

LEARNING LABEL DISTRIBUTION WITH SUBTASKS

Anonymous authors

Paper under double-blind review

ABSTRACT

Label distribution learning (LDL) is a novel learning paradigm that emulates label polysemy by assigning label distributions over the label space. However, recent LDL work seems to exhibit a notable contradiction: 1) some existing LDL methods employ auxiliary tasks to enhance performance, which narrows their focus to specific domains, thereby lacking generalization capability; 2) conversely, LDL methods without auxiliary tasks rely on losses tailored solely to label distributions of the primary task, lacking additional supervised information to guide the learning process. In this paper, we propose $\mathcal{S}$ -LDL, a novel and minimalist solution that partitions the label distribution of the primary task into subtask label distributions, i.e., a form of pseudo-supervised information, to reconcile the above contradiction. $\mathcal{S}$ -LDL encompasses two key aspects: 1) an algorithm capable of generating subtasks without any extra knowledge, with subtasks deemed valid and reconstructable via our analysis; and 2) a plug-and-play framework seamlessly compatible with existing LDL methods, and even adaptable to derivative tasks of LDL. Experiments demonstrate that $\mathcal{S}$ -LDL is effective and efficient. To the best of our knowledge, this represents the first endeavor to address LDL via subtasks. The code will soon be available on GitHub to facilitate reproducible research.

1 INTRODUCTION

Multi-label learning (MLL) (Zhang and Zhou, 2013) handles label polysemy in a binary manner, whereas label distribution learning (LDL) (Geng, 2016) offers a more nuanced perspective by answering: “How much does each label y describe the instance x ?”. This is accomplished through the concept of a *label distribution* $\mathbf{d}$, which is a form of probability simplex that assigns a real value (i.e., *description degree* d_x^y) to each label of each instance. This form introduces a quantitative manner to address label polysemy and extends LDL’s practical applications to a wider range, e.g., counting (or grading) (Geng et al., 2013; Wu et al., 2019), sentiment analysis (Chen et al., 2020; Le et al., 2023), segmentation (Gao et al., 2017; Li et al., 2023b), etc. Concurrently, more and more derivative tasks of LDL (González et al., 2021b; Lu and Jia, 2022; Wang, Jing and Geng, Xin, 2019; Xu and Zhou, 2017; Xu et al., 2019) are emerging to offer assistance in various real-world dilemmas.

However, LDL encounters a spectrum of challenges: 1) label distributions are bound by two constraints, non-negativity (i.e., $d_x^y \geq 0$) and sum-to-one (i.e., $\sum_{y \in \mathcal{Y}} d_x^y = 1$), and are often formed from mixture distributions, posing significant hurdles for fitting, particularly when employing a maximum entropy model (Shen et al., 2017); 2) label distribution matrices are usually obtained via crowdsourcing, which is time-consuming and labor-intensive, so one often copes with scarce and low-quality datasets (Wang et al., 2023). These two key issues stand as formidable barriers to performance improvement in LDL.

With the widespread use of multi-task learning, some LDL work tries to compensate for performance from the perspective of *auxiliary tasks*, which are learned concurrently alongside the primary task, thereby refining its representations and ultimately boosting performance. Unfortunately, though these methods can exploit additional supervised information, they 1) do not address the first key issue mentioned above; and 2) require extra knowledge (e.g., facial characteristics (Chen et al., 2020), pathology criteria (Wu et al., 2019), emotion wheel theory in psychology (Yang et al., 2017a), etc.) or similar domain-specific data (Zhao et al., 2023b), limiting their generalization capability to those corresponding specific domains. Conversely, LDL methods that do not take advantage of auxiliary tasks, despite their efforts in loss function engineering and network structure design, they 1) do not address the second key issue mentioned above; and 2) focus solely on one aspect of label correlations

(e.g., correlation of local instances (Jia et al., 2019), ranking relation (Jia et al., 2023), suboptimal label (Wang, Jing and Geng, Xin, 2019), etc.), each with its own set of limitations.

The generalizability across various domains appears to conflict with the ability to exploit additional data, so benefiting from both simultaneously seems elusive. However, we can still see the light from some MLL methods, which partition the label space and apply operations on these subspaces (Tsoumakas et al., 2008; 2010). These methods construct *subtasks* without involving extra knowledge and exhibit applicability across various domains. Intuitively, in the context of LDL, reliable supervised information can be generated from these subtasks, which can eventually be aggregated and reconstructed to the information of the primary task via ensemble strategies. Although existing label distribution ensemble practices demonstrate promising performance (González et al., 2021a; Shen et al., 2017), they focus only on the supervised information of the primary task.

In this paper, we introduce $\mathcal{S}$ -LDL, a novel and minimalist label distribution learning algorithm that constructs and exploits subtasks, to reconcile the contradiction between the generalizability across various domains and the ability to exploit additional data. Serving as auxiliary tasks, subtasks 1) provide different views of the primary task distribution, rendering the mixture of distributions more traceable (i.e., the key issue one); 2) furnish additional supervised data to mitigate the scarcity and ambiguity inherent in LDL datasets (i.e., the key issue two); 3) require no extra knowledge from specific domains; and 4) emphasize various label correlations via partitioning of the label space.

The main contributions of this paper are outlined below: 1) we propose $\mathcal{S}$ -LDL, which is considered the first endeavor to address LDL via subtasks; 2) our analysis shows the validity and reconstructability of these subtasks; 3) we present a plug-and-play framework seamlessly compatible with existing LDL methods, and adaptable to derivative tasks of LDL; and 4) the code will be available on GitHub soon, facilitating reproducible research endeavors.

2 RELATED WORK

LDL Our work is mainly related to LDL. Initially employed to tackle age estimation (Geng et al., 2013), LDL has evolved into a novel machine learning paradigm (Geng, 2016), which is supported by theoretical underpinnings (Wang and Geng, 2019) and features various derivative tasks (González et al., 2021b; Lu and Jia, 2022; Wang, Jing and Geng, Xin, 2019; Xu and Zhou, 2017; Xu et al., 2019). Most methods focus on improving performance via loss function engineering (Jia et al., 2019; 2023; Ren et al., 2019; Wen et al., 2023) or efficient model structures (González et al., 2021a; Jin et al., 2024; Shen et al., 2017; Yang et al., 2017b), while some work is dedicated to practical application scenarios (Gao et al., 2017; Li et al., 2023a; Shirani et al., 2019; Wu et al., 2019). However, the scarcity of label distribution datasets and the complexity of the label distribution itself make it difficult to further improve performance, at which point one may think of leveraging auxiliary tasks.

LDL with auxiliary tasks While there are LDL methods that leverage auxiliary tasks to enhance performance, they often rely on knowledge from disparate domains, extending beyond the scope of the LDL task. For example, LDL-ALSG (Chen et al., 2020) designs auxiliary tasks dedicated to facial emotion recognition, necessitating the use of external tools to extract facial points and action units from human faces. Wu et al. (2019) exploit the Hayashi criterion, a rule for counting and grading in acne lesions, which results in their method being only applicable in a small branch of the dermatology field. Yang et al. (2017a) employ a multi-task framework for image emotion classification, designing constraints inspired by Mikel’s wheel, a psychological emotion model, which also suffers from similar limitations. As LDL methods of transfer learning, GLDL (Zhao et al., 2023b) utilizes data from one or more source domains, which is not easy to obtain in practical applications. The need for specific extra knowledge significantly narrows the application scenarios of these methods.

MLL with partitioning of the label space For reference, there exist MLL methods based on partitioning of the label space, which can construct multi-label subtasks without involving additional knowledge and can be widely used in various domains. The most classic related work is that of HOMER (Tsoumakas et al., 2008) and RAKEL (Tsoumakas et al., 2010), the former forms a hierarchy of label subspaces while the latter randomly selects label subspaces. Many subsequent papers have been inspired by them (Prabhu et al., 2018; Read et al., 2013; Wang et al., 2021). Read et al. (2014) present a general framework of label subspaces and provide some theoretical justification for it. Since

Table 1: Key notation and terminology in this paper

Symbol	Description	Example
$\mathcal{Y} = \{y_j\}_{j=1}^L$	Label space (L labels)	$\mathcal{Y} = \{\text{HA, SA, SU, AN, DI, FE}\}^1$
$\mathcal{Y}^\circ$	Subtask label spaces	$\{\dots, \mathcal{Y}^{(t)} = \{\text{HA, SA, SU, FE}\}, \dots\}$
$d_{\mathbf{x}_i}^{y_j}$	Description degree of $\mathbf{x}_i$ about y_j	$d_{\mathbf{x}_i}^{y_0} = 0.4$, i.e., HA describes $\mathbf{x}_i$ by 0.4
$\mathbf{d}_i = (d_{\mathbf{x}_i}^{y_j})_{j=1}^L$	Label distribution of $\mathbf{x}_i$	$\mathbf{d}_i = (0.4, 0.05, 0.3, 0.1, 0.1, 0.05)$
$\mathbf{D} = (\mathbf{d}_i)_{i=1}^N = (\mathbf{d}_{\bullet,j})_{j=1}^L$	Distribution matrix (N samples)	$(\dots, \mathbf{d}_i, \dots)$
$\mathbf{d}_i^{(t)} = (d_{\mathbf{x}_i}^{(t)y_j})_{j=1}^{ \mathcal{Y}^{(t)} }$	Subtask label distribution	$\mathbf{d}_i^{(t)} = (0.5, 0.0625, 0.375, 0.0625)$
$\mathcal{D}^\circ$	Subtask distribution matrices	$\{\dots, \mathbf{D}^{(t)}, \dots\}$
$\mathbf{M} = (\mathbf{m}_t)_{t=1}^T = (M_{tj})$	Mask matrix (T anticipated tasks)	$(\dots, \mathbf{m}_t = (1, 1, 1, 0, 0, 1), \dots)$

label distribution contains rich knowledge, we can follow the patterns of these methods to construct label distribution subtasks.

LDL with ensemble strategy It is imperative to aggregate the output of subtasks. Fortunately, ensemble-based LDL methods have demonstrated promising performance. For instance, LDLFs (Shen et al., 2017) learns different label distributions on the leaf nodes of differentiable decision trees and learns weights that aggregate these label distributions. DF-LDL (González et al., 2021a) aggregates the label distribution of output of multiple base models by simple averaging, while Zhai et al. (2018) focus on aggregating the results of various neural networks via a combining learner. However, 1) the above methods are not suitable for incomplete label spaces (i.e., subtask label spaces); and 2) *none* of them involve the partitioning of the label space, therefore no extra supervised information of label distributions is constructed.

Drawing from the analysis of the aforementioned related work, we introduce $\mathcal{S}$ -LDL, which leverages pseudo-supervised information from subtasks to eliminate reliance on additional knowledge from disparate domains, and facilitates the creation of a novel knowledge dimension in a generic framework.

3 SUBTASK CONSTRUCTION

3.1 PRELIMINARY

Notation Vectors are denoted by lowercase bold letters, e.g., $\mathbf{v}$, and the corresponding regular letter with subscript i , i.e., v_i , indicates its i -th element. Matrices are denoted by uppercase bold letters, e.g., $\mathbf{A}$. The row vector $\mathbf{a}_i$ indicates its i -th row and the column vector $\mathbf{a}_{\bullet,j}$ indicates its j -th column. A_{ij} is the element in i -th row and j -th column of $\mathbf{A}$. The superscript (t) indicates that a symbol corresponds to the t -th subtask. Table 1 outlines the key notation in this paper.

Problem definition Let $\mathbf{x} \in \mathbb{R}^P$ denote the feature of the instance and $\mathbf{d} \in \Delta^{L-1}$ denote the label distribution, where $\Delta^{k-1} \triangleq \{\mathbf{v} \in \mathbb{R}^k \mid \mathbf{1}\mathbf{v}^T = 1, \mathbf{v} \geq 0\}$ is the $(k-1)$ -dimensional probability simplex. The goal of LDL is to find a mapping $\zeta : \mathbf{x} \mapsto \mathbf{d}$. In this paper, we partition the label space $\mathcal{Y}$ corresponding to $\mathbf{d}$ to obtain the subtask label space set $\mathcal{Y}^\circ$, then accordingly generate pseudo-supervised information, i.e., subtask distribution matrix set $\mathcal{D}^\circ$, to guide the learning of ζ .

Technical challenges Our first challenge arises from *the exponential growth in partitions* as the number of labels increases (Tsoumakas et al., 2010). When generating T tasks from a label space with L labels, the number of unique partitions is given by $(2^L - L - 2)! / (T!(2^L - L - 2 - T)!)!$. This makes it impractical to calculate metric for each case to select subtasks. We tackle this challenge in a mask matrix learning manner. The second challenge lies in *discerning reasonable partitions*. Since the label distribution matrix is usually imbalanced in average description degree (Zhao et al., 2023a), some partitions exhibit unreasonable local ignorance. As a result, the corresponding spaces struggle to handle the majority of instances, because 1) theoretically, there is no objective standard for the degree of negative correlation; and 2) empirically, weakly or negatively correlated information is

¹HA, SA, SU, AN, DI, FE, and NE represent the seven common emotions in sentiment analysis datasets, namely happiness, sadness, surprise, anger, disgust, fear, and neutral, respectively.

