

Feature-Induced Labeling Information Enrichment for Multi-Label Learning

Qian-Wen Zhang , Yun Zhong, Min-Ling Zhang

AAAI-2018

Contents



- Background
- Introduction
- Approach
- Experiments
- Conclusion

Background



Multi-class Learning



Classifier

$$\begin{bmatrix} 1 \\ -1 \\ -1 \\ -1 \\ -1 \\ \vdots \end{bmatrix} \begin{matrix} \text{horse} \\ \text{grassland} \\ \text{sky} \\ \text{sea} \\ \text{cow} \\ \dots \end{matrix}$$

Multi-label Learning



Classifier

$$\begin{bmatrix} 1 \\ 1 \\ 1 \\ -1 \\ -1 \\ \vdots \end{bmatrix} \begin{matrix} \text{horse} \\ \text{grassland} \\ \text{sky} \\ \text{sea} \\ \text{cow} \\ \dots \end{matrix}$$

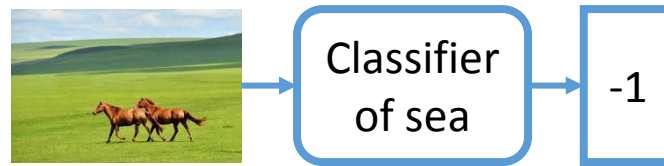
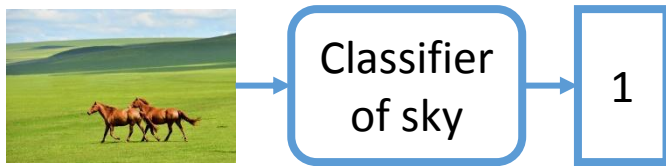
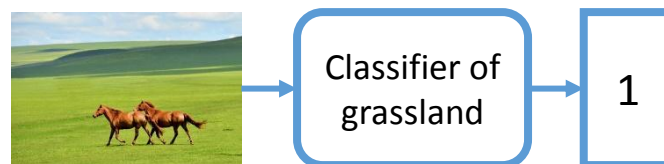
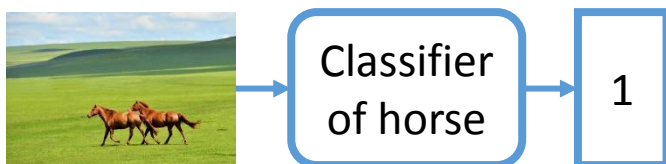


Background

Formalization

- Feature space: d-dimensional $\mathcal{X} = \mathbb{R}^d$
- Label space: q class labels $\mathcal{Y} = \{y_1, y_2, \dots, y_q\}$
- Training set: p instances $\mathcal{D} = \{(\mathbf{x}_i, Y_i) \mid 1 \leq i \leq p\}$
- Task: learn a predictive model $h : \mathcal{X} \rightarrow 2^{\mathcal{Y}}$

Cross Training



..... $[1 \ 1 \ 1 \ -1 \ \dots]^T$

Ignoring the co-existence of other labels.

