
Efficient and Consistent Adversarial Bipartite Matching

Rizal Fathony^{*1} Sima Behpour^{*1} Xinhua Zhang¹ Brian D. Ziebart¹

Abstract

Many important structured prediction problems, including learning to rank items, correspondence-based natural language processing, and multi-object tracking, can be formulated as weighted bipartite matching optimizations. Existing structured prediction approaches have significant drawbacks when applied under the constraints of perfect bipartite matchings. Exponential family probabilistic models, such as the conditional random field (CRF), provide statistical consistency guarantees, but suffer computationally from the need to compute the normalization term of its distribution over matchings, which is a #P-hard matrix permanent computation. In contrast, the structured support vector machine (SSVM) provides computational efficiency, but lacks Fisher consistency, meaning that there are distributions of data for which it cannot learn the optimal matching even under ideal learning conditions (i.e., given the true distribution and selecting from all measurable potential functions). We propose adversarial bipartite matching to avoid both of these limitations. We develop this approach algorithmically, establish its computational efficiency and Fisher consistency properties, and apply it to matching problems that demonstrate its empirical benefits.

1 Introduction

How can the elements from two sets be paired one-to-one to have the largest sum of pairwise utilities? This *maximum weighted perfect bipartite matching* problem is a classical combinatorial optimization problem in computer science. It can be formulated and efficiently solved in polynomial time as a linear program or using more specialized *Hungarian algorithm* techniques (Kuhn, 1955). This has made it an

attractive formalism for posing a wide range of problems, including recognizing correspondences in similar images (Belongie et al., 2002; Liu et al., 2008; Zhu et al., 2008; Rui et al., 2007), finding word alignments in text (Chan & Ng, 2008), and providing ranked lists of items for information retrieval tasks (Amini et al., 2008).

Machine learning methods seek to estimate the pairwise utilities of bipartite graphs so that the maximum weighted complete matching is most compatible with the (distribution of) ground truth matchings of training data. When these utilities are learned abstractly, they can be employed to make predictive matchings for test samples. Unfortunately, important measures of incompatibility (e.g., the Hamming loss) are often non-continuous with many local optima in the predictors' parameter spaces, making direct minimization intractable. Given this difficulty, two natural desiderata for any predictor are:

- **Efficiency:** learning from training data and making predictions must be computed efficiently in (low-degree) polynomial time; and
- **Consistency:** the predictor's training objectives must also minimize the underlying Hamming loss, at least under ideal learning conditions (given the true distribution and fully expressive model parameters).

Existing methods for learning bipartite matchings fail in one or the other of these desiderata; exponentiated potential fields models (Lafferty et al., 2001; Petterson et al., 2009) are intractable for large sets of items, while maximum margin methods based on the hinge loss surrogate (Taskar et al., 2005a; Tsochantaridis et al., 2005) lack Fisher consistency (Tewari & Bartlett, 2007; Liu, 2007). We discuss these limitations formally in Section 2.

Given the deficiencies of the existing methods, we contribute the first approach for learning bipartite matchings that is both computationally efficient and Fisher consistent. Our approach is based on an adversarial formulation for learning (Topsøe, 1979; Grünwald & Dawid, 2004; Asif et al., 2015) that poses prediction-making as a data-constrained zero-sum game between a player seeking to minimize the expected loss and an adversarial data approximator seeking to maximize the expected loss. We present two approaches for solving the corresponding zero-sum game arising from our formulation: (1) using the double

^{*}Equal contribution ¹Department of Computer Science, University of Illinois at Chicago. Correspondence to: Rizal Fathony <rfatho2@uic.edu>, Sima Behpour <sbehpo2@uic.edu>.

oracle method of constraint generation to find a sparsely-supported equilibrium for the zero-sum game; and (2) decomposing the game’s solution into marginal probabilities and optimizes these marginal probabilities directly to obtain an equilibrium saddle point for the game. We then establish the computational efficiency and consistency of this approach and demonstrate its benefits experimentally.

2 Previous Inefficiency and Inconsistency

2.1 Bipartite Matching Task

Given two sets of elements A and B of equal size ($|A| = |B|$), a maximum weighted bipartite matching π is the one-to-one mapping (e.g., Figure 1) from each element in B to each element in A that maximizes the sum of potentials: $\max_{\pi \in \Pi} \psi(\pi) = \max_{\pi \in \Pi} \sum_i \psi_i(\pi_i)$. Here $\pi_i \in [n] := \{1, 2, \dots, n\}$ is the entry in B that is matched with the i -th entry of A . The set of possible solutions Π is simply all permutation of $[n]$. Many machine learning tasks pose prediction as the solution to this problem, including: word alignment for natural language processing tasks (Taskar et al., 2005b; Padó & Lapata, 2006; MacCartney et al., 2008); learning correspondences between images in computer vision applications (Belongie et al., 2002; Dellaert et al., 2003); protein structure analysis in computational biology (Taylor, 2002; Wang et al., 2004); and learning to rank a set of items for information retrieval tasks (Dwork et al., 2001; Le & Smola, 2007). Thus, learning appropriate weights $\psi_i(\cdot)$ for bipartite graph matchings is a key problem for many application areas.

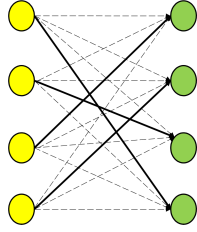


Figure 1. $n = 4$ bipartite matching task.

2.2 Performance Evaluation and Fisher Consistency

Given a predicted permutation, π' , and the “ground truth” permutation, π , the **Hamming loss** counts the number of mistaken pairings: $\text{loss}_{\text{Ham}}(\pi, \pi') = \sum_{i=1}^n 1(\pi'_i \neq \pi_i)$, where $1(\cdot) = 1$ if $\cdot$ is true and 0 otherwise. When the “ground truth” is a distribution over permutations, $P(\pi)$, rather than a single permutation, the (set of) **Bayes optimal** prediction(s) is: $\text{argmin}_{\pi'} \sum_{\pi} P(\pi) \text{loss}_{\text{Ham}}(\pi, \pi')$. For a predictor to be **Fisher consistent**, it must provide a Bayes optimal prediction for any possible distribution $P(\pi)$ when trained from that exact distribution using the predictor’s most general possible parameterization (e.g., all measurable functions ψ for potential-based models).

2.3 Exponential Family Random Field Approach

A probabilistic approach to learning bipartite graphs uses an exponential family distribution over permutations,

$P_{\psi}(\pi) = e^{\sum_{i=1}^n \psi_i(\pi_i)} / Z_{\psi}$, trained by maximizing training data likelihood. This provides certain statistical consistency guarantees for its marginal probability estimates (Peterson et al., 2009). Specifically, if the potentials ψ are chosen from the space of all measurable functions to maximize the likelihood of the true distribution of permutations $P(\pi)$, then $P_{\psi}(\pi)$ will match the marginal probabilities of the true distribution: $\forall i, j, P_{\psi}(\pi_i = j) = P(\pi_i = j)$. This implies Fisher consistency because the MAP estimate under this distribution, which can be obtained as a maximum weighted bipartite matching, is Bayes optimal.

The key challenge with this approach is its computational complexity. The normalization term, Z_{ψ} , is the permanent of a matrix defined in terms of exponentiated potential terms: $Z_{\psi} = \sum_{\pi} \prod_{i=1}^n e^{\psi_i(\pi_i)} = \text{perm}(\mathbf{M})$ where $M_{i,j} = e^{\psi_i(j)}$. For sets of small size (e.g., $n = 5$), enumerating the permutations is tractable and learning using the exponential random field model incurs a run-time cost that is acceptable in practice (Peterson et al., 2009). However, the matrix permanent computation is a $\#P$ -hard problem to compute exactly (Valiant, 1979). Monte Carlo sampling approaches are used instead of permutation enumeration to maximize the data likelihood (Peterson et al., 2009; Volkovs & Zemel, 2012). Though exact samples can be generated efficiently in polynomial time (Huber & Law, 2008), the number of samples needed for reliable likelihood or gradient estimates makes this approach infeasible for applications with even modestly-sized sets of $n = 20$ elements (Peterson et al., 2009).