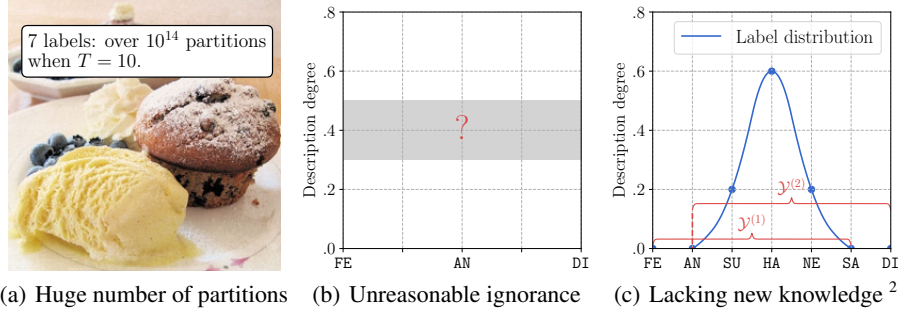


Figure 1: (a) is sourced from the emotion6 (Yang et al., 2017b) dataset, which has only 7 labels, but the number of potential partitions is huge. (b) exemplifies a subtask label space $\{FE, AN, DI\}$, which is challenging to describe (a). (c)’s two subtask label spaces encompass all descriptive information, meaning no new knowledge is generated about (a). This example vividly illustrates the limitations that may result from local ignorance and lack of diversity in subtask label spaces.

easily overlooked by human annotators in crowdsourced datasets. To mitigate this, we incorporate the description degree as a metric for the reliability of supervised information in guiding the generation of subtask masks. The third challenge is *avoiding analogous pseudo-supervised information*, i.e., to generate label distributions containing new knowledge (González et al., 2021b). This necessitates fostering richness and diversity in both subtask label spaces and distributions. To achieve this objective, we 1) minimize pairwise similarity among subtask masks; and 2) normalize each subtask label distribution to yield brand new insights distinct from the label distribution of the primary task. Fig. 1 portrays an illustrative example of these challenges.

3.2 LEARNING SUBTASK MASKS

Let $\mathbf{M} \in \{0, 1\}^{T \times L}$ denote the subtask mask matrix, where T represents the number of anticipated tasks. To ensure that the subtask label spaces contain as reliable information as possible, the learning of the subtask mask matrix can be converted into this problem: $\arg \max_{\mathbf{M}} \|\mathbf{D}\mathbf{M}^T\|_F$.

Obviously, a senseless solution is $\mathbf{m}_t = \mathbf{1}$ where $t = 1, \dots, T$, i.e., all pseudo-supervised information is equivalent to the primary task information. Therefore, solving the above problem alone is inappropriate. To address this, we consider pairwise similarity among subtask masks. We also employ exponential tricks to convert maximization into minimization. Finally, $\mathbf{M}$ is calculated as

$$\mathbf{M}^* = \arg \min_{\mathbf{M}} \left(\frac{1}{NT} \sum_{t=1}^T \sum_{i=1}^N \exp(-\mathbf{d}_i \mathbf{m}_t^\top) + \frac{2\lambda}{(T(T-1))} \sum_{i,j,i \neq j} \frac{\mathbf{m}_i \mathbf{m}_j^\top}{\|\mathbf{m}_i\| \|\mathbf{m}_j\|} \right), \quad (1)$$

s.t. $M_{tj} \in \{0, 1\}; t = 1, \dots, T; j = 1, \dots, L,$

where λ is a trade-off parameter. Eq. (1) is slightly more complicated than conventional integer programming. For convenience, we solve it using the stochastic gradient descent (SGD) method, with its constraint enforced via sigmoid (a conversion threshold is set, where outputs greater than it are set to 1, while those below are set to 0). Refer to Section 4.1 for an analysis of the validity of Eq. (1).

3.3 GENERATING SUBTASK DISTRIBUTIONS

We slice the label distribution matrix according to the subtask label space. To generate diversified subtask label distributions, we perform normalization on each subtask distribution with

$$[\mathcal{N}_{\text{SUM}}(\mathbf{v})]_j = \frac{v_j}{\sum_{i=1}^{|\mathbf{v}|} v_i}. \quad (2)$$

²Despite the discrete label space, in the field of LDL, the label distribution is intentionally plotted as a curve, to distinguish it from the logical labels.

Algorithm 1 Subtask construction

Input: Input matrix D , trade-off parameter λ , anticipated number of subtasks T .
Output: Subtask distribution matrices $\mathcal{D}^\circ$ (with corresponding subtask label spaces $\mathcal{Y}^\circ$).

- 1: Initialization: $\mathcal{Y}^\circ = \{\emptyset\}$, $\mathcal{D}^\circ = \{\emptyset\}$;
- 2: Calculate M using SGD; $\triangleright$ (Eq. (1))
- 3: **for** $t = 1$ **to** T **do**
- 4: $\mathcal{Y}^{(t)} \leftarrow \{y_j\}$ **if** $M_{tj} = 1$;
- 5: **if** $|\mathcal{Y}^{(t)}| = L$ **or** $|\mathcal{Y}^{(t)}| \leq 1$ **then**
- 6: **continue**; /* Ignore invalid subtask masks. */
- 7: **end if**
- 8: /* The clip(x, a, b) function limits x to be within $[a, b]$. */
- 9: $D^{(t)} \leftarrow \text{clip}(d_{\bullet j}, \varepsilon, 1)$ **if** $y_j \in \mathcal{Y}^{(t)}$; /* ε is a very small positive number. */
- 10: **for** $i = 1$ **to** N **do**
- 11: Normalization: $d_i^{(t)} \leftarrow \mathcal{N}_{\text{SUM}}(d_i^{(t)})$; $\triangleright$ (Eq. (2))
- 12: **end for**
- 13: $\mathcal{Y}^\circ = \mathcal{Y}^{(t)} \cup \mathcal{Y}^\circ$, $\mathcal{D}^\circ = D^{(t)} \cup \mathcal{D}^\circ$;
- 14: **end for**

Algorithm 2 S-LDL (shallow regime)

Input: Feature matrix X , label distribution matrix D , testing instance x' .
Output: Predicted label distribution d' for instance x' .

- 1: Initialize parameter of each estimator;
- 2: $\mathcal{D}^\circ \leftarrow \text{SC}(D)$;
- 3: **for** $t = 1$ **to** $|\mathcal{D}^\circ|$ **do**
- 4: Fit an estimator $f^{(t)}$ on dataset $\{X, D^{(t)}\}$;
- 5: $d^{(t)'} \leftarrow f^{(t)}(x')$;
- 6: **end for**
- 7: Concatenate X and all of the $D^{(t)}$ s to get Z , where $t = 1, \dots, |\mathcal{D}^\circ|$;
- 8: Fit an estimator f on dataset $\{Z, D\}$;
- 9: Concatenate x' and all of the $d^{(t)'}$ s to get z' , where $t = 1, \dots, |\mathcal{D}^\circ|$;
- 10: $d' \leftarrow f(z')$;

The rationale for utilizing $\mathcal{N}_{\text{SUM}}$ as the normalization function can be found in Section 4.2. The overall subtask construction process, denoted by SC, is illustrated in Alg. 1. Then, one can naturally come up with an adaptive LDL pipeline based on the shallow regime, as depicted in Alg. 2.

4 ANALYSIS ABOUT SUBTASK CONSTRUCTION

In this section, we analyze the subtask construction algorithm SC by studying the following questions:

- Q_1 : Are the subtask spaces provided by Eq. (1) valid for performance improvement? Can one configure λ and T in Eq. (1) without any prior knowledge?
- Q_2 : Are the subtask label distributions provided by Eq. (2) reconstructable? Can one replace Eq. (2) with other normalization functions?
- Q_3 : What is the overall time complexity of SC? Is it practical for large-scale datasets?

The validity, reconstructability, and complexity analysis are conducted for Q_1 , Q_2 , and Q_3 , respectively.

4.1 VALIDITY ANALYSIS

Eq. (1) manages the intricate task of selecting subtask spaces via λ and T . On the one hand, we strive to explain that it is useful for performance improvement to suppress local ignorance and increase diversity of each subtask label space simultaneously. On the other hand, we seek to determine the appropriate λ and T without any prior knowledge. To this end, we design the following two metrics.

Definition 1 (Information rate). *We call it informative if $M_{tj} = 1$ where $t = 1, \dots, T$ and $j = 1, \dots, L$. Let I be the summation of information; we define the information rate as $I/(TL)$.*

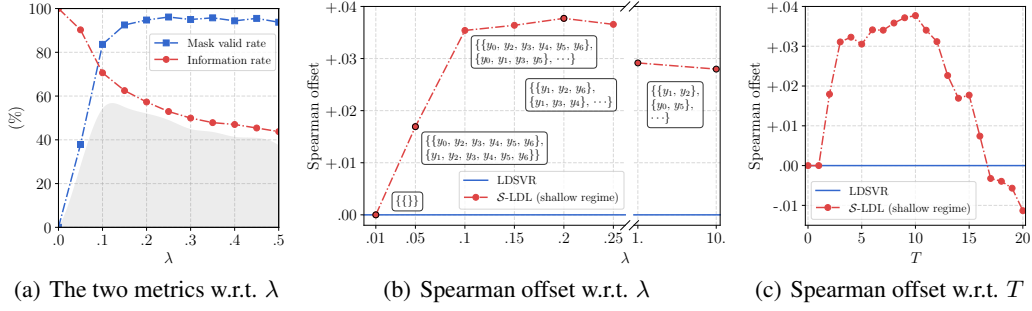


Figure 2: Visualized results of the validity analysis. Results of (a) are the average of experiments on all datasets (introduced in the appendix), while results of (b) and (c) are on the emotion6 dataset. The blue lines in (b) and (c) represent performance without auxiliary tasks.

Definition 2 (Mask valid rate). Let $\delta(\cdot, \cdot)$ be the Kronecker delta function. For all $t = 1, \dots, T$, the following are considered counting of invalid masks: 1) $\delta(\mathbf{m}_t, \mathbf{1})$, or 2) $\delta(|\mathbf{m}_t|!, 1)$, or 3) excluding masks in cases 1) and 2), for any remaining mask index i , $\delta(\mathbf{m}_i, \mathbf{m}_j)$, where $j = 1, \dots, i$.³ Let S be the summation of all invalid subtasks; we define the mask valid rate as $1 - S/T$.

While it is unknown which metric is more important, we intuitively claim that higher values for both metrics are likely to lead to better performance. First, we calculate the average of these two metrics for all datasets with varying λ , results of which are shown in Fig. 2(a). The grey area depicts the average of the two metrics, implying that the appropriate value of λ may be greater than 0.1.

It is important to highlight that Alg. 2 relies on naive concatenation operations and is not tied to representation learning. Consequently, any performance improvement over the base estimator is solely attributed to the effects of the subtask label distributions. Therefore, we employ Alg. 2 as the “scaffolding” for our analysis. With λ varying and T fixed at 10, we conduct ten-fold experiments repeated 10 times on the emotion6 dataset using Alg. 2. Here, $f^{(t)}$ s and f are implemented by a representative LDL method, LDSVR (Geng and Hou, 2015). We record the average Spearman’s coefficient (the higher the better). The results, which are shown in Fig. 2(b), support our claim. The panels from left to right display examples of subtask label spaces when the λ is 0.01, 0.05, 0.2, 1, and 10, respectively. When λ is suitable, label spaces are diverse and do not have excessive local ignorance; as λ decreases, label spaces tend to be homogeneous, and invalid masks account for the majority; as λ increases, the local ignorance of each label space becomes significant.

Besides, we also study the parameter sensitivity of T with λ fixed at 0.2. Results are shown in Fig. 2(c), illustrating that having a plethora of auxiliary tasks are detrimental to performance, which may be due to overfitting.

The validity analysis demonstrates that simultaneously avoiding local ignorance and homogeneity can lead to more efficient subtask label spaces, thereby improving performance. Without any prior knowledge, λ and T are recommended to be set to 0.2 and 10, respectively. Since the validity of Eq. (1) is ensured, one may wonder about the rationality and necessity of Eq. (2).

4.2 RECONSTRUCTABILITY ANALYSIS

We strive to choose a normalization function so that subtask label distributions retain more information, even efficacious enough to reconstruct the label distribution of the primary task. Theorem 1 illustrates that Eq. (2) is the only possibility.

Theorem 1. Let each subtask label space form a connected graph with its each label as a node. Then merge these graphs according to their respective labels to form $\mathcal{G}$. If and only if N_{SUM} is used for normalization, the primary label distribution can be reconstructed from these subtask label distributions, when the following conditions are satisfied: 1) $\mathcal{G}$ is connected; 2) $\mathcal{G}$ covers all labels in the label space, and 3) corresponding description degrees of all cut vertices of $\mathcal{G}$ are not zero.