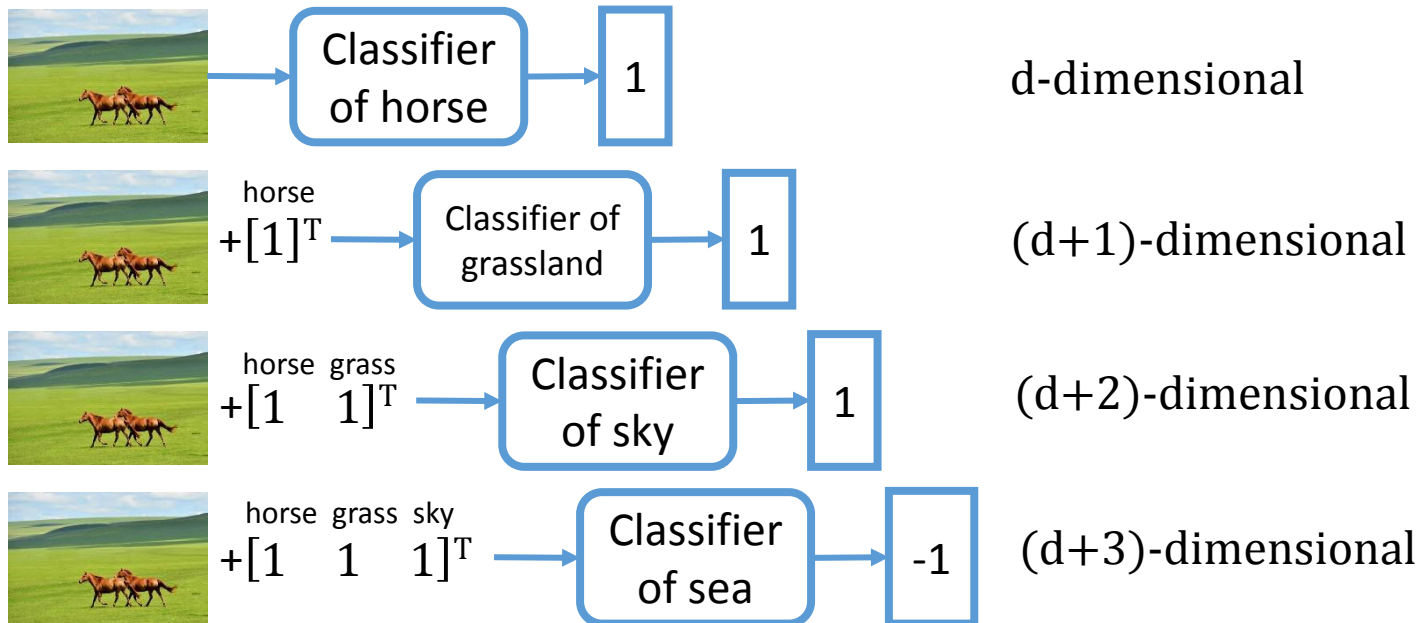
Background



Label Correlation

- If we know the image already has label **horse**, the probability of it having label **grassland** would be **high**.

Classifier Chain



Contents



- Background
- **Introduction**
- Approach
- Experiments
- Conclusion

Introduction



Motivation

- Categorical labeling information is actually a **simplification of the rich semantics** encoded by multi-label training examples.



20% sky, 20% horses and 60% grassland

特朗普就职典礼人数少过奥巴马 摄影师被要求P图

2018-09-08 19:08:51 来源: 观察网(上海)

△ 辛报

103

易信

微信

QQ空间

微博

更多

(原标题: 英媒: 特朗普干预就职典礼照片 摄影师被要求修图展现人多)

【文/观察网】特朗普去年上任伊始, 就和媒体在就职典礼人数上吵了几个回合。近两年过去, 这个问题如今又被提了出来。

英国《卫报》6日独家消息称, 特朗普当时因不满参与自己就职典礼的人数看起来比奥巴马少, 和时任白宫发言人肖恩·斯派塞 (Sean Spicer) 亲自致电照片提供方进行干预, 结果摄影师只能重新修图, 裁去空白部分, 让人群看起来多一些。

《卫报》根据《信息自由法案》从美国内政部获得的记录文件显示, 特朗普就职后的第一天, 即2017年1月21日清晨, 他就与发布就职典礼照片的美国国家公园管理局 (NPS) 的代理主管迈克尔·雷诺兹 (Michael Reynolds) 通了电话。

80% Trump, 20% Obama

- MLFE: **Enrich the labeling information** to induce multi-label predictive model with **strong generalization performance**.

Introduction



Basic Strategy

Leveraging the structural information in the feature space.

- The **underlying structure** of feature space is characterized by the **sparse reconstruction relationships** among training examples.
- The reconstruction information is utilized to guide the enrichment process of turning categorical labeling information into **numerical labeling information**.
- The desired **multi-label predictive model** is learned from training examples with enriched labeling information based on tailored multivariate regression techniques.



Contents

- Background
- Introduction
- Approach
 - Structural Information Discovery
 - Labeling Information Enrichment
 - Predictive Model Induction
- Experiments
- Conclusion

Approach



Intuition

- To characterize the underlying structure of feature space, MLFE works by **modeling the relationship between one example and all the other examples** via **sparse reconstruction**.



Approach



Structural Information Discovery

- Feature space: d-dimensional $\mathcal{X} = \mathbb{R}^d$
- Label space: q class labels $\mathcal{Y} = \{y_1, y_2, \dots, y_q\}$
- Training set: p instances $\mathcal{D} = \{(\mathbf{x}_i, Y_i) \mid 1 \leq i \leq p\}$

- To characterize the underlying structure of feature space, MLFE works by **modeling the relationship between one example and all the other examples** via **sparse reconstruction**.

