for each parameter update, and in Fig.1 the update interval is set to 1. Generally speaking: the smaller the update interval the higher chances are updated, and the larger the update interval the more experience could be learned. The update interval of the parameters could be tuned. We will show more details in Section III-B and V-A.

Remark 1. The implementation of the two levels of meta-learning is indispensable. As the proposed MLAM learns in an unsupervised way, the training process of ground-level meta-learning could be also regarded as solving when training on one problem. Therefore, MLAM can be directly applied to a specific problem by only carrying on the ground-level meta-learning [19]. In this paper, the proposed approach works with both of the two levels of the meta-learning.

The meta-learning implemented on the proposed MLAM approach mainly contributes to two aspects. (i) Recalling to the upper-level meta-learning, the training strategy of the MetaNet follows the meta-learning strategy, thus improving the adaptation of the learned parameters on solving different non-convex optimization problems $\{F_i(\mathbf{W}, \mathbf{X})\}_{i=1}^N$. Specifically, the meta-learning behavior significantly enhance the capacity of solving different tasks by leveraging the experience of solving a series of $\{F_i(\mathbf{W}, \mathbf{X})\}_{i=1}^N$. (ii) The optimization performance of solving each specific task is dramatically improved by achieving a dynamic and adaptive gradient based updating rule through the implemented ground-level meta-learning. The solution strategy of solving each task $F_i(\mathbf{W}, \mathbf{X})$ is meta-learned over the alternating iterations. In this way, the solution strategy is leveraged by the meta-learning behavior to allow the sub-problems $\{f_{\mathbf{W}_i}(\mathbf{X}), f_{\mathbf{W}_i}(\mathbf{X})\}_{i=1}^T$ to be solved in a less greedy but more effective way. Essentially, benefit of the meta-learning, the MLAM is capacity of leveraging the experience of solving a series of sub-problems $\{f_{\mathbf{W}_i}(\mathbf{X}), f_{\mathbf{W}_i}(\mathbf{X})\}_{i=1}^T$ to learn a gradient based strategy that provides better convergence on solving the current task $F(\mathbf{W}, \mathbf{X})$.

B. MLAM with LSTM-based MetaNet

In this section, we introduce the implementation details of using recurrent neural networks (RNNs) as the MetaNet in the proposed MLAM algorithm. Specifically, the Long Short-Term Memory (LSTM) network [38] is adopted as the variable update function at the ground-level in the MLAM model.

The proposed LSTM-MLAM updates the parameters of the MetaNet with respect to the accumulated global losses.

RNN has a sequentially processing chain structure to achieve the capacity of “memory” on sequential data. LSTM [38] is one of the most well-known RNNs and is able to memorize and forget different sequential information. Memory is the most important feature of LSTM (RNN). It stores the status information of previous iterations and allows for the information to flow along the entire chain process. In this way, LSTM can integrate previous information with the current step input. Mathematically, the output of LSTM at the i -th iteration $\mathbf{H}^{(i)}$ is determined by the current gradient $\nabla f(\mathbf{x}_i)$ and the last cell state $\mathbf{C}^{(i)}$ in the following forms:

$$\begin{aligned} \mathbf{H}^{(i)} &= \text{LSTM}(\nabla f(\mathbf{x}_i), \mathbf{C}^{(i)}, \theta^{(i)}), \\ \mathbf{C}^{(i+1)} &= \mathbf{z}_f \odot \mathbf{C}^{(i)} + \mathbf{z}_i \odot \tilde{\mathbf{C}}^{(i-1)}, \end{aligned} \quad (7)$$

where $\text{LSTM}(\cdot, \cdot, \theta)$ denotes the LSTM network with parameters θ , $\odot$ denotes Hadamard Product, $\mathbf{z}_f$ and $\mathbf{z}_i$ are the vectors of intermediate conditions inside LSTM, and $\tilde{\mathbf{C}}^{(i-1)} = \nabla_{\theta^{(i-1)}} \mathcal{L}^{(i)}$ is the candidate cell state, referring to the gradient of current loss $\mathcal{L}^{(i)}$ over the last parameters $\theta^{(i-1)}$ in our problem.

We adopt two LSTM networks $\text{LSTM}_{\mathbf{X}}$ and $\text{LSTM}_{\mathbf{W}}$ as MetaNet to generate variable update functions, recalling $\phi(\cdot)$ in (3), for solving sub-problems corresponding to variables $\mathbf{X}$ and $\mathbf{W}$, respectively. We also denote $\theta_{\mathbf{X}}$ and $\theta_{\mathbf{W}}$ as the parameters of $\text{LSTM}_{\mathbf{X}}$ and $\text{LSTM}_{\mathbf{W}}$, and denote $\mathbf{C}_{\mathbf{X}}$ and $\mathbf{C}_{\mathbf{W}}$ as their cell states. The inputs of the LSTM are the gradient of local loss function and the sequential knowledge of variables, represented by cell state $\mathbf{C}$. Then the LSTM outputs the variable update term that integrates step size and direction together. Denoting the inner loop update steps $i-1$ and $j-1$ at superscripts, and outer loops steps $t-1$ and t at subscripts for each sub-problem, the variables are updated in the following forms:

$$\begin{aligned} \mathbf{H}_{\mathbf{X}}^{(i-1)} &= \text{LSTM}_{\mathbf{X}}(\nabla f_{\mathbf{W}_{t-1}}(\mathbf{X}^{(i-1)}), \mathbf{C}_{\mathbf{X}}^{(i-1)}, \theta_{\mathbf{X}}), \\ \mathbf{X}^{(i)} &= \mathbf{X}^{(i-1)} + \mathbf{H}_{\mathbf{X}}^{(i-1)}, \end{aligned} \quad (8)$$

and

$$\begin{aligned} \mathbf{H}_{\mathbf{W}}^{(j-1)} &= \text{LSTM}_{\mathbf{W}}(\nabla f_{\mathbf{X}_t}(\mathbf{W}^{(j-1)}), \mathbf{C}_{\mathbf{W}}^{(j-1)}, \theta_{\mathbf{W}}), \\ \mathbf{W}^{(j)} &= \mathbf{W}^{(j-1)} + \mathbf{H}_{\mathbf{W}}^{(j-1)}. \end{aligned} \quad (9)$$

As mentioned in Section III-A, at the inner loops, the parameters $\theta_{\mathbf{X}}$ and $\theta_{\mathbf{W}}$ are frozen, and are used to generate the update steps $\mathbf{H}_{\mathbf{X}}$ and $\mathbf{H}_{\mathbf{W}}$ for variables with frozen iteration numbers; therefore, the update strategy is essentially determined by the parameters $\theta_{\mathbf{X}}$ and $\theta_{\mathbf{W}}$. At the outer loops, we leverage the accumulated global losses to guide the parameter update for better optimization strategy through backpropagation. Let t_{out} denote the update interval, the accumulated global loss is given by

$$\mathcal{L}_F^s = \frac{1}{t_{out}} \sum_{t_s=(s-1)t_{out}+1}^{st_{out}} \omega_{t_s} F(\mathbf{W}_{t_s}, \mathbf{X}_{t_s}), \quad (10)$$

where $\omega_t \in \mathbb{R}_{\geq 0}$ denotes the weight associated with each outer step, and $s = 1, 2, \dots, S$, with $S = T/t_{out}$ being the

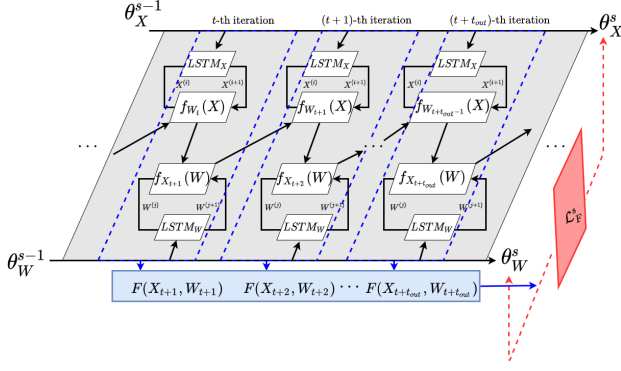


Fig. 2. The outer loop iterations on sub-problems are placed on the grey plane. Each iteration contains two sub-problems (dashed blue rhombus), and each sub-problem optimizes the corresponding variable at inner loop along with an LSTM MetaNet. The global loss $F(X, W)$ is computed after each iteration. For every t_{out} number of iterations, the accumulated global losses $\mathcal{L}_F^s$ (rhombus in red) is calculated and back-propagated to update the parameters θ_X and θ_W based on gradient descent manners (referred by the red dashed arrows).

maximum update number for LSTM networks, and T being the maximum outer steps. For every t_{out} outer loop iteration, the accumulated global losses $\mathcal{L}_F^s$ is computed and used to update θ_X and θ_W as follows

$$\begin{aligned}\theta_X^{s+1} &= \theta_X^s + \alpha_X \cdot \text{Adam}(\theta_X^s, \nabla_{\theta_X^s} \mathcal{L}_F^s), \\ \theta_W^{s+1} &= \theta_W^s + \alpha_W \cdot \text{Adam}(\theta_W^s, \nabla_{\theta_W^s} \mathcal{L}_F^s),\end{aligned}\quad (11)$$

where α_X and α_W denote the learning rates for the meta-networks LSTM_X and LSTM_W , respectively. The parameters of LSTMs are updated by the Adam method [37].

