

Semi-Supervised AUC Optimization without Guessing Labels of Unlabeled Data

Zheng Xie and Ming Li
AAAI, 2018

Introduction

□ Semi-supervised Learning

- Obtain the labels for the collected data is expensive
- Semi-supervised learning are based on certain distributional assumption, e.g. **cluster assumption** and **manifold assumption**
- Estimate the labels of unlabeled instances

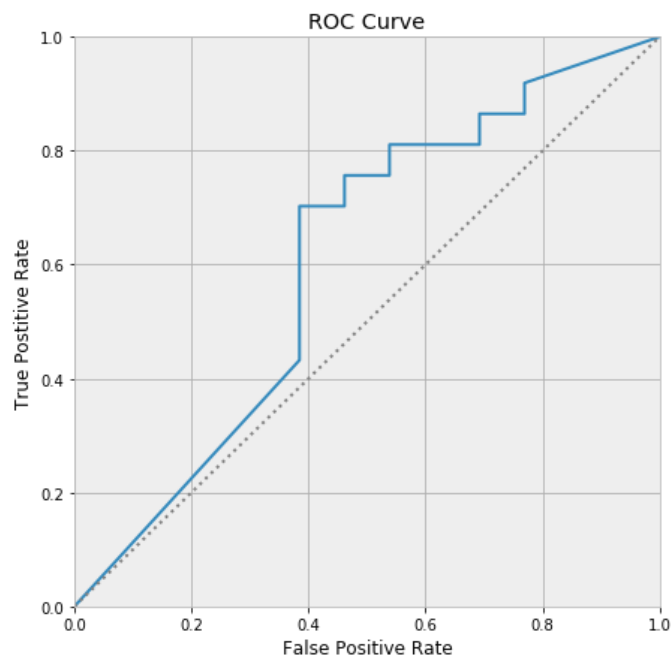
□ AUC Optimization for SSL

- Rely on the **distributional assumptions**
- Require **prior probability**
 - Sakai, Niu, and Sugiyama (2017) proposed an unbiased semi-supervised AUC optimization method based on positive-unlabeled learning

□ Proposed Method

- Do not depend on any distributional assumptions
- Do not need to know prior probability

AUC



$$TPR = \frac{TP}{TP + FN} \longrightarrow P$$

$$FPR = \frac{FP}{TN + FP} \longrightarrow N$$

The AUC is the probability the model will score a randomly chosen positive class higher than a randomly chosen negative class.

AUC Optimization

- Supervised learning

$$\begin{aligned} \mathcal{X}_P &:= \{\mathbf{x}_i\}_{i=1}^{n_P} \stackrel{\text{i.i.d.}}{\sim} p_P(\mathbf{x}) := p(\mathbf{x} \mid y = +1) \\ \mathcal{X}_N &:= \{\mathbf{x}'_j\}_{j=1}^{n_N} \stackrel{\text{i.i.d.}}{\sim} p_N(\mathbf{x}) := p(\mathbf{x} \mid y = -1) \end{aligned} \quad \Rightarrow \quad f(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}$$

- Since AUC is equivalent to the probability of a randomly drawn positive instance being ranked before a randomly drawn negative instance, it can be formulated as

$$\text{AUC} = 1 - \frac{\mathbb{E}_{\mathbf{x} \in \mathcal{X}_P} [\mathbb{E}_{\mathbf{x}' \in \mathcal{X}_N} [\ell_{01}(\mathbf{w}^\top (\mathbf{x} - \mathbf{x}'))]]}{R_{PN}}$$

AUC Optimization in SSL

- Rely on the distributional assumptions
 - These assumptions may be violated in many real-world applications
- Require an accurate estimation of the prior probability
 - It is usually difficult especially when the number of labeled data is extremely small.

We can achieve unbiased semi-supervised AUC optimization without distributional assumptions or prior knowledge about the distribution or class prior probabilities.

