

Article

An Online Semi-Definite Programming with a Generalized Log-Determinant Regularizer and Its Applications [†]

Yaxiong Liu ^{1,2,*} , Ken-ichiro Moridomi ³, Kohei Hatano ^{1,2}  and Eiji Takimoto ¹ 

¹ Department of Informatics, Kyushu University, Fukuoka 819-0395, Japan; hatano@inf.kyushu-u.ac.jp (K.H.); eiji@inf.kyushu-u.ac.jp (E.T.)

² AIP RIKEN, Tokyo 103-0027, Japan

³ SMN Corporation, Tokyo 141-0032, Japan; kenichiro_moridomi@so-netmedia.jp

* Correspondence: yaxiong.liu@inf.kyushu-u.ac.jp

[†] This paper is an extended version of our paper published in 2021 Proceedings of the 13th Asian Conference on Machine Learning, 17–19 November 2021, PMLR 157:1113–1128.

Abstract: We consider a variant of the online semi-definite programming problem (OSDP). Specifically, in our problem, the setting of the decision space is a set of positive semi-definite matrices constrained by two norms in parallel: the L_∞ norm to the diagonal entries and the Γ -trace norm, which is a generalized trace norm with a positive definite matrix Γ . Our setting recovers the original one when Γ is an identity matrix. To solve this problem, we design a follow-the-regularized-leader algorithm with a Γ -dependent regularizer, which also generalizes the log-determinant function. Next, we focus on online binary matrix completion (OBMC) with side information and online similarity prediction with side information. By reducing to the OSDP framework and applying our proposed algorithm, we remove the logarithmic factors in the previous mistake bound of the above two problems. In particular, for OBMC, our bound is optimal. Furthermore, our result implies a better offline generalization bound for the algorithm, which is similar to those of SVMs with the best kernel, if the side information is involved in advance.

Keywords: online semi-definite programming; log-determinant; sparse loss matrix; side information; online binary matrix completion

MSC: 68W27



Citation: Liu, Y.; Moridomi, K.-i.; Hatano, K.; Takimoto, E. An Online Semi-Definite Programming with a Generalized Log-Determinant Regularizer and Its Applications.

Mathematics **2022**, *10*, 1055. <https://doi.org/10.3390/math10071055>

Academic Editors: Adrian Deaconu, Petru Adrian Cotfas and Daniel Tudor Cotfas

Received: 21 February 2022

Accepted: 24 March 2022

Published: 25 March 2022

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2022 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Online binary matrix completion (OBMC) is standing on the frontier of research to online matrix completion, which is currently an active field in the machine learning community [1–4]. Intuitively, the OBMC problem is a sequential game of predicting the given entries from an unknown $m \times n$ target binary matrix. More specifically, the problem can be formulated as a repeated game between the algorithm and the adversarial environment as described below: on each round t , (i) the environment confirms the location of an entry in the target matrix $(i_t, j_t) \in [m] \times [n]$, (ii) the algorithm predicts $\hat{y}_t \in \{-1, 1\}$, and then (iii) the environment reveals the true label $y_t \in \{-1, 1\}$. The goal of the algorithm is to minimize the total number of mistakes $\sum_{t=1}^T \mathbb{I}_{\hat{y}_t \neq y_t}$. This OBMC model is widely applied in the real world, such as “Netflix Challenges” [5], where the rating matrix (target matrix) is composed of rows representing the viewers and the columns corresponding to the movies. The entry (i, j) is the rating of the viewer i to the movie j , concretely.

For convenience, we define an underlying matrix U , the comparator matrix, as a good enough approximation matrix to the unknown target $m \times n$ matrix. To be more precise, assume that $U \in \mathbb{R}^{m \times n}$ can be factorized into $U = PQ^\top$ for some matrices $P \in \mathbb{R}^{n \times d}$ and $Q \in \mathbb{R}^{m \times d}$ for some $d \geq 1$. Without loss of generality, we further assume that the rows of P and Q are normalized such that $\|P_i\| = \|Q_j\| = 1$ for all i and j , where P_i is the i -th row vector of P (interpreted as a linear classifier associated with row i of U) and Q_j is the j -th

row vector of $\mathbf{Q}$ (interpreted as a feature vector associated with column j of $\mathbf{U}$). Hence, the sign of $\mathbf{U}_{i,j}$ can be viewed as a classification of the classifier $\mathbf{P}_i$ to the feature $\mathbf{Q}_j$. Next, to quantify the predictiveness of the comparator matrix $\mathbf{U}$, we involve the hinge loss as $[1 - z_t/\gamma]_+$ for a given margin parameter $\gamma > 0$, where $[x]_+$ is x if $x > 0$ and 0 otherwise. Note that $z_t = \mathbf{y}_t \mathbf{U}_{i_t, j_t}$ can be considered as the margin of the labeled instance $(\mathbf{Q}_{j_t}, \mathbf{y}_t)$ with respect to a hyperplane $\mathbf{P}_{i_t}$. On the other hand, the hinge loss converges to the 0–1 loss function, if γ converges to 0.

Recently, Herbster et al. explored the OMBC problem by adding side information in advance to the algorithm [2]. The side information brings some prior knowledge about the target matrix, or more generally, about a comparator matrix $\mathbf{U}$. Moreover, side information is formally represented according to the columns and the rows of the comparator matrix as two symmetric positive definite matrices $\mathbf{M} \in \mathbb{R}^{m \times m}$ and $\mathbf{N} \in \mathbb{R}^{n \times n}$. To measure the quality of the side information, Herbster et al. involved a concept *quasi-dimension* of a comparator matrix $\mathbf{U}$, defined as the minimum of $\mathcal{D} = \text{Tr}(\mathbf{P}^\top \mathbf{M} \mathbf{P}) + \text{Tr}(\mathbf{Q}^\top \mathbf{N} \mathbf{Q})$ over all the factorizations of $\mathbf{U}$ such that $\mathbf{U} = \gamma \mathbf{P} \mathbf{Q}^\top$. Then, they proved a mistake bound given by the total hinge loss of $\mathbf{U}$ with an additional term expressed in terms of γ , m , n , and $\mathcal{D}$. In particular, Herbster et al. offered a bound $O(\mathcal{D} \ln(m+n)/\gamma^2)$, if the total hinge loss of $\mathbf{U}$ is zero (in the *realizable case*). Moreover, they obtained a mistake bound $O(kl \ln(m+n))$, when $\mathbf{U}$ has a (k, l) -biclustered structure (see Appendix A for details) and side information $\mathbf{M}, \mathbf{N}$ are in accordance with this particular structure. Unfortunately, however, there still remains a logarithmic gap from a lower bound of $\Omega(kl)$ [1].

In this paper, unlike the definition of quasi-dimension introduced by Herbster et al., we simplify the *quasi-dimension* in the following part as $\mathcal{D} = \text{Tr}(\mathbf{P}^\top \mathbf{M} \mathbf{P}) + \text{Tr}(\mathbf{Q}^\top \mathbf{N} \mathbf{Q})$, only the sum of the trace norms. Then, we obtain a mistake bound $O(\mathcal{D}/\gamma^2)$, by improving a logarithmic factor $\ln(m+n)$ in the mistake bound of Herbster et al.; further, our bound recovers the lower bound at most by a constant when the comparator matrix has a (k, l) -biclustered structure. The basic idea is to reduce the OBMC problem with side information to an online semi-definite programming (OSDP) problem specified by $\mathbf{\Gamma}$. In particular, the symmetric positive definite matrix $\mathbf{\Gamma}$ is transformed from the side information $(\mathbf{M}, \mathbf{N})$ in our reduction. Thus, the reduced OSDP problem is proposed as a repeated game based on a sparse loss matrices space and a decision space, which is a set of symmetric and positive semi-definite matrices $\mathbf{W}$. Note that our decision set is constrained by the L_∞ -norm of the diagonal entries of $\mathbf{W}$ and $\mathbf{\Gamma}$ -trace norm, $\text{Tr}(\mathbf{\Gamma} \mathbf{W} \mathbf{\Gamma})$, simultaneously. Actually, our reduced OSDP problem is a generalization of the standard OSDP problem [6,7], where in the standard form, $\mathbf{\Gamma} = \mathbf{E}$. We design and analyze our algorithm for the generalized OSDP problem under the follow-the-regularized-leader (FTRL) framework (see, e.g., [8–10]). Note that to guarantee good performance of the proposed algorithm, we choose a specialized regularizer as stated later.

The OSDP framework/problem solved by the FTRL approaching is a classical method for various problems of online matrix prediction, such as online gambling [6,11], online collaborative filtering [12–14], online similarity prediction [15], and especially a *non-binary version* of online matrix completion with no side information [6,7]. To measure the performance of the algorithm for the above problems, a new concept, *regret*, the difference between the cumulative loss of the algorithm and the global optimal comparator matrix in hindsight, is always involved.