Remark 2. The formulation (10) can be extended to accommodate prior information as follows,

$$\mathcal{L}' = \mathcal{L}_F^s + \omega_w \left\| \mathbf{W}_{t_s} - \tilde{\mathbf{W}} \right\|_F^2 + \omega_x \left\| \mathbf{X}_{t_s} - \tilde{\mathbf{X}} \right\|_F^2, \quad (12)$$

where ω_w and ω_x denote the weights of the prior knowledge, $\{\tilde{\mathbf{W}}\}$ and $\{\tilde{\mathbf{X}}\}$ are the available paired training samples from historical data.

Hence, θ_X and θ_W successfully build a connection between the variable update functions and the global losses. At inner loops, θ_X and θ_W convey an extra global loss information from the outer loops for the update functions. The accumulated global losses allow the LSTMs to be updated with respect to the mutual knowledge of dealing with different sub-problems. Therefore, in the ground-level meta-learning, the learnt algorithm has better adaptability to the new sub-problem in the sequence. The accumulated global losses based ground-level meta-learning is depicted in Fig.2.

There are two merits of adopting $\mathcal{L}_F^s$ to update the parameters. First, the update with a small update interval, e.g., $t_{out} = 1$, will lead to severe fluctuation. An appropriate update interval with accumulated global losses is able to effectively relieve the factor of the outliers in training process. We shall present the related simulation results in Section V-A. Second, the leveraged global losses can provide more mutual knowledge than one global loss value. The parameters are

Algorithm 1: General Structure of MLAM algorithm for solving one problem

```

1 Input: global loss function  $F(W, X)$ , local loss
   functions  $f_W(X)$  and  $f_X(W)$ , random initialization
    $W_0$ , number of outer loops  $T$ , and number of inner
   loops  $I$  and  $J$ .
2 Output: Estimated variables  $W_T, X_T$ .
3 for  $t \leftarrow 1, \dots, T$  do
4   for  $i \leftarrow 0, \dots, I-1$  do
5      $\Delta X = \text{LSTM}_X(\nabla f_{W_t}(X^{(i)}), C_X^{(i)}, \theta_X^s)$ 
6      $X^{(i+1)} \leftarrow X^{(i)} + \Delta X$ ;
7   end
8    $X_t \leftarrow X^{(i)}$ ;
9   Update local loss function  $f_{X_t}(W)$ ;
10  for  $j \leftarrow 0, \dots, J-1$  do
11     $\Delta W = \text{LSTM}_W(\nabla f_{X_t}(W^{(j)}), C_W^{(j)}, \theta_W^s)$ 
12     $W^{(j+1)} \leftarrow W^{(j)} + \Delta W$ ;
13  end
14   $W_t \leftarrow W^{(j)}$ ;
15  Update local loss function  $f_{W_t}(X)$ ;
16  Update global loss function  $F(W_t, X_t)$ ;
17  for  $s \leftarrow 1, \dots, t/t_{out}$  do
18     $\mathcal{L}_F^s = \frac{1}{t_{out}} \sum_{t_s=(s-1)t_{out}+1}^{st_{out}} \omega_{t_s} F(W_{t_s}, X_{t_s})$ 
19     $\theta_X^{s+1} = \theta_X^s - \alpha_X \nabla_{\theta_X^s} \mathcal{L}_F^s$ 
20     $\theta_W^{s+1} = \theta_W^s - \alpha_W \nabla_{\theta_W^s} \mathcal{L}_F^s$ 
21  end
22 end

```

updated with the objective to minimize a partial trajectory of the global losses. Therefore, it allows the MetaNet to learn a non-monotonic solution, where the global loss could increase at the beginning iterations but quickly decrease to better optimums on the global scope. We will discuss in details with a practical example in Section IV-B.

At this stage, the proposed MLAM algorithm is summarized in Algorithm 1. Specifically, the algorithm starts from a random initialization W_0 at the beginning of an outer loop. At the t -th outer loop, it contains two inner loops for updating X and W , respectively. Each inner loop starts from a random initialization $X^{(0)}$ and $W^{(0)}$ and updates variables based on equations (8) and (9) and then repeats for I and J times, respectively. In this paper, we set $I = J = t_{in}$ in which t_{in} indicates the maximum iteration number at the inner loops. The choices for t_{out} and t_{in} will be discussed in Section V-A. In practice, it is reasonable to set different values for I and J according to the demand of the objective in different problems. At the end of each inner loop, the output of this inner loop is regarded as X_t or W_t at the t -th outer loop, and is then assigned to generate sub-problem f_{X_t} and f_{W_t} , respectively. For every t_{out} step at outer loops, the parameters θ_X and θ_W are updated following equations (11). Finally, our MLAM algorithm will stop when t reaches the maximum outer loop number T .

Comparing to the deep unfolding algorithms, the fundamental differences of MLAM method are mainly three-fold: 1)

replacing the variable update function within each iteration by a meta-network instead of mapping the whole procedure at each iteration by black-box based network operator, thus the inner loop at each iteration is interpretable; 2) focusing on learning a new strategy for the whole iterations instead of following the basics of the original strategy with end-to-end network behavior on trainable parameters, leading to a better explainability on the optimization manner; and 3) is feasible for unsupervised learning while the deep unfolding methods are mostly supervised learning. Therefore, the proposed MLAM method highlights advances on the model explainability and the improvements of combining the advantages of learning-based and model-based approaches. We believe this is a further step ahead that makes a higher level of learning than the deep unfolding, where the learning objective is no longer the trainable parameters but the whole optimization strategy.

In a summary, a hierarchical LSTM-based MLAM model is proposed in this section. It contains two (or more for multiple variables if needed) LSTM networks which perform optimization at inner loops with frozen parameters; their parameters are updated during outer loops with respect to minimizing accumulated global losses. Therefore, the original structure of algorithms well-established in the field is maintained while the performance is improved by the meta-learned algorithmic rule.

IV. APPLICATIONS IN TYPICAL PROBLEMS

In this section, we shall apply the proposed MLAM approach to a bi-linear inverse problem and a non-linear problem. For those problems, the model-based methods perform less effective and learning-based methods typically do not work due to the lack of sufficient labelled data. Specifically, matrix completion and Gaussian mixture model (GMM) problems are analyzed in this paper as two representatives.

A. Bi-linear Inverse Problem: Matrix Completion

The bi-linear inverse problem is a typical optimization problem whose variables are within an intersection of two sets. Many non-convex optimization problems can be cast as bi-linear inverse problems, including low-rank matrix recovery [39]–[43], dictionary learning [44]–[46], and blind deconvolution [47]–[49]. For example, in blind deconvolution and in the matrix completion, $F(\cdot, \cdot)$ in (1) represents circular convolution and matrix product, respectively.

Next we focus on the matrix completion [41]–[43], [50], [51] as a representative bi-linear inverse problem to demonstrate how to apply MLAM to solve this type of problems. Matrix completion is a class of tasks that aims to recover missing entries in a data matrix [50]–[52] and has been widely applied in practice including recommending system [43] and collaborative filtering [53]. Typically, it is formulated as a low-rank matrix recovery problem in which the matrix to be completed is assumed to be low-rank with given rank information.

As for optimization, the low-rank matrix completion problem is usually formulated as the multiplication of two matrices and then converted into two corresponding sub-problems

which are generally strongly convex [39], [40], [54]. Then, gradient descent based methods, such as alternating least square (ALS) [55] and stochastic gradient descent (SGD) [56], are often applied on solving each sub-problem and can achieve good performance under certain conditions. However, satisfactory results are not universally guaranteed. On the one hand, the performance depends on a set of factors, including initialization strategies, the parameter setting of gradient descent algorithms, the sparsity level of the low-rank matrix, etc. On the other hand, the assumptions behind these constraints are stringent, such as the feature bias on variables for sparse subspace clustering mechanism [52], [57], which are not common in practice.

Therefore, the bi-linear inverse problem is basically ill-posed and thus has always been a challenging task. Consequently, the performance of the traditional model-based methods is highly limited by the prior knowledge on models and structural constraints. Meanwhile, the strong non-convexity and constraints also limit the application of the deep learning and unfolding techniques in this type of problems.

In this section, we further consider several more realistic and therefore challenging scenarios, including high-rank matrix completion, matrix completion without given rank knowledge, and mixed rank matrix completion problems, where the existing model-based and even learning-based approaches fail.