³These three cases correspond to 1) masks that are exactly the same as the primary task; 2) masks that fail to form label distributions; and 3) duplicate masks among the remaining masks, respectively.

Proof. We solely discuss the extreme case where two subtask label spaces overlap with just one label. Further specialized cases can be deduced by the reader via induction. With a little bit of symbol abuse, let the general normalization function be defined as $\mathcal{N}(\mathbf{v}) \triangleq p(\mathbf{v})/q(\mathbf{v})$. Assume that there is a label distribution $\mathbf{d} = (d_1, \dots, d_L)$ and its corresponding label space is $\mathcal{Y} = \{y_1, \dots, y_L\}$. The two decompositions of $\mathcal{Y}$ are $\mathcal{Y}_a = \{y_1, \dots, y_k\}$ and $\mathcal{Y}_b = \{y_k, \dots, y_L\}$, respectively. It is obvious that $\mathcal{Y}_a \cup \mathcal{Y}_b = \mathcal{Y}$ and $\mathcal{Y}_a \cap \mathcal{Y}_b = \{y_k\}$. Let the subspace label distribution corresponding to these two decompositions be $\mathbf{a} = (a_1, \dots, a_k)$ and $\mathbf{b} = (b_k, \dots, b_L)$. According to our assumptions, $d_k \neq 0$. Then, for any integer $j \in [1, k]$, we have

$$\frac{a_j}{a_k} = \frac{[\mathcal{N}(\mathbf{d})]_j}{[\mathcal{N}(\mathbf{d})]_k} = \frac{[p(\mathbf{d})]_j}{[q(\mathbf{d})]_j} \frac{[q(\mathbf{d})]_k}{[p(\mathbf{d})]_k}. \quad (3)$$

Typically, for most normalization functions, $q(\cdot)$ is a normalizing constant, i.e., $[q(\mathbf{d})]_j = [q(\mathbf{d})]_k$. Thus Eq. (3) can be rewritten into $a_j[p(\mathbf{d})]_k = a_k[p(\mathbf{d})]_j$. Plug it into $\sum_{j=1}^k a_j = 1$, and do the same for $\mathbf{b}$ as well, and get

$$\frac{a_k \sum_{j=1}^k [p(\mathbf{d})]_j}{[p(\mathbf{d})]_k} = 1, \quad \frac{b_k \sum_{j=k}^L [p(\mathbf{d})]_j}{[p(\mathbf{d})]_k} = 1. \quad (4)$$

Add these two equations together, we have

$$[p(\mathbf{d})]_k + \sum_{j=1}^L [p(\mathbf{d})]_j = \frac{[p(\mathbf{d})]_k}{a_k} + \frac{[p(\mathbf{d})]_k}{b_k}. \quad (5)$$

Eq. (5) implies that $\sum_{j=1}^L [p(\mathbf{d})]_j$ must be given, and $[p(\mathbf{d})]_k$ is related to d_k , and only d_k . To make it possible, the only thing we can exploit is the sum-to-one constraint of $\mathbf{d}$, i.e., $\sum_{j=1}^L d_j = 1$. Therefore $[p(\mathbf{v})]_j = v_j$. Since $\sum_i^{|\mathcal{V}|} [\mathcal{N}(\mathbf{v})]_i = 1$, we have $q(\mathbf{v}) = \sum_i^{|\mathcal{V}|} v_i$, i.e., the finally deduced normalization function is Eq. (2). In this case, for any integer $j \in [1, L]$, we have

$$d_j = \begin{cases} \frac{a_j b_k}{a_k + b_k - a_k b_k}, & j = 1, \dots, k \\ \frac{a_k b_j}{a_k + b_k - a_k b_k}, & j = k + 1, \dots, L \end{cases}, \quad (6)$$

which illustrates that the original label distribution $\mathbf{d}$ can be reconstructed by subtask label distributions $\mathbf{a}$ and $\mathbf{b}$. This is possible thanks to the use of $\mathcal{N}_{\text{SUM}}$. $\square$

Theorem 1 also states that it is not appropriate to replace Eq. (2) with the min-max or softmax function because doing so destroys the reconstruction information.

4.3 COMPLEXITY ANALYSIS

The overall time cost of SC is primarily influenced by the calculation of $\mathbf{M}$ and the normalization process. The time complexity of computing and updating $\mathbf{M}$ are $\mathcal{O}(L(TN + T^2))$ and $\mathcal{O}(LT)$, respectively. The time complexity of the normalization process is $\mathcal{O}(LTN)$. The overall time complexity of each iteration of SC is $\mathcal{O}(L(TN + T^2))$, which is linear with respect to the number of instances and labels. Therefore, it is clear that SC can be applied to large-scale datasets.

5 S-LDL OF THE DEEP REGIME

The aforementioned analysis has exposed the problems of the shallow regime: 1) shallow methods as base estimators have low potential in themselves; 2) there is a training gap between the primary task and subtasks, i.e., no representation learning is involved. Therefore, it is necessary to introduce our proposed S-LDL of the deep regime, the overview of which is illustrated in Fig. 3. We illustrate our framework by introducing the learnable parts one by one.

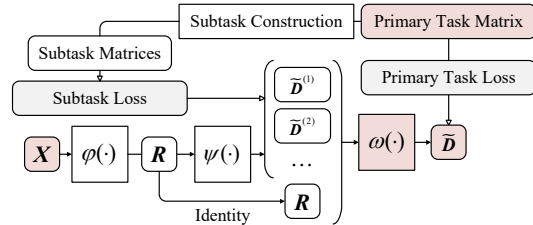


Figure 3: The overview of S-LDL (deep regime). White, red and gray highlight our proposed, existing methods, and loss functions, respectively.

Table 2: Modifications of different task adaptations

Type	Subtask construction	ℓ_{PRI}	ℓ_{SUB}
Vanilla LDL	$(\mathbf{D}^{(1)}, \dots) \leftarrow \text{SC}(\mathbf{D})$	$\ell(\mathbf{D}, \tilde{\mathbf{D}}) \in \mathcal{L}_{\text{LDL}}$	$\ell_{\text{SUB}}(\mathbf{D}^{(1)}, \dots; \tilde{\mathbf{D}}^{(1)}, \dots)$
LDL4C	$(\mathbf{D}^{(1)}, \dots) \leftarrow \text{SC}(\mathbf{D})$	$\ell(\mathbf{D}, \tilde{\mathbf{D}}, \tilde{\mathbf{D}}) \in \mathcal{L}_{\text{LDL4C}}$	$\ell_{\text{SUB}}(\mathbf{D}^{(1)}, \dots; \tilde{\mathbf{D}}^{(1)}, \dots)$
IncomLDL	$(\mathbf{D}_{\Omega}^{(1)}, \dots) \leftarrow \text{SC}(\mathcal{R}_{\Omega}(\mathbf{D}))$	$\ell(\mathcal{R}_{\Omega}(\mathbf{D}), \mathcal{R}_{\Omega}(\tilde{\mathbf{D}})) \in \mathcal{L}_{\text{IncomLDL}}$	$\ell_{\text{SUB}}(\mathbf{D}_{\Omega}^{(1)}, \dots; \mathcal{R}_{\Omega}(\tilde{\mathbf{D}}^{(1)}), \dots)$
LE	$(\mathbf{L}^{(1)}, \dots) \leftarrow \text{SC}(\mathbf{L})$	$\ell(\mathbf{L}, \tilde{\mathbf{D}}) \in \mathcal{L}_{\text{LE}}$	$\ell_{\text{SUB}}(\mathbf{L}^{(1)}, \dots; \tilde{\mathbf{D}}^{(1)}, \dots)$

- $\varphi(\cdot)$ is guided by subtasks to learn a powerful representation, i.e., $\mathbf{R} = \varphi(\mathbf{X})$.
- $\psi(\cdot)$ is responsible for predicting subtask label distributions, i.e., $(\tilde{\mathbf{D}}^{(1)}, \dots) = \psi(\mathbf{R})$. To ensure the precise prediction of subtask label distributions for reconstruction, we employ the mean absolute error function for subtask learning. The loss is weighted by the summation of the description degrees corresponding to the primary tasks, allowing more reliable label spaces to receive more attention. The subtask learning loss has the following form:

$$\ell_{\text{SUB}}(\mathcal{D}^{\circ}; \tilde{\mathcal{D}}^{\circ}) = \frac{1}{N |\mathcal{Y}^{\circ}|} \sum_{\mathcal{Y}^{(t)} \in \mathcal{Y}^{\circ}} \sum_{i=1}^N \left(\sum_{y_k \in \mathcal{Y}^{(t)}} d_{\mathbf{x}_i}^{y_k} \right) \sum_{j=1}^{|\mathcal{Y}^{(t)}|} \left| d_{\mathbf{x}_i}^{(t)y_j} - \tilde{d}_{\mathbf{x}_i}^{(t)y_j} \right|. \quad (7)$$

- $\omega(\cdot)$ can be any existing method that can be expressed as a network structure theoretically. Since the concatenation of the representation and subtask label distributions, we have $\mathbf{Z} = (\mathbf{R}, \psi(\mathbf{R}))$ and $\tilde{\mathbf{D}} = \omega(\mathbf{Z})$. In the case of the primary task being vanilla LDL, the primary task loss ℓ_{PRI} can be

$$\ell_{\text{KL}}(\mathbf{D}, \tilde{\mathbf{D}}) = \frac{1}{N} \sum_{i=1}^N \sum_{j=1}^L d_{\mathbf{x}_i}^{y_j} \ln \frac{d_{\mathbf{x}_i}^{y_j}}{\tilde{d}_{\mathbf{x}_i}^{y_j}}, \quad \ell_{\text{KL}} \in \mathcal{L}_{\text{LDL}}. \quad (8)$$

Note that ℓ_{PRI} changes as the primary task changes. Finally, we can learn the model parameters Θ by

$$\Theta^* = \arg \min_{\Theta} (\ell_{\text{PRI}} + \alpha \ell_{\text{SUB}}), \quad (9)$$

where α is a trade-off parameter. Compared with the shallow regime, $\mathcal{S}$ -LDL of the deep regime has the following advantages: 1) There is no two-stage training gap, which makes the representation contain insights from both the primary task and the subtasks; 2) the framework not only serves LDL, but can also be directly applied to derivative tasks of LDL, e.g., LDL for classification (LDL4C) (Wang, Jing and Geng, Xin, 2019), incomplete LDL (IncomLDL) (Xu and Zhou, 2017), label enhancement (LE) (Xu et al., 2019). The modifications involved are shown in Table 2, where $\mathcal{L}_X$ indicates the set of losses for adaptable methods in the task of type “X”. Special mathematical procedures of LDL4C and IncomLDL are defined as

$$\left[\tilde{\mathbf{D}}^{(t)} \right]_{ij} \triangleq \begin{cases} 1, & \text{if } y_j = \arg \max_{\bar{y} \in \mathcal{Y}^{(t)}} d_{\mathbf{x}_i}^{\bar{y}}, \\ 0, & \text{otherwise} \end{cases}, \quad [\mathcal{R}_{\Omega}(\mathbf{D})]_{ij} \triangleq \begin{cases} [\mathbf{D}]_{ij}, & \text{if } (i, y_j) \in \Omega \\ 0, & \text{otherwise} \end{cases}, \quad (10)$$

respectively, where $[\cdot]_{ij}$ represents the element in i -th row of the matrix corresponding to label y_j , and Ω represents observed elements sampled uniformly at random from $\mathbf{D}$ in IncomLDL. Such modifications are rational since: 1) targets of LDL4C and IncomLDL, i.e., $\tilde{\mathbf{D}}$ and $\mathcal{R}_{\Omega}(\mathbf{D})$, are essentially different forms of degradation of the label distribution matrix; and 2) the target of LE is a logical label matrix $\mathbf{L}$, the same as the target of MLL, which is actually a special case of LDL.

6 EXPERIMENTS

In this section, we evaluate $\mathcal{S}$ -LDL of the deep regime. Due to page limitations, datasets, comparison methods, and their parameter settings are introduced in the appendix.

Metrics For LDL, we use the same metrics suggested by Jia et al. (2023). Due to page limitations, we only present results on Cheby. $\downarrow$ (Chebyshev distance), Clark $\downarrow$ (Clark distance), Cosine $\uparrow$ (cosine similarity), and Spear. $\uparrow$ (Spearman’s coefficient) in the main paper, where $\downarrow$ ($\uparrow$) indicates “the lower (higher) the better”. Note that these metrics are *not* as intuitive as accuracy or error rate, i.e., *small changes can mean large performance differences*. For LDL4C, objective of which is different from LDL, we use 0/1 loss $\downarrow$ (zero one loss) and Err. prob. $\downarrow$ (error probability) as metrics (Wang, Jing and Geng, Xin, 2019).