For the aforementioned results about non-binary online matrix completion with no side information, Hazan et al. [6] firstly proposed a reduction to the standard OSDP problem, and then, they utilized the FTRL algorithm with an *entropic regularizer*, obtaining a sub-optimal regret bound. Moridomi et al. [7] improved the regret bound by deploying a *log-determinant regularizer*, while they noticed that the loss matrices are sparse in the reduction.

Next, for the OMBC problem with side information, Herbster et al. [2] reduced this problem to another variant of the OSDP problem with different or fewer constraints to the decision space than ours. Then they followed a similar FTRL-based algorithm with the entropic regularizer as [6]. Note that instead of a general regret analysis for their OSDP

problem, they only gave a particular analysis to the OSDP problem instance obtained from the reduction. Due to the inspiration of the work of [7], the gap of the mistake bound is from the choice of the entropic regularizer. As same as the case mentioned in [7], in our reduction, the loss matrices are sparse, which implies that the log-determinant regularizer can lead to better performance.

As mentioned previously, the OBMC problem with side information is reduced to an OSDP problem, where the decision space is parameterized with the side information. It requests a new form of the log-determinant regularizer since the standard log-determinant regularizer $R(W) = -\ln \det(W + \epsilon E)$ performs unsatisfactorily from our examination. Meanwhile, we attempt to reduce our problem to the standard OSDP problem in a straightforward and natural way. Unfortunately, this reduction fails, as we will show a counterexample in the latter section since this trivial reduction damages the sparsity of the loss matrices and the bound to the diagonal entries of the decision space. Therefore, to solve our reduced OSDP problem, a generalized log-determinant regularizer, $R(W) = -\ln \det(\Gamma W \Gamma + \epsilon E)$, is required and this specified regularizer assists a successful regret bound. Conclusively, our reduction and solution not only show the power of the not well explored log-determinant regularizer compared with the entropic or Frobenius-norm regularizer in the OSDP framework, but also demonstrate that an appropriate choice of the regularizer, which is dependent on the form of the decision set, further on the side information, and the loss space can effectively improve the performance of the algorithm in theory. Note that although our derivation is similar to the analysis of Moridomi et al. [7], it is in fact a non-trivial generalization.

Furthermore, we apply our online algorithm in the statistical (batch) learning setting by the standard online-to-batch conversion framework (see, for example, the work of Mohri et al. [16]) and derive a generalization error bound with side information. Our generalized error bound is similar to the known margin-based bound of SVMs (e.g., Mohri et al. [16]) with the best kernel when the side information is vacuous. It is remarkable that we can not only obtain such a bound without knowing the best kernel, but it also implies that the error bound in the batch learning setting can be improved when the side information is given to the learner in advance.

Our main contribution is summarized as follows:

1. Firstly, we generalize the OSDP problem by parameterizing a symmetric and positive definite matrix Γ in the decision set. This generalization definitely extends the standard OSDP problem and offers a more widely applicable framework. Next, we design an FTRL-based algorithm with a generalized log-determinant regularizer depending on the matrix Γ from the decision set. Our result recovers the previously known bound [7] in the case where Γ is the identity matrix.
2. We obtain refined mistake bounds to the OBMC problem with side information and the online similarity prediction with side information under the umbrella of the above results. We reduce these problems to the OSDP framework and parameterize the side information into a symmetric positive definite matrix in the setting of decision space. Compared with the analysis of Herbster et al. [2], our reduction is explicit and easy to follow. Due to the achievement of the generalized OSDP problems, we improve the previously known mistake bounds by logarithmic factors for both of the problems. In particular, for the former problem, our mistake bound is optimal.
3. In addition, we added the online–offline conversion method to the OBMC problem with side information compared with the preliminary version [17]. We offer a standard online-to-batch conversion framework, which guarantees that our online algorithm can perform not worse than the traditional offline algorithm (SVM with the best kernel). As we demonstrated in the following section, our error bound recovers the best margin-based bound when the side information is vacuous. With the assistance of the ideal side information, our proposed algorithm performs even better than the previous one. On the one hand, we improve the error bound to the OBMC with side information in batch setting; on the other hand, our result implies that there might be

a more effective algorithm for batch setting if the algorithm can make good use of the side information.

This paper is organized as follows. In Section 2, we give some basic notations and formally formulate the generalized OSDP. Then with a toy example, we show that a naive reduction from our problem to the standard OSDP can even lead to a worse regret bound. It yields the necessity of our generalized log-determinant regularizer. The main algorithm with its regret bound for the generalized OSDP is given in Section 3. In Section 4, we give the reductions of the OBMC problem and online similarity prediction with side information to our demonstrated OSDP problems, respectively. Moreover, we show that the mistake bound for the OBMC problem is optimal in the realizable case where the comparator matrix has a biclustered structure. In Section 5, we derive the batch setting to the OBMC problem with side information. In Appendix A.1, we describe some necessary lemmata for proof of our results. Further, we define the (k, l) -biclustered structure in Appendix A.2.

2. Preliminaries

For any positive integer T , a subset of $\mathbb{N}$, $\{1, 2, \dots, T\} \subseteq \mathbb{N}$ is denoted by $[T]$. Let $\mathbb{S}^{N \times N}$, $\mathbb{S}_+^{N \times N}$ and $\mathbb{S}_{++}^{N \times N}$ denote the sets of $N \times N$ symmetric matrices, symmetric positive semi-definite matrices and symmetric strictly positive definite matrices, respectively. The identity matrix is denoted as E . For an $m \times n$ matrix $X \in \mathbb{R}^{m \times n}$ and $(i, j) \in [m] \times [n]$, we denote the i -th row vector of X and the (i, j) entry of X by X_i and $X_{i,j}$ respectively. Furthermore, $\text{vec}(X)$ is an mn -dimensional vector according to X and arranges all entries $X_{i,j}$, $i \in [m]$ and $j \in [n]$, in some order. For any matrices $X, Y \in \mathbb{R}^{m \times n}$, $X \bullet Y = \text{Tr}(X^\top Y) = \text{vec}(X)^\top \text{vec}(Y)$ denotes the Frobenius inner product of them. The trace norm of matrix $X \in \mathbb{S}_+^{N \times N}$ is defined as $\text{Tr}(X) = \sum_{i=1}^N |\lambda_i(X)|$, where $\lambda_i(X)$ denotes the i -th largest eigenvalue of X . Meanwhile, $\text{Tr}(X) = \sum_{i=1}^N X_{i,i}$. In addition, we generalize the trace norm of matrix X to the Γ -trace norm of X as $\text{Tr}(\Gamma X \Gamma)$, for some $\Gamma \in \mathbb{S}_{++}^{N \times N}$. Note that the Γ -trace norm recovers the trace norm when $\Gamma = E$. For a vector x , the L_p -norm of x is denoted by $\|x\|_p$.

2.1. Generalized OSDP Problem with Bounded Γ -Trace Norm

Our generalized OSDP problem specified with a symmetric positive definite matrix $\Gamma \in \mathbb{S}_{++}^{N \times N}$ is formulated by a pair $(\mathcal{K}, \mathcal{L})$, where

$$\mathcal{K} = \{W \in \mathbb{S}_+^{N \times N} : \text{Tr}(\Gamma W \Gamma) \leq \tau \wedge \|\text{vec}(W)\|_\infty \leq \beta\} \quad (1)$$

is called the decision space/set, and

$$\mathcal{L} = \{L \in \mathbb{S}^{N \times N} : \|\text{vec}(L)\|_1 \leq g\} \quad (2)$$

is called the loss space, where $\tau > 0$, $\beta > 0$ and $g > 0$ are parameters. The generalized OSDP problem $(\mathcal{K}, \mathcal{L})$ is a repeated game between the algorithm and the adversarial environment as described below: On each round $t \in [T]$, we have the following:

1. The algorithm predicts a matrix $W_t \in \mathcal{K}$.
2. The algorithm receives a loss matrix $L_t \in \mathcal{L}$, returning from the environment.
3. The algorithm incurs the loss on round t : $W_t \bullet L_t$.

The goal of the algorithm is to minimize the following regret

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) = \sum_{t=1}^T W_t \bullet L_t - \min_{W \in \mathcal{K}} \sum_{t=1}^T W \bullet L_t. \quad (3)$$

Note that the standard OSDP problem corresponds to the special case of our setting that $\Gamma = E$.