$\mathbf{A}_i = [\mathbf{x}_1, \dots, \mathbf{x}_{i-1}, \mathbf{x}_{i+1}, \dots, \mathbf{x}_p]$ $d \times (p-1)$ matrix
all training instances other than $\mathbf{x}_i$

$\mathbf{v}_i = [w_{1i}, \dots, w_{i-1,i}, w_{i+1,i}, \dots, w_{pi}]^\top$ $(p-1)$ -dimensional
reconstruction coefficients

- The coefficient vector $\mathbf{v}_i$ is learned by solving the following optimization problem:

$$\min_{\mathbf{v}_i} \underbrace{\|\mathbf{A}_i \mathbf{v}_i - \mathbf{x}_i\|_2^2}_{\text{linear reconstruction error}} + \underbrace{\lambda \|\mathbf{v}_i\|_1}_{\text{sparsity of the co-efficients}} \quad \text{Solved by ADMM}$$



Contents

- Background
- Introduction
- Approach
 - Structural Information Discovery
 - Labeling Information Enrichment
 - Predictive Model Induction
- Experiments
- Conclusion



Approach

Labeling Information Enrichment

- binary labeling vector $\mathbf{t}_i = (t_{i1}, t_{i2}, \dots, t_{iq})^\top$ $\mathbf{t}_i \in \{1, -1\}^q$



- numerical labeling vector $\mathbf{u}_i = (u_{i1}, u_{i2}, \dots, u_{iq})^\top \in \mathbb{R}^q$

$$\min_{\mathbf{v}_i} \underbrace{\|\mathbf{A}_i \mathbf{v}_i - \mathbf{x}_i\|_2^2}_{\text{linear reconstruction error}} + \lambda \underbrace{\|\mathbf{v}_i\|_1}_{\text{sparsity of the co-efficients}}$$

- reconstruction error over the training set $E(\mathbf{W}) = \sum_{i=1}^p \left\| \mathbf{x}_i - \sum_{j=1}^p w_{ji} \mathbf{x}_j \right\|_2^2$

Suppose that the structural relationship specified in the feature space also holds in the output space.

- reconstruction error in the label space $\Phi(\mathbf{U}) = \sum_{i=1}^p \left\| \mathbf{u}_i - \sum_{j=1}^p w_{ji} \mathbf{u}_j \right\|_2^2$
 $\mathbf{U} = [\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_p]$

Approach



Labeling Information Enrichment

- *Suppose that the structural relationship specified in the feature space also holds in the output space.*

$$\Phi(\mathbf{U}) = \sum_{i=1}^p \|\mathbf{u}_i - \sum_{j=1}^p w_{ji} \mathbf{u}_j\|_2^2 \quad \mathbf{U} = [\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_p]$$

- The enriched labeling information is generated by solving the following optimization problem:

$$\min_{\mathbf{U}} \sum_{i=1}^p \|\mathbf{u}_i - \sum_{j=1}^p w_{ji} \mathbf{u}_j\|_2^2 \quad \text{A standard QP problem}$$

$$\text{s.t. } c_1 \leq t_{ij} u_{ij} \leq c_2 \quad (1 \leq i \leq p, 1 \leq j \leq q)$$

- The constraint ensures that the numerical label possesses the same sign with the binary label and takes value with reasonable magnitude.



Contents

- Background
- Introduction
- **Approach**
 - Structural Information Discovery
 - Labeling Information Enrichment
 - **Predictive Model Induction**
- Experiments
- Conclusion

Approach



Predictive Model Induction

- The original multi-label training set can be transformed into its enriched version $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, \mathbf{u}_i) \mid 1 \leq i \leq p\}$
- Induce the predictive model by employing **multi-output regression techniques**.

$$\Omega(\Theta, \mathbf{b}) = \frac{1}{2} \sum_{j=1}^q \|\boldsymbol{\theta}_j\|_2^2 + \beta_1 \sum_{i=1}^p \Omega_1(u_i) + \beta_2 \sum_{i=1}^p \sum_{j=1}^q \Omega_2(o_{ij}) + \beta_3 \sum_{i=1}^p \sum_{j=1}^q \Omega_3(r_{ij})$$

To minimize the objective function $\Omega(\Theta, \mathbf{b})$, MLFE employs the quasi-Newton iterative method named Iterative Re-Weighted Least Square (IRWLS)

$$\Omega_1(u_i) = \begin{cases} 0, & u_i < \epsilon \\ (u_i - \epsilon)^2, & u_i \geq \epsilon \end{cases}$$

ϵ -insensitive cost

$$\Omega_2(o_{ij}) = \begin{cases} 0, & o_{ij} > 0 \\ -o_{ij}, & o_{ij} \leq 0 \end{cases}$$

penalize the case where the sign of predictive output is different to that of original binary label

$$o_{ij} = t_{ij} (\boldsymbol{\theta}_j^\top \boldsymbol{\varphi}(\mathbf{x}_i) + b_j)$$

$$\Omega_3(r_{ij}) = \begin{cases} 1, & r_{ij} > 0 \\ 0, & r_{ij} \leq 0 \end{cases}$$

penalize the case where the predictive model yields large number of relevant labels for the training example

$$r_{ij} = \boldsymbol{\theta}_j^\top \boldsymbol{\varphi}(\mathbf{x}_i) + b_j.$$

Contents



- Background
- Introduction
- Approach
- **Experiments**
- Conclusion

Experiments



Table 2: Predictive performance of each comparing algorithm (mean $\pm$ std. deviation) on the regular-scale data sets.