Theorem

In AUC optimization,

$$\min R_{PN}$$



$$\min R_{PU}$$

$$\min R_{NU}$$

$$R_{PU} = \mathbb{E}_{\mathbf{x} \in \mathcal{X}_P} [\mathbb{E}_{\mathbf{x}'' \in \mathcal{X}_U} [\ell_{01}(\mathbf{w}^\top (\mathbf{x} - \mathbf{x}''))]] ,$$

$$R_{NU} = \mathbb{E}_{\mathbf{x}'' \in \mathcal{X}_U} [\mathbb{E}_{\mathbf{x}' \in \mathcal{X}_N} [\ell_{01}(\mathbf{w}^\top (\mathbf{x}'' - \mathbf{x}'))]] .$$

Proof:

$$\begin{aligned} R_{PU} &= \mathbb{E}_{\mathbf{x} \in \mathcal{X}_P} [\mathbb{E}_{\mathbf{x}'' \in \mathcal{X}_U} [\ell_{01}(\mathbf{w}^\top (\mathbf{x} - \mathbf{x}''))]] \\ &= \mathbb{E}_{\mathbf{x} \in \mathcal{X}_P} [\theta_P \mathbb{E}_{\bar{\mathbf{x}} \in \mathcal{X}_P} [\ell_{01}(\mathbf{w}^\top (\mathbf{x} - \bar{\mathbf{x}}))] + \theta_N \mathbb{E}_{\mathbf{x}' \in \mathcal{X}_N} [\ell_{01}(\mathbf{w}^\top (\mathbf{x} - \mathbf{x}'))]] \\ &= \frac{1}{2} \theta_P + \theta_N \mathbb{E}_{\mathbf{x} \in \mathcal{X}_P} [\mathbb{E}_{\mathbf{x}' \in \mathcal{X}_N} [\ell_{01}(\mathbf{w}^\top (\mathbf{x} - \mathbf{x}'))]] \end{aligned}$$

$$\mathcal{X}_U := \{\mathbf{x}''_k\}_{k=1}^{n_U} \stackrel{\text{i.i.d.}}{\sim} p(\mathbf{x}) = \theta_P p_P(\mathbf{x}) + \theta_N p_N(\mathbf{x})$$

$$\frac{1}{2} [l_{01}(w(x_1 - x_2)) + l_{01}(w(x_2 - x_1))] = \frac{1}{2}$$

Unbiased Estimation

$$\theta_P + \theta_N = 1$$

$$R_{PU} = \theta_N R_{PN} + \frac{1}{2} \theta_P$$

$$R_{NU} = \theta_P R_{PN} + \frac{1}{2} \theta_N$$

$$R_{PU} + R_{NU} - \frac{1}{2} = R_{PN}$$

Conduct unbiased AUC risk estimation without knowing the class prior probabilities θ_P and θ_N .

Alg-1. SAMULT

Semi-supervised AUC Maximization by treating the UnLabeled data in Two ways

Objective: $\min_{\mathbf{w}} \gamma \hat{R}_{\text{PN}} + (1 - \gamma) \left(\hat{R}_{\text{PU}} + \hat{R}_{\text{NU}} - \frac{1}{2} \right) + \lambda \|\mathbf{w}\|^2$

$\hat{R}_{\text{PN}}$ supervised AUC risk

$\hat{R}_{\text{PU}} + \hat{R}_{\text{NU}} - \frac{1}{2}$ unbiased semi-supervised AUC risk