Due to the definition of the OSDP above, this problem is categorized to online linear optimization framework, since the convexity of the decision space and the Frobenius production is linear. Therefore, as Moridomi et al. [7] did for the standard OSDP problem,

we can apply a standard FTRL algorithm. The FTRL algorithm outputs a matrix $\mathbf{W}_t$ according to

$$\mathbf{W}_t = \arg \min_{\mathbf{W} \in \mathcal{K}} \left(R(\mathbf{W}) + \eta \sum_{s=1}^{t-1} \mathbf{L}_s \bullet \mathbf{W} \right), \quad (4)$$

where $R : \mathcal{K} \rightarrow \mathbb{R}$ is a strongly convex function, and called *regularizer*. Entropy and Euclidean norm are classical regularizers in online linear optimization [9]. In particular, Moridomi et al. chose the log-determinant regularizer, which is not well studied, defined as

$$R(\mathbf{W}) = -\ln \det(\mathbf{W} + \epsilon \mathbf{E}), \quad (5)$$

where $\epsilon > 0$ is a parameter and derived the following regret bound for the standard OSDP problem.

Theorem 1 ([7]). *For the standard OSDP problem $(\mathcal{K}, \mathcal{L})$ with $\mathbf{\Gamma} = \mathbf{E}$, The FTRL algorithm with the log-determinant regularizer achieves*

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) = O(g\sqrt{\tau\beta T}). \quad (6)$$

In the next sub-section, we show that the standard log-determinant regularizer performs unsatisfactorily to our generalized OSDP, even if our generalized OSDP can be reduced to the standard OSDP.

2.2. A Naive Reduction

Define a generalized OSDP $(\mathcal{K}, \mathcal{L})$ as in Equations (1) and (2), naturally, there is a reduction to a standard OSDP problem $(\tilde{\mathcal{K}}, \tilde{\mathcal{L}})$ where

$$\tilde{\mathcal{K}} = \{\tilde{\mathbf{W}} \in \mathbb{S}_+^{N \times N} : \text{Tr}(\tilde{\mathbf{W}}) \leq \tau \wedge \|\text{vec}(\tilde{\mathbf{W}})\|_\infty \leq \beta'\}, \quad \tilde{\mathcal{L}} = \{\tilde{\mathbf{L}} \in \mathbb{S}^{N \times N} : \|\text{vec}(\tilde{\mathbf{L}})\|_1 \leq g'\},$$

for some parameters $\beta' > 0$ and $g' > 0$. For convenience, we denote $\mathcal{A}$ as a FTRL-based algorithm with the log-determinant regularizer.

The reduction consists of two transformations: one is to transform the decision matrix $\tilde{\mathbf{W}}_t \in \tilde{\mathcal{K}}$ produced from $\mathcal{A}$ to the decision matrix $\mathbf{W}_t = \mathbf{\Gamma}^{-1} \tilde{\mathbf{W}}_t \mathbf{\Gamma}^{-1}$ for OSDP $(\mathcal{K}, \mathcal{L})$. The other one is to transform the loss matrices $\mathbf{L}_t \in \mathcal{L}$ from the environment of $(\mathcal{K}, \mathcal{L})$ to $\tilde{\mathbf{L}}_t = \mathbf{\Gamma}^{-1} \mathbf{L}_t \mathbf{\Gamma}^{-1}$, which is fed to the algorithm $\mathcal{A}$. Note that the loss is preserved under this reduction, that is, $\mathbf{W}_t \bullet \mathbf{L}_t = \text{Tr}(\mathbf{W}_t \mathbf{L}_t) = \text{Tr}(\mathbf{\Gamma}^{-1} \tilde{\mathbf{W}}_t \mathbf{\Gamma}^{-1} \mathbf{\Gamma} \tilde{\mathbf{L}}_t \mathbf{\Gamma}) = \text{Tr}(\tilde{\mathbf{W}}_t \tilde{\mathbf{L}}_t) = \tilde{\mathbf{W}}_t \bullet \tilde{\mathbf{L}}_t$. Moreover, the $\mathbf{\Gamma}$ -trace norm of $\mathbf{W}_t$ is the trace norm of $\tilde{\mathbf{W}}_t$, i.e., $\text{Tr}(\mathbf{\Gamma} \mathbf{W}_t \mathbf{\Gamma}) = \text{Tr}(\tilde{\mathbf{W}}_t)$. Therefore, setting β' and g' appropriately such that for any $\mathbf{W} \in \mathcal{K}$ and $\mathbf{L} \in \mathcal{L}$, $\mathbf{\Gamma} \mathbf{W} \mathbf{\Gamma} \in \tilde{\mathcal{K}}$ and $\mathbf{\Gamma}^{-1} \mathbf{L} \mathbf{\Gamma}^{-1} \in \tilde{\mathcal{L}}$, respectively, we have that $\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) \leq \text{Regret}_{\text{OSDP}}(T, \tilde{\mathcal{K}}, \tilde{\mathcal{L}})$.

Hence, according to the regret bound to $(\tilde{\mathcal{K}}, \tilde{\mathcal{L}})$ with log-determinant regularizer [7], we immediately have

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) = O(g'\sqrt{\tau\beta'T}).$$

by Theorem 1.

In the following part, we give an example, which implies that the above reduction yields a worse regret bound than our proposed algorithm.

Example 1. Define $\mathbf{\Gamma} \in \mathbb{S}_{++}^{N \times N}$ as

$$\mathbf{\Gamma} = \begin{bmatrix} N & -1 & \cdots & -1 \\ -1 & N & \cdots & -1 \\ \vdots & \vdots & \ddots & \vdots \\ -1 & -1 & \cdots & N \end{bmatrix} \quad \text{with} \quad \mathbf{\Gamma}^{-1} = \begin{bmatrix} \frac{2}{N+1} & \frac{1}{N+1} & \cdots & \frac{1}{N+1} \\ \frac{1}{N+1} & \frac{2}{N+1} & \cdots & \frac{1}{N+1} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{1}{N+1} & \frac{1}{N+1} & \cdots & \frac{2}{N+1} \end{bmatrix}$$

and let $\tau = N^3 + N^2 - N$, $\beta = 1$ and $g = 4$ so that $\mathbf{E} \in \mathcal{K}$. Next we define loss matrix $\mathbf{L} \in \mathcal{L}$, where $L_{i,j} \leq 1$ if $(i, j) \in \{1, N\} \times \{1, N\}$ and 0 otherwise, as follows:

$$\mathbf{L} = \begin{bmatrix} L_{1,1} & 0 & \cdots & L_{1,N} \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ L_{N,1} & 0 & \cdots & L_{N,N} \end{bmatrix}.$$

Then, with a simple calculation, we obtain $|\mathbf{\Gamma E \Gamma}|_{i,i} = N^2 + N - 1$ for all $i \in [N]$, which implies that we need $\beta' \geq N^2 + N - 1$. Meanwhile, we have that $\|\text{vec}(\mathbf{\Gamma}^{-1} \mathbf{L} \mathbf{\Gamma}^{-1})\|_1 = L_{1,1} + L_{1,N} + L_{N,1} + L_{N,N} \leq 4$, which suggests that $g' \geq 4$. In other words, the regret bound we obtained must be larger than the order of $4N\sqrt{\tau T}$, if we only process the a naive reduction. Parallel, the regret bound to above example is $O(\sqrt{\tau T})$, if we directly utilize our proposed algorithm in the next section, since $\rho = \max_{i,j} |(\mathbf{\Gamma}^{-1} \mathbf{\Gamma}^{-1})_{i,j}| \leq 1$. Thus, our algorithm can improve the FTRL-based algorithm with the standard log-determinant regularizer significantly.

3. Algorithm for the Generalized OSDP Problem

In this section, we give the main algorithm and regret bound to the generalized OSDP problem $(\mathcal{K}, \mathcal{L})$ specified by (1) and (2), with respect to some $\mathbf{\Gamma} \in \mathbb{S}_{++}^{N \times N}$. We propose the FTRL algorithm (4) with the $\mathbf{\Gamma}$ -calibrated log-determinant regularizer:

$$R(\mathbf{W}) = -\ln \det(\mathbf{\Gamma W \Gamma} + \epsilon \mathbf{E}), \quad (7)$$

where $\epsilon > 0$ is a parameter.

The following theorem gives a regret bound of our algorithm.