Comparing algorithms	One-error $\downarrow$							
	cal500	emotions	medical	llog	msra	image	scene	yeast
MLFE	0.129 $\pm$ 0.047	0.260 $\pm$ 0.030	0.113$\pm$0.041	0.669$\pm$0.044	0.040$\pm$0.008	0.258$\pm$0.020	0.149$\pm$0.023	0.231 $\pm$ 0.015
LIFT	0.116$\pm$0.040	0.255$\pm$0.037	0.152 $\pm$ 0.039	0.705 $\pm$ 0.049	0.057 $\pm$ 0.014	0.277 $\pm$ 0.023	0.162 $\pm$ 0.023	0.229$\pm$0.013
RELIAB	0.159 $\pm$ 0.062	0.277 $\pm$ 0.036	0.168 $\pm$ 0.039	0.746 $\pm$ 0.028	0.066 $\pm$ 0.017	0.350 $\pm$ 0.023	0.270 $\pm$ 0.023	0.247 $\pm$ 0.017
ML ²	0.166 $\pm$ 0.078	0.268 $\pm$ 0.032	0.129 $\pm$ 0.028	0.699 $\pm$ 0.038	0.040$\pm$0.013	0.267 $\pm$ 0.022	0.156 $\pm$ 0.029	0.254 $\pm$ 0.027
CLR	0.269 $\pm$ 0.061	0.322 $\pm$ 0.032	0.360 $\pm$ 0.170	0.830 $\pm$ 0.058	0.144 $\pm$ 0.027	0.437 $\pm$ 0.019	0.344 $\pm$ 0.027	0.241 $\pm$ 0.012
RAKEL	0.611 $\pm$ 0.084	0.315 $\pm$ 0.074	0.246 $\pm$ 0.038	0.879 $\pm$ 0.026	0.239 $\pm$ 0.031	0.412 $\pm$ 0.029	0.339 $\pm$ 0.027	0.280 $\pm$ 0.016
Comparing algorithms	Coverage $\downarrow$							
	cal500	emotions	medical	llog	msra	image	scene	yeast
MLFE	0.764 $\pm$ 0.013	0.278$\pm$0.022	0.031$\pm$0.012	0.227 $\pm$ 0.027	0.518$\pm$0.010	0.163$\pm$0.014	0.016$\pm$0.009	0.452$\pm$0.010
LIFT	0.759 $\pm$ 0.020	0.280 $\pm$ 0.025	0.048 $\pm$ 0.022	0.173 $\pm$ 0.020	0.539 $\pm$ 0.011	0.172 $\pm$ 0.014	0.020 $\pm$ 0.009	0.459 $\pm$ 0.010
RELIAB	0.738$\pm$0.013	0.299 $\pm$ 0.029	0.047 $\pm$ 0.016	0.161$\pm$0.020	0.541 $\pm$ 0.011	0.199 $\pm$ 0.012	0.107 $\pm$ 0.011	0.457 $\pm$ 0.004
ML ²	0.805 $\pm$ 0.013	0.279 $\pm$ 0.022	0.031$\pm$0.011	0.181 $\pm$ 0.019	0.522 $\pm$ 0.011	0.165 $\pm$ 0.012	0.018 $\pm$ 0.009	0.455 $\pm$ 0.011
CLR	0.795 $\pm$ 0.008	0.334 $\pm$ 0.020	0.080 $\pm$ 0.068	0.186 $\pm$ 0.044	0.618 $\pm$ 0.013	0.247 $\pm$ 0.016	0.137 $\pm$ 0.017	0.480 $\pm$ 0.008
RAKEL	0.964 $\pm$ 0.006	0.348 $\pm$ 0.021	0.089 $\pm$ 0.019	0.340 $\pm$ 0.023	0.670 $\pm$ 0.010	0.253 $\pm$ 0.009	0.174 $\pm$ 0.015	0.564 $\pm$ 0.008
Comparing algorithms	Ranking loss $\downarrow$							
	cal500	emotions	medical	llog	msra	image	scene	yeast
MLFE	0.185 $\pm$ 0.008	0.142$\pm$0.020	0.020 $\pm$ 0.010	0.233 $\pm$ 0.027	0.118$\pm$0.006	0.135$\pm$0.014	0.052$\pm$0.010	0.167$\pm$0.007
LIFT	0.182 $\pm$ 0.004	0.142$\pm$0.025	0.033 $\pm$ 0.017	0.157 $\pm$ 0.021	0.125 $\pm$ 0.005	0.147 $\pm$ 0.015	0.056 $\pm$ 0.009	0.170 $\pm$ 0.006
RELIAB	0.177$\pm$0.005	0.161 $\pm$ 0.031	0.031 $\pm$ 0.012	0.128 $\pm$ 0.018	0.131 $\pm$ 0.005	0.181 $\pm$ 0.013	0.090 $\pm$ 0.010	0.180 $\pm$ 0.008
ML ²	0.210 $\pm$ 0.009	0.144 $\pm$ 0.023	0.019$\pm$0.008	0.170 $\pm$ 0.020	0.119 $\pm$ 0.006	0.138 $\pm$ 0.011	0.055 $\pm$ 0.011	0.172 $\pm$ 0.009
CLR	0.224 $\pm$ 0.008	0.199 $\pm$ 0.024	0.065 $\pm$ 0.059	0.152$\pm$0.039	0.190 $\pm$ 0.009	0.243 $\pm$ 0.018	0.119 $\pm$ 0.016	0.187 $\pm$ 0.005
RAKEL	0.364 $\pm$ 0.014	0.217 $\pm$ 0.026	0.067 $\pm$ 0.015	0.292 $\pm$ 0.028	0.232 $\pm$ 0.011	0.250 $\pm$ 0.012	0.154 $\pm$ 0.014	0.250 $\pm$ 0.005