Where,

$$\hat{R}_{\text{PN}} = \frac{1}{n_{\text{P}} n_{\text{N}}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{P}}} \sum_{\mathbf{x}' \in \mathcal{X}_{\text{N}}} \ell(\mathbf{w}^{\top}(\mathbf{x} - \mathbf{x}'))$$

$$\hat{R}_{\text{PU}} = \frac{1}{n_{\text{P}} n_{\text{U}}} \sum_{\mathbf{x} \in \mathcal{X}_{\text{P}}} \sum_{\mathbf{x}'' \in \mathcal{X}_{\text{U}}} \ell(\mathbf{w}^{\top}(\mathbf{x} - \mathbf{x}''))$$

$$\hat{R}_{\text{NU}} = \frac{1}{n_{\text{U}} n_{\text{N}}} \sum_{\mathbf{x}'' \in \mathcal{X}_{\text{U}}} \sum_{\mathbf{x}' \in \mathcal{X}_{\text{N}}} \ell(\mathbf{w}^{\top}(\mathbf{x}'' - \mathbf{x}'))$$

$$\gamma \in [0, 1]$$

$$\lambda \geq 0$$

SAMULT does not determine the decision boundary.

Analytical Solution

$$\hat{\mathbf{w}} = (\gamma \mathbf{H}_{\text{PN}} + (1 - \gamma)(\mathbf{H}_{\text{PU}} + \mathbf{H}_{\text{NU}}) + \lambda \mathbf{I}_d)^{-1} \\ (\gamma \mathbf{h}_{\text{PN}} + (1 - \gamma)(\mathbf{h}_{\text{PU}} + \mathbf{h}_{\text{NU}})) ,$$

where

$$\begin{aligned} \mathbf{h}_{\text{PN}} &= \frac{1}{n_{\text{P}}} \mathbf{X}_{\text{P}}^{\top} \mathbf{1}_{n_{\text{P}}} - \frac{1}{n_{\text{N}}} \mathbf{X}_{\text{N}}^{\top} \mathbf{1}_{n_{\text{N}}} , \\ \mathbf{h}_{\text{PU}} &= \frac{1}{n_{\text{P}}} \mathbf{X}_{\text{P}}^{\top} \mathbf{1}_{n_{\text{P}}} - \frac{1}{n_{\text{U}}} \mathbf{X}_{\text{U}}^{\top} \mathbf{1}_{n_{\text{U}}} , \\ \mathbf{h}_{\text{NU}} &= \frac{1}{n_{\text{U}}} \mathbf{X}_{\text{U}}^{\top} \mathbf{1}_{n_{\text{U}}} - \frac{1}{n_{\text{N}}} \mathbf{X}_{\text{N}}^{\top} \mathbf{1}_{n_{\text{N}}} , \\ \mathbf{H}_{\text{PN}} &= \frac{1}{n_{\text{P}}} \mathbf{X}_{\text{P}}^{\top} \mathbf{X}_{\text{P}} - \frac{1}{n_{\text{P}} n_{\text{N}}} \mathbf{X}_{\text{P}}^{\top} \mathbf{1}_{n_{\text{P}}} \mathbf{1}_{n_{\text{N}}}^{\top} \mathbf{X}_{\text{N}} \\ &\quad - \frac{1}{n_{\text{P}} n_{\text{N}}} \mathbf{X}_{\text{N}}^{\top} \mathbf{1}_{n_{\text{N}}} \mathbf{1}_{n_{\text{P}}}^{\top} \mathbf{X}_{\text{P}} + \frac{1}{n_{\text{N}}} \mathbf{X}_{\text{N}}^{\top} \mathbf{X}_{\text{N}} , \\ \mathbf{H}_{\text{PU}} &= \frac{1}{n_{\text{P}}} \mathbf{X}_{\text{P}}^{\top} \mathbf{X}_{\text{P}} - \frac{1}{n_{\text{P}} n_{\text{U}}} \mathbf{X}_{\text{P}}^{\top} \mathbf{1}_{n_{\text{P}}} \mathbf{1}_{n_{\text{U}}}^{\top} \mathbf{X}_{\text{U}} \\ &\quad - \frac{1}{n_{\text{P}} n_{\text{U}}} \mathbf{X}_{\text{U}}^{\top} \mathbf{1}_{n_{\text{U}}} \mathbf{1}_{n_{\text{P}}}^{\top} \mathbf{X}_{\text{P}} + \frac{1}{n_{\text{U}}} \mathbf{X}_{\text{U}}^{\top} \mathbf{X}_{\text{U}} , \\ \mathbf{H}_{\text{NU}} &= \frac{1}{n_{\text{U}}} \mathbf{X}_{\text{U}}^{\top} \mathbf{X}_{\text{U}} - \frac{1}{n_{\text{U}} n_{\text{N}}} \mathbf{X}_{\text{U}}^{\top} \mathbf{1}_{n_{\text{U}}} \mathbf{1}_{n_{\text{N}}}^{\top} \mathbf{X}_{\text{N}} \\ &\quad - \frac{1}{n_{\text{U}} n_{\text{N}}} \mathbf{X}_{\text{N}}^{\top} \mathbf{1}_{n_{\text{N}}} \mathbf{1}_{n_{\text{U}}}^{\top} \mathbf{X}_{\text{U}} + \frac{1}{n_{\text{N}}} \mathbf{X}_{\text{N}}^{\top} \mathbf{X}_{\text{N}} , \end{aligned}$$