Theorem 2 (Main Theorem). Given $\mathbf{\Gamma} \in \mathbb{S}_{++}^{N \times N}$. For the generalized OSDP problem specified with Equations (1) and (2), denoting $\rho = \max_{i,j} |(\mathbf{\Gamma}^{-1} \mathbf{\Gamma}^{-1})_{i,j}|$, running the FTRL algorithm with the $\mathbf{\Gamma}$ -calibrated log-determinant regularizer for T times, the regret is bounded as follows:

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) = O\left(g^2(\beta + \rho\epsilon)^2 T \eta + \frac{\tau}{\epsilon \eta}\right).$$

In particular, letting $\eta = \sqrt{\frac{\tau}{g^2(\beta + \rho\epsilon)^2 \epsilon T}}$ and $\epsilon = \beta/\rho$, we have

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) = O\left(g\sqrt{\beta\rho\tau T}\right). \quad (8)$$

Note that we can recover the same regret bound of Theorem 1, when $\mathbf{\Gamma} = \mathbf{E}$.

Before we give the proof of our main theorem, we need to introduce strong convexity, which will play a central role in our proof. The definition of strong convexity is as follows.

Definition 1. For a decision space $\mathcal{K}$ and a real number $s \geq 0$, a regularizer $R : \mathcal{K} \rightarrow \mathbb{R}$ is said to be s -strongly convex with respect to the loss space $\mathcal{L}$ if for any $\alpha \in [0, 1]$, any $\mathbf{X}, \mathbf{Y} \in \mathcal{K}$ and any $\mathbf{L} \in \mathcal{L}$, the following holds

$$R(\alpha \mathbf{X} + (1 - \alpha) \mathbf{Y}) \leq \alpha R(\mathbf{X}) + (1 - \alpha) R(\mathbf{Y}) - \frac{s}{2} \alpha (1 - \alpha) |\mathbf{L} \bullet (\mathbf{X} - \mathbf{Y})|^2. \quad (9)$$

This is equivalent to the following condition: for any $\mathbf{X}, \mathbf{Y} \in \mathcal{K}$ and $\mathbf{L} \in \mathcal{L}$,

$$R(\mathbf{X}) \geq R(\mathbf{Y}) + \nabla R(\mathbf{Y}) \bullet (\mathbf{X} - \mathbf{Y}) + \frac{s}{2} |\mathbf{L} \bullet (\mathbf{X} - \mathbf{Y})|^2. \quad (10)$$

Note that the notion of strong convexity defined above is quite different from the standard one: usually, the strong convexity is defined with respect to some norm $\|\cdot\|$ in decision set [9], but now it is defined with respect to the decision space and the loss space.

A direct reason is that our decision set is constrained by two norms simultaneously, and the same definition is involved by [18].

The following lemma from [7] states a general form between the regret bound to the OSDP problem and the strongly convex regularizer of the FTRL-based algorithm.

Lemma 1 ([7]). *Let $R : \mathcal{K} \rightarrow \mathbb{R}$ be an s -strongly convex regularizer with respect to a decision space $\mathcal{L}$ for a decision space $\mathcal{K}$. Then the FTRL with the regularizer R applied to $(\mathcal{K}, \mathcal{L})$ achieves*

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}) \leq \frac{H_0}{\eta} + \frac{\eta}{s} T, \quad (11)$$

where $H_0 = \max_{W, W' \in \mathcal{K}} (R(W) - R(W'))$.

Due to the lemma above, it suffices to analyze the strong convexity of our Γ -calibrated log-determinant regularizer with respect to our decision space (1) and loss space (2). We show the result in the main proposition.

Proposition 1 (Main proposition). *The Γ -calibrated log-determinant regularizer $R(W) = -\ln \det(\Gamma W \Gamma + \epsilon E)$ is s -strongly convex with respect to $\mathcal{L}$ for $\mathcal{K}$ with $s = 1/(576\sqrt{\epsilon}(\beta + \rho\epsilon)^2 g^2)$, where $\rho = \max_{i,j} |(\Gamma^{-1}\Gamma^{-1})_{i,j}|$.*

We prove this proposition in the next sub-section. Based on this proposition, we firstly give a proof sketch of our main Theorem. The details are in the next sub-section.

Proof Sketch of Theorem 2. According to Proposition 1 and Lemma 1, the only gap of proving this theorem is the bound of H_0 . As we show in the following subsection, $H_0 \leq \frac{\tau}{\epsilon}$, due to the definition of R . Obviously, the size of the matrix N has no effect on our regret bound in Theorem 2. $\square$

Proof for Main Proposition and Theorem

Before we prove Theorem 2, we need to involve some lemmata and notations.

Given a distribution P over $\mathbb{R}^N$, the negative entropy function in respect of P is defined by $H(P) = \mathbb{E}_{x \sim P} [\ln(P(x))]$. We define the characteristic function of P by $\phi(u) = \mathbb{E}_{x \sim P} [e^{iu^T x}]$ where i is the imaginary unit. For two distributions P and Q over $\mathbb{R}^N$, $\frac{1}{2} \int_{\mathbb{R}^N} |P(x) - Q(x)| dx$ denotes the total variation distance between P and Q .

Lemma 2. *Let G_1 and G_2 be two zero mean Gaussian distributions with covariance matrix $\Gamma \mathbf{X} \Gamma$ and $\Gamma \mathbf{Y} \Gamma$, where $\mathbf{X}, \mathbf{Y}, \Gamma \in \mathbb{S}_{++}^{N \times N}$, respectively. If there exists (i, j) such that*

$$|\mathbf{X}_{i,j} - \mathbf{Y}_{i,j}| \geq \delta(\mathbf{X}_{i,i} + \mathbf{Y}_{i,i} + \mathbf{X}_{j,j} + \mathbf{Y}_{j,j}), \quad (12)$$

for some $\delta > 0$, then the total variation distance between G_1 and G_2 is at least $\frac{1}{12e^{1/4}} \delta$.

Proof. Given $\phi_1(u)$ and $\phi_2(u)$ as characteristic functions of G_1 and G_2 , respectively, due to Lemma A1, we have

$$\int_{\mathbb{R}^N} |G_1(x) - G_2(x)| dx \geq \max_{u \in \mathbb{R}^N} |\phi_1(u) - \phi_2(u)|, \quad (13)$$

so we only need to show the lower bound of $\max_{u \in \mathbb{R}^N} |\phi_1(u) - \phi_2(u)|$.

Then we set that characteristic function of G_1 and G_2 are $\phi_1(u) = e^{\frac{-1}{2} u^T \Gamma^T \mathbf{X} \Gamma u}$ and $\phi_2(u) = e^{\frac{-1}{2} u^T \Gamma^T \mathbf{Y} \Gamma u}$, respectively. Let that $\alpha_1 = (\Gamma v)^T \mathbf{X} (\Gamma v)$, $\alpha_2 = (\Gamma v)^T \mathbf{Y} (\Gamma v)$ and $\Gamma u = \frac{\Gamma v}{\sqrt{\alpha_1 + \alpha_2}}$ for some $v \in \mathbb{R}^N$. Moreover we define that given Γ , for any $\bar{v} \in \mathbb{R}^N$ such that $\bar{v} = \Gamma v$, there exists $v \in \mathbb{R}^N$. $\bar{u} = \Gamma u$ in the same way.

Now, let us show the lower bound of $\max_{u \in \mathbb{R}^N} |\phi_1(u) - \phi_2(u)|$.

$$\begin{aligned}
 & \max_{u \in \mathbb{R}^N} |\phi_1(u) - \phi_2(u)| \\
 &= \max_{u \in \mathbb{R}^N} \left| e^{\frac{-1}{2} u^T \Gamma X \Gamma u} - e^{\frac{-1}{2} u^T \Gamma Y \Gamma u} \right| \\
 &= \max_{u \in \mathbb{R}^N} \left| e^{\frac{-1}{2} (\Gamma u)^T X (\Gamma u)} - e^{\frac{-1}{2} (\Gamma u)^T Y (\Gamma u)} \right| \\
 &\geq \max_{\bar{v} \in \mathbb{R}^N} \left| e^{\frac{-\alpha_1}{2(\alpha_1 + \alpha_2)}} - e^{\frac{-\alpha_2}{2(\alpha_1 + \alpha_2)}} \right| \\
 &\geq \max_{\bar{v} \in \mathbb{R}^N} \left| \frac{1}{2e^{1/4}} \frac{\alpha_1 - \alpha_2}{\alpha_1 + \alpha_2} \right|.
 \end{aligned} \tag{14}$$

Then second inequality is due to Lemma A4, since $\min\{\frac{\alpha_1}{\alpha_1 + \alpha_2}, \frac{\alpha_2}{\alpha_1 + \alpha_2}\} \in (0, \frac{1}{2}]$. Due to the assumption in the Lemma, we obtain for some (i, j) that

$$\begin{aligned}
 & \delta(X_{i,i} + Y_{i,i} + X_{j,j} + Y_{j,j}) \leq |X_{i,j} - Y_{i,j}| \\
 &= \frac{1}{2} |(e_i + e_j)^T (X - Y)(e_i + e_j) - e_i^T (X - Y)e_i - e_j^T (X - Y)e_j|
 \end{aligned} \tag{15}$$

It implies that one of $(e_i + e_j)^T (X - Y)(e_i + e_j)$, $e_i^T (X - Y)e_i$ and $e_j^T (X - Y)e_j$ has absolute value greater than $\frac{2\delta}{3}(X_{i,i} + Y_{i,i} + X_{j,j} + Y_{j,j})$.