2.4 Maximum Margin Approach

Maximum margin methods for structured prediction seek potentials ψ that minimize the training sample hinge loss:

$$\min_{\psi} \mathbb{E}_{\pi \sim \tilde{P}} \left[\max_{\pi'} \{ \text{loss}(\pi, \pi') + \psi(\pi') \} - \psi(\pi) \right], \quad (1)$$

where $\tilde{P}$ is the empirical distribution. Finding the optimal ψ is a convex optimization problem (Boyd & Vandenberghe, 2004) that can generally be tractably solved using constraint generation methods as long as the maximizing assignments can be found efficiently. In the case of permutation learning, finding the permutation π' with highest hinge loss reduces to a maximum weighted bipartite matching problem and can therefore be solved efficiently.

Though computationally efficient, maximum margin approaches for learning to make perfect bipartite matches lack **Fisher consistency**, which requires the prediction $\pi^* = \text{argmax}_{\pi} \psi(\pi)$ resulting from Equation (1) to minimize the expected risk, $\mathbb{E}_{\pi \sim \tilde{P}} [\text{loss}(\pi, \pi^*)]$, for all distributions $\tilde{P}$. We consider a distribution over permutations that is an extension of a counterexample for multiclass classification consistency analysis with no majority label (Liu,

2007): $P(\pi = [1 \ 2 \ 3]) = 0.4$; $P(\pi = [2 \ 3 \ 1]) = 0.3$; and $P(\pi = [3 \ 1 \ 2]) = 0.3$. The potential function $\psi_i(j) = 1$ if $i = j$ and 0 otherwise, provides a Bayes optimal permutation prediction for this distribution and an expected hinge loss of $3.6 = 0.4(3 - 3) + 0.3(3 + 3) + 0.3(3 + 3)$. However, the expected hinge loss is optimally minimized with a value of 3 when $\psi_i(j) = 0, \forall i, j$, which is indifferent between all permutations and is not Bayes optimal. Thus, Fisher consistency is not guaranteed.

3 Approach

To overcome the computational inefficiency of exponential random field methods and the Fisher inconsistency of maximum margin methods, we formulate the task of learning for bipartite matching problems as an adversarial structured prediction task. We present two approaches for efficiently solving the resulting game over permutations.

3.1 Permutation Mixture Formulation

The training data for bipartite matching consists of triplets (A, B, π) where A and B are two sets of nodes with equal size and π is the assignment. To simplify the notation, we denote x as the bipartite graph containing the nodes A and B . We also denote $\phi(x, \pi)$ as a vector that enumerates the joint feature representations based on the bipartite graph x and the matching assignment π . This joint feature is defined additively over each node assignment, i.e., $\phi(x, \pi) = \sum_{i=1}^n \phi_i(x, \pi_i)$.

Our approach seeks a predictor that robustly minimizes the Hamming loss against the worst-case permutation mixture probability that is consistent with the statistics of the training data. In this setting, a predictor makes a probabilistic prediction over the set of all possible assignments (denoted as $\hat{P}$). Instead of evaluating the predictor with the empirical distribution, the predictor is pitted against an adversary that also makes a probabilistic prediction (denoted as $\tilde{P}$). The predictor's objective is to minimize the expected loss function calculated from the predictor's and adversary's probabilistic predictions, while the adversary seeks to maximize the loss. The adversary (and only the adversary) is constrained to select a probabilistic prediction that matches the statistical summaries of the empirical training distribution (denoted as $\tilde{P}$) via moment matching constraints on joint features $\phi(x, \pi)$. Formally, we write our formulation as:

$$\min_{\hat{P}(\hat{\pi}|x)} \max_{\tilde{P}(\tilde{\pi}|x)} \mathbb{E}_{x \sim \tilde{P}; \hat{\pi}|x \sim \hat{P}; \tilde{\pi}|x \sim \tilde{P}} [\text{loss}(\hat{\pi}, \tilde{\pi})] \text{ s.t.} \quad (2)$$

$$\mathbb{E}_{x \sim \tilde{P}; \tilde{\pi}|x \sim \tilde{P}} \left[\sum_{i=1}^n \phi_i(x, \tilde{\pi}_i) \right] = \mathbb{E}_{(x, \pi) \sim \tilde{P}} \left[\sum_{i=1}^n \phi_i(x, \pi_i) \right].$$

This follows a recent line of work for adversarial classification under additive (Asif et al., 2015) and non-additive

(Wang et al., 2015) loss functions that has been employed for chain-structured prediction (Li et al., 2016), object detection (Behpour et al., 2017), and robust cut learning (Behpour et al., 2018). Using the method of Lagrangian multipliers and strong duality for convex-concave saddle point problems (Von Neumann & Morgenstern, 1945; Sion, 1958), The optimization in Eq. (2) can be equivalently solved in the dual formulation:

$$\min_{\theta} \mathbb{E}_{x, \pi \sim \tilde{P}} \min_{\hat{P}(\hat{\pi}|x)} \max_{\tilde{P}(\tilde{\pi}|x)} \mathbb{E}_{\hat{\pi}|x \sim \hat{P}} \left[\text{loss}(\hat{\pi}, \tilde{\pi}) + \theta \cdot \sum_{i=1}^n (\phi_i(x, \tilde{\pi}_i) - \phi_i(x, \pi_i)) \right] \quad (3)$$

where θ is the Lagrange dual variable for the moment matching constraints. We refer the reader to Appendix A in the supplementary materials for a more detailed explanation of this construction (i.e., the transformation from Eq. (2) to Eq. (3)). In this paper, we use Hamming distance, $\text{loss}(\hat{\pi}, \tilde{\pi}) = \sum_{i=1}^n 1(\hat{\pi}_i \neq \tilde{\pi}_i)$, as the loss function.

Table 1 shows the payoff matrix for the game of size $n = 3$ with $3!$ actions (permutations) for the predictor player $\hat{\pi}$ and for the adversarial approximation player $\tilde{\pi}$. Here, we define the difference between the Lagrangian potential of the adversary's action and the ground truth permutation as $\delta_{\tilde{\pi}} = \psi(\tilde{\pi}) - \psi(\pi) = \theta \cdot \sum_{i=1}^n (\phi_i(x, \tilde{\pi}_i) - \phi_i(x, \pi_i))$.

Table 1. Augmented Hamming loss matrix for $n=3$ permutations.

	$\tilde{\pi} = 123$	$\tilde{\pi} = 132$	$\tilde{\pi} = 213$	$\tilde{\pi} = 231$	$\tilde{\pi} = 312$	$\tilde{\pi} = 321$
$\hat{\pi} = 123$	$0 + \delta_{123}$	$2 + \delta_{132}$	$2 + \delta_{213}$	$3 + \delta_{231}$	$3 + \delta_{312}$	$2 + \delta_{321}$
$\hat{\pi} = 132$	$2 + \delta_{123}$	$0 + \delta_{132}$	$3 + \delta_{213}$	$2 + \delta_{231}$	$2 + \delta_{312}$	$3 + \delta_{321}$
$\hat{\pi} = 213$	$2 + \delta_{123}$	$3 + \delta_{132}$	$0 + \delta_{213}$	$2 + \delta_{231}$	$2 + \delta_{312}$	$3 + \delta_{321}$
$\hat{\pi} = 231$	$3 + \delta_{123}$	$2 + \delta_{132}$	$2 + \delta_{213}$	$0 + \delta_{231}$	$3 + \delta_{312}$	$2 + \delta_{321}$
$\hat{\pi} = 312$	$3 + \delta_{123}$	$2 + \delta_{132}$	$2 + \delta_{213}$	$3 + \delta_{231}$	$0 + \delta_{312}$	$2 + \delta_{321}$
$\hat{\pi} = 321$	$2 + \delta_{123}$	$3 + \delta_{132}$	$3 + \delta_{213}$	$2 + \delta_{231}$	$2 + \delta_{312}$	$0 + \delta_{321}$

Unfortunately, the number of permutations, π , grows factorially ($\mathcal{O}(n!)$) with the number of elements being matched (n). This makes explicit construction of the Lagrangian minimax game intractable for modestly-sized problems.