The matrix completion problem is then formulated as [58],

$$\min_{\mathbf{U}, \mathbf{V}} F(\mathbf{U}, \mathbf{V}) := \frac{1}{2} \|\mathcal{P}_{\Omega}(\mathbf{R} - \mathbf{U}\mathbf{V}^T)\|_F^2 + \frac{\lambda}{2} (\|\mathbf{U}\|_F^2 + \|\mathbf{V}\|_F^2), \quad (13)$$

where the projection $\mathcal{P}_{\Omega}(\cdot)$ preserves the observed elements defined by Ω and replaces the missing entries with 0, and λ is the weight parameter of the regularizers. The matrix completion problem is typically formulated as a low-rank matrix recovery problem, which parameterizes a low-rank matrix $\mathbf{R} \in \mathbb{R}^{z \times q}$ as a multiplication of two matrices $\mathbf{U}\mathbf{V}^T$ with $\mathbf{U} \in \mathbb{R}^{z \times p}$, $\mathbf{V} \in \mathbb{R}^{q \times p}$ and $p \leq \min(z, q)$.

Remark 3. When taking the matrix multiplication $\mathbf{R} = \mathbf{U}\mathbf{V}^T$, the parameter p is set to the rank of $\mathbf{R}$, which is typically known. If the rank is not provided, the problem will be much more difficult and the existing methods would be unworkable. In Section V-A, we will verify that the proposed MLAM still works properly when p is unknown.

Here we define (13) as the overall problem, and $F(\mathbf{U}, \mathbf{V})$ as the global loss function. It is obvious that the overall problem is not convex in terms of $\mathbf{U}$ and $\mathbf{V}$, but the sub-problems are convex when fixing one variable and updating the other. Therefore, we split (13) into two sub-problems in quadratic form, fixing one variable in (13) and updating the other one, referring to $f_{\mathbf{U}}(\mathbf{V})$ and $f_{\mathbf{V}}(\mathbf{U})$. The ALS and SGD methods can be applied by iteratively minimizing the two sub-problems.

Meanwhile, according to Algorithm 1, our MLAM method can be directly applied to solve the problem (13) using two LSTM networks $\text{LSTM}_{\mathbf{U}}$ and $\text{LSTM}_{\mathbf{V}}$ with parameters $\theta_{\mathbf{U}}$ and $\theta_{\mathbf{V}}$, to optimize matrices $\mathbf{U}$ and $\mathbf{V}$, respectively. The variable update equations are given by

$$\begin{aligned} \mathbf{U}^{(i)} &= \mathbf{U}^{(i-1)} + \mathbf{H}_{\mathbf{U}}^{(i-1)}, \\ \mathbf{V}^{(j)} &= \mathbf{V}^{(j-1)} + \mathbf{H}_{\mathbf{V}}^{(j-1)}, \end{aligned} \quad (14)$$

where $\mathbf{H}_U$ and $\mathbf{H}_V$ are the outputs of LSTM_U and LSTM_V , respectively. Two LSTM networks are updated for every t_{out} outer loop steps, via back-propagating accumulated global losses $\mathcal{L}_F^s$ according to equation (10). In this way the updating rule of the learned algorithm could be adjusted to find gradient descent steps $\mathbf{H}_U$ and $\mathbf{H}_V$ for minimizing $F(\mathbf{U}, \mathbf{V})$ adaptively.

The advantages of our MLAM model for the matrix completion problem may be explained by the replacement of updating step function ϕ by the neural networks. In the AM methods, when a local optimum $(\mathbf{U}', \mathbf{V}')$ is reached, the gradient of the variables equals to zero, e.g., $\frac{\partial f_{\mathbf{V}'}(\mathbf{U})}{\partial \mathbf{U}}|_{\mathbf{U}=\mathbf{U}'} = \mathbf{0}$. At this stage, the update function, which is determined by the gradients of variables, will be stuck at local optimum. While in our MLAM model, our update functions are determined by their parameters, θ_U and θ_V , which are further determined by leveraging on partial global loss trajectories across the outer loop steps. This major difference mainly brings two benefits to our MLAM method. One advantage is that even at local optimum points, our MLAM can still provide a certain step update on variables. This can be understood that even when one of the input $\frac{\partial f_{\mathbf{V}'}(\mathbf{U})}{\partial \mathbf{U}}|_{\mathbf{U}=\mathbf{U}'} = \mathbf{0}$, LSTM_U still obtains some non-zero outputs given a non-linear function of zero input, cell state, and parameters. Another advantage is that the leveraged global loss leads to a smooth optimization on a global loss landscape, which allows the learned algorithm to be essentially guided by the inductive bias from a smoother transform of the global loss landscape. In Section V, besides traditional low rank matrix completion problem, we further evaluate MLAM in matrix completion problems in the case of high rank, unknown rank and mixed rank (a set of matrices with different ranks, from low rank to full rank).

B. Non-linear Problem: Gaussian Mixture Model

Different from the bi-linear inverse problems, optimization over the intersection of two variables that has a non-linear mapping between variables and observed samples also plays a significant role in statistical machine learning, including Bayesian model [59], graphic model [60], [61], and finite mixture model [62]. The sub-problem in non-linear problem typically has no closed-form solution and the overall problem possesses many local optimums. Therefore, it is difficult to guarantee the convergence of the model-based approaches, especially for the global optima, as well as intractable to obtain labelled data for the existing learning based methods.

GMM [63]–[65] is one of the most important probabilistic models in machine learning. GMM problem, consisting of a set of Gaussian distributions in the form of weighted Gaussian density components, is usually estimated by maximum log-likelihood method [65]. The variables in GMM problem possess a non-linear mapping to the observations. Many methods have been proposed to solve the GMM problem, i.e., maximising the likelihood of GMM problems, such as conjugate gradients, quasi-Newton and Newton [66]. However, these methods typically perform inferior to the one called expectation-maximization (EM) algorithm [62], [63]. One possible reason is due to the non-convexity and non-

linearity of the GMM problem that requires a sophisticated step descent strategy to find a good stationary point. On the other hand, EM algorithm omits the hyperparameter related to the step size by converting the origin estimation problem of maximising likelihood into a relaxed problem where a lower bound is maximized monotonically and analytically. Even though the EM algorithm has been widely applied in the GMM problem, it also suffers from the aforementioned non-convexity and non-linearity. When the non-convexity is high (referring to some real-world scenarios: the number of observation is not sufficient, dealing with high-dimensional data [67] and large number of clusters), the convergence is not guaranteed and the performance significantly degrades. Many works attempt to replace EM algorithm through reformulating GMM as adopting matrix manifold optimization [64], [68], and also learning based method in high dimensions [67].

Detailed descriptions of GMM problems can be found in [69]. Given a set of G i.i.d samples $\mathbf{X} = \{\mathbf{x}_g\}_{g=1}^G$, each entry $\mathbf{x}_g$ is a D -dimensional data vector. Then, a typical optimization when using the GMM to model the samples is to maximize the log-likelihood (MLL) [65], which is equivalent to minimize the Kullback–Leibler divergence from the empirical distribution. The parameters of the GMM can then be optimized as follows,

$$\max_{\{\pi_k, \mu_k, \Sigma_k\}_{k=1}^K} \log p(\mathbf{X}) = \sum_{g=1}^G \log \sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{x}_g | \mu_k, \Sigma_k), \quad (15)$$

where

$$\mathcal{N}(\mathbf{x}_g | \mu_k, \Sigma_k) = \frac{\exp\{-\frac{1}{2}(\mathbf{x}_g - \mu_k)^T \Sigma_k^{-1}(\mathbf{x}_g - \mu_k)\}}{(2\pi)^{D/2} |\Sigma_k|^{1/2}}. \quad (16)$$

In (15), for the k -th Gaussian component, μ_k is the mean vector and defines the cluster centre, covariance Σ_k denotes the cluster scatter, π_k represents mixing proportion with $\sum_{k=1}^K \pi_k = 1$, and $|\Sigma_k|$ represents the determinant of Σ_k .

However, it is intractable to directly obtain a closed-form solution that maximizes $\log p(\mathbf{X})$ in (15). The key difficulty is that by differentiating $\log p(\mathbf{X})$ (summation of logarithmic summation) and equalizing it to 0, each parameter is intertwined with each other. Gradient descent methods in an AM manner can alternatively solve (15), but they typically perform inferior to the EM algorithm [62], [63].

We should also point out that the EM algorithm still updates the GMM parameters in an AM manner, i.e., iterating to optimize over each parameter at which its gradient equals to 0 with the other parameters being frozen. Furthermore, the EM algorithm, together with other first-order methods, has been proved to converge to arbitrary bad local optimum almost surely [70]. As we have discussed before, this AM strategy can be improved by replacing the frozen updating rule that searches optimum in a local landscape, by a less-greedy rule that updates variables up to global scope knowledge on the loss landscape of the global objective function.