Due to the positive definiteness of X and Y , we have that for all $v \in \{e_i + e_j, e_i, e_j\}$

$$v^T (X + Y)v \leq 2(X + Y)_{i,i} + 2(X + Y)_{j,j}. \tag{16}$$

and therefore we have that

$$\max_{\bar{v} \in \mathbb{R}^N} \left| \frac{1}{2e^{1/4}} \frac{\alpha_1 - \alpha_2}{\alpha_1 + \alpha_2} \right| \geq \max_{\bar{v} \in \{e_i + e_j, e_i, e_j\}} \left| \frac{1}{2e^{1/4}} \frac{v^T (X - Y)v}{v^T (X + Y)v} \right| \geq \frac{\delta}{6e^{1/4}} \tag{17}$$

□

Lemma 3. Let $X, Y \in \mathbb{S}_+^{N \times N}$ be such that

$$|X_{i,j} - Y_{i,j}| \geq \delta(X_{i,i} + Y_{i,i} + X_{j,j} + Y_{j,j}), \tag{18}$$

and Γ is a symmetric strictly positive definite matrix. Then the following inequality holds that

$$\begin{aligned}
 & -\ln \det(\alpha \Gamma X \Gamma + (1 - \alpha) \Gamma Y \Gamma) \\
 & \leq -\alpha \ln \det(\Gamma X \Gamma) - (1 - \alpha) \ln \det(\Gamma Y \Gamma) - \frac{\alpha(1 - \alpha)}{2} \frac{\delta^2}{36e^{1/2}}.
 \end{aligned} \tag{19}$$

Proof. Let G_1 and G_2 be zero mean Gaussian distributions with covariance matrix $\Gamma X \Gamma$ and $\Gamma Y \Gamma$. The total variation distance between G_1 and G_2 is lower bounded by $\frac{\delta}{12e^{1/4}}$, since the assumption in this Lemma and the result in Lemma 2. Consider the entropy of the following probability distribution of v with probability α that $v \sim G_1$ and $v \sim G_2$ otherwise. Its covariance matrix is $\alpha \Gamma X \Gamma + (1 - \alpha) \Gamma Y \Gamma$.

Due to Lemma A3, we obtain that

$$-\ln \det(\alpha \Gamma X \Gamma + (1 - \alpha) \Gamma Y \Gamma) \leq 2H(\alpha G_1 + (1 - \alpha) G_2) + \ln(2\pi e)^V.$$

By Lemma A2, we further have that

$$\begin{aligned} & 2H(\alpha G_1 + (1 - \alpha)G_2) + \ln(2\pi e)^V \\ & \leq 2\alpha H(G_1) + 2(1 - \alpha)H(G_2) + \ln(2\pi e)^V - 2\alpha(1 - \alpha)\frac{\delta^2}{144e^{1/2}} \\ & \leq 2\alpha H(G_1) + 2(1 - \alpha)H(G_2) + \ln(2\pi e)^V - \alpha(1 - \alpha)\frac{\delta^2}{72e^{1/2}}, \end{aligned}$$

where the second inequality is due the fact that $\delta \geq 0$. At last, from the assumption in this Lemma, G_1 and G_2 are Gaussian distributions, due to the statement in Lemma A3, we have that

$$\begin{aligned} & 2\alpha H(G_1) + 2(1 - \alpha)H(G_2) + \ln(2\pi e)^V - \alpha(1 - \alpha)\frac{\delta^2}{72e^{1/2}} \\ & = -\alpha \ln \det(\Gamma \Sigma \Gamma) - (1 - \alpha) \ln \det(\Gamma \Theta \Gamma) - \alpha(1 - \alpha)\frac{\delta^2}{72e^{1/2}}. \end{aligned} \quad (20)$$

Thus, we have our conclusion. $\square$

Lemma 4 (Lemma 5.4 [7]). Let $\mathbf{X}, \mathbf{Y} \in \mathbb{S}_{++}^{N \times N}$ be such that for all $i \in [N]$ $|\mathbf{X}_{i,i}| \leq \beta'$ and $|\mathbf{Y}_{i,i}| \leq \beta'$. Then for any $\mathbf{L} \in \mathcal{L} = \{\mathbf{L} \in \mathbb{S}_{+}^{N \times N} : \|\text{vec}(\mathbf{L})\|_1 \leq g\}$ there exists that

$$|\mathbf{X}_{i,j} - \mathbf{Y}_{i,j}| \geq \frac{|\mathbf{L} \bullet (\mathbf{X} - \mathbf{Y})|}{4\beta'g} (\mathbf{X}_{i,i} + \mathbf{Y}_{i,i} + \mathbf{X}_{j,j} + \mathbf{Y}_{j,j}). \quad (21)$$

Proposition 2 (Main proposition). For $\Gamma \in \mathbb{S}_{++}^{N \times N}$, the generalized log-determinant regularizer $R(\mathbf{X}) = -\ln \det(\Gamma \mathbf{X} \Gamma + \epsilon \mathbf{E})$ is s -strongly convex with respect to $\mathcal{L}$ for $\mathcal{K}$ with $s = 1/(576\sqrt{e}(\beta + \rho\epsilon)^2g^2)$. Here $\mathbf{E}$ is identity matrix.

Proof. Since the assumption of the proposition that Γ is positive definite then we have that $\Gamma \mathbf{X} \Gamma + \epsilon \mathbf{E} = \Gamma(\mathbf{X} + \Gamma^{-1}\epsilon \mathbf{E} \Gamma^{-1})\Gamma$.

For any $\mathbf{X} \in \mathcal{K}$, we obtain that

$$\max_{i,j} |(\mathbf{X} + \Gamma^{-1}\epsilon \mathbf{E} \Gamma^{-1})_{i,j}| \leq \max_{i,j} |\mathbf{X}_{i,j}| + \epsilon\rho = \beta + \epsilon\rho,$$

where $\rho = \max_{i,j} |(\Gamma^{-1}\Gamma^{-1})_{i,j}|$.

Setting $\beta' = \beta + \epsilon\rho$ in Lemma 4, combining Lemma 3 and Definition 1, our conclusion follows. $\square$

The proof of the main theorem is given as follows:

Proof of Theorem 2. Due to Lemma 1 (in main part) we obtain that

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}, \mathbf{W}^*) \leq \frac{H_0}{\eta} + \frac{\eta}{s}T. \quad (22)$$

Due to the main proposition in main part we know that $s = 1/(576(\beta + \rho\epsilon)^2\sqrt{e}g^2)$.

Thus we need only to show $H_0 \leq \frac{\tau}{\epsilon}$. Denoting $\mathbf{W}_0$ and $\mathbf{W}_1$ as the minimizer and maximizer of R , respectively, then we obtain that

$$\begin{aligned} \max_{\mathbf{W}, \mathbf{W}' \in \mathcal{K}} (R(\mathbf{W}) - R(\mathbf{W}')) &= R(\mathbf{W}_1) - R(\mathbf{W}_0) \\ &= -\ln \det(\Gamma \mathbf{W}_1 \Gamma + \epsilon \mathbf{E}) + \ln \det(\Gamma \mathbf{W}_0 \Gamma + \epsilon \mathbf{E}) \\ &= \sum_{i=1}^N \ln \frac{\lambda_i(\Gamma \mathbf{W}_0 \Gamma) + \epsilon}{\lambda_i(\Gamma \mathbf{W}_1 \Gamma) + \epsilon}, \end{aligned}$$

where the last equality is due to the fact that $\det(A) = \prod_{i=1}^N \lambda_i(A)$, for any $A \in \mathbb{S}_{++}^{N \times N}$. Further, we have that

$$\begin{aligned} \sum_{i=1}^N \ln \frac{\lambda_i(\Gamma W_0 \Gamma) + \epsilon}{\lambda_i(\Gamma W_1 \Gamma) + \epsilon} &= \sum_{i=1}^N \ln \left(\frac{\lambda_i(\Gamma W_0 \Gamma)}{\lambda_i(\Gamma W_1 \Gamma) + \epsilon} + \frac{\epsilon}{\lambda_i(\Gamma W_1 \Gamma) + \epsilon} \right) \\ &\leq \sum_{i=1}^N \ln \left(\frac{\lambda_i(\Gamma W_0 \Gamma)}{\epsilon} + 1 \right) \\ &\leq \sum_{i=1}^N \frac{\lambda_i(\Gamma W_0 \Gamma)}{\epsilon} = \frac{\text{Tr}(\Gamma W_0 \Gamma)}{\epsilon} \leq \frac{\tau}{\epsilon}, \end{aligned}$$

where the first inequality is from the inequality $\ln(1+x) \leq x$. Plugging s , we obtain that

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}, W^*) = O\left(g^2(\beta + \rho\epsilon)^2 T\eta + \frac{\tau}{\epsilon\eta}\right). \quad (23)$$

□

4. Application to OBMC with Side Information

We demonstrate an explicit reduction from OBMC with side information to our aforementioned OSDP problem $(\mathcal{K}, \mathcal{L})$ in the following part. The reduction is two-fold. In the first step, we reduce OBMC with side information to an online matrix prediction (OMP) problem with side information by involving hinge loss function and mistake-driven technique. In the second step, we reduce OMP with side information to the OSDP problem, while representing the side information into the Γ -trace norm.