3.2 Optimization by Constraint Generation

Our first approach for taming the factorial computational complexity of explicitly constructing games for matching tasks is a constraint-generation approach known as the *double oracle method* (McMahan et al., 2003). It obtains the equilibrium solution to the adversarial prediction game without explicitly constructing the entire game matrix (Table 1). Based on the key observation that the equilibrium of the zero-sum game is typically supported by a relatively small number of permutations, it seeks to efficiently uncover this sparse set of permutations for each player.

Algorithm 1 Double Oracle Algorithm for Adversarial Bipartite Matching Equilibria.

Input: Lagrangian potentials $\Psi(\cdot)$; Initial label π_{initial}
Output: The (sparse) Nash equilibrium $(\hat{\mathcal{S}}, \hat{\mathcal{S}}, \hat{P}, \hat{P})$

- 1: $\hat{\mathcal{S}} \leftarrow \hat{\mathcal{S}} \leftarrow \{\pi_{\text{initial}}\}$
- 2: **repeat**
- 3: $(\hat{P}, \check{P}, \check{V}) \leftarrow \text{solveGame}(\Psi(\hat{\mathcal{S}}), \text{loss}_{\text{Ham}}(\hat{\mathcal{S}}, \hat{\mathcal{S}}))$
- 4: $(\tilde{\pi}_{\text{new}}, V_{\text{max}}) \leftarrow \arg\max_{\tilde{\pi}} \mathbb{E}_{\tilde{\pi} \sim \hat{P}} [\text{loss}_{\text{Ham}}(\hat{\pi}, \tilde{\pi}) + \Psi(\tilde{\pi})]$
- 5: **if** $(V \neq V_{\text{max}})$ **then** $\hat{\mathcal{S}} \leftarrow \hat{\mathcal{S}} \cup \tilde{\pi}_{\text{new}}$
- 6: $(\hat{P}, \check{P}, \check{V}) \leftarrow \text{solveGame}(\Psi(\hat{\mathcal{S}}), \text{loss}_{\text{Ham}}(\hat{\mathcal{S}}, \hat{\mathcal{S}}))$
- 7: $(\hat{\pi}_{\text{new}}, V_{\text{min}}) \leftarrow \arg\min_{\hat{\pi}} \mathbb{E}_{\hat{\pi} \sim \check{P}} [\text{loss}_{\text{Ham}}(\hat{\pi}, \tilde{\pi})]$
- 8: **if** $(V \neq V_{\text{min}})$ **then** $\hat{\mathcal{S}} \leftarrow \hat{\mathcal{S}} \cup \hat{\pi}_{\text{new}}$
- 9: **until** $\check{V} = V_{\text{max}} = \hat{V} = V_{\text{min}}$
- 10: **return** $(\hat{\mathcal{S}}, \hat{\mathcal{S}}, \hat{P}, \hat{P})$

Algorithm 1 produces this set of “active” permutations for each player, $\hat{\mathcal{S}}$ and $\check{\mathcal{S}}$ (subsets of rows and columns in Table 1), and the associated Nash equilibrium $(\hat{P}, \check{P})$. Starting from an initial permutation, π_{initial} (Line 1), it repeatedly obtains the Nash equilibrium solution $(\hat{P}, \check{P})$ with value $\hat{V}$ or $\check{V}$ for the zero-sum game defined only by permutations in $\hat{\mathcal{S}}$ and $\check{\mathcal{S}}$ (Lines 3 and 6). This is efficiently accomplished using a linear program (Von Neumann & Morgenstern, 1945). The algorithm then obtains the other player’s best response to either $\hat{P}$ or $\check{P}$ (Lines 4 and 7) with values V_{max} and V_{min} using the Kuhn-Munkres (Hungarian) algorithm in $\mathcal{O}(n^3)$ time for sets of size n . These best responses, $\tilde{\pi}_{\text{new}}$ and $\hat{\pi}_{\text{new}}$, are added to the set of active permutations (i.e., new rows or columns in the game matrix) if they have better values than the previous equilibrium values (Lines 5 and 8). This is repeated until no game value improvement exists for either player (Line 9), at which point a Nash equilibrium for the full game has been obtained.

We solve the convex optimization of Lagrange parameters θ in Eq. (3) using the results of Algorithm 1. We employ AdaGrad (Duchi et al., 2011) with the gradient calculated as the difference between expected features under the adversary’s distribution and the empirical training data: $\mathbb{E}_{x \sim \check{P}; \tilde{\pi} \sim \check{P}} [\sum_{i=1}^n \phi_i(x, \tilde{\pi}_i)] - \mathbb{E}_{x, \pi \sim \hat{P}} [\sum_{i=1}^n \phi_i(x, \pi_i)]$.

In contrast with SSVM, which compute the hinge loss for each training instance using only a single run of the Hungarian algorithm, our double oracle method must solve this problem repeatedly to find the equilibrium. Though in practice the total number of active permutations is much smaller than the $n!$ possibilities, no formal polynomial bound is known—and, consequentially, the run time of the approach as a whole cannot be characterized as polynomial.

3.3 Marginal Distribution Formulation

Our second approach, which significantly improves the efficiency of solving the adversarial bipartite matching game,

leverages the key insight that all quantities of interest for evaluating the loss and satisfying the constraints depend only on marginal probabilities of the permutation’s value assignments. Based on this, we employ a marginal distribution decomposition of the game.

Table 2. Doubly stochastic matrices $\mathbf{P}$ and $\mathbf{Q}$ for the marginal decompositions of each player’s mixture of permutations.

	1	2	3		1	2	3
$\hat{\pi}_1$	$p_{1,1}$	$p_{1,2}$	$p_{1,3}$	$\tilde{\pi}_1$	$q_{1,1}$	$q_{1,2}$	$q_{1,3}$
$\hat{\pi}_2$	$p_{2,1}$	$p_{2,2}$	$p_{2,3}$	$\tilde{\pi}_2$	$q_{2,1}$	$q_{2,2}$	$q_{2,3}$
$\hat{\pi}_3$	$p_{3,1}$	$p_{3,2}$	$p_{3,3}$	$\tilde{\pi}_3$	$q_{3,1}$	$q_{3,2}$	$q_{3,3}$

We begin this reformulation by first defining a matrix representation of permutation π as $\mathbf{Y}(\pi) \in \mathbb{R}^{n \times n}$ (or simply $\mathbf{Y}$) where the value of its cell $Y_{i,j}$ is 1 when $\pi_i = j$ and 0 otherwise. To be a valid complete bipartite matching or permutation, each column and row of $\mathbf{Y}$ can only have one entry of 1. For each feature function $\phi_i^{(k)}(x, \pi_i)$, we also denote its matrix representation as $\mathbf{X}_k$ whose (i, j) -th cell represents the k -th entry of $\phi_i(x, j)$. For a given distribution of permutations, $P(\pi)$, we denote the marginal probabilities of matching i with j as $p_{i,j} \triangleq P(\pi_i = j)$. We let $\mathbf{P} = \sum_{\pi} P(\pi) \mathbf{Y}(\pi)$ be the predictor’s marginal probability matrix where its (i, j) cell represents $\hat{P}(\hat{\pi}_i = j)$, and similarly let $\mathbf{Q}$ be the adversary’s marginal probability matrix (based on $\check{P}$), as shown in Table 2.