Experiments



Table 2: Predictive performance of each comparing algorithm (mean $\pm$ std. deviation) on the regular-scale data sets.

Comparing algorithms	Average precision $\uparrow$							
	cal500	emotions	medical	llog	msra	image	scene	yeast
MLFE	0.490 $\pm$ 0.025	0.815 $\pm$ 0.020	0.914$\pm$0.024	0.393 $\pm$ 0.036	0.827 $\pm$ 0.008	0.833$\pm$0.014	0.917$\pm$0.014	0.767$\pm$0.010
LIFT	0.503 $\pm$ 0.015	0.817$\pm$0.027	0.874 $\pm$ 0.029	0.402$\pm$0.039	0.830 $\pm$ 0.008	0.819 $\pm$ 0.015	0.909 $\pm$ 0.013	0.762 $\pm$ 0.008
RELIAB	0.507$\pm$0.019	0.797 $\pm$ 0.028	0.869 $\pm$ 0.028	0.393 $\pm$ 0.034	0.818 $\pm$ 0.009	0.776 $\pm$ 0.013	0.840 $\pm$ 0.014	0.744 $\pm$ 0.011
ML ²	0.456 $\pm$ 0.027	0.811 $\pm$ 0.022	0.903 $\pm$ 0.021	0.396 $\pm$ 0.036	0.836$\pm$0.007	0.827 $\pm$ 0.014	0.913 $\pm$ 0.016	0.757 $\pm$ 0.014
CLR	0.436 $\pm$ 0.019	0.762 $\pm$ 0.024	0.687 $\pm$ 0.192	0.295 $\pm$ 0.075	0.741 $\pm$ 0.013	0.718 $\pm$ 0.014	0.795 $\pm$ 0.018	0.745 $\pm$ 0.008
RAKEL	0.332 $\pm$ 0.019	0.766 $\pm$ 0.031	0.802 $\pm$ 0.027	0.233 $\pm$ 0.026	0.694 $\pm$ 0.014	0.725 $\pm$ 0.013	0.780 $\pm$ 0.018	0.710 $\pm$ 0.009
Comparing algorithms	Macro-averaging F1 $\uparrow$							
	cal500	emotions	medical	llog	msra	image	scene	yeast
MLFE	0.237 $\pm$ 0.022	0.670$\pm$0.052	0.720$\pm$0.073	0.461$\pm$0.062	0.553$\pm$0.015	0.658$\pm$0.024	0.819$\pm$0.026	0.425 $\pm$ 0.023
LIFT	0.176 $\pm$ 0.021	0.630 $\pm$ 0.042	0.690 $\pm$ 0.079	0.399 $\pm$ 0.057	0.516 $\pm$ 0.017	0.621 $\pm$ 0.035	0.797 $\pm$ 0.015	0.388 $\pm$ 0.022