Before we show the reductions, we need define some necessary notations and the OBMC problem with side information formally.

4.1. The Problem Statement

In principle, our problem statement simplifies the settings in work of Herbster et al. [2].

Given $m, n \in \mathbb{N}_+$, let the pair (M, N) be the side information given to the algorithm, where $M \in \mathbb{S}_{++}^{m \times m}$ and $N \in \mathbb{S}_{++}^{n \times n}$.

The online binary matrix completion (OBMC) problem is a repeated game between the the algorithm and the adversarial environment formulated as follows: On each round $t \in [T]$:

1. The environment selects $(i_t, j_t) \in [m] \times [n]$.
2. The algorithm returns a prediction $\hat{y}_t \in \{-1, +1\}$.
3. The environment reveals the true label $y_t \in \{-1, 1\}$.

The target of the algorithm is to minimize the mistake number during the whole learning process $M = \sum_{t=1}^T \mathbb{I}_{y_t \neq \hat{y}_t}$. Particularly, with the assistance of the side information (M, N) , the mistake bound is supposed to be refined in the case that the side information is constructive.

Equivalently, if we describe the selections and true labels from the environment as a sequence $\mathcal{S} = \{(i_t, j_t), y_t\}_{t=1}^T \subseteq ([m] \times [n] \times \{-1, 1\})^T$. The problem can be seen as a sequential prediction of the entries in an underlying $m \times n$ target matrix. On each single round t , the environment confirms the location of the entry (i_t, j_t) and the algorithm is required to predict the label of the entry. However, we release the consistence of such unknown target matrix, that is, it can happen $y_t \neq y_{t'}$ even if $(i_t, j_t) = (i_{t'}, j_{t'})$.

Instead of the 0–1 loss in the OBMC problem, we firstly need introduce a convex surrogate loss function for the FTRL framework. Specifically, for a positive parameter γ , we define a hinge loss $h_\gamma : \mathbb{R} \rightarrow \mathbb{R}$ with respect to γ as follows:

$$h_\gamma(x) = \begin{cases} 0 & \text{if } \gamma \leq x, \\ 1 - x/\gamma & \text{otherwise,} \end{cases}$$

where γ is also named the margin parameter.

Then for the sequence $\mathcal{S}$ we factorize the *comparator matrix* with $PQ^\top \in \mathbb{R}^{m \times n}$, where $P \in \mathbb{R}^{m \times d}$ and $Q \in \mathbb{R}^{n \times d}$ for some d . Combining the definition of the hinge loss, we define the hinge loss of the sequence $\mathcal{S}$ as

$$\text{hloss}(\mathcal{S}, (P, Q), \gamma) = \sum_{t=1}^T h_\gamma \left(\frac{y_t P_{i_t} Q_{j_t}^\top}{\|P_{i_t}\|_2 \|Q_{j_t}\|_2} \right), \quad (24)$$

in terms of the factorization pair (P, Q) and margin parameter γ . This hinge loss can be interpreted as a measurement for the predictions in terms of the *comparator matrix* $PQ^\top$ to the true label y_t . In the following part, we can, without loss of generality, assume that each row of P and Q is normalized as $\|P_i\|_2 = \|Q_j\|_2 = 1$ for every $(i, j) \in [m] \times [n]$. Moreover, we sometimes call the pair (P, Q) the comparator matrix. Moreover, in the following part, we involve the row normalized matrix $\bar{X}$, according to any matrix X , such that

$$\bar{X} = \text{diag} \left(\frac{1}{\|X_1\|_2}, \dots, \frac{1}{\|X_k\|_2} \right) X.$$

Now for each pair of the comparator matrix, we define the *quasi-dimension*, measuring the quality of the side information. Specifically, the quasi-dimension of a comparator matrix (P, Q) with respect to the side information (M, N) , is defined as

$$\mathcal{D}_{M,N}(P, Q) = \alpha_M \text{Tr}(P^\top M P) + \alpha_N \text{Tr}(Q^\top N Q),$$

where $\alpha_M = \max_{i \in [m]} (M^{-1})_{i,i}$ and $\alpha_N = \max_{j \in [n]} (N^{-1})_{j,j}$. Note that M and N are set as identity matrices, when the side information is empty. In this case, the quasi-dimension is $m + n$ for any comparator matrix. Nevertheless, the quasi-dimension will be smaller, if the rows of P and/or the columns of Q are correlated to M and/or N , which implies that the side information is appropriate and reflects the useful information to the comparator matrix.

Note that the notion of quasi-dimension is defined in a different way in [2].

4.2. Reduction from OBMC with Side Information to an Online Matrix Prediction (OMP)

We formulate an OMP problem, to which our problem is firstly reduced. The OMP problem is specified by a decision space $\mathcal{X} \subseteq [-1, 1]^{m \times n}$ and a margin parameter $\gamma > 0$, and again it is described as a repeated game between the algorithm and adversary. On each round $t \in [T]$,

1. The algorithm predicts a matrix $X_t \in \mathbb{R}^{m \times n}$.
2. The adversary returns a triple $(i_t, j_t, y_t) \in [m] \times [n] \times \{-1, 1\}$, and
3. the algorithm suffers the loss defined by $h_\gamma(y_t X_{t,(i_t,j_t)})$.

The goal of the algorithm is to minimize the regret:

$$\text{Regret}_{\text{OMP}}(T, \mathcal{X}, X^*) = \sum_{t=1}^T h_\gamma(y_t X_{t,(i_t,j_t)}) - \min_{X^* \in \mathcal{X}} \sum_{t=1}^T h_\gamma(y_t X_{t,(i_t,j_t)}^*),$$

$\sup_{X^* \in \mathcal{X}} \text{Regret}_{\text{OMP}}(T, \mathcal{X}, X^*) = \text{Regret}_{\text{OMP}}(T, \mathcal{X})$. Note that unlike the standard setting of online prediction, we do not require $X_t \in \mathcal{X}$.

Below we reduce the OBMC problem with side information (M, N) to the OMP problem with the following decision space:

$$\mathcal{X} = \{\bar{P}\bar{Q}^\top : PQ^\top \in \mathbb{R}^{m \times n}, \mathcal{D}_{M,N}(\bar{P}, \bar{Q}) \leq \hat{D}\},$$

where $\hat{D}$ is an arbitrary parameter. Assume that we have an algorithm $\mathcal{A}$ for the OMP problem $(\mathcal{X}, \gamma)$. In this reduction, we involve the mistake-driven technique, that is the

algorithm $\mathcal{A}$ will be only launched when the prediction error appears. The reduced algorithm is as follows.

Run the algorithm $\mathcal{A}$ and receive the first prediction matrix $\mathbf{X}_1$ from $\mathcal{A}$. Then, in each round $t \in [T]$, we do the following:

1. Observe an index pair $(i_t, j_t) \in [m] \times [n]$.
2. Predict $\hat{y}_t = \text{sgn}(\mathbf{X}_{t,(i_t,j_t)})$.
3. Observe a true label $y_t \in \{-1, 1\}$.
4. If $\hat{y}_t = y_t$ then $\mathbf{X}_{t+1} = \mathbf{X}_t$, and if $\hat{y}_t \neq y_t$, then feed (i_t, j_t, y_t) to $\mathcal{A}$ to let it proceed and receive $\mathbf{X}_{t+1}$.

Note that due to the mistake-driven technique, we run the algorithm $\mathcal{A}$ for at most $M = \sum_{t=1}^T \mathbb{I}_{\hat{y}_t \neq y_t}$ rounds, where M is the number of mistakes of the reduction algorithm above.

The next lemma shows the performance of the reduction.