The Birkhoff–von Neumann theorem (Birkhoff, 1946; Von Neumann, 1953) states that the convex hull of the set of $n \times n$ permutation matrices forms a convex polytope in $\mathbb{R}^{n^2}$ (known as the Birkhoff polytope B_n) in which points are doubly stochastic matrices, i.e., the $n \times n$ matrices with non-negative elements where each row and column must sum to one. This implies that both marginal probability matrices $\mathbf{P}$ and $\mathbf{Q}$ are doubly stochastic matrices. In contrast to the space of distributions over permutation of n objects, which grows factorially ($\mathcal{O}(n!)$ with $n! - 1$ free parameters), the size of this marginal matrices grows only quadratically ($\mathcal{O}(n^2)$ with $n^2 - 2n$ free parameters). This provides a significant benefit in terms of the optimization.

Starting with the minimax over $\hat{P}(\hat{\pi})$ and $\check{P}(\tilde{\pi})$ in the permutation mixture formulation, and using the matrix notation above, we rewrite Eq. (3) as a minimax over marginal probability matrices $\mathbf{P}$ and $\mathbf{Q}$ with additional constraints that both $\mathbf{P}$ and $\mathbf{Q}$ are doubly-stochastic matrices, i.e., $\mathbf{P} \geq \mathbf{0}$ (elementwise), $\mathbf{Q} \geq \mathbf{0}$, $\mathbf{P}\mathbf{1} = \mathbf{P}^\top \mathbf{1} = \mathbf{Q}\mathbf{1} = \mathbf{Q}^\top \mathbf{1} = \mathbf{1}$ where $\mathbf{1} = (1, \dots, 1)^\top$. That is:

$$\begin{aligned} \min_{\theta} \mathbb{E}_{\mathbf{X}, \mathbf{Y} \sim \hat{P}} \min_{\mathbf{P} \geq \mathbf{0}} \max_{\mathbf{Q} \geq \mathbf{0}} [n - \langle \mathbf{P}, \mathbf{Q} \rangle + \langle \mathbf{Q} - \mathbf{Y}, \sum_k \theta_k \mathbf{X}_k \rangle] \\ \text{s.t. : } \mathbf{P}\mathbf{1} = \mathbf{P}^\top \mathbf{1} = \mathbf{Q}\mathbf{1} = \mathbf{Q}^\top \mathbf{1} = \mathbf{1}, \end{aligned} \quad (4)$$

where $\langle \cdot, \cdot \rangle$ denotes the Frobenius inner product between two matrices, i.e., $\langle A, B \rangle = \sum_{i,j} A_{i,j} B_{i,j}$.

3.3.1 OPTIMIZATION

We reduce the computational costs of the optimization in Eq. (4) by focusing on optimizing the adversary's marginal probability $\mathbf{Q}$. By strong duality, we then push the maximization over $\mathbf{Q}$ in the formulation above to the outermost level of Eq. (4). Note that the objective above is a non-smooth function (i.e., piece-wise linear). For the purpose of smoothing the objective, we add a small amount of strongly convex prox-functions to both $\mathbf{P}$ and $\mathbf{Q}$. We also add a regularization penalty to the parameter θ to improve the generalizability of our model. We unfold Eq. (4) by replacing the empirical expectation with an average over all training examples, resulting in the following optimization:

$$\begin{aligned} \max_{\mathbf{Q} \geq 0} \min_{\theta} \frac{1}{m} \sum_{i=1}^m \min_{\mathbf{P}_i \geq 0} & \left[\langle \mathbf{Q}_i - \mathbf{Y}_i, \sum_k \theta_k \mathbf{X}_{i,k} \rangle - \langle \mathbf{P}_i, \mathbf{Q}_i \rangle \right. \\ & \left. + \frac{\mu}{2} \|\mathbf{P}_i\|_F^2 - \frac{\mu}{2} \|\mathbf{Q}_i\|_F^2 \right] + \frac{\lambda}{2} \|\theta\|_2^2 \\ \text{s.t. : } & \mathbf{P}_i \mathbf{1} = \mathbf{P}_i^\top \mathbf{1} = \mathbf{Q}_i \mathbf{1} = \mathbf{Q}_i^\top \mathbf{1} = \mathbf{1}, \quad \forall i, \end{aligned} \quad (5)$$

where m is the number of bipartite matching problems in the training set, λ is the regularization penalty parameter, μ is the smoothing penalty parameter, and $\|A\|_F$ denotes the Frobenius norm of matrix A . The subscript i in $\mathbf{P}_i$, $\mathbf{Q}_i$, $\mathbf{X}_i$, and $\mathbf{Y}_i$ refers to the i -th example in the training set.

In the formulation above, given a fixed $\mathbf{Q}$, the inner minimization over θ and $\mathbf{P}$ can then be solved separately. The optimal θ in the inner minimization admits a closed-form solution, in which the k -th element of θ^* is:

$$\theta_k^* = -\frac{1}{\lambda m} \sum_{i=1}^m \langle \mathbf{Q}_i - \mathbf{Y}_i, \mathbf{X}_{i,k} \rangle. \quad (6)$$

The inner minimization over $\mathbf{P}$ can be solved independently for each training example. Given the adversary's marginal probability matrix $\mathbf{Q}_i$ for the i -th example, the optimal $\mathbf{P}_i$ can be formulated as:

$$\mathbf{P}_i^* = \underset{\{\mathbf{P}_i \geq 0 | \mathbf{P}_i \mathbf{1} = \mathbf{P}_i^\top \mathbf{1} = \mathbf{1}\}}{\operatorname{argmin}} \quad \frac{\mu}{2} \|\mathbf{P}_i\|_F^2 - \langle \mathbf{P}_i, \mathbf{Q}_i \rangle \quad (7)$$

$$= \underset{\{\mathbf{P}_i \geq 0 | \mathbf{P}_i \mathbf{1} = \mathbf{P}_i^\top \mathbf{1} = \mathbf{1}\}}{\operatorname{argmin}} \quad \|\mathbf{P}_i - \frac{1}{\mu} \mathbf{Q}_i\|_F^2. \quad (8)$$

We can interpret this minimization as projecting the matrix $\frac{1}{\mu} \mathbf{Q}_i$ to the set of doubly-stochastic matrices. We will discuss our projection technique in the upcoming subsection.

For solving the outer optimization over $\mathbf{Q}$ with the doubly-stochastic constraints, we employ a projected Quasi-Newton algorithm (Schmidt et al., 2009). Each iteration of the algorithm optimizes the quadratic approximation of the objective function (using limited-memory Quasi-Newton) over the the convex set. In each update step, a projection to the set of doubly-stochastic matrices is needed, akin to the inner minimization of $\mathbf{P}$ in Eq. (8).

The optimization above provides the adversary's optimal marginal probability $\mathbf{Q}^*$. To achieve our learning goal, we recover θ^* using Eq. (6) computed over the optimal $\mathbf{Q}^*$. We use the θ^* that our model learns from this optimization to construct a weighted bipartite graph for making predictions for test examples.