SAMULT^{P+U}

When only the positive and unlabeled data are available

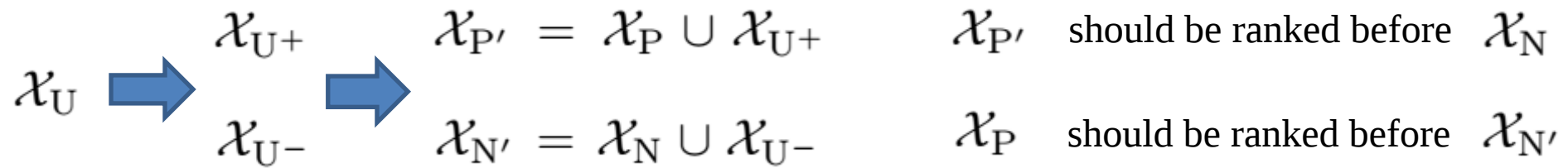
$$\min_{\boldsymbol{w}} \quad \gamma \hat{R}_{\text{PN}} + (1 - \gamma) \left(\hat{R}_{\text{PU}} + \hat{R}_{\text{NU}} - \frac{1}{2} \right) + \lambda ||\boldsymbol{w}||^2$$



$$\min_{\boldsymbol{w}} \quad \hat{R}_{\text{PU}}(\boldsymbol{w}) + \lambda ||\boldsymbol{w}||^2$$

Alg-2. SAMPURA

Semi-supervised AUC Maximization by **P**artitioning **U**nabeled data at **R**andom



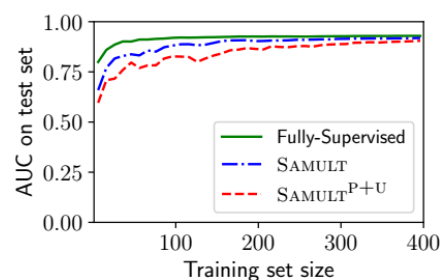
$$\min_{\mathbf{w}} \quad \hat{R}_{P'N} + \hat{R}_{PN'} + \lambda ||\mathbf{w}||^2$$

take the average of those $\mathbf{w}$ s to construct the final ensemble

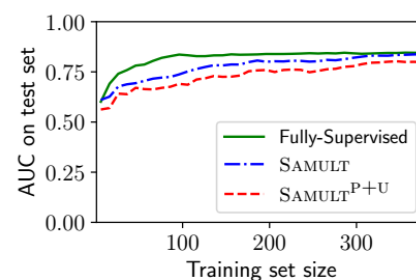
Experiment - 1

The model that minimizes the PU-AUC risk or the PNU-AUC risk converges to the supervised case