Therefore, we propose to solve GMMs by adopting our MLAM method, which directly applies a learning-based gradient descent algorithm to the original ML problem (15) without any extra constraints. In this scenario, we consider the GMM

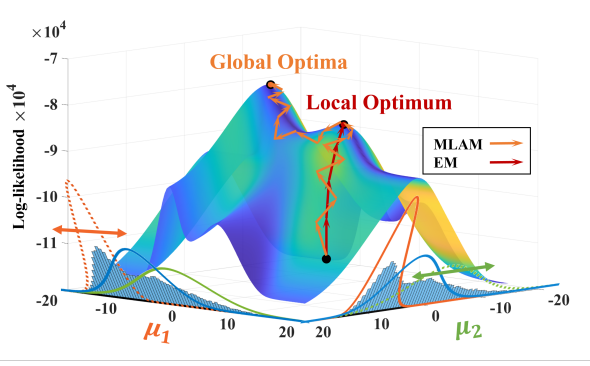


Fig. 3. An illustration of converge trajectories achieved by the proposed MLAM and EM methods in a simple GMM problem. In this GMM, the samples were obtained by three Gaussian distributions of parameters $\{\mu_1 = -3, \mu_2 = 3, \mu_3 = 0\}$, $\{\sigma_1 = 1, \sigma_2 = 10, \sigma_3 = 5\}$, and $\{\pi_1 = \pi_2 = \pi_3 = \frac{1}{3}\}$. For the convenience of visualization, we here show the surface of log-likelihood and converge trajectories with regard to μ_1 and μ_2 , whereas setting the other parameters to the ground-truth when optimizing GMMs; this means the global optima exists at $\mu_1 = -3$ and $\mu_2 = 3$.

problem with covariance Σ being given; hence the MLL estimation of GMM is presented in terms of negative log-likelihood as follows

$$\min_{\{\pi_k, \mu_k\}_{k=1}^K} F_{ne}(\pi, \mu) = - \sum_{g=1}^G \log \sum_{k=1}^K \pi_k \mathcal{N}(x_g | \mu_k, \Sigma_k). \quad (17)$$

In this case, we treat the (17) as our overall problem, and split it into two sub-problems $f_\pi(\mu)$ and $f_\mu(\pi)$. Define vector $\pi \in \mathbb{R}^K$ and matrix $\mu \in \mathbb{R}^{K \times D}$, where K is the maximum number of clusters and D is the dimensionality of samples. In our MLAM framework, we build two MetaNets LSTM $_\pi$ and LSTM $_\mu$ with parameters θ_π and θ_μ to update π and μ as follows

$$\begin{aligned} \pi^{(i)} &= \pi^{(i-1)} + H_\pi^{(i-1)}, \\ \mu^{(j)} &= \mu^{(j-1)} + H_\mu^{(j-1)}, \end{aligned} \quad (18)$$

where H_π and H_μ are the outputs of the two LSTM neural networks respectively.

Similar to applying MLAM in matrix completion, we start from a random initialization μ^0 and π^0 . During our MLAM procedure given in Algorithm 1, π_k and μ_k are updated based on equation (18) for t_{in} steps in inner loops respectively. Meanwhile, we back-forward the accumulated global losses with respect to the negative log-likelihood of GMM, referring to $\mathcal{L}_F^s$, for every t_{out} steps on the outer loops. In this way, the algorithm is updated based on the global scope knowledge about the global losses across outer loop iterations.

Considering that the existing numerical gradient-based solutions typically perform less effectively and accurately than the EM algorithm [63], we focus on comparing our MLAM method and the EM algorithm. EM algorithm converts maximization of the log-likelihood into maximization on its lower bound; hence it has closed-form formulations to implement an AM strategy on its maximization step for variable updating. Nevertheless, our MLAM method directly optimizes the cost function, i.e., the log-likelihood, and the variables are updated constantly up to the global scope knowledge extracted by

LSTMs. In Fig.3 we show the convergence trajectory of our MLAM and EM algorithm on the geometry of a GMM problem with three clusters. As mentioned in Section III-B, the proposed MLAM with accumulated global losses allows the learnt algorithm converges in a non-monotonic way. It can be seen that the EM algorithm (red arrows) quickly converges to the local optimum and stops at it. In contrast, the MLAM approach (orange arrows), first going to the local optimum yet, is able to escape from the local optimum and converges to the global optima. Specifically, when encountering local optimum, the proposed MLAM first goes down on the geometry, and then moves towards to the global optima, thus escaping from the local optimum. This reveals that the learnt algorithm does not request each step moving towards to the most ascending direction, but the global losses on a partial trajectory should be minimized, recalling the accumulated global losses minimization $\mathcal{L}_F^s$ in equation (10). Therefore, the MLAM enjoys significant freedom to learn a non-monotonic algorithm for convergence in terms of the update strategy on $\mathcal{L}_F^s$. We also shall point out that this also results in the fluctuations on the trajectory, which could increase the needed iterations when the geometry is smooth and benign.

V. EXPERIMENTS

In this section, we will present simulation results on the aforementioned matrix completion and GMM problems to validate the effectiveness and efficiency of our MLAM method.

For the experimental settings, our MetaNets employ two-layer LSTM networks with 500 hidden units in each layer. Each network is trained by minimizing the accumulated loss functions according to equation (10) via truncating backpropagation through time (BPTT) [71], which is a typical training algorithm to update weights in RNNs including LSTMs. The weights of LSTM are updated by Adam [37], and the learning rate is set to 10^{-3} . In all simulations, we set $\omega_{t_s} = 1$ for simplicity. The parameters of LSTM networks are randomly initialized and continuously updated through the whole training process. For evaluation, we fix the parameters of our MLAM model and evaluate the performance on the testing datasets.

A. Numerical Results on Matrix Completion

In this subsection, we consider learning to optimize synthetic D -dimensional matrix completion problems. We take $D = 10$ and $D = 100$ to evaluate algorithms for both small and large scale cases. For each matrix completion problem, the ground truth matrix $\mathbf{R} \in \mathbb{R}^{D \times D}$ is randomly and synthetically generated with rank p . Meanwhile, the observation $\mathbf{R}_S = \mathcal{P}_\Omega(\mathbf{R})$ is generated by randomly setting a certain percentage of entries in $\mathbf{R}$ to be zeros, and the non-zero fraction of entries is the observation rate. Matrix completion for $\mathbf{R}$ is then achieved by solving the low-rank matrix recovery on $\mathbf{R}_S$ in the form of equation (13) in Section IV-A. The two factorized low-dimensional matrices $\mathbf{U} \in \mathbb{R}^{D \times p}$ and $\mathbf{V}^{D \times p}$ are then used to generate reconstruction of the ground truth

matrix, denoted by $\hat{\mathbf{R}} = \mathbf{UV}^T$. The evaluation criterion is given by Relative Mean Square Error (RMSE) of

$$\text{RMSE} = \frac{\|\mathbf{R} - \hat{\mathbf{R}}\|_F}{\|\mathbf{R}\|_F}.$$

As aforementioned, the classical AM-based ALS [55] and SGD [56] approaches, as well as the learning-based deep matrix factorization (DMF) [5] and unfolding matrix factorization (UMF) methods [11], have been adopted for comparisons. The parameters of all compared methods are carefully adjusted to present their best performances.

Different simulation scenarios on matrix completion are evaluated comprehensively. The detailed simulation settings include: i) each simulation contains a set of 200 matrix completion problems, half of which are employed as training samples for parameter update on LSTM, while the remaining 100 matrix completion problems are used to evaluate the performance as testing samples; ii) we set $T = 100$ as total alternating steps for each problem; iii) the averaged RMSE over 100 testing samples is used for evaluation; iv) in all the training and testing processes, the ground truth matrix $\mathbf{R}$ is not given, and is only used to evaluate performance after the optimizing processes.

1) *Parameter setting*: The numbers of variable update steps on inner loops t_{in} and update interval t_{out} are the two most important hyper-parameters. Different settings on t_{in} and t_{out} are thus tested at first to provide a brief guidance on the choices of t_{in} and t_{out} .

Empirically, there is a trade-off between performance and efficiency. Here we set t_{in} and t_{out} to both vary from 1 to 20 with 20% observation rate, and there are therefore $20 \times 20 = 400$ different parameter combinations to be evaluated for rank-5 matrix completion problems. The performance of these 400 simulations are shown in Fig. 4. We can see that for all choices of t_{out} , the increase on t_{in} leads to significant improvements on performance when $t_{in} \leq 10$, however, further increasing on t_{in} witnesses little gain on performances. At the same time, the variations of t_{out} have less impact on RMSE results, while larger choices ($t_{out} \geq 10$) bring improved stability (there are less fluctuations when $t_{in} \geq 10$, and $t_{out} \geq 10$). Fig. 4 indicates that a sufficient number of update steps for inner and outer loops play a significant role in these optimization processes. We can see that when either t_{in} or t_{out} is small (less than 5), the optimization does not perform well enough. A larger choice of t_{in} , however, typically brings higher computational cost. Thus in the rest of this paper, we set $t_{in} = 10$ and $t_{out} = 10$ as the default parameter setting.