Table 3: Experimental results of LDL on JAFFE and Yeast_diau formatted as (mean $\pm$ std(rank))

Algorithms	JAFFE (Lyons et al., 1998)		Algorithms	Yeast_diau (Geng, 2016)	
	Clark $\downarrow$	Cosine $\uparrow$		Cheby. $\downarrow$	Spear. $\uparrow$
LDSVR (Geng and Hou, 2015)	.3280 $\pm$.027 (6)	.9549 $\pm$.010 (7)	CPNN (Geng et al., 2013)	.0385 $\pm$.001 (9)	.2962 $\pm$.034 (10)
AA- k NN (Geng, 2016)	.3483 $\pm$.032 (8)	.9497 $\pm$.010 (9)	AA- k NN	.0385 $\pm$.001 (9)	.3674 $\pm$.029 (9)
LDLFs (Shen et al., 2017)	.3637 $\pm$.032 (10)	.9494 $\pm$.009 (10)	LDLFs	.0371 $\pm$.001 (8)	.4088 $\pm$.021 (8)
DF-BFGS (González et al., 2021a)	.3062 $\pm$.025 (3)	.9633 $\pm$.007 (2)	DF-BFGS	.0368 $\pm$.001 (5)	.4161 $\pm$.027 (5)
KLD (Geng, 2016) •	.3608 $\pm$.031 (9)	.9538 $\pm$.008 (8)	LRR •	.0370 $\pm$.001 (7)	.4154 $\pm$.023 (6)
$\mathcal{S}$ -KLD	.3007 $\pm$.032 (2)	.9625 $\pm$.009 (3)	$\mathcal{S}$ -LRR	.0366 $\pm$.001 (1)	.4198 $\pm$.023 (2)
SCL (Jia et al., 2019) •	.3358 $\pm$.024 (7)	.9592 $\pm$.006 (6)	QFD ² (Wen et al., 2023) •	.0369 $\pm$.001 (6)	.4118 $\pm$.025 (7)
$\mathcal{S}$ -SCL	.3184 $\pm$.025 (4)	.9604 $\pm$.008 (5)	$\mathcal{S}$ -QFD ²	.0366 $\pm$.001 (1)	.4203 $\pm$.021 (1)
LRR (Jia et al., 2023) •	.3230 $\pm$.027 (5)	.9616 $\pm$.008 (4)	CJS (Wen et al., 2023)	.0367 $\pm$.001 (4)	.4164 $\pm$.025 (4)
$\mathcal{S}$ -LRR	.2934 $\pm$.028 (1)	.9635 $\pm$.008 (1)	$\mathcal{S}$ -CJS	.0366 $\pm$.001 (1)	.4198 $\pm$.024 (2)

Table 4: Experimental results of LDL4C on sBU_3DFE and Flickr formatted as (mean $\pm$ std(rank))

Algorithms	sBU_3DFE (Geng, 2016)		Algorithms	Flickr (Yang et al., 2017b)	
	0/1 loss $\downarrow$	Err. prob. $\downarrow$		0/1 loss $\downarrow$	Err. prob. $\downarrow$
LDL4C (Wang, Jing and Geng, Xin, 2019)	.5578 $\pm$.028 (6)	.7671 $\pm$.007 (5)	LDL4C	.8971 $\pm$.008 (6)	.8884 $\pm$.004 (6)
$\mathcal{S}$ -LDL4C	.5526 $\pm$.025 (5)	.7686 $\pm$.006 (6)	$\mathcal{S}$ -LDL4C	.8705 $\pm$.138 (5)	.8702 $\pm$.100 (5)
HR (Wang and Geng, 2021a) •	.5167 $\pm$.027 (3)	.7596 $\pm$.006 (2)	HR •	.4513 $\pm$.015 (4)	.5823 $\pm$.007 (4)
$\mathcal{S}$ -HR	.5069 $\pm$.025 (2)	.7598 $\pm$.006 (3)	$\mathcal{S}$ -HR	.4219 $\pm$.015 (1)	.5639 $\pm$.007 (1)
LDLM (Wang and Geng, 2021b) •	.5258 $\pm$.034 (4)	.7619 $\pm$.009 (4)	LDLM •	.4384 $\pm$.014 (3)	.5740 $\pm$.007 (3)
$\mathcal{S}$ -LDLM	.4809 $\pm$.024 (1)	.7524 $\pm$.005 (1)	$\mathcal{S}$ -LDLM	.4321 $\pm$.016 (2)	.5667 $\pm$.007 (2)

Results and discussion We apply $\mathcal{S}$ -LDL to existing methods to demonstrate performance improvements. For each dataset we conduct ten-fold experiments repeated 10 times, and the average performance is recorded. Tables 3 to 4 show representative results and the remainder are in the appendix, where • (o) indicates that more than half of the metrics support that “ $\mathcal{S}$ -X” is statistically superior (inferior) to the corresponding methods “X” (pairwise t -test at 0.05 significance level); there is no significant if neither • nor o is shown. LRR focuses on the label ranking relationship, which is also emphasized by each subtask. We believe this is why $\mathcal{S}$ -LDL and LRR fit so well. Note that our method has the least improvement in SCL, which may be attributed to its reliance on shallow regime methods in the prediction phase. It is also worth noting that the improvement in vanilla KLD is considerable, which just illustrates the limitations of loss function engineering that considers label correlation one-sidedly. QFD² and CJS are tailored for ordinary LDL, and may have better results than LRR on this regard. Powered by $\mathcal{S}$ -LDL, these methods can all achieve better level. For LDL4C, $\mathcal{S}$ -LDL significantly improves both HR and LDLM. However, it can be observed that $\mathcal{S}$ -LDL4C is unstable on the Flickr dataset, which is not surprising since LDL4C itself fails on it. We believe this is caused by the combined effect of the sparsity of the dataset and the information entropy operation involved in LDL4C.

Parameter sensitivity We check the sensitivity of the trade-off parameter α on the LDL task with the Natural_Scene dataset by varying the parameter in $\{0.01, 0.05, 0.1, 0.5, 1, 5\}$. Results are shown in Fig. 4. Spearman’s coefficient of $\mathcal{S}$ -LDL first increases and then decreases as α varies, demonstrating a desirable bell-shaped curve. This justifies our motivation of jointly learning the primary task and subtasks, as a good trade-off between them can enhance the performance.

Ablation study Here we are interested in the importance of each part of $\mathcal{S}$ -LDL, thus an ablation study is performed with $\mathcal{S}$ -KLD: 1) we replace $\mathcal{N}_{\text{SUM}}$ in SC with the min-max function to examine the importance of the subtask distribution reconstruction, and this model is denoted as $\mathcal{S}$ -KLD (min-max); 2) we remove the identity mapping in Fig. 3 to examine the importance of the prediction via subtask representation, and this model is denoted as $\mathcal{S}$ -KLD w/o id.; 3) we train without the term of ℓ_{SUB} (i.e.,

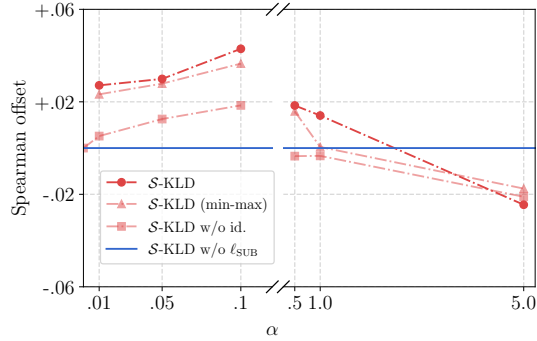


Figure 4: Visualized results of the ablation study and the parameter sensitivity analysis, which is on the Natural_Scene (Geng, 2016) dataset.

setting $\alpha = 0$) to examine the importance of subtask learning, and this model is denoted as $\mathcal{S}$ -KLD w/o ℓ_{SUB} . Results are also shown in Fig. 4, which confirms that each part of $\mathcal{S}$ -LDL contributes as long as there is a good trade-off.

7 LIMITATIONS AND CONCLUSION

Limitations First, $\mathcal{S}$ -LDL of the shallow regime is proposed out of intuition, and in Section 5, we have discussed its limitations, which are addressed via the designing of $\mathcal{S}$ -LDL of the deep regime. Second, when the label space is large, especially when labels are continuous and result in unimodal label distributions (e.g., age estimation), our proposed cannot be rationally applied. Fortunately, one possible workaround is to use a binning tricks for preprocessing, and then construct subtasks.

Conclusion We propose $\mathcal{S}$ -LDL, a subtask learning framework nested into LDL. $\mathcal{S}$ -LDL is generic: it generates pseudo-supervised information via subtask construction without any extra knowledge; $\mathcal{S}$ -LDL is minimalist: it can be attached to existing methods and handle derivative tasks; $\mathcal{S}$ -LDL is efficient: it captures a wide variety of label correlations. The analysis shows the validity and reconstructability of subtasks, and experiments show the superiority of our framework.

REFERENCES

- Shikai Chen, Jianfeng Wang, Yuedong Chen, Zhongchao Shi, Xin Geng, and Yong Rui. Label distribution learning on auxiliary label space graphs for facial expression recognition. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13984–13993, 2020.
- Bin-Bin Gao, Chao Xing, Chen-Wei Xie, Jianxin Wu, and Xin Geng. Deep label distribution learning with label ambiguity. *IEEE Transactions on Image Processing*, 26(6):2825–2838, 2017.
- Xin Geng. Label distribution learning. *IEEE Transactions on Knowledge and Data Engineering*, 28(7):1734–1748, 2016.
- Xin Geng and Peng Hou. Pre-release prediction of crowd opinion on movies by label distribution learning. In *Proceedings of the 24th International Joint Conference on Artificial Intelligence*, pages 3511–3517, 2015.
- Xin Geng, Chao Yin, and Zhi-Hua Zhou. Facial age estimation by learning from label distributions. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 35(10):2401–2412, 2013.
- Manuel González, Germán González-Almagro, Isaac Triguero, José-Ramón Cano, and Salvador García. Decomposition-fusion for label distribution learning. *Information Fusion*, 66:64–75, 2021a.
- Manuel González, Julián Luengo, José-Ramón Cano, and Salvador García. Synthetic sample generation for label distribution learning. *Information Sciences*, 544:197–213, 2021b.
- Xiuyi Jia, Zechao Li, Xiang Zheng, Weiwei Li, and Sheng-Jun Huang. Label distribution learning with label correlations on local samples. *IEEE Transactions on Knowledge and Data Engineering*, 33(4):1619–1631, 2019.
- Xiuyi Jia, Xiaoxia Shen, Weiwei Li, Yunan Lu, and Jihua Zhu. Label distribution learning by maintaining label ranking relation. *IEEE Transactions on Knowledge and Data Engineering*, 35(02):1695–1707, 2023.
- Yufei Jin, Richard Gao, Yi He, and Xingquan Zhu. Gldl: Graph label distribution learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 12965–12974, 2024.
- Diederik P. Kingma and Jimmy Lei Ba. Adam: A method for stochastic optimization. In *Proceedings of the 3rd International Conference on Learning Representations*, 2015.
- Nhat Le, Khanh Nguyen, Quang Tran, Erman Tjiputra, Bac Le, and Anh Nguyen. Uncertainty-aware label distribution learning for facial expression recognition. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*, pages 6088–6097, 2023.

- Xiangyu Li, Xinjie Liang, Gongning Luo, Wei Wang, Kuanquan Wang, and Shuo Li. Ambiguity-aware breast tumor cellularity estimation via self-ensemble label distribution learning. *Medical Image Analysis*, 90:102944, 2023a.
- Xiangyu Li, Gongning Luo, Wei Wang, Kuanquan Wang, and Shuo Li. Curriculum label distribution learning for imbalanced medical image segmentation. *Medical Image Analysis*, 89:102911, 2023b.
- Lingyu Liang, LuoJun Lin, Lianwen Jin, Duorui Xie, and Mengru Li. Scut-fbp5500: A diverse benchmark dataset for multi-paradigm facial beauty prediction. In *Proceedings of the 24th International Conference on Pattern Recognition*, pages 1598–1603, 2018.
- Yunan Lu and Xiuyi Jia. Predicting label distribution from multi-label ranking. In *Proceedings of the 36th Annual Conference on Neural Information Processing Systems*, pages 36931–36943, 2022.
- Michael Lyons, Shigeru Akamatsu, Miyuki Kamachi, and Jiro Gyoba. Coding facial expressions with gabor wavelets. In *Proceedings of the 3rd IEEE International Conference on Automatic Face and Gesture Recognition*, pages 200–205, 1998.
- Yashoteja Prabhu, Anil Kag, Shrutendra Harsola, Rahul Agrawal, and Manik Varma. Parabel: Partitioned label trees for extreme classification with application to dynamic search advertising. In *Proceedings of the 2018 World Wide Web Conference*, pages 993–1002, 2018.
- Jesse Read, Concha Bielza, and Pedro Larrañaga. Multi-dimensional classification with super-classes. *IEEE Transactions on Knowledge and Data Engineering*, 26(7):1720–1733, 2013.
- Jesse Read, Antti Puurula, and Albert Bifet. Multi-label classification with meta-labels. In *Proceedings of the 2014 IEEE International Conference on Data Mining*, pages 941–946, 2014.
- Tingting Ren, Xiuyi Jia, Weiwei Li, Lei Chen, and Zechao Li. Label distribution learning with label-specific features. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, pages 3318–3324, 2019.
- Wei Shen, Kai Zhao, Yilu Guo, and Alan Yuille. Label distribution learning forests. In *Proceedings of the 31st Annual Conference on Neural Information Processing Systems*, pages 834–843, 2017.
- Amirreza Shirani, Franck Dernoncourt, Paul Asente, Nedim Lipka, Seokhwan Kim, Jose Echevarria, and Thamar Solorio. Learning emphasis selection for written text in visual media from crowd-sourced label distributions. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 1167–1172, 2019.
- Grigorios Tsoumakas, Ioannis Katakis, and Ioannis Vlahavas. Effective and efficient multilabel classification in domains with large number of labels. In *ECML/PKDD 2008 Workshop on Mining Multidimensional Data*, volume 21, pages 53–59, 2008.
- Grigorios Tsoumakas, Ioannis Katakis, and Ioannis Vlahavas. Random k-labelsets for multilabel classification. *IEEE Transactions on Knowledge and Data Engineering*, 23(7):1079–1089, 2010.
- Jing Wang and Xin Geng. Theoretical analysis of label distribution learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 5256–5263, 2019.
- Jing Wang and Xin Geng. Learn the highest label and rest label description degrees. In *Proceedings of the 30th International Joint Conference on Artificial Intelligence*, pages 3097–3103, 2021a.
- Jing Wang and Xin Geng. Label distribution learning machine. In *Proceedings of the 38th International Conference on Machine Learning*, pages 10749–10759. PMLR, 2021b.
- Ke Wang, Ning Xu, Miaogen Ling, and Xin Geng. Fast label enhancement for label distribution learning. *IEEE Transactions on Knowledge and Data Engineering*, 35(02):1502–1514, 2023.
- Ran Wang, Sam Kwong, Xu Wang, and Yuheng Jia. Active k-labelsets ensemble for multi-label classification. *Pattern Recognition*, 109:107583, 2021.
- Wang, Jing and Geng, Xin. Classification with label distribution learning. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, pages 3712–3718, 2019.