3.3.2 DOUBLY-STOCHASTIC MATRIX PROJECTION

The projection from an arbitrary matrix $\mathbf{R}$ to the set of doubly-stochastic matrices can be formulated as:

$$\min_{\mathbf{P} \geq 0} \|\mathbf{P} - \mathbf{R}\|_F^2, \quad \text{s.t. : } \mathbf{P} \mathbf{1} = \mathbf{P}^\top \mathbf{1} = \mathbf{1}. \quad (9)$$

We employ the alternating direction method of multipliers (ADMM) technique (Douglas & Rachford, 1956; Glowinski & Marroco, 1975; Boyd et al., 2011) to solve the optimization problem above. We divide the doubly-stochastic matrix constraint into two sets of constraints $C_1 : \mathbf{P} \mathbf{1} = \mathbf{1}$ and $\mathbf{P} \geq 0$, and $C_2 : \mathbf{P}^\top \mathbf{1} = \mathbf{1}$ and $\mathbf{P} \geq 0$. Using this construction, we convert the optimization above into ADMM form as follows:

$$\begin{aligned} \min_{\mathbf{P}, \mathbf{S}} & \frac{1}{2} \|\mathbf{P} - \mathbf{R}\|_F^2 + \frac{1}{2} \|\mathbf{S} - \mathbf{R}\|_F^2 + \mathcal{I}_{C_1}(\mathbf{P}) + \mathcal{I}_{C_2}(\mathbf{S}) \\ \text{s.t. : } & \mathbf{P} - \mathbf{S} = 0. \end{aligned} \quad (10)$$

The augmented Lagrangian for this optimization is:

$$\begin{aligned} \mathcal{L}_\rho(\mathbf{P}, \mathbf{S}, \mathbf{W}) = & \frac{1}{2} \|\mathbf{P} - \mathbf{R}\|_F^2 + \frac{1}{2} \|\mathbf{S} - \mathbf{R}\|_F^2 + \mathcal{I}_{C_1}(\mathbf{P}) \\ & + \mathcal{I}_{C_2}(\mathbf{S}) + \frac{\rho}{2} \|\mathbf{P} - \mathbf{S} + \mathbf{W}\|_F^2, \end{aligned} \quad (11)$$

where ρ is the ADMM penalty parameter and $\mathbf{W}$ is the scaled dual variable. From the augmented Lagrangian, we compute the update for $\mathbf{P}$ as:

$$\begin{aligned} \mathbf{P}^{t+1} = & \underset{\mathbf{P}}{\operatorname{argmin}} \mathcal{L}_\rho(\mathbf{P}, \mathbf{S}^t, \mathbf{W}^t) \\ = & \underset{\{\mathbf{P} \geq 0 | \mathbf{P} \mathbf{1} = \mathbf{1}\}}{\operatorname{argmin}} \quad \frac{1}{2} \|\mathbf{P} - \mathbf{R}\|_F^2 + \frac{\rho}{2} \|\mathbf{P} - \mathbf{S}^t + \mathbf{W}^t\|_F^2 \\ = & \underset{\{\mathbf{P} \geq 0 | \mathbf{P} \mathbf{1} = \mathbf{1}\}}{\operatorname{argmin}} \quad \|\mathbf{P} - \frac{1}{1+\rho} (\mathbf{R} + \rho (\mathbf{S}^t - \mathbf{W}^t))\|_F^2. \end{aligned} \quad (12)$$

The minimization above can be interpreted as a projection to the set $\{\mathbf{P} \geq 0 | \mathbf{P} \mathbf{1} = \mathbf{1}\}$ which can be realized by projecting to the probability simplex independently for each row of the matrix $\frac{1}{1+\rho} (\mathbf{R} + \rho (\mathbf{S}^t - \mathbf{W}^t))$. Similarly, the ADMM update for $\mathbf{S}$ can also be formulated as a column-wise probability simplex projection. The technique for projecting a point to the probability simplex has been studied previously, e.g., by Duchi et al. (2008). Therefore, our ADMM algorithm consists of the following updates:

$$\mathbf{P}^{t+1} = \operatorname{Proj}_{C_1} \left(\frac{1}{1+\rho} (\mathbf{R} + \rho (\mathbf{S}^t - \mathbf{W}^t)) \right) \quad (13)$$

$$\mathbf{S}^{t+1} = \operatorname{Proj}_{C_2} \left(\frac{1}{1+\rho} (\mathbf{R} + \rho (\mathbf{P}^{t+1} + \mathbf{W}^t)) \right) \quad (14)$$

$$\mathbf{W}^{t+1} = \mathbf{W}^t + \mathbf{P}^{t+1} - \mathbf{S}^{t+1}. \quad (15)$$

We run this series of updates until the stopping conditions are met. Our stopping conditions are based on the primal and dual residual optimality as described in Boyd et al. (2011). In our overall algorithm, this ADMM projection algorithm is used both in the projected Quasi-Newton algorithm for optimizing $\mathbf{Q}$ (Eq. (5)) and in the inner optimization for minimizing $\mathbf{P}_i$ (Eq. (8)).

3.3.3 CONVERGENCE PROPERTY

The convergence rate of ADMM is $\mathcal{O}(\log \frac{1}{\epsilon})$ thanks to the strong convexity of the objective (Deng & Yin, 2016). Each step inside ADMM is simply a projection to a simplex, hence costing $\tilde{\mathcal{O}}(n)$ computations (Duchi et al., 2008).

In terms of optimization on $\mathbf{Q}$, since no explicit rates of convergence are available for the projected Quasi-Newton algorithm (Schmidt et al., 2009) that finely characterize the dependency on the condition numbers, we simply illustrate the $\sqrt{L/\mu} \log \frac{1}{\epsilon}$ rate using Nesterov’s accelerated gradient algorithm (Nesterov, 2003), where L is the Lipschitz continuous constant of the gradient. In our case, $L = \frac{1}{m^2\lambda} \sum_k \sum_{i=1}^m \|\mathbf{X}_{i,k}\|_F^2 + 1/\mu$.

Comparison with Structured SVM (SSVM) Conventional SSVMs for learning bipartite matchings have only $\mathcal{O}(1/\epsilon)$ rates due to the lack of smoothness (Joachims et al., 2009; Teo et al., 2010). If smoothing is added, then similar linear convergence rates can be achieved with similar condition numbers. However, it is noteworthy that at each iteration we need to apply ADMM to solve a projection problem to the doubly stochastic matrix set (Eq. (9)), while SSVMs (without smoothing) solves a matching problem with the Hungarian algorithm, incurring $\mathcal{O}(n^3)$ time.

3.4 Consistency Analysis

Despite its apparent differences from standard empirical risk minimization (ERM), adversarial loss minimization (Eq. (3)) can be equivalently recast as an ERM:

$$\min_{\theta} \mathbb{E}_{x \sim P} \left[AL_{f_{\theta}}^{\text{perm}}(x, \pi) \right] \text{ where } AL_{f_{\theta}}^{\text{perm}}(x, \pi) \triangleq \min_{\hat{P}(\hat{\pi}|x)} \max_{\tilde{P}(\tilde{\pi}|x)} \mathbb{E}_{\hat{\pi} \sim \hat{P}} \left[\text{loss}(\hat{\pi}, \tilde{\pi}) + f_{\theta}(x, \tilde{\pi}) - f_{\theta}(x, \pi) \right]$$

and $f_{\theta}(x, \pi) = \theta \cdot \sum_{i=1}^n \phi(x, \pi_i)$ is the Lagrangian potential function. Here we consider f_{θ} as the linear discriminant function for a proposed permutation π , using parameter value θ . $AL_{f_{\theta}}^{\text{perm}}(x, \pi)$ is then the surrogate loss for input x and permutation π .

As described in Section 2.2, Fisher consistency is an important property for a surrogate loss L . It requires that under the true distribution $P(x, \pi)$, the hypothesis that minimizes L is Bayes optimal (Tewari & Bartlett, 2007; Liu,

2007). For the cases of multiclass classification and ordinal regression, Fisher consistency for adversarial surrogate loss has been established by Fathony et al. (2016; 2017). In our setting, the Fisher consistency of AL_f^{perm} can be written as:

$$f^* \in \mathcal{F}^* \triangleq \underset{f}{\operatorname{argmin}} \mathbb{E}_{\pi|x \sim P} \left[AL_f^{\text{perm}}(x, \pi) \right] \quad (16)$$

$$\Rightarrow \operatorname{argmax}_{\pi} f^*(x, \pi) \subseteq \Pi^{\diamond} \triangleq \operatorname{argmin}_{\pi} \mathbb{E}_{\pi|x \sim P} [\text{loss}(\pi, \bar{\pi})].$$

Note that in Eq. (16) we allow f to be optimized over the set of all measurable functions on the input space (x, π) . In our formulation, we have restricted f to be additively decomposable over individual elements of permutation, $f(x, \pi) = \sum_i g_i(x, \pi_i)$. In the sequel, we will show that the condition in Eq. (16) also holds for this restricted set provided that g is allowed to be optimized over the set of all measurable functions on the space of individual input (x, π_i) . We start by establishing Fisher consistency for the case of singleton loss minimizing sets $\Pi^{\diamond}$ in Theorem 1 and then for more general cases in Theorem 2.