It is understandable that t_{in} directly determines the number of variable update steps within each inner loops. When t_{in} is small, the learned updating rule needs to optimize variables in a few steps, however this could be intractable in general. Meanwhile, the accumulated global losses depend on t_{out} at outer loops which indicates the length of the trajectory of global losses at outer loops. Therefore, a large enough t_{out} could provide sufficient global losses trajectory for parameter update. According to the two-levels meta-learning in our MLAM model, sufficient update steps at inner and outer loops can ensure that each level of meta-learning works well.

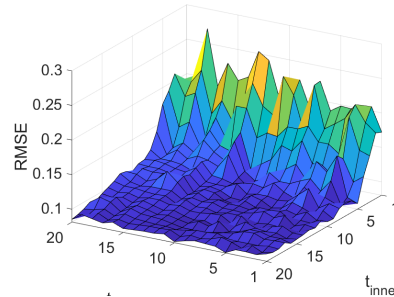


Fig. 4. Performance of the proposed MLAM method on rank-5 matrix completion problem with different parameter combinations of t_{in} and t_{out} which range from 1 to 20 with step size 1.

We infer that small t_{in} may limit the ground-level meta-learning corresponding to inner loops, making it unable to extract effective sub-problems across knowledge with merely few update steps, and small t_{out} could cause that the upper-level meta-learning corresponding to outer loops becomes less stable due to the lack of updates on parameters.

2) *Standard matrix completion*: In this part, we will first evaluate all methods on low-rank matrix completion tasks and then apply them on high rank matrix completion tasks. In these simulations the rank of the matrices is given. In this subsection, the main goal is the proof of the concept of the proposed MLAM. We evaluate four existing approaches, including conventional model-based methods and the state-of-the-art learning-based methods for comparison.

In Table I, average RMSE of the five evaluated methods on four sets of 100-dimensional rank-5 matrix completion problems have been reported. Different sets have different observation rates, including 20%, 40%, 60% and 80%. From Table I, it is clear that our MLAM algorithm significantly outperforms all the existing methods, especially on high observation rate scenarios. It is also noticeable that when the observation rate is 20%, both of the model-based methods (ALS and SGD) and learning-based methods (DMF and UMF) work not well, while our MLAM method achieves good reconstruction with $\text{RMSE} < 0.1$.

In Fig.5, we present the RMSE variance of the five tested methods on 100 trails of rank-5 matrix completion tasks with 20% observation rate. It can be observed that the MLAM obtains the best performance while keep the variance relatively small. It is noticeably that the ALS approach fluctuates severely, while the SGD and DMF approaches present relatively large variance comparing to the MLAM. Though the UMF method gains robust results with small variance, the accuracy is significantly larger than the MLAM.

From these simulation results, we can conclude that in the classic low-rank matrix completion problem, our MLAM method has shown better performances than the comparison methods, especially when the observation rate is low.

We then consider more challenging matrix completion problems. Generally, the matrix completion problem is assumed to be solved as a low-rank matrix completion problem. Here we are aiming to solving high-rank or even full-rank matrix completion problems without adding any further assumptions.

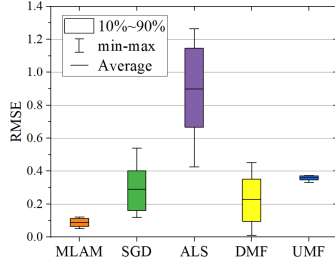


Fig. 5. The RMSE variance of the evaluated methods on rank-5 matrix completion tasks with 20% observation rate.

TABLE I

RMSE OF RANK-5 MATRIX COMPLETION WITH DIFFERENT OBSERVATION RATES Θ .

Methods $\backslash \Theta$	0.2	0.4	0.6	0.8
MLAM	0.089	0.057	0.045	0.003
SGD	0.292	0.213	0.136	0.072
ALS	0.864	0.717	0.423	0.228
DMF	0.228	0.091	0.057	0.025
UMF	0.363	0.359	0.335	0.312

In this case, we test all the five approaches on six set of 100-dimensional matrix completion problems with 20% observation rate, whose ranks are ranged from 5 to 100.

The averaged RMSE results of matrix completion problems with variant rank are listed in Table II. There is a clear trend of decreasing performances of the comparison methods when the rank of matrices increases. Noticeably, ALS method is no longer workable when the rank is larger than 10. SGD, DMF and UMF methods fail after the rank reaches 40. This can be understood that these approaches are typically assumed specific penalty terms recalling to the low-rank property. However, the proposed MLAM method shows different results: the higher the rank, the smaller the RMSEs. We can conclude that our algorithm is capable of solving the matrix completion problem in high-rank, even full-rank scenarios, without any extra constraints, while standard methods typically fail.

The execution time of the evaluated methods is compared in Table IV. It can be seen that the proposed MLAM approach has larger time cost than the compared methods, but the difference is within second-level. This is caused by the alternatively utilizing network at each iteration. There is a trade off between the time consuming and performance: the higher the accuracy, the larger the time consumption. In this paper, the main goal is the proof of the concept that the MLAM approach could provide better performance on solving multi-variable non-convex optimization problem.

3) *Blind matrix completion*: In this part, we consider more challenging cases where the rank information is not known. We consider these are blind matrix completion problems which are typically intractable using previous methods. We first test our MLAM method and standard methods in the case where a set of matrices completion problems have the same but unknown rank. Then we will further test our MLAM method by more difficult case which has a set of matrices completion problems with different unknown ranks.

In the first case, we test our MLAM method and two stan-

TABLE II
RMSE OF MATRIX COMPLETION WITH DIFFERENT RANKS.

Methods $\backslash$ Rank	5	10	20	40	80	100
MLAM	0.089	0.085	0.078	0.072	0.053	0.052
SGD	0.292	0.405	0.689	>1	>1	>1
ALS	0.864	0.942	>1	>1	>1	>1
DMF	0.228	0.263	0.392	0.643	>1	>1
UMF	0.363	0.376	0.417	0.582	>1	>1

TABLE III

RMSE OF RANK-10 MATRIX COMPLETION WITH VARIANT FACTORIZED MATRIX DIMENSION p .

Methods $\backslash p$	10	20	40	80	100
MLAM	0.09	0.12	0.19	0.30	0.36
SGD	0.31	0.78	>1	>1	>1
ALS	0.89	>1	>1	>1	>1
DMF	0.26	0.53	0.81	>1	>1
UMF	0.37	0.68	0.87	>1	>1

dard methods on a set of rank-10 matrix completion problems with variant factorized matrix dimension p . In Table III we report the results of applying different p for reconstructing a rank-10 matrix through the five tested approaches. We can see that the proposed MLAM method is more robust when there is a mismatch between p and the true rank compared to the other methods.

When we take a large gap such as $p = 80$ or $p = 100$ to reconstruct a rank-10 matrix, the MLAM method still achieves acceptable performances. Meanwhile, the rest of the tested methods quickly degrade with the increase of the mismatch between p and the rank values. The RMSEs of the existing methods quickly arise to more than 100% when $p \geq 40$. In contrast, our MLAM method shows a good tolerance on the increase of the difference between the real rank and p . Thus, the MLAM method does not require an accurate rank information to achieve a successful matrix completion.

Then, we test our MLAM method on matrix completion problems with different ranks for the matrix R . This means that in the training and testing sample sets, there are mixed matrices with different ranks. Thus the learned algorithm is required to adapt on matrix completion problems across different ranks.

In Table V, RMSE results of three sets of matrix completion problems are listed. For each set, 100 training samples and 100 testing samples are generated by setting their ground truth rank uniformly distributed between 10 and 100 with stepsize 10. We take three different choices of p for each set, and calculate the mean RMSE on the samples of each rank respectively.

The proposed MLAM method in the three cases performs generally well on samples with different ranks. The comparison of the three cases further reveals that our method of choosing $p = 10$ achieves the best performance in all samples, especially on low-rank scenarios where it provides significantly better results than the others. For the other two choices of p , our method also has comparable performances on high-rank samples, and achieves sufficiently good accuracy on the majority of the samples.

The major computational cost in each iteration of MLAM

TABLE IV
TIME COMPLEXITY OF DIFFERENT METHODS.

Methods	MLAM	DMF	UMF	ALS	SGD
Time(ms)	4618	3704	2243	1272	1518

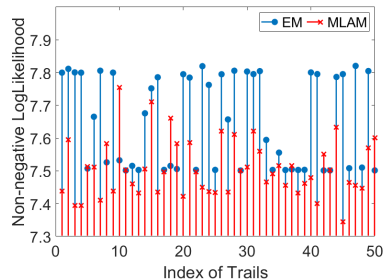


Fig. 6. The performance of EM and MLAM over 50 testing samples. Averaged non-negative log-likelihood of 50 trails through EM algorithm is 7.77, while the result through MLAM model is 7.56.

is from the network-based computation of the gradient update term at inner loops. In this instance, due to the implementation of the networks at each inner loop iteration, the computational complexity of the MLAM approach is significantly regarded to the number of the inner and outer loops.