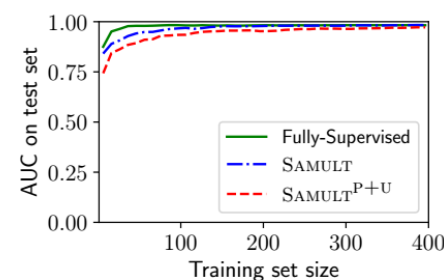
AUC



(a) *australian*



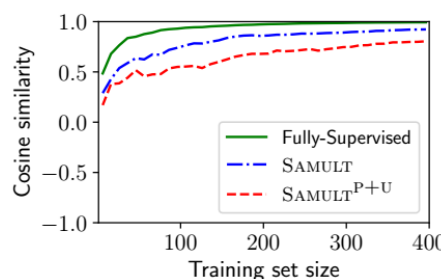
(b) *clean1*



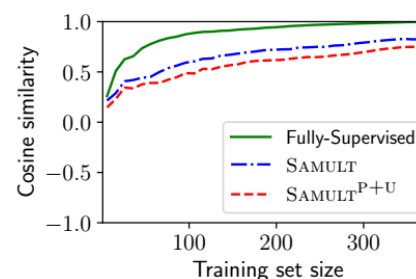
(c) *wdbc*

10% data in training set is labeled

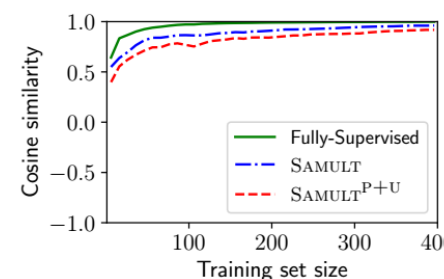
$\cos\langle \hat{\mathbf{w}}, \hat{\mathbf{w}}^* \rangle$



(a) *australian*



(b) *clean1*



(c) *wdbc*

$\hat{\mathbf{w}}^*$ is learned from all available data with label

Experiment - 2

Compare our methods to the state-of-the-art semi-supervised AUC optimization methods.

- **SAMULT**
- **SAMPURA**
- **SSRankBoost** (Amini, Truong, and Goutte 2008), a boosting based algorithm for learning a bipartite ranking function with partially labeled data
- **PNU-AUC** (Sakai, Niu, and Sugiyama 2017), which is a semi-supervised AUC optimization method based on positive-unlabeled learning.
- **Supervised AUC optimization**
- **Logistic regression**

Experiment - 2

Dataset	Supervised	Log. Reg.	SSRankBoost	PNU-AUC	SAMULT	SAMPURA
<i>australian</i>	.879±.029	.860±.027	.886±.013	.903±.009	.903±.009	.903±.009
<i>breast</i>	.655±.097	.625±.095	.647±.065	.701±.029	.701±.029	.704±.026
<i>breastw</i>	.987±.009	.980±.006	.984±.013	.992±.001	.996±.001	.996±.001
<i>clean1</i>	.760±.062	.725±.060	.737±.050	.767±.042	.777±.039	.782±.038
<i>colic</i>	.829±.112	.818±.074	.721±.062	.858±.013	.869±.013	.870±.013
<i>colic.orig</i>	.647±.093	.645±.076	.612±.081	.644±.048	.658±.049	.663±.044
<i>credit-a</i>	.893±.024	.886±.023	.885±.013	.906±.008	.906±.007	.906±.008
<i>credit-g</i>	.719±.043	.709±.030	.665±.027	.748±.018	.748±.018	.761±.017
<i>fourclass</i>	.825±.023	.826±.026	.692±.029	.827±.008	.827±.008	.828±.006
<i>german</i>	.683±.057	.672±.048	.709±.025	.727±.019	.727±.019	.729±.017
<i>haberman</i>	.551±.086	.530±.075	.582±.067	.547±.051	.551±.045	.556±.043
<i>heart</i>	.857±.065	.842±.060	.823±.042	.876±.025	.876±.025	.878±.024
<i>house</i>	.975±.038	.961±.015	.972±.034	.979±.012	.979±.012	.979±.011
<i>ijcnn1</i>	.912±.003	.901±.004	.902±.002	.904±.009	.913±.005	.915±.004
<i>madelon</i>	.510±.037	.541±.020	.571±.023	.528±.029	.517±.027	.530±.022
<i>parkinsons</i>	.848±.129	.826±.082	.799±.051	.860±.023	.860±.023	.863±.011
<i>phishing</i>	.975±.097	.972±.001	.983±.003	.974±.002	.976±.002	.985±.002
<i>vehicle</i>	.932±.038	.922±.022	.912±.039	.965±.020	.965±.020	.970±.014
<i>vote</i>	.965±.038	.951±.015	.972±.034	.979±.012	.979±.012	.979±.011
<i>wdbc</i>	.971±.014	.963±.006	.964±.016	.983±.006	.983±.006	.983±.005
# Best/Comp.	2	0	4	11	15	18