Theorem 1. Suppose $\text{loss}(\pi, \bar{\pi}) = \text{loss}(\bar{\pi}, \pi)$ (symmetry) and $\text{loss}(\pi, \pi) < \text{loss}(\bar{\pi}, \pi)$ for all $\bar{\pi} \neq \pi$. Then the adversarial permutation loss AL_f^{perm} is Fisher consistent if f is over all measurable functions and $\Pi^{\diamond}$ is a singleton.

Theorem 2. Suppose $\text{loss}(\pi, \bar{\pi}) = \text{loss}(\bar{\pi}, \pi)$ (symmetry) and $\text{loss}(\pi, \pi) < \text{loss}(\bar{\pi}, \pi)$ for all $\bar{\pi} \neq \pi$. Furthermore if f is over all measurable functions, then:

- (a) there exists $f^* \in \mathcal{F}^*$ such that $\operatorname{argmax}_{\pi} f^*(x, \pi) \subseteq \Pi^{\diamond}$ (i.e., satisfies the Fisher consistency requirement). In fact, all elements in $\Pi^{\diamond}$ can be recovered by some $f^* \in \mathcal{F}^*$.
- (b) if $\operatorname{argmin}_{\pi} \sum_{\pi' \in \Pi^{\diamond}} \alpha_{\pi'} \text{loss}(\pi', \pi) \subseteq \Pi^{\diamond}$ for all $\alpha_{(\cdot)} \geq 0$; $\sum_{\pi' \in \Pi^{\diamond}} \alpha_{\pi'} = 1$, then $\operatorname{argmax}_{\pi} f^*(x, \pi) \subseteq \Pi^{\diamond}$ for all $f^* \in \mathcal{F}^*$. In this case, all $f^* \in \mathcal{F}^*$ satisfy the Fisher consistency requirement.

These assumptions of loss functions in the theorems above are quite mild, requiring only that wrong predictions suffer higher loss than correct ones. We refer the reader to Appendix B for the detailed proofs of theorems. The key to the proofs is the observation that for the optimal potential function f^* , $f^*(x, \pi) + \text{loss}(\pi, \pi^{\diamond})$ is invariant to π when $\Pi^{\diamond} = \{\pi^{\diamond}\}$. We refer to this as the *loss reflective property*. Note that this generalizes the observation for the case of ordinal regression loss (Fathony et al., 2017) into matching loss functions, subject to the mild pre-conditions assumed by the theorem.

Theorem 3. Suppose the loss is Hamming loss, and the potential function $f(x, \pi)$ decomposes additively by $\sum_i g_i(x, \pi_i)$. Then, the adversarial permutation loss AL_f^{perm} is Fisher consistent provided that g_i is allowed to

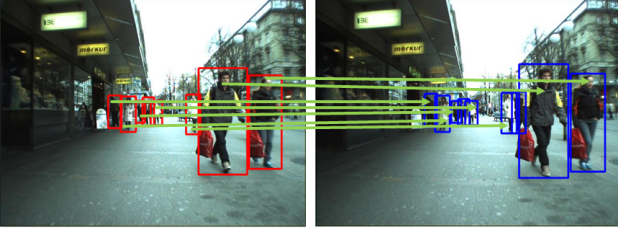


Figure 2. An example of bipartite matching in video tracking.

be optimized over the set of all measurable functions on the space of individual inputs (x, π_i) .

Proof. Simply choose g_i such that for each sample x in the population, $g_i(x, \pi_i) = -(\pi_i \neq \pi_i^\circ)$. This renders the loss reflective property under the Hamming loss. $\square$

4 Experimental Evaluation

To evaluate our approach, we apply our adversarial bipartite matching model to video tracking tasks using public benchmark datasets (Leal-Taixé et al., 2015). In this problem, we are given a set of images (video frames) and a list of objects in each image. We are also given the correspondence matching between objects in frame t and objects in frame $t + 1$. Figure 2 shows an example of the problem setup. It is important to note that the number of objects are not the same in every frames. Some of the objects may enter, leave, or remain in the consecutive frames. To handle the this issue, we setup our experiment as follows. Let k_t be the number of objects in frame t and k^* be the maximum number of objects a frame can have, i.e., $k^* = \max_{t \in T} k_t$. Starting from k^* nodes to represent the objects, we add k^* more nodes as “invisible” nodes to allow new objects to enter and existing objects to leave. As a result, the total number of nodes in each frame doubles to $n = 2k^*$.

4.1 Feature Representation

We define the features for pairs of bounding boxes (i.e., $\phi_i(x, j)$ for pairing bounding box i with bounding box j) in two consecutive video frames so that we can compute the associative feature vectors, $\phi(x, \pi) = \sum_{i=1}^n \phi_i(x, \pi_i)$, for each possible matching π . To define the feature vector $\phi_i(\cdot, \cdot)$, we follow the feature representation reported by Kim et al. (2012) using six different types of features:

- Intersection over union (IoU) overlap ratio between bounding boxes, $\text{area}(\text{BB}_i^t \cap \text{BB}_j^{t+1}) / \text{area}(\text{BB}_i^t \cup \text{BB}_j^{t+1})$, where BB_i^t denotes the bounding box of object i at time frame t ;
- Euclidean distance between object centers;
- 21 color histogram distance features (RGB) from the Bhattacharyaa distance, $\frac{1}{4} \ln \left(\frac{1}{4} \left(\frac{\sigma_p^2}{\sigma_q^2} + \frac{\sigma_q^2}{\sigma_p^2} + 2 \right) \right) +$

Table 3. Dataset properties

DATASET	# ELEMENTS	# EXAMPLES
TUD-CAMPUS	12	70
TUD-STADTMITTE	16	178
ETH-SUNNYDAY	18	353
ETH-BAHNHOF	34	999
ETH-PEDCROSS2	30	836

$\frac{1}{4} \left(\frac{(\mu_p - \mu_q)^2}{\mu_p^2 + \mu_q^2} \right)$, between distributions from the histograms of 7×3 blocks, in which p and q are two different distributions of the blocks at time frames t and $t + 1$, μ and σ^2 are the mean and the variance of the distribution respectively;

- 21 local binary pattern (LBP) features from similar Bhattacharyaa distances and bounding box blocks;
- Optical flow (motion) between bounding boxes; and
- Three indicator variables (for *entering*, *leaving*, and *staying invisible*).

We explain this feature representation in more detail in Appendix C.

4.2 Experimental Setup

We compare our approach with the Structured SVM (SSVM) model (Taskar et al., 2005a; Tsochantaridis et al., 2005) implemented based on Kim et al. (2012) using SVM-Struct (Joachims, 2008; Vedaldi, 2011). We implement our marginal version of adversarial bipartite matching using minConf (Schmidt, 2008) for performing projected Quasi-Newton optimization.

We consider two different groups of datasets in our experiment: TUD datasets and ETH datasets. Each dataset contains different numbers of elements (i.e., the number of pedestrian bounding box in the frame plus the number of extra nodes to indicate entering or leaving) and different numbers of examples (i.e., pairs of two consecutive frames that we want to match). Table 3 contains the detailed information about the datasets.

To avoid having test examples that are too similar with the training set, we train the models on one dataset and test the model on another dataset that has similar characteristics. In particular, we perform evaluations for every pair of datasets in TUD and ETH collections. This results in eight pairs of training/test datasets, as shown in Table 4.