In the future work, we will further improve the computational speed of the MLAM to reduce the time cost for better practical applications.

In summary, we have been shown that our MLAM method is capable of solving the matrix completion problem without any prior information. Typically, these problems have underlying complicated landscapes geometries. Therefore, it is hard for standard gradient descent based methods to achieve satisfactory results. The proposed MLAM model makes it possible to find good solutions through learning an appropriate updating rule that is dynamic and adaptive.

B. Numerical Results on Gaussian Mixture Model

In this section, we apply MLAM method to GMM problems. Given the data set $\mathbf{X} = \{\mathbf{x}_g\}_{g=1}^G$ (G denotes the number of observation samples), we optimize the mean cluster centre μ_k and mixing proportion π_k whilst keeping the covariance Σ_k frozen, which is similar to the optimization of the k -means problem. We consider one GMM problem as one sample in the training and testing set. For solving each GMM problem, we set the maximum alternating steps as 100, and record the negative log-likelihood F_{ne} trajectory according to (17). Similar to the settings in the matrix completion simulation, we also choose $t_{in} = 10$ and $t_{out} = 10$ for all the simulation scenarios in the GMM problems.

EM algorithm has been adopted for comparison. The stopping criterion for EM algorithm is $|F_{ne}(t-1) - F_{ne}(t)| < 10^{-4}$. EM algorithm and our MLAM method all start from the same random initialization for μ_k and π_k .

We first start from the 2-dimensional GMM problems, whose data vector $\mathbf{x}_g \in \mathbb{R}^2$. Given four clusters ($K = 4$), $G = 500$ data points in $\mathbf{X} = \{\mathbf{x}_g\}_{g=1}^G$, we show 50 tested results with random initialization in Fig 6. We only

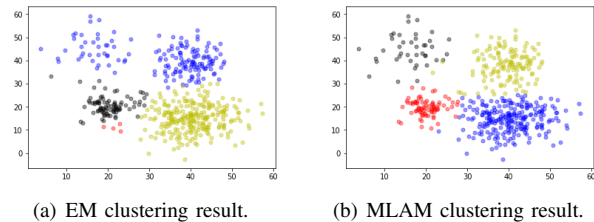


Fig. 7. Clustering results of one GMM problem with random initialization.

provide 500 data points here to increase the difficulty for optimizing these GMM problems due to the limited number of observations. From Fig 6, it can be seen that on these 50 tests, the MLAM method outperforms the EM algorithm in most cases. The mean negative log-likelihood F_{ne} of MLAM method on these 50 samples is 7.52 while that of the EM algorithm is 7.75. More specifically, we randomly select one clustering result among these 50 trails which is shown in Fig.7, which clearly shows that the MLAM method obtains much better clustering results than the EM algorithm.

Furthermore, we conduct simulations on a flower-shaped synthetic data ($G = 10,000$) with random initializations. Each cluster in the flower-shaped data is composed by Gaussian-distributed samples. This typically is a hard problem as the anisotropic clusters lead to extensive local optimum. 10 randomly selected optimization can be found in Fig 8, in which our algorithm consistently achieves nearly optimal clustering but the results from EM are highly biased. Although being illustrative in 2 dimensions, the data in Fig. 8 consists of 8 anisotropic clusters, which might be the main reason that the EM method converges to bad local optimum. As has been proved in [70], when solving GMMs with more than 2 clusters ($K > 2$), the EM method is highly likely to converge to arbitrarily bad local optimum under random initialization. More specifically, The EM method is a special case of Bayesian variational inference, and provides a tractable solver to maximize the lower bound of the log-likelihood of GMM problems. More importantly, maximizing the log-likelihood is equivalent of using the Kullback-Leibler (KL) divergence in estimating two distributions [63]. Thus, the performance of the EM method is basically limited to the intrinsic nature of the KL divergence, which can output infinitely large values with gradient vanishing issues when two distributions are well-separated [70], [72]. In other words, the EM method is highly sensitive to initialization and may converge to bad local optimum from random initialization, the phenomenon that has been studied in [68]. In contrast, the proposed MLAM method consistently achieves the global estimation, verifying the global optimization nature of our method. Therefore, our MLAM method obtains significant improvements on accuracy, even for some challenging scenarios such as insufficient observations and anisotropic clusters in these illustrative evaluations.

We further evaluate our method in estimating high dimensional GMM problems. Several sets of high dimensional synthetic data ($G = 500$, $K = 4$) are also randomly generated for the evaluation, with dimensions 4, 8, 16, 32 and 64. The

TABLE V
RMSES OF MIXED MATRIX COMPLETION WITH DIFFERENT CHOICE OF p .

$p \backslash$ Rank	10	20	30	40	50	60	70	80	90	100
10	0.120	0.081	0.065	0.055	0.049	0.046	0.043	0.042	0.042	0.039
50	0.250	0.130	0.095	0.075	0.067	0.055	0.052	0.048	0.046	0.045
100	0.450	0.210	0.140	0.100	0.081	0.068	0.061	0.056	0.052	0.048

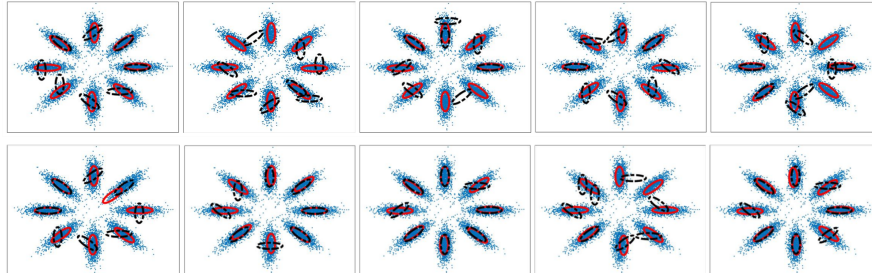


Fig. 8. EM and MLAM performance on 10 flower-shaped data. Black and red dash lines denote EM and MLAM clustering results, respectively

TABLE VI
HIGH DIMENSIONAL GMM SIMULATION RESULTS.

Methods \ Dimension	4	8	16	32	64
EM	6.52	12.64	23.83	47.10	96.88
MLAM	6.44	12.42	23.05	46.93	95.72

averaged negative log-likelihoods of 100 tests for each testing dimension are reported in Table VI. It is clear that from Table VI our MLAM method outperforms the EM algorithm on all the high dimensional sample sets and the variation on dimensions does not affect the performance of our method, while this typically decreases the performance of EM algorithm in general.

In this stage, it has been verified that our proposed method is able to solve non-convex multi-variable problem with non-linear relationships between variables and observation data. Even without closed-form landscapes in these problems, our MLAM method still successfully finds good solutions whilst EM algorithm fails.

VI. CONCLUSION

We have proposed a meta-learning based alternating minimization (MLAM) method for solving non-convex optimization problems. The learned algorithm has been verified to have a faster convergence speed and better performances than existing alternating minimization (AM)-based methods. To achieve that, our MLAM method has employed LSTM-based meta-learners to build an interaction between variable updates and the global loss. In this way, the variables are updated by the LSTM networks with frozen parameters at inner loops, which are then updated by minimizing accumulated global losses at outer loops. This paper is just a proof of the concept that a less greedy and learning-based solution for non-convex problem could surpass both of the traditional model-based and learning-based methods. More importantly, these concepts: optimize each independent solution step in-exhaustively by globally learning the optimization strategy, and integrate these

independent steps by the bi-level meta-learning optimization model, spark a new direction of improving the optimization-inspired solutions for advanced performance. It reveals that applying neural network based behavior to assist the well-established algorithmic principle, instead of replacing it by black-box network behavior, could bring significant advances.

For the future work, we plan to apply the proposed MLAM method to more challenging and practical non-convex problems, such as dictionary learning in compress sensing, blind image super resolution, and MIMO beamforming. We would also like to extend some theoretical analysis on our MLAM model. Moreover, the MLAM can be directly apply to solve non-convex problem online without pre-training. It is deserved to discover the application of applying MLAM as an algorithm without training for non-convex problems.