- Changsong Wen, Xin Zhang, Xingxu Yao, and Jufeng Yang. Ordinal label distribution learning. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 23481–23491, 2023.
- Xiaoping Wu, Ni Wen, Jie Liang, Yukun Lai, Dongyu She, Mingming Cheng, and Jufeng Yang. Joint acne image grading and counting via label distribution learning. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 10642–10651, 2019.
- Miao Xu and Zhi-Hua Zhou. Incomplete label distribution learning. In *Proceedings of the 26th International Joint Conference on Artificial Intelligence*, pages 3175–3181, 2017.
- Ning Xu, Yunpeng Liu, and Xin Geng. Label enhancement for label distribution learning. *IEEE Transactions on Knowledge and Data Engineering*, 33(4):1632–1643, 2019.
- Ning Xu, Jun Shu, Renyi Zheng, Xin Geng, Deyu Meng, and Min-Ling Zhang. Variational label enhancement. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(5):6537–6551, 2023.
- Jufeng Yang, Dongyu She, and Ming Sun. Joint image emotion classification and distribution learning via deep convolutional neural network. In *Proceedings of the 26th International Joint Conference on Artificial Intelligence*, pages 3266–3272, 2017a.
- Jufeng Yang, Ming Sun, and Xiaoxiao Sun. Learning visual sentiment distributions via augmented conditional probability neural network. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 224–230, 2017b.
- Yansheng Zhai, Jianhua Dai, and Hong Shi. Label distribution learning based on ensemble neural networks. In *Proceedings of the 25th International Conference on Neural Information Processing*, pages 593–602, 2018.
- Min-Ling Zhang and Zhi-Hua Zhou. A review on multi-label learning algorithms. *IEEE Transactions on Knowledge and Data Engineering*, 26(8):1819–1837, 2013.
- Xingyu Zhao, Yuexuan An, Ning Xu, Jing Wang, and Xin Geng. Imbalanced label distribution learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 11336–11344, 2023a.
- Xingyu Zhao, Lei Qi, Yuexuan An, and Xin Geng. Generalizable label distribution learning. In *Proceedings of the 31st ACM International Conference on Multimedia*, pages 8932–8941, 2023b.
- Qinghai Zheng, Jihua Zhu, and Haoyu Tang. Label information bottleneck for label enhancement. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 7497–7506, 2023.

A APPENDIX: DETAILS OF EXPERIMENTS

Here we try our best to provide as much information as possible for reproducible research.

A.1 METRICS

For LDL, IncomLDL, and LE, we use the same metrics suggested by Geng (2016), which are Cheby. $\downarrow$ (Chebyshev distance), Clark $\downarrow$ (Clark distance), Can. $\downarrow$ (Canberra distance), KLD $\downarrow$ (Kullback-Leibler divergence), Cosine $\uparrow$ (cosine similarity), and Int. $\uparrow$ (intersection similarity), respectively. Here $\downarrow$ ($\uparrow$) indicates “the lower (higher) the better”. For LDL and IncomLDL, we additionally use two ranking metrics: Spear. $\uparrow$ (Spearman’s coefficient) and Ken. $\uparrow$ (Kendall’s coefficient) (Jia et al., 2023). Note that these metrics are *not* as intuitive as accuracy or error rate, i.e., *small changes can mean large performance differences*. For LDL4C, objective of which is different from LDL, we use 0/1 loss $\downarrow$ (zero one loss) and Err. prob. $\downarrow$ (error probability) as metrics (Wang, Jing and Geng, Xin, 2019). Let the real distribution be denoted by $\mathbf{u} = \{u_j\}_{j=1}^L$, and the predicted distribution be denoted by $\mathbf{v} = \{v_j\}_{j=1}^L$, then the above metrics can be summarized in Table 5, where $\rho(\cdot)$ and $\delta(\cdot, \cdot)$ are the ranking function and the Kronecker delta function, respectively.

Table 5: Summary of the metrics

Name	Formula	Name	Formula
Cheby. ↓	$\text{Dis}_1(\mathbf{u}, \mathbf{v}) = \max_j u_j - v_j $	Cosine ↑	$\text{Sim}_1(\mathbf{u}, \mathbf{v}) = \frac{\sum_{j=1}^L u_j v_j}{\sqrt{\sum_{j=1}^L u_j^2} \sqrt{\sum_{j=1}^L v_j^2}}$
Clark ↓	$\text{Dis}_2(\mathbf{u}, \mathbf{v}) = \sqrt{\sum_{j=1}^L \frac{(u_j - v_j)^2}{(u_j + v_j)^2}}$	Int. ↑	$\text{Sim}_2(\mathbf{u}, \mathbf{v}) = \sum_{j=1}^L \min(u_j, v_j)$
Can. ↓	$\text{Dis}_3(\mathbf{u}, \mathbf{v}) = \sum_{j=1}^L \frac{ u_j - v_j }{u_j + v_j}$	Spear. ↑	$\text{Rnk}_1(\mathbf{u}, \mathbf{v}) = 1 - \frac{6 \sum_{j=1}^L (\rho(u_j) - \rho(v_j))^2}{L(L^2 - 1)}$
KLD ↓	$\text{Dis}_4(\mathbf{u}, \mathbf{v}) = \sum_{j=1}^L u_j \ln \frac{u_j}{v_j}$	Ken. ↑	$\text{Rnk}_2(\mathbf{u}, \mathbf{v}) = \frac{2 \sum_{j < k} \text{sgn}(u_j - u_k) \text{sgn}(v_j - v_k)}{L(L-1)}$
0/1 loss ↓	$C_1(\mathbf{u}, \mathbf{v}) = \delta(\arg \max(\mathbf{u}), \arg \max(\mathbf{v}))$	Err. prob. ↓	$C_2(\mathbf{u}, \mathbf{v}) = 1 - u_{\arg \max(\mathbf{v})}$

A.2 DATASETS

We adopt several widely used label distribution datasets, including: JAFFE (Lyons et al., 1998);⁴ fbp5500 (Liang et al., 2018);⁵ sBU_3DFE, Movie, Natural_Scene, Yeast_heat, Yeast_diau, Yeast_cold, and Yeast_dtt provided by Geng (2016);⁶ emotion6, Twitter, and Flickr provided by Yang et al. (2017b).⁷ The information of these datasets are summarized in Table 6.

A.3 COMPARISON METHODS

Table 6: Summary of datasets

Dataset	# Instances N	# Features P	# Labels L
JAFFE	213	243	6
sBU_3DFE	2500	243	6
Movie	7755	1869	5
Nature_Scene	2000	294	9
fbp5500	5500	512	5
Yeast_heat	2465	24	6
Yeast_diau	2465	24	7
Yeast_cold	2465	24	4
Yeast_dtt	2465	24	4
emotion6	1980	168	7
Twitter	10045	168	8
Flickr	11150	168	8

On the one hand, we apply our proposed $\mathcal{S}$ -LDL to existing methods to demonstrate performance improvements in the LDL task (denoted by the “ $\mathcal{S}$ ” prefix). These methods are BFGS-LLD (KLD) (Geng, 2016), SCL (Jia et al., 2019), LRR (Jia et al., 2023), QFD² (Wen et al., 2023), and CJS (Wen et al., 2023) (the losses of these methods constitute the set $\mathcal{L}_{\text{LDL}}$). On the other hand, we compare $\mathcal{S}$ -LDL with methods that have specialized structure, which our proposed cannot directly adapt to. These methods are CPNN (Geng et al., 2013), LDSVR (Geng and Hou, 2015), AA- k NN (Geng, 2016), LDLFs (Shen et al., 2017), and DF-LDL (denoted by DF-BFGS since we use BFGS-LLDs as base estimators) (González et al., 2021a). Moreover, we apply our proposed to derivative tasks of LDL (i.e., LDL4C, IncomLDL, and LE) and the comparison methods involved are LDL4C (Wang, Jing and Geng, Xin, 2019), HR (Wang and Geng, 2021a), LDLM (Wang and Geng, 2021b), IncomLDL (Xu and Zhou, 2017), LP (Xu et al., 2019), GLLE (Xu et al., 2019), LEVI (Xu et al., 2023), and LIBLE (Zheng et al., 2023).

A.4 PARAMETER SETTINGS AND EXPERIMENTAL ENVIRONMENT

The parameter settings of the proposed $\mathcal{S}$ -LDL and comparison algorithms are summarized in Table 7. Note that DF-LDL is parameter-free, and we use BFGS-LLDs as its base estimators, parameter settings of which are the same as BFGS-LLD as the comparison algorithm. We use Adam (Kingma and Ba, 2015) for the optimization of $\mathcal{S}$ -LDL. For all methods of the deep regime, the learning rate is chosen among $\{1, 2, 5\} \times 10^{\{-4, -3, -2\}}$, and the selection of the number of epochs is nested into

⁴<https://zenodo.org/records/3451524>

⁵<https://github.com/HCIILAB/SCUT-FBP5500-Database-Release>

⁶<https://palm.seu.edu.cn/xgeng/LDL/download.htm>

⁷<https://cv.nankai.edu.cn/projects>

Table 7: Summary of algorithms and parameter settings

Algorithms	Parameter	Value (Range)
AA- k NN	k : # Neighbors	5
LDLFs	# Estimators (trees)	5
	Depth	6
	Latent units (leaves)	64
BFGS-LLD	ε : Convergence criterion	10^{-6}
	Max iteration	600
SCL	m : # Clusters	5
	$\lambda_1, \lambda_2, \lambda_3$: Trade-off	$10^{-3}, 10^{-3}, 0.1$
LRR	λ : Trade-off (ranking loss)	$10^{\{-5, -4, -3, -2, -1\}}$
	β : Trade-off (regularization)	$10^{\{-3, -2, -1, 0, 1, 2\}}$
LDL4C	C_1, C_2 : Balance coefficients	$10^{-2}, 10^{-6}$
	ρ : Margin	10^{-2}
HR	$\lambda_1, \lambda_2, \lambda_3$: Trade-off	$10^{-2}, 10^{-6}$
	ρ : Margin	10^{-2}
LDLM	$\lambda_1, \lambda_2, \lambda_3$: Trade-off	$10^{-6}, 10^{\{-3, -2, -1\}}, 10^{\{-3, -2, -1\}}$
	ρ : Margin	10^{-2}
IncomLDL	ε : Convergence criterion	10^{-6}
	γ : Factor of Lipschitz constant	2
	λ : Trade-off	1
LP	α : Balance coefficient	0.5
GLLE	λ_1, λ_2 : Trade-off	$10^{-2}, 10^{-4}$
	σ : Width parameter for similarity calculation	10
LEVI	λ : Trade-off	1
LIBLE	α, β : Trade-off	$10^{\{-3, -2, -1, 0, 1, 2\}}$
S-LDL	α, λ, T	0.1, 0.2, 10

a ten-fold cross validation. All the results are obtained on a Linux workstation with Intel Core i9 (3.70GHz), NVIDIA GeForce RTX 3090 (24GB), and 32GB memory.