To tune the regularization parameter (λ in adversarial matching, and C in SSVM), we perform 5-fold cross validation based on the training dataset only. The resulting best regularization parameter is used to train the model over all training examples to obtain parameters θ , which we then use to predict the matching for the testing data. For SSVM and the marginal version of adversarial matching, the pre-

Table 4. The mean and standard deviation (in parenthesis) of the average accuracy (1 - the average Hamming loss) for the adversarial bipartite matching model compared with Structured-SVM.

TRAINING/ TESTING	ADV DO	ADV MARG.	SSVM	ADV DO #PERM.
CAMPUS/ STADTMITTE	0.662 (0.09)	0.662 (0.08)	0.662 (0.08)	11.4
STADTMITTE/ CAMPUS	0.672 (0.12)	0.667 (0.11)	0.660 (0.12)	7.4
BAHNHOF/ SUNNYDAY	0.758 (0.12)	0.754 (0.10)	0.729 (0.15)	5.8
PEDCROSS2/ SUNNYDAY	0.760 (0.08)	0.750 (0.10)	0.736 (0.13)	8.2
SUNNYDAY/ BAHNHOF	0.755 (0.20)	0.751 (0.18)	0.739 (0.20)	9.8
PEDCROSS2/ BAHNHOF	0.760 (0.12)	0.763 (0.16)	0.731 (0.21)	10.8
BAHNHOF/ PEDCROSS2	0.718 (0.16)	0.714 (0.16)	0.701 (0.18)	8.5
SUNNYDAY/ PEDCROSS2	0.719 (0.18)	0.712 (0.17)	0.700 (0.18)	14.4

diction is done by finding the bipartite matching that maximizes the potential value, i.e., $\operatorname{argmax}_{\mathbf{Y}} \langle \mathbf{Y}, \sum_k \theta_k \mathbf{X}_k \rangle$ which can be solved using the Hungarian algorithm. The double oracle version of adversarial matching makes predictions by finding the most likely permutation from the predictor’s strategy in the equilibrium.

4.3 Results

We report the average accuracy, which in this case is defined as (1 - the average Hamming loss) over all examples in the testing dataset. Table 4 shows the mean and the standard deviation of our metric across different dataset pairs. We report the results for both the double-oracle (DO) and marginal (MARG) versions of the adversarial model. Our experiment indicates that both methods result in very similar values of θ . The slight advantage of the double-oracle version is caused by the difference in prediction techniques between the double-oracle (argmax over predictor’s equilibrium strategy) and marginal version (argmax over potentials). We also observe that the double-oracle approach requires only a small number of augmenting permutations to converge as shown in the last column (the average number of permutations) of Table 4. This indicates the sparseness of the set of permutations that support the equilibrium.

To compare with SSVM, we highlight (using bold font) the cases in which our result is better with statistical significance (under paired t-test with $\alpha < 0.05$) in Table 4. Compared with SSVM, our proposed adversarial matching outperforms SSVM in all pairs of datasets—with statistical

Table 5. Running time (in seconds) of the model for various number of elements n with fixed number of samples ($m = 50$)

DATASET	# ELEMENTS	ADV MARG.	SSVM
CAMPUS	12	1.96	0.22
STADTMITTE	16	2.46	0.25
SUNNYDAY	18	2.75	0.15
PEDCROSS2	30	8.18	0.26
BAHNHOF	34	9.79	0.31

significance on all six pairs of the ETH datasets and slightly better than SSVM on the TUD datasets. This suggests that our adversarial bipartite matching model benefits from its Fisher consistency property.

In terms of the running time, Table 5 shows that the marginal version of adversarial method is relatively fast. It only takes a few seconds to train until convergence in the case of 50 examples, with the number of elements varied up to 34. The running time grows roughly quadratically in the number of elements, which is natural since the size of the marginal probability matrices $\mathbf{P}$ and $\mathbf{Q}$ also grow quadratically in the number of elements. This shows that our approach is much more efficient than the CRF approach, which has a running time that is impractical even for small problems with 20 elements. The training time of SSVM is faster than the adversarial methods due to two different factors: (1) the inner optimization of SSVM can be solved using a single execution of the Hungarian algorithm compared with the inner optimization of adversarial method which requires ADMM optimization for projection to doubly stochastic matrix set; (2) different tools for implementation, i.e., C++ for SSVM and MATLAB for our method, which benefits the running time of SSVM. In addition, though the game size is relatively small, as indicated by the final column in Table 4, the double oracle version of adversarial method takes much longer to train compared to the marginal version.

5 Conclusions and Future Work

In this paper, we have presented an adversarial approach for learning bipartite matchings that is not only computationally efficient to employ but also provides Fisher consistency guarantees. We showed that these theoretical advantages translate into better empirical performance for our model compared with previous approaches. Our future work will explore matching problems with different loss functions and other graphical structures.

Acknowledgements

We thank our anonymous reviewers for their useful feedback and suggestions. This research was supported in part by NSF Grants RI-#1526379 and CAREER-#1652530.

References

- Ahonen, T., Hadid, A., and Pietikainen, M. Face description with local binary patterns: Application to face recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 28(12):2037–2041, 2006.
- Amini, M. R., Truong, T. V., and Goutte, C. A boosting algorithm for learning bipartite ranking functions with partially labeled data. In *Proceedings of the International ACM SIGIR Conference*, pp. 99–106. ACM, 2008.
- Asif, K., Xing, W., Behpour, S., and Ziebart, B. D. Adversarial cost-sensitive classification. In *Proceedings of the Conference on Uncertainty in Artificial Intelligence*, pp. 92–101, 2015.
- Beauchemin, S. S. and Barron, J. L. The computation of optical flow. *ACM Computing Surveys (CSUR)*, 27(3): 433–466, 1995.
- Behpour, S., Kitani, K. M., and Ziebart, B. D. Adversarially optimizing intersection over union for object localization tasks. *CoRR*, abs/1710.07735, 2017.
- Behpour, S., Xing, W., and Ziebart, B. D. ARC: Adversarial robust cuts for semi-supervised and multi-label classification. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pp. 2704–2711, 2018.
- Belongie, S., Malik, J., and Puzicha, J. Shape matching and object recognition using shape contexts. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 24(4):509–522, 2002.
- Birkhoff, G. Three observations on linear algebra. *Univ. Nac. Tacuman, Rev. Ser. A*, 5:147–151, 1946.
- Boyd, S. and Vandenberghe, L. *Convex Optimization*. Cambridge University Press, 2004.
- Boyd, S., Parikh, N., Chu, E., Peleato, B., Eckstein, J., et al. Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends® in Machine Learning*, 3(1):1–122, 2011.
- Chan, Y. S. and Ng, H. T. Maxsim: A maximum similarity metric for machine translation evaluation. *Proceedings of ACL-08: HLT*, pp. 55–62, 2008.
- Dellaert, F., Seitz, S. M., Thorpe, C. E., and Thrun, S. EM, MCMC, and chain flipping for structure from motion with unknown correspondence. *Machine Learning*, 50(1-2):45–71, 2003.
- Deng, W. and Yin, W. On the global and linear convergence of the generalized alternating direction method of multipliers. *Journal of Scientific Computing*, 66(3), 2016.
- Douglas, J. and Rachford, H. H. On the numerical solution of heat conduction problems in two and three space variables. *Transactions of the American Mathematical Society*, 82(2):421–439, 1956.
- Duchi, J., Shalev-Shwartz, S., Singer, Y., and Chandra, T. Efficient projections onto the l_1 -ball for learning in high dimensions. In *Proceedings of the International Conference on Machine Learning*, pp. 272–279. ACM, 2008.
- Duchi, J., Hazan, E., and Singer, Y. Adaptive subgradient methods for online learning and stochastic optimization. *Journal of Machine Learning Research*, 12(Jul):2121–2159, 2011.
- Dwork, C., Kumar, R., Naor, M., and Sivakumar, D. Rank aggregation methods for the web. In *Proceedings of the International Conference on World Wide Web*, pp. 613–622. ACM, 2001.
- Fathony, R., Liu, A., Asif, K., and Ziebart, B. Adversarial multiclass classification: A risk minimization perspective. In *Advances in Neural Information Processing Systems*, pp. 559–567, 2016.
- Fathony, R., Bashiri, M. A., and Ziebart, B. Adversarial surrogate losses for ordinal regression. In *Advances in Neural Information Processing Systems*, pp. 563–573, 2017.
- Glowinski, R. and Marroco, A. Sur l’approximation, par éléments finis d’ordre un, et la résolution, par pénalisation-dualité d’une classe de problèmes de Dirichlet non linéaires. *Revue française d’automatique, informatique, recherche opérationnelle. Analyse numérique*, 9(R2):41–76, 1975.
- Grünwald, P. D. and Dawid, A. P. Game theory, maximum entropy, minimum discrepancy, and robust Bayesian decision theory. *Annals of Statistics*, 32:1367–1433, 2004.
- Huber, M. and Law, J. Fast approximation of the permanent for very dense problems. In *Proceedings of the nineteenth annual ACM-SIAM symposium on Discrete algorithms*, pp. 681–689. Society for Industrial and Applied Mathematics, 2008.
- Joachims, T. SVM-struct: Support vector machine for complex outputs. http://www.cs.cornell.edu/People/tj/svm_light/svm_struct.html, 2008.
- Joachims, T., Finley, T., and Yu, C.-N. Cutting-plane training of structural SVMs. *Machine Learning*, 77(1):27–59, 2009.
- Kim, S., Kwak, S., Feyereisl, J., and Han, B. Online multi-target tracking by large margin structured learning. In *Asian Conference on Computer Vision*, pp. 98–111. Springer, 2012.
- Kuhn, H. W. The hungarian method for the assignment problem. *Naval Research Logistics*, 2(1-2):83–97, 1955.
- Lafferty, J., McCallum, A., and Pereira, F. Conditional random fields: Probabilistic models for segmenting and labeling sequence data. In *Proc. of the International Conference on Machine Learning*, pp. 282–289, 2001.
- Le, Q. and Smola, A. Direct optimization of ranking measures. *arXiv preprint arXiv:0704.3359*, 2007.
- Leal-Taixé, L., Milan, A., Reid, I., Roth, S., and Schindler,