REFERENCES

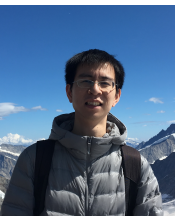
- [1] M. J. Powell, "On search directions for minimization algorithms," *Mathematical programming*, vol. 4, no. 1, pp. 193–201, 1973.
- [2] Q. Hu, Y. Cai, Q. Shi, K. Xu, G. Yu, and Z. Ding, "Iterative algorithm induced deep-unfolding neural networks: Precoding design for multi-user mimo systems," *IEEE Transactions on Wireless Communications*, vol. 20, no. 2, pp. 1394–1410, 2020.
- [3] D. Ren, K. Zhang, Q. Wang, Q. Hu, and W. Zuo, "Neural blind deconvolution using deep priors," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2020, pp. 3341–3350.
- [4] J. Bernstein, Y.-X. Wang, K. Azizzadenesheli, and A. Anandkumar, "signsgd: Compressed optimisation for non-convex problems," in *International Conference on Machine Learning*. PMLR, 2018, pp. 560–569.
- [5] J. Fan and J. Cheng, "Matrix completion by deep matrix factorization," *Neural Networks*, vol. 98, pp. 34–41, 2018.
- [6] K. Gregor and Y. LeCun, "Learning fast approximations of sparse coding," in *Proceedings of the 27th international conference on international conference on machine learning*, 2010, pp. 399–406.
- [7] D. You, J. Xie, and J. Zhang, "Ista-net++: Flexible deep unfolding network for compressive sensing," in *2021 IEEE International Conference on Multimedia and Expo (ICME)*. IEEE, 2021, pp. 1–6.
- [8] R. Fu, V. Monardo, T. Huang, and Y. Liu, "Deep unfolding network for block-sparse signal recovery," in *ICASSP 2021-2021 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2021, pp. 2880–2884.
- [9] Z. Luo, Y. Huang, S. Li, L. Wang, and T. Tan, "Unfolding the alternating optimization for blind super resolution," *arXiv preprint arXiv:2010.02631*, 2020.

- [10] L. Pellaco, M. Bengtsson, and J. Jaldén, “Deep unfolding of the weighted mmse beamforming algorithm,” *arXiv preprint arXiv:2006.08448*, 2020.
- [11] T. T. N. Mai, E. Y. Lam, and C. Lee, “Ghost-free hdr imaging via unrolling low-rank matrix completion,” in *2021 IEEE International Conference on Image Processing (ICIP)*. IEEE, 2021, pp. 2928–2932.
- [12] R. Li, S. Zhang, C. Zhang, Y. Liu, and X. Li, “Deep learning approach for sparse aperture isar imaging and autofocusing based on complex-valued admm-net,” *IEEE Sensors Journal*, vol. 21, no. 3, pp. 3437–3451, 2020.
- [13] S. Ravi and H. Larochelle, “Optimization as a model for few-shot learning,” 2016.
- [14] J. Li and M. Hu, “Continuous model adaptation using online meta-learning for smart grid application,” *IEEE Transactions on Neural Networks and Learning Systems*, 2020.
- [15] L. Liu, B. Wang, Z. Kuang, J.-H. Xue, Y. Chen, W. Yang, Q. Liao, and W. Zhang, “Gendet: Meta learning to generate detectors from few shots,” *IEEE Transactions on Neural Networks and Learning Systems*, 2021.
- [16] L. Chen, B. Hu, Z.-H. Guan, L. Zhao, and X. Shen, “Multiagent meta-reinforcement learning for adaptive multipath routing optimization,” *IEEE Transactions on Neural Networks and Learning Systems*, 2021.
- [17] C. Finn and S. Levine, “Meta-learning and universality: Deep representations and gradient descent can approximate any learning algorithm,” *arXiv preprint arXiv:1710.11622*, 2017.
- [18] K. Li and J. Malik, “Learning to optimize,” *arXiv preprint arXiv:1606.01885*, 2016.
- [19] J. Xia and G. Deniz, “Meta-learning based beamforming design for miso downlink,” *arXiv preprint arXiv:2103.11978*, 2021.
- [20] C. Finn, P. Abbeel, and S. Levine, “Model-agnostic meta-learning for fast adaptation of deep networks,” in *Proceedings of the 34th International Conference on Machine Learning-Volume 70*. JMLR. org, 2017, pp. 1126–1135.
- [21] A. Li, W. Huang, X. Lan, J. Feng, Z. Li, and L. Wang, “Boosting few-shot learning with adaptive margin loss,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2020, pp. 12 576–12 584.
- [22] M. Andrychowicz, M. Denil, S. Gomez, M. W. Hoffman, D. Pfau, T. Schaul, B. Shillingford, and N. De Freitas, “Learning to learn by gradient descent by gradient descent,” in *Advances in neural information processing systems*, 2016, pp. 3981–3989.
- [23] C. L. Byrne, “Alternating minimization and alternating projection algorithms: A tutorial,” *Sciences New York*, 2011.
- [24] W. Ha and R. F. Barber, “Alternating minimization and alternating descent over nonconvex sets,” *arXiv preprint arXiv:1709.04451*, 2017.
- [25] U. Niesen, D. Shah, and G. Wornell, “Adaptive alternating minimization algorithms,” in *2007 IEEE International Symposium on Information Theory*. IEEE, 2007, pp. 1641–1645.
- [26] T. Sun, D. Li, H. Jiang, and Z. Quan, “Iteratively reweighted penalty alternating minimization methods with continuation for image deblurring,” in *ICASSP 2019-2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2019, pp. 3757–3761.
- [27] D. Goldfarb, S. Ma, and K. Scheinberg, “Fast alternating linearization methods for minimizing the sum of two convex functions,” *Mathematical Programming*, 2013.
- [28] S. Guminov, P. Dvurechensky, and A. Gasnikov, “Accelerated alternating minimization,” *arXiv preprint arXiv:1906.03622*, 2019.
- [29] X. Zhang, W. Jiang, K. Huo, Y. Liu, and X. Li, “Robust adaptive beamforming based on linearly modified atomic-norm minimization with target contaminated data,” *IEEE Transactions on Signal Processing*, vol. 68, pp. 5138–5151, 2020.
- [30] J. Sui, Z. Liu, L. Liu, B. Peng, T. Liu, and X. Li, “Online non-cooperative radar emitter classification from evolving and imbalanced pulse streams,” *IEEE Sensors Journal*, vol. 20, no. 14, pp. 7721–7730, 2020.
- [31] J. Sui, Z. Liu, L. Liu, A. Jung, and X. Li, “Dynamic sparse subspace clustering for evolving high-dimensional data streams,” *IEEE Transactions on Cybernetics*, 2020.
- [32] C. Dong, C. C. Loy, K. He, and X. Tang, “Learning a deep convolutional network for image super-resolution,” in *European conference on computer vision*. Springer, 2014, pp. 184–199.
- [33] C. Dong, C. C. Loy, and X. Tang, “Accelerating the super-resolution convolutional neural network,” in *European conference on computer vision*. Springer, 2016, pp. 391–407.
- [34] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2016, pp. 770–778.
- [35] W. Samek, T. Wiegand, and K.-R. Müller, “Explainable artificial intelligence: Understanding, visualizing and interpreting deep learning models,” *arXiv preprint arXiv:1708.08296*, 2017.
- [36] K. Zhang, L. V. Gool, and R. Timofte, “Deep unfolding network for image super-resolution,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2020, pp. 3217–3226.
- [37] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” *arXiv preprint arXiv:1412.6980*, 2014.
- [38] J. Schmidhuber and S. Hochreiter, “Long short-term memory,” *Neural Comput*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [39] P. Jain, P. Netrapalli, and S. Sanghavi, “Low-rank matrix completion using alternating minimization,” in *Proceedings of the forty-fifth annual ACM symposium on Theory of computing*. ACM, 2013.
- [40] D. Gross, “Recovering low-rank matrices from few coefficients in any basis,” *IEEE Transactions on Information Theory*, vol. 57, no. 3, pp. 1548–1566, 2011.
- [41] Z. Wen, W. Yin, and Y. Zhang, “Solving a low-rank factorization model for matrix completion by a nonlinear successive over-relaxation algorithm,” *Mathematical Programming Computation*, vol. 4, no. 4, pp. 333–361, 2012.
- [42] M. Hardt, “Understanding alternating minimization for matrix completion,” in *2014 IEEE 55th Annual Symposium on Foundations of Computer Science*. IEEE, 2014, pp. 651–660.
- [43] Y. Koren, R. Bell, and C. Volinsky, “Matrix factorization techniques for recommender systems,” *Computer*, pp. 30–37, 2009.
- [44] I. Tosic and P. Frossard, “Dictionary learning: What is the right representation for my signal?” *IEEE Signal Processing Magazine*, vol. 28, no. ARTICLE, pp. 27–38, 2011.
- [45] M. Aharon, M. Elad, and A. Bruckstein, “K-svd: An algorithm for designing overcomplete dictionaries for sparse representation,” *IEEE Transactions on signal processing*, vol. 54, no. 11, pp. 4311–4322, 2006.
- [46] Q. Yu, W. Dai, Z. Cvetkovic, and J. Zhu, “Bilinear dictionary update via linear least squares,” in *ICASSP 2019-2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2019, pp. 7923–7927.
- [47] Y. Zhang, Y. Lau, H.-w. Kuo, S. Cheung, A. Pasupathy, and J. Wright, “On the global geometry of sphere-constrained sparse blind deconvolution,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 4894–4902.
- [48] T. G. Stockham, T. M. Cannon, and R. B. Ingebreetsen, “Blind deconvolution through digital signal processing,” *Proceedings of the IEEE*, vol. 63, no. 4, pp. 678–692, 1975.
- [49] J. R. Hopgood and P. J. Rayner, “Blind single channel deconvolution using nonstationary signal processing,” *IEEE Transactions on Speech and Audio Processing*, vol. 11, no. 5, pp. 476–488, 2003.
- [50] X. Lu, T. Gong, P. Yan, Y. Yuan, and X. Li, “Robust alternative minimization for matrix completion,” *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, vol. 42, no. 3, pp. 939–949, 2012.
- [51] E. J. Candès and B. Recht, “Exact matrix completion via convex optimization,” *Foundations of Computational mathematics*, vol. 9, no. 6, p. 717, 2009.
- [52] J. Fan and T. W. Chow, “Sparse subspace clustering for data with missing entries and high-rank matrix completion,” *Neural Networks*, vol. 93, pp. 36–44, 2017.
- [53] Y. Zhou, D. Wilkinson, R. Schreiber, and R. Pan, “Large-scale parallel collaborative filtering for the netflix prize,” in *International conference on algorithmic applications in management*. Springer, 2008, pp. 337–348.
- [54] S. Choudhary and U. Mitra, “On identifiability in bilinear inverse problems,” in *2013 IEEE International Conference on Acoustics, Speech and Signal Processing*. IEEE, 2013, pp. 4325–4329.
- [55] P. M. Kroonenberg and J. De Leeuw, “Principal component analysis of three-mode data by means of alternating least squares algorithms,” *Psychometrika*, vol. 45, no. 1, pp. 69–97, 1980.
- [56] R. Gemulla, E. Nijkamp, P. J. Haas, and Y. Sismanis, “Large-scale matrix factorization with distributed stochastic gradient descent,” in *Proceedings of the 17th ACM SIGKDD international conference on Knowledge discovery and data mining*. ACM, 2011, pp. 69–77.
- [57] C. Yang, D. Robinson, and R. Vidal, “Sparse subspace clustering with missing entries,” in *International Conference on Machine Learning*, 2015, pp. 2463–2472.
- [58] J. D. M. Rennie and N. Srebro, “Fast maximum margin matrix factorization for collaborative prediction,” in *Machine Learning, Proceedings of the Twenty-Second International Conference (ICML 2005), Bonn, Germany, August 7-11, 2005*, 2005.