A.5 FULL EXPERIMENTAL RESULTS

Here we provide complete results of all conducted experiments. Tables 8 to 19 are results on the LDL task with different datasets. For IncomLDL, we follow *the incomplete settings* (Xu and Zhou, 2017) and vary the observed rate $\omega\%$ from 20% to 40%. Tables 20 to 21 are results on the IncomLDL task. Tables 22 to 23 are on the LDL4C task. For LE, we follow *the settings of the recovery experiment* (Xu et al., 2019). Tables 24 to 25 show results on the LE task.

Table 8: Experimental results of LDL on the JAFFE dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
LDSVR	.0959 $\pm$.013	.3280 $\pm$.027	.6778 $\pm$.058	.0476 $\pm$.011	.9549 $\pm$.010	.8838 $\pm$.012	.5175 $\pm$.102	.4508 $\pm$.086
AA- k NN	.0978 $\pm$.012	.3483 $\pm$.032	.7164 $\pm$.066	.0527 $\pm$.011	.9497 $\pm$.010	.8766 $\pm$.012	.4111 $\pm$.083	.3514 $\pm$.070
LDLFs	.0940 $\pm$.010	.3637 $\pm$.032	.7355 $\pm$.066	.0550 $\pm$.009	.9494 $\pm$.009	.8766 $\pm$.011	.4364 $\pm$.108	.3749 $\pm$.093
DF-BFGS	.0827 $\pm$.009	.3062 $\pm$.025	.6239 $\pm$.052	.0388 $\pm$.007	.9633 $\pm$.007	.8944 $\pm$.010	.5244 $\pm$.087	.4493 $\pm$.077
KLD $\bullet$	.0925 $\pm$.010	.3608 $\pm$.031	.7363 $\pm$.064	.0508 $\pm$.009	.9538 $\pm$.008	.8777 $\pm$.011	.4572 $\pm$.097	.3873 $\pm$.084
S-KLD	.0818 $\pm$.011	.3007 $\pm$.032	.6132 $\pm$.067	.0395 $\pm$.010	.9625 $\pm$.009	.8960 $\pm$.012	.5461 $\pm$.105	.4769 $\pm$.096
SCL $\bullet$	.0873 $\pm$.008	.3358 $\pm$.024	.6874 $\pm$.051	.0439 $\pm$.006	.9592 $\pm$.006	.8851 $\pm$.009	.4744 $\pm$.092	.4020 $\pm$.080
S-SCL	.0854 $\pm$.010	.3184 $\pm$.025	.6526 $\pm$.053	.0420 $\pm$.008	.9604 $\pm$.008	.8896 $\pm$.010	.5110 $\pm$.095	.4388 $\pm$.084
LRR $\bullet$	.0853 $\pm$.010	.3230 $\pm$.027	.6560 $\pm$.055	.0412 $\pm$.008	.9616 $\pm$.008	.8906 $\pm$.010	.5117 $\pm$.094	.4420 $\pm$.084
S-LRR	.0804 $\pm$.009	.2934 $\pm$.028	.5989 $\pm$.059	.0383 $\pm$.009	.9635 $\pm$.008	.8981 $\pm$.011	.5448 $\pm$.092	.4819 $\pm$.084

Table 9: Experimental results of LDL on the sBU_3DFE dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
LDSVR	.1250 $\pm$.005	.3710 $\pm$.010	.8009 $\pm$.021	.0720 $\pm$.004	.9298 $\pm$.004	.8559 $\pm$.004	.3524 $\pm$.031	.3011 $\pm$.026
AA-kNN	.1272 $\pm$.004	.4001 $\pm$.009	.8281 $\pm$.020	.0801 $\pm$.004	.9217 $\pm$.004	.8488 $\pm$.004	.2053 $\pm$.030	.1767 $\pm$.026
LDLFs	.1016 $\pm$.003	.3262 $\pm$.008	.6841 $\pm$.017	.0504 $\pm$.003	.9499 $\pm$.003	.8776 $\pm$.003	.4212 $\pm$.023	.3620 $\pm$.019
DF-BFGS	.1146 $\pm$.004	.3616 $\pm$.008	.7627 $\pm$.019	.0618 $\pm$.003	.9388 $\pm$.003	.8626 $\pm$.004	.3026 $\pm$.031	.2621 $\pm$.026
KLD $\bullet$	.1147 $\pm$.004	.3697 $\pm$.008	.7804 $\pm$.019	.0624 $\pm$.003	.9387 $\pm$.003	.8604 $\pm$.003	.3021 $\pm$.026	.2643 $\pm$.022
S-KLD	.1014 $\pm$.004	.3203 $\pm$.009	.6736 $\pm$.018	.0514 $\pm$.003	.9487 $\pm$.003	.8789 $\pm$.004	.4334 $\pm$.025	.3729 $\pm$.022
SCL $\bullet$	.1145 $\pm$.004	.3648 $\pm$.008	.7748 $\pm$.018	.0605 $\pm$.003	.9404 $\pm$.003	.8614 $\pm$.003	.3091 $\pm$.026	.2701 $\pm$.021
S-SCL	.1041 $\pm$.004	.3301 $\pm$.009	.6936 $\pm$.019	.0535 $\pm$.003	.9468 $\pm$.003	.8754 $\pm$.004	.3956 $\pm$.030	.3381 $\pm$.027
LRR $\bullet$	.1067 $\pm$.003	.3476 $\pm$.008	.7320 $\pm$.017	.0543 $\pm$.003	.9465 $\pm$.003	.8695 $\pm$.003	.3626 $\pm$.026	.3123 $\pm$.022
S-LRR	.0996 $\pm$.004	.3157 $\pm$.008	.6610 $\pm$.017	.0499 $\pm$.003	.9502 $\pm$.003	.8812 $\pm$.003	.4455 $\pm$.026	.3837 $\pm$.023

Table 10: Experimental results of LDL on the Yeast_heat dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
CPNN	.0419 $\pm$.001	.1818 $\pm$.005	.3633 $\pm$.009	.0125 $\pm$.001	.9881 $\pm$.001	.9404 $\pm$.001	.1507 $\pm$.034	.1221 $\pm$.028
AA-kNN	.0441 $\pm$.001	.1913 $\pm$.005	.3840 $\pm$.010	.0140 $\pm$.001	.9867 $\pm$.001	.9370 $\pm$.002	.1678 $\pm$.031	.1384 $\pm$.026
LDLFs	.0420 $\pm$.001	.1818 $\pm$.005	.3627 $\pm$.009	.0125 $\pm$.001	.9881 $\pm$.001	.9405 $\pm$.001	.1731 $\pm$.032	.1409 $\pm$.026
DF-BFGS	.0420 $\pm$.001	.1816 $\pm$.005	.3624 $\pm$.009	.0125 $\pm$.001	.9881 $\pm$.001	.9405 $\pm$.001	.1964 $\pm$.034	.1624 $\pm$.028
LRR $\bullet$	.0423 $\pm$.001	.1828 $\pm$.005	.3644 $\pm$.009	.0126 $\pm$.001	.9880 $\pm$.001	.9402 $\pm$.001	.1655 $\pm$.033	.1351 $\pm$.028
S-LRR	.0417 $\pm$.001	.1806 $\pm$.005	.3609 $\pm$.009	.0124 $\pm$.001	.9882 $\pm$.001	.9408 $\pm$.001	.1882 $\pm$.034	.1548 $\pm$.028
QFD ² $\bullet$	.0423 $\pm$.001	.1827 $\pm$.005	.3644 $\pm$.009	.0126 $\pm$.001	.9880 $\pm$.001	.9402 $\pm$.001	.1677 $\pm$.032	.1351 $\pm$.027
S-QFD ²	.0417 $\pm$.001	.1808 $\pm$.005	.3611 $\pm$.009	.0124 $\pm$.001	.9882 $\pm$.001	.9408 $\pm$.001	.1880 $\pm$.032	.1544 $\pm$.027
CJS $\bullet$	.0423 $\pm$.001	.1827 $\pm$.005	.3643 $\pm$.009	.0126 $\pm$.001	.9880 $\pm$.001	.9402 $\pm$.001	.1632 $\pm$.032	.1329 $\pm$.027
S-CJS	.0417 $\pm$.001	.1804 $\pm$.005	.3603 $\pm$.009	.0124 $\pm$.001	.9882 $\pm$.001	.9409 $\pm$.001	.1940 $\pm$.030	.1589 $\pm$.025

Table 11: Experimental results of LDL on the Yeast_diau dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
CPNN	.0385 $\pm$.001	.2069 $\pm$.006	.4439 $\pm$.012	.0138 $\pm$.001	.9872 $\pm$.001	.9383 $\pm$.002	.2962 $\pm$.034	.2427 $\pm$.027
AA-kNN	.0385 $\pm$.001	.2085 $\pm$.006	.4487 $\pm$.014	.0145 $\pm$.001	.9867 $\pm$.001	.9377 $\pm$.002	.3674 $\pm$.029	.2976 $\pm$.024
LDLFs	.0371 $\pm$.001	.2014 $\pm$.006	.4324 $\pm$.012	.0132 $\pm$.001	.9879 $\pm$.001	.9401 $\pm$.002	.4088 $\pm$.021	.3254 $\pm$.018
DF-BFGS	.0368 $\pm$.001	.1999 $\pm$.006	.4294 $\pm$.013	.0131 $\pm$.001	.9879 $\pm$.001	.9405 $\pm$.002	.4161 $\pm$.027	.3404 $\pm$.022
LRR $\bullet$	.0370 $\pm$.001	.2007 $\pm$.006	.4307 $\pm$.012	.0131 $\pm$.001	.9879 $\pm$.001	.9403 $\pm$.002	.4154 $\pm$.023	.3343 $\pm$.020
S-LRR	.0366 $\pm$.001	.1983 $\pm$.006	.4257 $\pm$.012	.0129 $\pm$.001	.9881 $\pm$.001	.9410 $\pm$.002	.4198 $\pm$.023	.3389 $\pm$.019
QFD ² $\bullet$	.0369 $\pm$.001	.2000 $\pm$.006	.4296 $\pm$.012	.0131 $\pm$.001	.9879 $\pm$.001	.9404 $\pm$.002	.4118 $\pm$.025	.3326 $\pm$.021
S-QFD ²	.0366 $\pm$.001	.1985 $\pm$.006	.4261 $\pm$.012	.0129 $\pm$.001	.9881 $\pm$.001	.9409 $\pm$.002	.4203 $\pm$.021	.3387 $\pm$.018
CJS	.0367 $\pm$.001	.1989 $\pm$.006	.4272 $\pm$.012	.0130 $\pm$.001	.9880 $\pm$.001	.9408 $\pm$.002	.4164 $\pm$.025	.3366 $\pm$.021
S-CJS	.0366 $\pm$.001	.1984 $\pm$.006	.4260 $\pm$.012	.0130 $\pm$.001	.9881 $\pm$.001	.9409 $\pm$.002	.4198 $\pm$.024	.3392 $\pm$.019

Table 12: Experimental results of LDL on the Yeast_cold dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
CPNN	.0510 $\pm$.002	.1392 $\pm$.005	.2396 $\pm$.008	.0121 $\pm$.001	.9886 $\pm$.001	.9410 $\pm$.002	.2651 $\pm$.036	.2263 $\pm$.032
AA-kNN	.0542 $\pm$.002	.1476 $\pm$.005	.2549 $\pm$.008	.0135 $\pm$.001	.9872 $\pm$.001	.9371 $\pm$.002	.2189 $\pm$.035	.1866 $\pm$.031
LDLFs	.0511 $\pm$.002	.1396 $\pm$.005	.2404 $\pm$.009	.0122 $\pm$.001	.9885 $\pm$.001	.9408 $\pm$.002	.2482 $\pm$.038	.2112 $\pm$.033
DF-BFGS	.0514 $\pm$.002	.1404 $\pm$.005	.2424 $\pm$.008	.0123 $\pm$.001	.9885 $\pm$.001	.9403 $\pm$.002	.2581 $\pm$.036	.2190 $\pm$.030
LRR	.0511 $\pm$.002	.1395 $\pm$.005	.2402 $\pm$.009	.0122 $\pm$.001	.9886 $\pm$.001	.9408 $\pm$.002	.2490 $\pm$.035	.2111 $\pm$.030
S-LRR	.0510 $\pm$.002	.1391 $\pm$.005	.2395 $\pm$.009	.0121 $\pm$.001	.9886 $\pm$.001	.9410 $\pm$.002	.2618 $\pm$.037	.2238 $\pm$.032
QFD ²	.0513 $\pm$.002	.1401 $\pm$.005	.2413 $\pm$.009	.0123 $\pm$.001	.9885 $\pm$.001	.9405 $\pm$.002	.2534 $\pm$.037	.2158 $\pm$.032
S-QFD ²	.0510 $\pm$.002	.1391 $\pm$.005	.2396 $\pm$.008	.0121 $\pm$.001	.9886 $\pm$.001	.9410 $\pm$.002	.2571 $\pm$.039	.2197 $\pm$.033
CJS	.0513 $\pm$.002	.1401 $\pm$.005	.2412 $\pm$.008	.0123 $\pm$.001	.9884 $\pm$.001	.9406 $\pm$.002	.2535 $\pm$.038	.2152 $\pm$.032
S-CJS	.0510 $\pm$.002	.1392 $\pm$.005	.2396 $\pm$.009	.0121 $\pm$.001	.9886 $\pm$.001	.9410 $\pm$.002	.2621 $\pm$.037	.2241 $\pm$.031