RELIAB	0.301$\pm$0.022	0.650 $\pm$ 0.039	0.712 $\pm$ 0.053	0.392 $\pm$ 0.058	0.546 $\pm$ 0.014	0.556 $\pm$ 0.035	0.665 $\pm$ 0.025	0.405 $\pm$ 0.024
ML ²	0.236 $\pm$ 0.021	0.650 $\pm$ 0.047	0.674 $\pm$ 0.061	0.370 $\pm$ 0.060	0.548 $\pm$ 0.012	0.646 $\pm$ 0.029	0.799 $\pm$ 0.029	0.443$\pm$0.025
CLR	0.211 $\pm$ 0.025	0.601 $\pm$ 0.038	0.600 $\pm$ 0.129	0.395 $\pm$ 0.062	0.499 $\pm$ 0.017	0.525 $\pm$ 0.022	0.620 $\pm$ 0.025	0.400 $\pm$ 0.018
RAKEL	0.187 $\pm$ 0.020	0.618 $\pm$ 0.036	0.672 $\pm$ 0.058	0.366 $\pm$ 0.051	0.492 $\pm$ 0.020	0.540 $\pm$ 0.012	0.644 $\pm$ 0.019	0.430 $\pm$ 0.014
Comparing algorithms	Micro-averaging F1 $\uparrow$							
	cal500	emotions	medical	llog	msra	image	scene	yeast
MLFE	0.374 $\pm$ 0.024	0.684$\pm$0.043	0.816$\pm$0.032	0.211$\pm$0.042	0.725$\pm$0.009	0.657$\pm$0.021	0.810$\pm$0.026	0.649$\pm$0.014
LIFT	0.323 $\pm$ 0.016	0.659 $\pm$ 0.024	0.775 $\pm$ 0.036	0.177 $\pm$ 0.037	0.716 $\pm$ 0.015	0.622 $\pm$ 0.031	0.787 $\pm$ 0.015	0.645 $\pm$ 0.011
RELIAB	0.483$\pm$0.012	0.647 $\pm$ 0.036	0.725 $\pm$ 0.029	0.181 $\pm$ 0.037	0.714 $\pm$ 0.007	0.560 $\pm$ 0.030	0.654 $\pm$ 0.025	0.637 $\pm$ 0.011
ML ²	0.358 $\pm$ 0.027	0.661 $\pm$ 0.039	0.782 $\pm$ 0.026	0.057 $\pm$ 0.021	0.722 $\pm$ 0.008	0.645 $\pm$ 0.028	0.788 $\pm$ 0.029	0.641 $\pm$ 0.014
CLR	0.326 $\pm$ 0.019	0.614 $\pm$ 0.037	0.598 $\pm$ 0.157	0.176 $\pm$ 0.049	0.624 $\pm$ 0.010	0.525 $\pm$ 0.019	0.612 $\pm$ 0.026	0.628 $\pm$ 0.012
RAKEL	0.355 $\pm$ 0.018	0.634 $\pm$ 0.031	0.685 $\pm$ 0.031	0.148 $\pm$ 0.027	0.613 $\pm$ 0.015	0.540 $\pm$ 0.011	0.636 $\pm$ 0.023	0.632 $\pm$ 0.009