- [59] S. Li, W. Ying, L. Jie, and X. Gao, "Multispectral image classification using a new bayesian approach with weighted markov random fields," in *CCF Chinese Conference on Computer Vision*, 2015.
- [60] V. Chandrasekaran, P. A. Parrilo, and A. S. Willsky, "Rejoinder: Latent variable graphical model selection via convex optimization," 2011.
- [61] L. Stankovic, D. Mandic, M. Dakovic, M. Brajovic, B. Scalzo, S. Li, and A. G. Constantinides, "Graph signal processing—part iii: Machine learning on graphs, from graph topology to applications," *arXiv preprint arXiv:2001.00426*, 2020.
- [62] A. P. Dempster, N. M. Laird, and D. B. Rubin, "Maximum likelihood from incomplete data via the em algorithm," *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 39, no. 1, pp. 1–22, 1977.
- [63] L. Xu and M. I. Jordan, "On convergence properties of the em algorithm for gaussian mixtures," *Neural computation*, vol. 8, no. 1, pp. 129–151, 1996.
- [64] R. Hosseini and S. Sra, "Matrix manifold optimization for gaussian mixtures," in *Advances in Neural Information Processing Systems*, 2015, pp. 910–918.
- [65] D. A. Reynolds, "Gaussian mixture models," *Encyclopedia of biometrics*, vol. 741, 2009.
- [66] R. A. Redner and H. F. Walker, "Mixture densities, maximum likelihood and the em algorithm," *SIAM review*, vol. 26, no. 2, pp. 195–239, 1984.
- [67] R. Ge, Q. Huang, and S. M. Kakade, "Learning mixtures of gaussians in high dimensions," in *Proceedings of the forty-seventh annual ACM symposium on Theory of computing*, 2015, pp. 761–770.
- [68] S. Li, Z. Yu, M. Xiang, and D. Mandic, "Solving general elliptical mixture models through an approximate wasserstein manifold," vol. 34, no. 04, pp. 4658–4666, 2020.
- [69] G. J. McLachlan and D. Peel, *Finite mixture models*. John Wiley & Sons, 2004.
- [70] C. Jin, Y. Zhang, S. Balakrishnan, M. J. Wainwright, and M. I. Jordan, "Local maxima in the likelihood of Gaussian mixture models: Structural results and algorithmic consequences," in *Advances in neural information processing systems*, 2016, pp. 4116–4124.
- [71] P. J. Werbos, "Backpropagation through time: what it does and how to do it," *Proceedings of the IEEE*, vol. 78, no. 10, pp. 1550–1560, 1990.
- [72] J. Xu, D. J. Hsu, and A. Maleki, "Global analysis of expectation maximization for mixtures of two Gaussians," *Advances in Neural Information Processing Systems*, vol. 29, 2016.

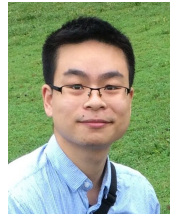


Jingyuan Xia received his B.Sc. and M.Sc. degree from the National University of Defense Technology (NUDT), Changsha, China, and received his Ph.D. degree from Imperial College London (ICL) in 2020. He has been a lecturer with the college of the Electronic Science at the National University of Defense Technology (NUDT) since 2020. His current research interests include machine learning, non-convex optimization, and representation learning.



Shengxi Li (S'14-M'22) received his Bachelor degree in Beihang University in Jul. 2014, the Master degree in Mar. 2016, also in Beihang University, and got his PhD degree in Aug. 2021, in EEE department of Imperial College London, under the supervision by Prof. Danilo Mandic. He is now an associate professor in Beihang University. His research interests include generative models, statistical signal processing, rate distortion theory, and perceptual video coding. He is the recipient of Imperial Lee Family Scholarship and Chinese Government Award

for Outstanding Self-financed Students Abroad.



Jun-Jie Huang (Member, IEEE) received the B.Eng. (Hons.) degree with First Class Honours in Electronic Engineering and the M.Phil. degree in Electronic and Information Engineering from The Hong Kong Polytechnic University, Hong Kong, China, in 2013 and 2015, respectively, and the Ph.D. degree from Imperial College London (ICL), London, U.K., in 2019. He is a Lecturer with College of Computer Science, National University of Defence Technology (NUDT), Changsha, China. During 2019 - 2021, he was a Postdoc with Communications and Signal Processing (CSP) Group, Electrical and Electronic Engineering Department, ICL, London, U.K. His research interests include the areas of computer vision, signal processing and deep learning.



Zhixiong Yang received the B.Sc. degree from the Northeastern University, China. He is now pursuing a M.Sc. degree at the College of Electronic Science, the National University of Defense Technology. His research interests include deep learning on signal processing and image processing.



Imad M. Jaimoukha received the B.Sc degree in Electrical Engineering from the University of Southampton, UK, in 1983, and the MSc and PhD degrees in control systems from Imperial College London, UK, in 1986 and 1990, respectively. He was a Research Fellow with the Centre for Process Systems Engineering at Imperial College London from 1990 to 1994. Since 1994, he has been with the Department of Electrical and Electronic Engineering, Imperial College London. His research interests include robust and fault tolerant control, system approximation, and global optimization.



Deniz Gündüz received his M.S. and Ph.D. degrees from NYU Tandon School of Engineering (formerly Polytechnic University) in 2004 and 2007, respectively. After his PhD, he served as a postdoctoral research associate at Princeton University, as a consulting assistant professor at Stanford University, and as a research associate at CTTC in Barcelona, Spain. In Sep. 2012, he joined the Electrical and Electronic Engineering Department of Imperial College London, UK, where he is currently a Professor of Information Processing, and serves as the deputy head of the Intelligent Systems and Networks Group. He is also a part-time faculty member at the University of Modena and Reggio Emilia, Italy, and has held visiting positions at University of Padova (2018-2020) and Princeton University (2009-2012). Dr. Gündüz is a Fellow of the IEEE, and a Distinguished Lecturer for the IEEE Information Theory Society (2020-22). He is the recipient of a Consolidator Grant of the European Research Council (ERC) in 2022, the IEEE Communications Society - Communication Theory Technical Committee (CTTC) Early Achievement Award in 2017, a Starting Grant of the European Research Council (ERC) in 2016, and several best paper awards.