Lemma 5. Let $\text{Regret}_{\text{OMP}}(M, \mathcal{X}, \mathbf{X}^*)$ denote the regret of the algorithm $\mathcal{A}$ in the reduction above for a competitor matrix $\mathbf{X}^* \in \mathcal{X}$, where $M = \sum_{t=1}^T \mathbb{I}(\hat{y}_t \neq y_t)$. Then,

$$\begin{aligned} M &\leq \inf_{\bar{\mathbf{P}}\bar{\mathbf{Q}}^T \in \mathcal{X}} (\text{Regret}_{\text{OMP}}(M, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^T) + \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma)) \\ &\leq \text{Regret}_{\text{OMP}}(M, \mathcal{X}) + \text{hloss}(\mathcal{S}, \gamma), \end{aligned} \quad (25)$$

where we define

$$\text{hloss}(\mathcal{S}, \gamma) = \min_{\bar{\mathbf{P}}\bar{\mathbf{Q}}^T \in \mathcal{X}} \text{hloss}(\mathcal{S}, (\bar{\mathbf{P}}, \bar{\mathbf{Q}}), \gamma). \quad (26)$$

Remark 1. If $\mathbf{M}$ and $\mathbf{N}$ are identity matrices, then we have $\mathcal{D}_{\mathbf{M}, \mathbf{N}}(\bar{\mathbf{P}}, \bar{\mathbf{Q}}) = m + n$, and thus the decision space is an unconstrained set $\mathcal{X} = \{\bar{\mathbf{P}}\bar{\mathbf{Q}}^T : \bar{\mathbf{P}}\bar{\mathbf{Q}}^T \in \mathbb{R}^{m \times n}\}$.

Proof. Let $\mathbf{P}$ and $\mathbf{Q}$ be arbitrary matrices such that $\bar{\mathbf{P}}\bar{\mathbf{Q}}^T \in \mathcal{X}$. Since $\mathbb{I}(\text{sgn}(x) \neq y) \leq h_\gamma(yx)$ for any $x \in \mathbb{R}$ and $y \in \{-1, 1\}$, we have

$$\begin{aligned} M &= \sum_{t=1}^T \mathbb{I}(\hat{y}_t \neq y_t) \leq \sum_{\{t: \hat{y}_t \neq y_t\}} h_\gamma(y_t \mathbf{X}_{t,(i_t,j_t)}) \\ &= \text{Regret}_{\text{OMP}}(M, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^T) + \sum_{\{t: \hat{y}_t \neq y_t\}} h_\gamma(y_t (\bar{\mathbf{P}}\bar{\mathbf{Q}}^T)_{i_t,j_t}) \\ &\leq \text{Regret}_{\text{OMP}}(M, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^T) + \sum_{t=1}^T h_\gamma(y_t (\bar{\mathbf{P}}\bar{\mathbf{Q}}^T)_{i_t,j_t}) \\ &= \text{Regret}_{\text{OMP}}(M, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^T) + \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma), \end{aligned}$$

where the second equality follows from the definition of regret, and the third equality follows from the fact that $(\bar{\mathbf{P}}\bar{\mathbf{Q}}^T)_{i,j} = \mathbf{P}_i \mathbf{Q}_j^T / (\|\mathbf{P}_i\|_2 \|\mathbf{Q}_j\|_2)$. Since the choice of $\mathbf{P}$ and $\mathbf{Q}$ is arbitrary, the following inequality holds:

$$M \leq \inf_{\bar{\mathbf{P}}\bar{\mathbf{Q}}^T \in \mathcal{X}} (\text{Regret}_{\text{OMP}}(M, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^T) + \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma)). \quad (27)$$

Now, let $\mathbf{P}$ and $\mathbf{Q}$ be the matrices that attain (26). Hence, our lemma follows from

$$M \leq \text{Regret}_{\text{OMP}}(M, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^T) + \text{hloss}(\mathcal{S}, \gamma) \leq \sup_{\mathbf{X}^* \in \mathcal{X}} \text{Regret}_{\text{OMP}}(M, \mathcal{X}, \mathbf{X}^*) + \text{hloss}(\mathcal{S}, \gamma).$$

□

4.3. Reduction from OMP to the Generalised OSDP Problem

Throughout this sub-section, we reduce the OMP with side information to the OSDP problem parameterized with Γ . Our reduction is similar to the work of [1,6]. For convenience, we denote $N = m + n$ in the following part.

First of all, we formulate the side information $M \in \mathbb{S}_{++}^{m \times n}, N \in \mathbb{S}_{++}^{n \times n}$ into $\Gamma \in \mathbb{S}_{++}^{N \times N}$ for our generalized OSDP problem in the following equation:

$$\Gamma = \begin{bmatrix} \sqrt{\alpha_M} M & 0 \\ 0 & \sqrt{\alpha_N} N \end{bmatrix}. \quad (28)$$

Next we define the decision space $\mathcal{K}$. For any comparator matrix (P, Q) such that $PQ^\top \in \mathbb{R}^{m \times n}$, we define

$$W_{P,Q} = \begin{bmatrix} \bar{P} \\ \bar{Q} \end{bmatrix} \begin{bmatrix} \bar{P}^\top & \bar{Q}^\top \end{bmatrix} = \begin{bmatrix} \bar{P}\bar{P}^\top & \bar{P}\bar{Q}^\top \\ \bar{Q}\bar{P}^\top & \bar{Q}\bar{Q}^\top \end{bmatrix}.$$

Trivially, $W_{P,Q}$ is an $N \times N$ symmetric and positive semi-definite matrix. Intuitively, any comparator matrix (P, Q) such that $\bar{P}\bar{Q}^\top \in \mathcal{X}$, is embedded into the upper right block in $W_{P,Q}$. In addition, since the normalization of $\bar{P}$ and $\bar{Q}$, $(W_{P,Q})_{i,i} \leq 1$ for all $i \in [N]$. Hence, we need to find a convex decision space $\mathcal{K} \subseteq \mathbb{S}_{++}^{N \times N}$ which satisfies

$$\mathcal{K} \supseteq \{W_{P,Q} : \bar{P}\bar{Q}^\top \in \mathcal{X}\}.$$

First, we involve the following Lemma:

Lemma 6 (Lemma 8 [2]). For any pair of side information matrices (M, N) , where $M \in \mathbb{S}_{++}^{m \times m}$ and $N \in \mathbb{S}_{++}^{n \times n}$, and Γ is induced as in Equation (28).

$$\text{Tr}(\Gamma W_{P,Q} \Gamma) = \alpha_M \text{Tr}(\bar{P}^\top M \bar{P}) + \alpha_N \text{Tr}(\bar{Q}^\top N \bar{Q}). \quad (29)$$

Due to this lemma and the definition of $\mathcal{X}$, we can directly define $\mathcal{K}$ as follows:

$$\mathcal{K} = \{W \in \mathbb{S}_{++}^{N \times N} : \|\text{vec}(W)\|_\infty \leq 1 \wedge \text{Tr}(\Gamma W \Gamma) \leq \hat{D}\} \supseteq \{W_{P,Q} : \bar{P}\bar{Q}^\top \in \mathcal{X}\}. \quad (30)$$

Then, we describe the loss matrix class $\mathcal{L}$. We first define a sparse matrix $Z\langle i, j \rangle \in \mathbb{S}_+^{N \times N}$ with any pair $(i, j) \in [m] \times [n]$ such that only entries $(i, m+j)$ and $(m+j, i)$ are 1, others are 0. More formally,

$$Z\langle i, j \rangle = \frac{1}{2} (e_i e_{m+j}^\top + e_{m+j} e_i^\top),$$

where e_k is the k -th orthogonal basis vector of $\mathbb{R}^N$. Note that due to the definition of the Frobenius product, we have that

$$W_{P,Q} \bullet Z\langle i, j \rangle = (\bar{P}\bar{Q}^\top)_{i,j},$$

which is what we focus on. Thus, $\mathcal{L}$ is defined as

$$\mathcal{L} = \{cZ\langle i, j \rangle : c \in \{-1/\gamma, 1/\gamma\}, i \in [m], j \in [n]\}. \quad (31)$$

Now we state the reduction from the OMP problem with side information to the OSDP problem $(\mathcal{K}, \mathcal{L})$ specified by Γ . Assume that there is an algorithm $\mathcal{A}$ for the OSDP problem.

Run the algorithm $\mathcal{A}$ and receive the first prediction matrix $W_1 \in \mathcal{K}$ from $\mathcal{A}$.