Table 13: Experimental results of LDL on the Yeast_dtt dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
CPNN	.0361 $\pm$.001	.0984 $\pm$.004	.1690 $\pm$.006	.0063 $\pm$.001	.9941 $\pm$.000	.9583 $\pm$.001	.1735 $\pm$.035	.1494 $\pm$.030
AA-kNN	.0386 $\pm$.001	.1047 $\pm$.004	.1797 $\pm$.006	.0071 $\pm$.001	.9933 $\pm$.000	.9556 $\pm$.001	.1591 $\pm$.033	.1399 $\pm$.030
LDLFs	.0360 $\pm$.001	.0981 $\pm$.004	.1689 $\pm$.006	.0063 $\pm$.001	.9941 $\pm$.000	.9583 $\pm$.001	.1986 $\pm$.038	.1727 $\pm$.034
DF-BFGS	.0365 $\pm$.001	.0995 $\pm$.004	.1712 $\pm$.006	.0064 $\pm$.001	.9939 $\pm$.000	.9578 $\pm$.001	.1804 $\pm$.033	.1592 $\pm$.030
LRR	.0360 $\pm$.001	.0982 $\pm$.004	.1690 $\pm$.006	.0063 $\pm$.001	.9941 $\pm$.000	.9583 $\pm$.001	.2016 $\pm$.037	.1738 $\pm$.032
S-LRR	.0359 $\pm$.001	.0977 $\pm$.004	.1680 $\pm$.006	.0062 $\pm$.001	.9941 $\pm$.000	.9585 $\pm$.001	.2068 $\pm$.035	.1811 $\pm$.031
QFD ²	.0362 $\pm$.001	.0986 $\pm$.004	.1696 $\pm$.006	.0063 $\pm$.001	.9940 $\pm$.000	.9582 $\pm$.001	.1917 $\pm$.035	.1665 $\pm$.031
S-QFD ²	.0359 $\pm$.001	.0977 $\pm$.004	.1681 $\pm$.006	.0062 $\pm$.001	.9941 $\pm$.000	.9585 $\pm$.001	.2086 $\pm$.036	.1822 $\pm$.032
CJS	.0361 $\pm$.001	.0984 $\pm$.004	.1692 $\pm$.006	.0063 $\pm$.001	.9941 $\pm$.000	.9582 $\pm$.001	.1975 $\pm$.040	.1722 $\pm$.035
S-CJS	.0359 $\pm$.001	.0978 $\pm$.004	.1682 $\pm$.006	.0062 $\pm$.001	.9941 $\pm$.000	.9585 $\pm$.001	.2080 $\pm$.035	.1804 $\pm$.031

Table 14: Experimental results of LDL on the emotion6 dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
LDSVR	.3152 $\pm$.010	1.8217 $\pm$.020	4.1452 $\pm$.064	1.0744 $\pm$.081	.6906 $\pm$.015	.5773 $\pm$.012	.3915 $\pm$.030	.3235 $\pm$.025
AA-kNN	.3288 $\pm$.011	1.7116 $\pm$.026	3.8757 $\pm$.076	.9512 $\pm$.115	.6632 $\pm$.013	.5564 $\pm$.010	.2920 $\pm$.027	.2401 $\pm$.022
LDLFs	.3120 $\pm$.010	1.6625 $\pm$.026	3.7330 $\pm$.075	.5871 $\pm$.024	.7143 $\pm$.011	.5802 $\pm$.010	.3631 $\pm$.029	.3025 $\pm$.024
DF-BFGS	.3026 $\pm$.010	1.6765 $\pm$.025	3.7675 $\pm$.071	.5805 $\pm$.026	.7206 $\pm$.013	.5909 $\pm$.010	.3940 $\pm$.027	.3256 $\pm$.022
KLD $\bullet$	.3037 $\pm$.010	1.6774 $\pm$.025	3.7729 $\pm$.074	.5863 $\pm$.027	.7191 $\pm$.013	.5897 $\pm$.011	.3959 $\pm$.028	.3259 $\pm$.023
S-KLD	.3024 $\pm$.010	1.6548 $\pm$.026	3.6984 $\pm$.075	.5631 $\pm$.024	.7282 $\pm$.012	.5926 $\pm$.010	.4063 $\pm$.027	.3361 $\pm$.023
SCL $\bullet$	.3020 $\pm$.010	1.6750 $\pm$.025	3.7642 $\pm$.073	.5803 $\pm$.027	.7219 $\pm$.013	.5917 $\pm$.011	.4003 $\pm$.028	.3299 $\pm$.023
S-SCL	.3018 $\pm$.010	1.6554 $\pm$.026	3.6993 $\pm$.076	.5631 $\pm$.025	.7281 $\pm$.012	.5936 $\pm$.010	.4089 $\pm$.027	.3383 $\pm$.023
LRR $\bullet$	.3030 $\pm$.010	1.6736 $\pm$.025	3.7601 $\pm$.073	.5804 $\pm$.026	.7212 $\pm$.013	.5899 $\pm$.010	.3941 $\pm$.027	.3243 $\pm$.023
S-LRR	.3028 $\pm$.009	1.6524 $\pm$.026	3.6923 $\pm$.074	.5607 $\pm$.023	.7299 $\pm$.011	.5923 $\pm$.010	.4078 $\pm$.027	.3373 $\pm$.022

Table 15: Experimental results of LDL on the Twitter dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
LDSVR	.4236 $\pm$.008	2.6722 $\pm$.002	7.3015 $\pm$.009	5.0018 $\pm$.115	.7627 $\pm$.008	.5761 $\pm$.008	.5237 $\pm$.008	.4246 $\pm$.007
AA-kNN	.3172 $\pm$.004	2.0142 $\pm$.012	4.5597 $\pm$.043	3.1429 $\pm$.148	.7926 $\pm$.006	.6024 $\pm$.005	.5014 $\pm$.009	.4432 $\pm$.008
LDLFs	.4035 $\pm$.014	2.5461 $\pm$.010	6.8269 $\pm$.040	1.6884 $\pm$.115	.6756 $\pm$.018	.5318 $\pm$.013	.4164 $\pm$.013	.3349 $\pm$.010
DF-BFGS	.2982 $\pm$.004	2.4025 $\pm$.005	6.2416 $\pm$.020	.6304 $\pm$.012	.8250 $\pm$.006	.6220 $\pm$.004	.5467 $\pm$.008	.4454 $\pm$.007
KLD $\circ$	.2966 $\pm$.004	2.4059 $\pm$.005	6.2558 $\pm$.020	.6307 $\pm$.013	.8243 $\pm$.006	.6249 $\pm$.005	.5470 $\pm$.008	.4456 $\pm$.007
S-KLD	.2995 $\pm$.005	2.4112 $\pm$.005	6.2883 $\pm$.020	.6491 $\pm$.013	.8203 $\pm$.006	.6205 $\pm$.005	.5384 $\pm$.009	.4385 $\pm$.008
SCL $\bullet$	.2977 $\pm$.004	2.4028 $\pm$.005	6.2435 $\pm$.021	.6262 $\pm$.013	.8256 $\pm$.006	.6233 $\pm$.005	.5488 $\pm$.008	.4471 $\pm$.007
S-SCL	.2940 $\pm$.005	2.4059 $\pm$.006	6.2589 $\pm$.023	.6203 $\pm$.013	.8268 $\pm$.006	.6281 $\pm$.006	.5518 $\pm$.008	.4497 $\pm$.007
LRR $\bullet$	.2984 $\pm$.004	2.4046 $\pm$.005	6.2525 $\pm$.019	.6351 $\pm$.012	.8232 $\pm$.006	.6220 $\pm$.004	.5443 $\pm$.008	.4636 $\pm$.007
S-LRR	.2937 $\pm$.004	2.4056 $\pm$.005	6.2589 $\pm$.019	.6189 $\pm$.013	.8271 $\pm$.006	.6283 $\pm$.005	.5519 $\pm$.008	.4498 $\pm$.007

Table 16: Experimental results of LDL on the Flickr dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
LDSVR	.5174 $\pm$.006	2.6364 $\pm$.002	7.2094 $\pm$.011	5.0366 $\pm$.086	.6636 $\pm$.008	.4683 $\pm$.006	.4622 $\pm$.009	.3811 $\pm$.008
AA-kNN	.3286 $\pm$.005	2.0685 $\pm$.009	4.9363 $\pm$.033	2.2172 $\pm$.107	.7200 $\pm$.006	.5582 $\pm$.005	.4265 $\pm$.009	.3465 $\pm$.007
LDLFs	.4051 $\pm$.011	2.4012 $\pm$.012	6.3262 $\pm$.050	1.4274 $\pm$.077	.6073 $\pm$.015	.4822 $\pm$.011	.3478 $\pm$.014	.2847 $\pm$.012
DF-BFGS	.3007 $\pm$.005	2.1995 $\pm$.007	5.4900 $\pm$.025	.6309 $\pm$.011	.7801 $\pm$.005	.5979 $\pm$.004	.5102 $\pm$.009	.4226 $\pm$.008
KLD $\circ$	.3015 $\pm$.005	2.2008 $\pm$.007	5.4969 $\pm$.025	.6348 $\pm$.012	.7787 $\pm$.005	.5973 $\pm$.004	.5113 $\pm$.009	.4234 $\pm$.008
S-KLD	.3052 $\pm$.005	2.2044 $\pm$.007	5.5222 $\pm$.026	.6485 $\pm$.012	.7720 $\pm$.005	.5926 $\pm$.004	.5030 $\pm$.009	.4166 $\pm$.008
SCL $\bullet$	.3280 $\pm$.013	2.2986 $\pm$.024	5.9247 $\pm$.099	.8301 $\pm$.055	.7268 $\pm$.018	.5713 $\pm$.015	.4566 $\pm$.024	.3748 $\pm$.022
S-SCL	.2929 $\pm$.005	2.2045 $\pm$.007	5.5289 $\pm$.028	.6113 $\pm$.012	.7862 $\pm$.005	.6070 $\pm$.005	.5265 $\pm$.008	.4373 $\pm$.007
LRR $\bullet$	.3057 $\pm$.005	2.1969 $\pm$.007	5.4763 $\pm$.025	.6431 $\pm$.012	.7752 $\pm$.006	.5929 $\pm$.004	.5047 $\pm$.009	.4229 $\pm$.008
S-LRR	.2938 $\pm$.005	2.2013 $\pm$.006	5.5138 $\pm$.023	.6105 $\pm$.012	.7864 $\pm$.005	.6058 $\pm$.005	.5261 $\pm$.009	.4369 $\pm$.008

Table 17: Experimental results of LDL on the Natural_Scene dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
LDSVR	.4899 $\pm$.016	2.0831 $\pm$.025	5.7724 $\pm$.092	2.0862 $\pm$.085	.5740 $\pm$.017	.4430 $\pm$.015	.4997 $\pm$.015	.3695 $\pm$.012
AA-kNN	.3113 $\pm$.014	1.9066 $\pm$.034	4.5413 $\pm$.110	1.0874 $\pm$.082	.7113 $\pm$.015	.5636 $\pm$.013	.4921 $\pm$.021	.3518 $\pm$.016
LDLFs	.2808 $\pm$.034	2.4329 $\pm$.024	6.6027 $\pm$.108	.6464 $\pm$.118	.7679 $\pm$.046	.5839 $\pm$.043	.5406 $\pm$.058	.4072 $\pm$.045
DF-BFGS	.3074 $\pm$.013	2.4126 $\pm$.017	6.5896 $\pm$.072	.7603 $\pm$.033	.7381 $\pm$.013	.5568 $\pm$.011	.5110 $\pm$.017	.3837 $\pm$.013
KLD $\bullet$	.3201 $\pm$.013	2.4242 $\pm$.017	6.6560 $\pm$.070	.8285 $\pm$.044	.7172 $\pm$.015	.5485 $\pm$.011	.4958 $\pm$.016	.3715 $\pm$.012
S-KLD	.2743 $\pm$.013	2.3866 $\pm$.020	6.4733 $\pm$.077	.6608 $\pm$.039	.7751 $\pm$.014	.6133 $\pm$.012	.5592 $\pm$.017	.4221 $\pm$.014
SCL $\bullet$	.3379 $\pm$.014	2.4800 $\pm$.018	6.8659 $\pm$.076	.8867 $\pm$.035	.7014 $\pm$.014	.4801 $\pm$.014	.4109 $\pm$.018	.3025 $\pm$.013
S-SCL	.2733 $\pm$.013	2.3734 $\pm$.018	6.4376 $\pm$.072	.6703 $\pm$.043	.7744 $\pm$.015	.6156 $\pm$.013	.5573 $\pm$.018	.4207 $\pm$.014
LRR $\bullet$	.3138 $\pm$.013	2.4469 $\pm$.018	6.7118 $\pm$.074	.7703 $\pm$.032	.7363 $\pm$.013	.5456 $\pm$.011	.5056 $\pm$.016	.3782 $\pm$.012
S-LRR	.2740 $\pm$.018	2.3461 $\pm$.023	6.3467 $\pm$.087	.6867 $\pm$.070	.7715 $\pm$.021	.6199 $\pm$.017	.5595 $\pm$.023	.4228 $\pm$.018