Experiments



Table 3: Predictive performance of each comparing algorithm (mean $\pm$ std. deviation) on the large-scale data sets.

Comparing algorithms	One-error $\downarrow$						
	slashdot	corel5k	rcv1subset1	rcv1subset2	rcv1subset3	rcv1subset4	rcv1subset5
MLFE	0.372 $\pm$ 0.020	0.646 $\pm$ 0.021	0.413 $\pm$ 0.013	0.463 $\pm$ 0.014	0.472 $\pm$ 0.023	0.453 $\pm$ 0.020	0.445 $\pm$ 0.015
LIFT	0.397 $\pm$ 0.026	0.669 $\pm$ 0.014	0.415 $\pm$ 0.019	0.455 $\pm$ 0.012	0.473 $\pm$ 0.020	0.457 $\pm$ 0.022	0.445 $\pm$ 0.017
RELIAB	0.516 $\pm$ 0.017	0.718 $\pm$ 0.015	0.467 $\pm$ 0.020	0.476 $\pm$ 0.011	0.477 $\pm$ 0.024	0.462 $\pm$ 0.015	0.467 $\pm$ 0.017
ML ²	0.363$\pm$0.018	0.627$\pm$0.023	0.396$\pm$0.021	0.448$\pm$0.014	0.461$\pm$0.020	0.438$\pm$0.017	0.437$\pm$0.017
CLR	0.979 $\pm$ 0.005	0.741 $\pm$ 0.018	0.501 $\pm$ 0.027	0.507 $\pm$ 0.019	0.533 $\pm$ 0.037	0.499 $\pm$ 0.017	0.503 $\pm$ 0.018
RAKEL	0.615 $\pm$ 0.020	0.872 $\pm$ 0.014	0.623 $\pm$ 0.023	0.592 $\pm$ 0.017	0.598 $\pm$ 0.018	0.592 $\pm$ 0.013	0.595 $\pm$ 0.021
Comparing algorithms	Coverage $\downarrow$						
	slashdot	corel5k	rcv1subset1	rcv1subset2	rcv1subset3	rcv1subset4	rcv1subset5
MLFE	0.117 $\pm$ 0.008	0.263$\pm$0.013	0.100$\pm$0.007	0.094$\pm$0.005	0.096$\pm$0.004	0.081$\pm$0.006	0.094$\pm$0.006
LIFT	0.107 $\pm$ 0.009	0.286 $\pm$ 0.013	0.132 $\pm$ 0.009	0.139 $\pm$ 0.006	0.142 $\pm$ 0.006	0.120 $\pm$ 0.009	0.134 $\pm$ 0.004
RELIAB	0.134 $\pm$ 0.005	0.300 $\pm$ 0.011	0.137 $\pm$ 0.009	0.121 $\pm$ 0.005	0.125 $\pm$ 0.006	0.113 $\pm$ 0.011	0.119 $\pm$ 0.006
ML ²	0.101$\pm$0.007	0.288 $\pm$ 0.012	0.110 $\pm$ 0.008	0.105 $\pm$ 0.007	0.109 $\pm$ 0.005	0.088 $\pm$ 0.007	0.108 $\pm$ 0.007
CLR	0.258 $\pm$ 0.009	0.287 $\pm$ 0.015	0.112 $\pm$ 0.008	0.105 $\pm$ 0.006	0.114 $\pm$ 0.024	0.095 $\pm$ 0.007	0.107 $\pm$ 0.006
RAKEL	0.218 $\pm$ 0.012	0.874 $\pm$ 0.012	0.417 $\pm$ 0.012	0.359 $\pm$ 0.022	0.369 $\pm$ 0.014	0.358 $\pm$ 0.020	0.363 $\pm$ 0.015
Comparing algorithms	Ranking loss $\downarrow$						
	slashdot	corel5k	rcv1subset1	rcv1subset2	rcv1subset3	rcv1subset4	rcv1subset5
MLFE	0.098 $\pm$ 0.008	0.109$\pm$0.006	0.039$\pm$0.003	0.039$\pm$0.002	0.040$\pm$0.002	0.034$\pm$0.003	0.038$\pm$0.002
LIFT	0.092 $\pm$ 0.008	0.122 $\pm$ 0.005	0.053 $\pm$ 0.003	0.059 $\pm$ 0.002	0.062 $\pm$ 0.002	0.051 $\pm$ 0.004	0.055 $\pm$ 0.002
RELIAB	0.118 $\pm$ 0.005	0.130 $\pm$ 0.005	0.058 $\pm$ 0.003	0.045 $\pm$ 0.002	0.052 $\pm$ 0.002	0.047 $\pm$ 0.004	0.048 $\pm$ 0.002
ML ²	0.084$\pm$0.006	0.163 $\pm$ 0.008	0.042 $\pm$ 0.003	0.043 $\pm$ 0.003	0.046 $\pm$ 0.003	0.037 $\pm$ 0.003	0.043 $\pm$ 0.003
CLR	0.245 $\pm$ 0.010	0.131 $\pm$ 0.008	0.048 $\pm$ 0.002	0.046 $\pm$ 0.002	0.054 $\pm$ 0.020	0.044 $\pm$ 0.002	0.046 $\pm$ 0.003
RAKEL	0.198 $\pm$ 0.013	0.586 $\pm$ 0.011	0.233 $\pm$ 0.007	0.209 $\pm$ 0.012	0.222 $\pm$ 0.009	0.224 $\pm$ 0.013	0.209 $\pm$ 0.012

Experiments



Table 3: Predictive performance of each comparing algorithm (mean $\pm$ std. deviation) on the large-scale data sets.