Experiment - 3

Dataset	PU-RSVM	PU-AUC	SAMULT ^{P+U}
<i>australian</i>	.844±.034	.898±.021	.900±.019
<i>breast</i>	.615±.104	.701±.077	.701±.077
<i>breastw</i>	.987±.009	.993±.002	.996±.002
<i>clean1</i>	.709±.072	.786±.050	.796±.050
<i>colic</i>	.807±.103	.877±.060	.877±.060
<i>colic.orig</i>	.621±.078	.650±.068	.670±.061
<i>credit-a</i>	.876±.028	.912±.015	.912±.015
<i>credit-g</i>	.688±.047	.755±.028	.757±.027
<i>fourclass</i>	.823±.031	.832±.026	.832±.025
<i>german</i>	.642±.045	.734±.034	.736±.034
<i>haberman</i>	.572±.083	.561±.081	.555±.080
<i>heart</i>	.835±.073	.883±.039	.883±.040
<i>house</i>	.945±.029	.980±.012	.983±.011
<i>ijcnn1</i>	.927±.005	.900±.011	.905±.012
<i>madelon</i>	.470±.015	.533±.031	.514±.032
<i>parkinsons</i>	.797±.089	.870±.033	.870±.032
<i>phishing</i>	.960±.008	.966±.005	.970±.005
<i>vehicle</i>	.942±.034	.959±.030	.966±.025
<i>vote</i>	.945±.029	.980±.012	.983±.011
<i>wdbc</i>	.967±.021	.984±.007	.985±.007
#Best/Com.	2	7	17

Degeneration to AUC Optimization for Positive and Unlabeled Data

PU-RSVM (Sellamanickam, Garg, and Selvaraj 2011):

A ranking SVM based method for positive and unlabeled data

PU-AUC (Sakai, Niu, and Sugiyama 2017):

A positive unlabeled AUC optimization method by optimizing an unbiased AUC risk estimator relies only on positive and unlabeled data

Do not need to know the prior.

Conclusion

- In semi-supervised AUC optimization, it is unnecessary to design sophisticated strategies to estimate the possible labels of the unlabeled data or the class prior probabilities.
- Proposed two semi-supervised AUC optimization methods: SAMULT and SAMPURA
- The positive-unlabeled AUC optimization problem can be addressed by a degenerated version of proposed method that simply treats the unlabeled data as negative.

Reference

- Zheng Xie, Ming Li. **Semi-Supervised AUC Optimization without Guessing Labels of Unlabeled Data.** AAAI, 2018.
- Sundararajan Sellamanickam, Priyanka Garg, Sathiya Keerthi Selvaraj. **A pairwise ranking based approach to learning with positive and unlabeled examples.** CIKM, 2011.
- Tomoya Sakai, Gang Niu, and Masashi Sugiyama. **Semi-Supervised AUC Optimization Based on Positive-Unlabeled Learning.** *Machine Learning*, 2017.
- Massih-Reza Amini, Tuong-Vinh Truong, Cyril Goutte. **A Boosting Algorithm for Learning Bipartite Ranking Functions with Partially Labeled Data.** SIGIR, 2008.