- K. Motchallenge 2015: Towards a benchmark for multi-target tracking. *arXiv preprint arXiv:1504.01942*, 2015.
- Li, J., Asif, K., Wang, H., Ziebart, B. D., and Berger-Wolf, T. Y. Adversarial sequence tagging. In *International Joint Conference on Artificial Intelligence*, 2016.
- Liu, L., Sun, L., Rui, Y., Shi, Y., and Yang, S. Web video topic discovery and tracking via bipartite graph reinforcement model. In *Proceedings of the 17th International Conference on World Wide Web*, pp. 1009–1018. ACM, 2008.
- Liu, Y. Fisher consistency of multicategory support vector machines. In *International Conference on Artificial Intelligence and Statistics*, pp. 291–298, 2007.
- MacCartney, B., Galley, M., and Manning, C. D. A phrase-based alignment model for natural language inference. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing*, pp. 802–811. Association for Computational Linguistics, 2008.
- McMahan, H. B., Gordon, G. J., and Blum, A. Planning in the presence of cost functions controlled by an adversary. In *Proceedings of the International Conference on Machine Learning*, pp. 536–543, 2003.
- Nesterov, Y. *Introductory Lectures on Convex Optimization: A Basic Course*. Springer, 2003.
- Padó, S. and Lapata, M. Optimal constituent alignment with edge covers for semantic projection. In *Proceedings of the International Conference on Computational Linguistics*, pp. 1161–1168. Association for Computational Linguistics, 2006.
- Petterson, J., Yu, J., McAuley, J. J., and Caetano, T. S. Exponential family graph matching and ranking. In *Advances in Neural Information Processing Systems*, pp. 1455–1463, 2009.
- Rui, X., Li, M., Li, Z., Ma, W.-Y., and Yu, N. Bipartite graph reinforcement model for web image annotation. In *Proceedings of the 15th ACM International Conference on Multimedia*, pp. 585–594. ACM, 2007.
- Schmidt, M. minConf: projection methods for optimization with simple constraints in Matlab. <http://www.cs.ubc.ca/~schmidtm/Software/minConf.html>, 2008.
- Schmidt, M., Berg, E., Friedlander, M., and Murphy, K. Optimizing costly functions with simple constraints: A limited-memory projected quasi-Newton algorithm. In *Artificial Intelligence and Statistics*, pp. 456–463, 2009.
- Sion, M. On general minimax theorems. *Pacific Journal of mathematics*, 8(1):171–176, 1958.
- Taskar, B., Chatalbashev, V., Koller, D., and Guestrin, C. Learning structured prediction models: A large margin approach. In *Proceedings of the International Conference on Machine Learning*, pp. 896–903. ACM, 2005a.
- Taskar, B., Lacoste-Julien, S., and Klein, D. A discriminative matching approach to word alignment. In *Conference on Human Language Technology and Empirical Methods in Natural Language Processing*, pp. 73–80. Association for Computational Linguistics, 2005b.
- Taylor, W. R. Protein structure comparison using bipartite graph matching and its application to protein structure classification. *Molecular & Cellular Proteomics*, 1(4): 334–339, 2002.
- Teo, C. H., Vishwanathan, S. V. N., Smola, A. J., and Le, Q. V. Bundle methods for regularized risk minimization. *Journal of Machine Learning Research*, 11:311–365, January 2010.
- Tewari, A. and Bartlett, P. On the consistency of multiclass classification methods. *Journal of Machine Learning Research*, 8:1007–1025, 2007.
- Topsøe, F. Information-theoretical optimization techniques. *Kybernetika*, 15(1):8–27, 1979.
- Tsochantaridis, I., Joachims, T., Hofmann, T., and Altun, Y. Large margin methods for structured and interdependent output variables. *Journal of Machine Learning Research*, 6(Sep):1453–1484, 2005.
- Valiant, L. G. The complexity of computing the permanent. *Theoretical Computer Science*, 8(2):189–201, 1979.
- Vedaldi, A. A MATLAB wrapper of SVM^{struct}. <http://www.vlfeat.org/~vedaldi/code/svm-struct-matlab.>, 2011.
- Volkovs, M. and Zemel, R. S. Efficient sampling for bipartite matching problems. In *Advances in Neural Information Processing Systems*, pp. 1313–1321, 2012.
- Von Neumann, J. A certain zero-sum two-person game equivalent to the optimal assignment problem. *Contributions to the Theory of Games*, 2:5–12, 1953.
- Von Neumann, J. and Morgenstern, O. Theory of games and economic behavior. *Bull. Amer. Math. Soc.*, 51(7): 498–504, 1945.
- Wang, H., Xing, W., Asif, K., and Ziebart, B. Adversarial prediction games for multivariate losses. In *Advances in Neural Information Processing Systems*, pp. 2710–2718, 2015.
- Wang, X.-Y., Wu, J.-F., and Yang, H.-Y. Robust image retrieval based on color histogram of local feature regions. *Multimedia Tools and Applications*, 49(2):323–345, 2010.
- Wang, Y., Makedon, F., Ford, J., and Huang, H. A bipartite graph matching framework for finding correspondences between structural elements in two proteins. In *International Conference of the IEEE Engineering in Medicine and Biology Society*, pp. 2972–2975, 2004.
- Zhu, J., Hoi, S. C., Lyu, M. R., and Yan, S. Near-duplicate keyframe retrieval by nonrigid image matching. In *Proceedings of the 16th ACM international conference on Multimedia*, pp. 41–50. ACM, 2008.