Comparing algorithms	Average precision $\uparrow$						
	slashdot	corel5k	rcv1subset1	rcv1subset2	rcv1subset3	rcv1subset4	rcv1subset5
MLFE	0.715 $\pm$ 0.015	0.316$\pm$0.012	0.615 $\pm$ 0.010	0.622 $\pm$ 0.006	0.614 $\pm$ 0.013	0.640 $\pm$ 0.013	0.626 $\pm$ 0.007
LIFT	0.695 $\pm$ 0.019	0.291 $\pm$ 0.010	0.582 $\pm$ 0.013	0.579 $\pm$ 0.008	0.569 $\pm$ 0.010	0.596 $\pm$ 0.010	0.586 $\pm$ 0.009
RELIAB	0.610 $\pm$ 0.012	0.269 $\pm$ 0.009	0.563 $\pm$ 0.010	0.588 $\pm$ 0.009	0.586 $\pm$ 0.013	0.611 $\pm$ 0.010	0.586 $\pm$ 0.009
ML ²	0.728$\pm$0.015	0.315 $\pm$ 0.013	0.629$\pm$0.013	0.630$\pm$0.008	0.622$\pm$0.011	0.647$\pm$0.014	0.631$\pm$0.006
CLR	0.261 $\pm$ 0.006	0.247 $\pm$ 0.009	0.564 $\pm$ 0.012	0.579 $\pm$ 0.011	0.554 $\pm$ 0.050	0.589 $\pm$ 0.013	0.576 $\pm$ 0.012
RAKEL	0.516 $\pm$ 0.012	0.120 $\pm$ 0.007	0.391 $\pm$ 0.009	0.429 $\pm$ 0.010	0.423 $\pm$ 0.010	0.431 $\pm$ 0.011	0.421 $\pm$ 0.012
Comparing algorithms	Macro-averaging F1 $\uparrow$						
	slashdot	corel5k	rcv1subset1	rcv1subset2	rcv1subset3	rcv1subset4	rcv1subset5
MLFE	0.455$\pm$0.048	0.338$\pm$0.014	0.282 $\pm$ 0.024	0.263 $\pm$ 0.023	0.243 $\pm$ 0.021	0.290 $\pm$ 0.036	0.265 $\pm$ 0.024
LIFT	0.427 $\pm$ 0.036	0.324 $\pm$ 0.014	0.219 $\pm$ 0.038	0.163 $\pm$ 0.020	0.151 $\pm$ 0.020	0.203 $\pm$ 0.034	0.165 $\pm$ 0.022
RELIAB	0.433 $\pm$ 0.047	0.303 $\pm$ 0.019	0.332$\pm$0.026	0.332$\pm$0.023	0.333$\pm$0.022	0.335$\pm$0.039	0.332$\pm$0.012
ML ²	0.424 $\pm$ 0.050	0.331 $\pm$ 0.015	0.228 $\pm$ 0.025	0.225 $\pm$ 0.020	0.216 $\pm$ 0.019	0.263 $\pm$ 0.037	0.232 $\pm$ 0.020
CLR	0.165 $\pm$ 0.035	0.276 $\pm$ 0.015	0.278 $\pm$ 0.028	0.269 $\pm$ 0.016	0.255 $\pm$ 0.035	0.297 $\pm$ 0.023	0.286 $\pm$ 0.013
RAKEL	0.363 $\pm$ 0.033	0.257 $\pm$ 0.013	0.266 $\pm$ 0.029	0.237 $\pm$ 0.024	0.243 $\pm$ 0.023	0.256 $\pm$ 0.020	0.255 $\pm$ 0.016
Comparing algorithms	Micro-averaging F1 $\uparrow$						
	slashdot	corel5k	rcv1subset1	rcv1subset2	rcv1subset3	rcv1subset4	rcv1subset5
MLFE	0.550$\pm$0.016	0.177 $\pm$ 0.015	0.411 $\pm$ 0.011	0.411 $\pm$ 0.011	0.397 $\pm$ 0.016	0.426 $\pm$ 0.013	0.414 $\pm$ 0.010
LIFT	0.509 $\pm$ 0.020	0.077 $\pm$ 0.011	0.311 $\pm$ 0.020	0.297 $\pm$ 0.013	0.289 $\pm$ 0.014	0.327 $\pm$ 0.019	0.314 $\pm$ 0.012
RELIAB	0.449 $\pm$ 0.014	0.226$\pm$0.010	0.417$\pm$0.007	0.468$\pm$0.011	0.431$\pm$0.013	0.485$\pm$0.009	0.466$\pm$0.009
ML ²	0.531 $\pm$ 0.015	0.126 $\pm$ 0.016	0.378 $\pm$ 0.016	0.383 $\pm$ 0.014	0.377 $\pm$ 0.017	0.411 $\pm$ 0.012	0.394 $\pm$ 0.015
CLR	0.008 $\pm$ 0.003	0.123 $\pm$ 0.019	0.361 $\pm$ 0.008	0.356 $\pm$ 0.015	0.338 $\pm$ 0.029	0.365 $\pm$ 0.017	0.368 $\pm$ 0.014
RAKEL	0.362 $\pm$ 0.014	0.135 $\pm$ 0.009	0.341 $\pm$ 0.008	0.337 $\pm$ 0.008	0.335 $\pm$ 0.012	0.349 $\pm$ 0.010	0.350 $\pm$ 0.010

Contents



- Background
- Introduction
- Approach
- Experiments
- Conclusion

Conclusion



- A novel approach is proposed to learning from multi-label data by **leveraging the structural information** in feature space.
- The key strategy is to convey the structural information modeled by **sparse reconstruction in feature space** to facilitate **generating enriched labeling information in output space**.
- The effectiveness of feature-induced labeling information enrichment is clearly validated with extensive experiments on benchmark multi-label data sets.