Table 18: Experimental results of LDL on the Movie dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
CPNN	.1337 $\pm$.003	.5639 $\pm$.010	1.0746 $\pm$.020	.1191 $\pm$.005	.9194 $\pm$.003	.8164 $\pm$.004	.6610 $\pm$.013	.7080 $\pm$.002
AA-kNN	.1223 $\pm$.002	.5451 $\pm$.009	1.0445 $\pm$.018	.1129 $\pm$.004	.9254 $\pm$.003	.8250 $\pm$.003	.6557 $\pm$.011	.5710 $\pm$.010
LDLFs	.1172 $\pm$.003	.5233 $\pm$.013	1.0134 $\pm$.026	.1086 $\pm$.006	.9305 $\pm$.003	.8324 $\pm$.004	.6929 $\pm$.013	.6051 $\pm$.012
DF-BFGS	.1210 $\pm$.002	.5282 $\pm$.009	1.0158 $\pm$.019	.1084 $\pm$.005	.9289 $\pm$.003	.8301 $\pm$.003	.6848 $\pm$.012	.5963 $\pm$.012
LRR $\bullet$	.1135 $\pm$.002	.5101 $\pm$.009	.9770 $\pm$.018	.0957 $\pm$.004	.9369 $\pm$.002	.8385 $\pm$.003	.7119 $\pm$.011	.6203 $\pm$.011
S-LRR	.1125 $\pm$.002	.5086 $\pm$.009	.9717 $\pm$.018	.0945 $\pm$.004	.9376 $\pm$.002	.8398 $\pm$.003	.7126 $\pm$.011	.6227 $\pm$.011
QFD ² $\bullet$	.1159 $\pm$.002	.5200 $\pm$.009	.9920 $\pm$.018	.0975 $\pm$.004	.9355 $\pm$.002	.8357 $\pm$.003	.7075 $\pm$.011	.6158 $\pm$.011
S-QFD ²	.1123 $\pm$.002	.5073 $\pm$.009	.9700 $\pm$.018	.0945 $\pm$.004	.9376 $\pm$.002	.8401 $\pm$.003	.7125 $\pm$.011	.6224 $\pm$.011
CJS $\bullet$	.1153 $\pm$.002	.5127 $\pm$.009	.9845 $\pm$.019	.0984 $\pm$.004	.9352 $\pm$.002	.8368 $\pm$.003	.7103 $\pm$.012	.6178 $\pm$.011
S-CJS	.1123 $\pm$.002	.5072 $\pm$.009	.9699 $\pm$.018	.0945 $\pm$.004	.9376 $\pm$.002	.8401 $\pm$.003	.7125 $\pm$.011	.6223 $\pm$.011

Table 19: Experimental results of LDL on the fbp5500 dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
CPNN	.1864 $\pm$.005	1.3367 $\pm$.009	2.3604 $\pm$.020	.1664 $\pm$.005	.9281 $\pm$.004	.7958 $\pm$.005	.8688 $\pm$.005	.7831 $\pm$.007
AA-kNN	.1515 $\pm$.004	1.0443 $\pm$.015	1.7295 $\pm$.031	.1846 $\pm$.016	.9419 $\pm$.004	.8317 $\pm$.005	.8865 $\pm$.006	.8123 $\pm$.008
LDLFs	.1307 $\pm$.003	1.2787 $\pm$.010	2.1703 $\pm$.024	.1002 $\pm$.005	.9575 $\pm$.003	.8552 $\pm$.004	.9060 $\pm$.005	.8352 $\pm$.007
DF-BFGS	.1341 $\pm$.003	1.2889 $\pm$.010	2.1982 $\pm$.023	.1050 $\pm$.005	.9551 $\pm$.003	.8523 $\pm$.004	.9047 $\pm$.005	.8337 $\pm$.007
LRR	.1312 $\pm$.003	1.2767 $\pm$.010	2.1655 $\pm$.024	.1004 $\pm$.004	.9575 $\pm$.002	.8547 $\pm$.003	.9059 $\pm$.004	.8350 $\pm$.006
S-LRR	.1302 $\pm$.003	1.2796 $\pm$.010	2.1717 $\pm$.024	.0997 $\pm$.005	.9576 $\pm$.002	.8558 $\pm$.003	.9063 $\pm$.004	.8425 $\pm$.006
QFD ² $\bullet$	.1380 $\pm$.003	1.2803 $\pm$.010	2.1858 $\pm$.024	.1084 $\pm$.005	.9535 $\pm$.003	.8476 $\pm$.004	.9021 $\pm$.004	.8297 $\pm$.006
S-QFD ²	.1321 $\pm$.004	1.2811 $\pm$.010	2.1779 $\pm$.024	.1027 $\pm$.006	.9561 $\pm$.003	.8537 $\pm$.004	.9044 $\pm$.005	.8330 $\pm$.007
CJS $\bullet$	.1343 $\pm$.003	1.3057 $\pm$.010	2.2374 $\pm$.024	.1084 $\pm$.005	.9544 $\pm$.003	.8527 $\pm$.004	.9044 $\pm$.004	.8334 $\pm$.006
S-CJS	.1302 $\pm$.003	1.2802 $\pm$.010	2.1731 $\pm$.024	.0997 $\pm$.005	.9575 $\pm$.002	.8559 $\pm$.003	.9066 $\pm$.004	.8429 $\pm$.007

Table 20: Experimental results of IncomLDL on the JAFFE dataset formatted as (mean $\pm$ std)

Algorithms	$\omega = 20\%$							
	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
IncomLDL $\bullet$	.0898 $\pm$.010	.3304 $\pm$.024	.6742 $\pm$.049	.0425 $\pm$.007	.9598 $\pm$.007	.8861 $\pm$.009	.4742 $\pm$.094	.4017 $\pm$.082
S-IncomLDL	.0863 $\pm$.013	.3179 $\pm$.036	.6525 $\pm$.077	.0433 $\pm$.012	.9590 $\pm$.011	.8893 $\pm$.014	.5034 $\pm$.114	.4401 $\pm$.100
Algorithms	$\omega = 40\%$							
	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$	Ken. $\uparrow$
IncomLDL $\bullet$	.0946 $\pm$.010	.3454 $\pm$.026	.7073 $\pm$.053	.0465 $\pm$.007	.9558 $\pm$.007	.8801 $\pm$.010	.4231 $\pm$.086	.3534 $\pm$.074
S-IncomLDL	.0868 $\pm$.013	.3211 $\pm$.038	.6568 $\pm$.081	.0439 $\pm$.014	.9585 $\pm$.012	.8886 $\pm$.015	.5075 $\pm$.115	.4434 $\pm$.098

Table 21: Experimental results of IncomLDL on the SBU_3DFE dataset formatted as (mean $\pm$ std)

Algorithms	$\omega = 20\%$						
	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$
IncomLDL •	.1088 $\pm$.003	.3586 $\pm$.008	.7586 $\pm$.017	.0574 $\pm$.003	.9439 $\pm$.003	.8655 $\pm$.003	.3171 $\pm$.027
<i>S</i> -IncomLDL	.1014 $\pm$.004	.3208 $\pm$.009	.6727 $\pm$.019	.0516 $\pm$.003	.9485 $\pm$.003	.8790 $\pm$.004	.4255 $\pm$.026

Algorithms	$\omega = 40\%$						
	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$	Spear. $\uparrow$
IncomLDL •	.1104 $\pm$.003	.3621 $\pm$.008	.7673 $\pm$.017	.0586 $\pm$.003	.9426 $\pm$.003	.8638 $\pm$.003	.3003 $\pm$.024
<i>S</i> -IncomLDL	.1016 $\pm$.004	.3213 $\pm$.009	.6748 $\pm$.019	.0516 $\pm$.003	.9485 $\pm$.003	.8786 $\pm$.004	.4232 $\pm$.025

Table 22: Experimental results of LDL4C on JAFFE and Twitter formatted as (mean $\pm$ std)

Algorithms	JAFFE		Algorithms	Twitter	
	0/1 loss $\downarrow$	Err. prob. $\downarrow$		0/1 loss $\downarrow$	Err. prob. $\downarrow$
LDL4C •	.4973 $\pm$.108	.7665 $\pm$.020	LDL4C	.9081 $\pm$.009	.8846 $\pm$.005
<i>S</i> -LDL4C	.4453 $\pm$.102	.7600 $\pm$.019	<i>S</i> -LDL4C	.8714 $\pm$.207	.8729 $\pm$.156
LDL-HR	.4786 $\pm$.097	.7676 $\pm$.020	LDL-HR •	.3656 $\pm$.017	.4928 $\pm$.011
<i>S</i> -HR	.4653 $\pm$.105	.7655 $\pm$.019	<i>S</i> -HR	.2753 $\pm$.013	.4250 $\pm$.008
LDLM	.4787 $\pm$.109	.7687 $\pm$.021	LDLM •	.2814 $\pm$.013	.4291 $\pm$.008
<i>S</i> -LDLM	.4737 $\pm$.097	.7689 $\pm$.019	<i>S</i> -LDLM	.2753 $\pm$.014	.4250 $\pm$.008

Table 23: Experimental results of LDL4C on sBU_3DFE and Flickr formatted as (mean $\pm$ std)

Algorithms	sBU_3DFE		Algorithms	Flickr	
	0/1 loss $\downarrow$	Err. prob. $\downarrow$		0/1 loss $\downarrow$	Err. prob. $\downarrow$
LDL4C	.5578 $\pm$.028	.7671 $\pm$.007	LDL4C	.8971 $\pm$.008	.8884 $\pm$.004
<i>S</i> -LDL4C	.5526 $\pm$.025	.7686 $\pm$.006	<i>S</i> -LDL4C	.8705 $\pm$.138	.8702 $\pm$.100
LDL-HR •	.5167 $\pm$.027	.7596 $\pm$.006	LDL-HR •	.4513 $\pm$.015	.5823 $\pm$.007
<i>S</i> -HR	.5069 $\pm$.025	.7598 $\pm$.006	<i>S</i> -HR	.4219 $\pm$.015	.5639 $\pm$.007
LDLM •	.5258 $\pm$.034	.7619 $\pm$.009	LDLM •	.4384 $\pm$.014	.5740 $\pm$.007
<i>S</i> -LDLM	.4809 $\pm$.024	.7524 $\pm$.005	<i>S</i> -LDLM	.4321 $\pm$.016	.5667 $\pm$.007

Table 24: Experimental results of LE on the JAFFE dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$
LP	.0812 $\pm$.001	.3446 $\pm$.002	.7125 $\pm$.005	.0424 $\pm$.001	.9618 $\pm$.001	.8808 $\pm$.001
GLLE	.0821 $\pm$.002	.3196 $\pm$.013	.6518 $\pm$.028	.0386 $\pm$.003	.9638 $\pm$.002	.8901 $\pm$.004
LEVI	.0787 $\pm$.003	.3316 $\pm$.013	.6864 $\pm$.028	.0391 $\pm$.003	.9649 $\pm$.002	.8860 $\pm$.004
LIBLE •	.0813 $\pm$.006	.3106 $\pm$.020	.6358 $\pm$.044	.0370 $\pm$.005	.9652 $\pm$.005	.8929 $\pm$.008
<i>S</i> -LIBLE	.0770 $\pm$.003	.2942 $\pm$.007	.5997 $\pm$.016	.0332 $\pm$.002	.9685 $\pm$.002	.8987 $\pm$.003

Table 25: Experimental results of LE on the Yeast_heat dataset formatted as (mean $\pm$ std)

Algorithms	Cheby. $\downarrow$	Clark $\downarrow$	Can. $\downarrow$	KLD $\downarrow$	Cosine $\uparrow$	Int. $\uparrow$
LP	.0421 $\pm$.000	.2148 $\pm$.000	.4711 $\pm$.001	.0153 $\pm$.000	.9860 $\pm$.000	.9235 $\pm$.000
GLLE	.0481 $\pm$.001	.2114 $\pm$.005	.4282 $\pm$.011	.0168 $\pm$.001	.9842 $\pm$.001	.9298 $\pm$.002
LEVI	.0494 $\pm$.007	.2125 $\pm$.027	.4307 $\pm$.056	.0169 $\pm$.004	.9838 $\pm$.004	.9289 $\pm$.009
LIBLE •	.0453 $\pm$.000	.1973 $\pm$.001	.3982 $\pm$.003	.0148 $\pm$.000	.9859 $\pm$.000	.9346 $\pm$.000
<i>S</i> -LIBLE	.0445 $\pm$.000	.1901 $\pm$.002	.3790 $\pm$.005	.0137 $\pm$.000	.9869 $\pm$.000	.9376 $\pm$.001