In each round t ,

1. let X_t be the upper right $m \times n$ component matrix of W_t .
// $X_{t,(i,j)} = W_t \bullet Z\langle i, j \rangle$

2. observe a triple $(i_t, j_t, y_t) \in [m] \times [n] \times \{-1, 1\}$,
3. suffer loss $\ell_t(\mathbf{W}_t)$ where $\ell_t : \mathbf{W} \mapsto h_\gamma(y_t(\mathbf{W} \bullet \mathbf{Z}\langle i_t, j_t \rangle))$,
4. let $\mathbf{L}_t = \nabla_{\mathbf{W}} \ell_t(\mathbf{W}_t) = \begin{cases} -\frac{y_t}{\gamma} \mathbf{Z}\langle i_t, j_t \rangle & \text{if } y_t \mathbf{X}_{t,(i_t, j_t)} \leq \gamma \\ 0 & \text{otherwise} \end{cases}$,
5. feed $\mathbf{L}_t$ to the algorithm $\mathcal{A}$ to let it proceed and receive $\mathbf{W}_{t+1}$.

Due to the convexity of ℓ_t , a standard linearization argument ([9]) gives

$$\ell_t(\mathbf{W}_t) - \ell_t(\mathbf{W}^*) \leq \mathbf{W}_t \bullet \mathbf{L}_t - \mathbf{W}^* \bullet \mathbf{L}_t$$

for any $\mathbf{W}^* \in \mathcal{K}$. Moreover, according to our reduction that $\ell_t(\mathbf{W}_t) = h_\gamma(y_t \mathbf{X}_{t,(i_t, j_t)})$ and $\ell_t(\mathbf{W}_{P,Q}) = h_\gamma(y_t(\bar{\mathbf{P}}\bar{\mathbf{Q}}^\top)_{i_t, j_t})$, the following lemma immediately follows.

Lemma 7. Let $\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}, \mathbf{W}_{P,Q}) = \sum_{t=1}^T (\mathbf{W}_t - \mathbf{W}_{P,Q}) \bullet \mathbf{L}_t$ denote the regret of the algorithm $\mathcal{A}$ in the reduction above for a competitor matrix $\mathbf{W}_{P,Q}$ and $\text{Regret}_{\text{OMP}}(T, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^\top) = \sum_{t=1}^T (h_\gamma(y_t \mathbf{X}_{t,(i_t, j_t)}) - h_\gamma(y_t(\bar{\mathbf{P}}\bar{\mathbf{Q}}^\top)_{i_t, j_t}))$ denote the regret of the reduction algorithm for $\bar{\mathbf{P}}\bar{\mathbf{Q}}^\top$. Then,

$$\text{Regret}_{\text{OMP}}(T, \mathcal{X}, \bar{\mathbf{P}}\bar{\mathbf{Q}}^\top) \leq \text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}, \mathbf{W}_{P,Q}).$$

Combining Lemmas 5 and 7, we have the following corollary.

Corollary 1. There exists an algorithm for the OBMC problem with side information with the following mistake bounds.

$$\begin{aligned} M &\leq \inf_{\bar{\mathbf{P}}\bar{\mathbf{Q}}^\top \in \mathcal{X}} (\text{Regret}_{\text{OSDP}}(M, \mathcal{K}, \mathcal{L}, \mathbf{W}_{P,Q}) + \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma)) \\ &\leq \text{Regret}_{\text{OSDP}}(M, \mathcal{K}, \mathcal{L}) + \text{hloss}(\mathcal{S}, \gamma). \end{aligned}$$

4.4. Application to Matrix Completion

Combining previous two reductions, our ultimate reduction immediately follows. We reduce OBMC with side information to OSDP specified with Γ . Compared with the analysis of Herbster et al. [2], our reduction is explicit and easy to follow. Note that in our reduction, the side information $\mathbf{M}$ and $\mathbf{N}$ for OBMC is parameterized by Γ . Again, the Γ is the identity matrix if the side information is vacuous. Finally, running our proposed FTRL-based algorithm with the Γ -calibrated log-determinant regularizer, we improve the logarithmic factor in the previous mistake bound [2]. Particularly, our mistake bound is optimal.

Remark 2. Since the definition of Γ in Equation (28), we have that $\rho = 1$.

According to our analysis, the parameters in the reduced OSDP problem $(\mathcal{K}, \mathcal{L})$ with Γ defined in Equation (28), can trivially be set as $g = 1/\gamma$, $\beta = \epsilon = \rho = 1$, $\tau = \widehat{D}$, then utilizing Theorem 2, we obtain the following result

$$\text{Regret}_{\text{OSDP}}(T, \mathcal{K}, \mathcal{L}, \mathbf{W}^*) = O\left(\frac{T\eta}{\gamma^2} + \frac{\widehat{D}}{\eta}\right). \quad (32)$$

Next, we give our main algorithm for the OBMC problem with side information $\mathbf{M}, \mathbf{N}$ in Algorithm 1. Putting the two reductions together and proceeding with the FTRL-based algorithm (4), the main algorithm is as follows:

Theorem 3. Running Algorithm 1 with parameter $\eta = \sqrt{\gamma^2 \widehat{D}}/T$, $\gamma \in (0, 1]$, the hinge loss of OBMC with side information is bounded as follows:

$$\sum_{t=1}^T h_{\gamma}(y_t \cdot \hat{y}_t) - \sum_{t=1}^T h_{\gamma}(y_t \cdot (\bar{\mathbf{P}}\bar{\mathbf{Q}}^{\top})_{i_t, j_t}) \leq O\left(\sqrt{\frac{\widehat{D}T}{\gamma^2}}\right). \quad (33)$$

Compared with [2], the logarithmic factor $\ln(m+n)$ is improved in our regret bound to hinge loss. Meanwhile, since the mistake-driven technique is involved in our reduction, the horizon T is replaced by M , the number of mistakes, which is unknown in advance. Then, by choosing η independent of M , we can derive a good mistake bound due to above theorem, resulting in Equation (32).

Algorithm 1 Online binary matrix completion with side information algorithm.

```

1: Parameters:  $\gamma > 0$ ,  $\eta > 0$ , side information matrices  $\mathbf{M} \in \mathbb{S}_{++}^{m \times m}$  and  $\mathbf{N} \in \mathbb{S}_{++}^{n \times n}$ . Quasi
   dimension estimator  $1 \leq \widehat{D}$ .  $\Gamma$  is composed as in Equation (28), and decision set  $\mathcal{K}$  is
   given as (30).
2: Initialize  $\forall \mathbf{W} \in \mathcal{K}$ , set  $\mathbf{W}_1 = \mathbf{W}$ .
3: for  $t = 1, 2, \dots, T$  do
4:   Receive  $(i_t, j_t) \in [m] \times [n]$ .
5:   Let  $\mathbf{Z}_t = \frac{1}{2}(\mathbf{e}_{i_t}\mathbf{e}_{m+j_t}^T + \mathbf{e}_{m+j_t}\mathbf{e}_{i_t}^T)$ .
6:   Predict  $\hat{y}_t = \text{sgn}(\mathbf{W}_t \bullet \mathbf{Z}_t)$  and receive  $y_t \in \{-1, 1\}$ .
7:   if  $\hat{y}_t \neq y_t$  then
8:     Let  $\mathbf{L}_t = \frac{-y_t}{\gamma}\mathbf{Z}_t$  and  $\mathbf{W}_{t+1} = \arg \min_{\mathbf{W} \in \mathcal{K}} -\ln \det(\Gamma \mathbf{W} \Gamma + \mathbf{E}) + \eta \sum_{s=1}^t \mathbf{W} \bullet \mathbf{L}_s$ .
9:   else
10:    Let  $\mathbf{L}_t = 0$  and  $\mathbf{W}_{t+1} = \mathbf{W}_t$ .
11:   end if
12: end for
```

Theorem 4. Algorithm 1 with $\eta = c\gamma^2$ for some $c > 0$ achieves

$$M = \sum_{t=1}^T \mathbb{I}_{\hat{y}_t \neq y_t} = O\left(\frac{\widehat{D}}{\gamma^2}\right) + 2\text{hloss}(\mathcal{S}, \gamma). \quad (34)$$

Proof. Combining Corollary 1 and the regret bound (32), we have

$$M = O\left(\frac{M\eta}{\gamma^2} + \frac{\widehat{D}}{\eta}\right) + \text{hloss}(\mathcal{S}, \gamma).$$

Choosing $\eta = c\gamma^2$ for sufficiently small constant c , we get

$$M \leq \frac{M}{2} + O\left(\frac{\widehat{D}}{\gamma^2}\right) + \text{hloss}(\mathcal{S}, \gamma),$$

from which (34) follows. $\square$

Again if the side information is vacuous, which means that $\mathbf{M}, \mathbf{N}$ are identity matrices, from Remark 1 and Theorem 4, we can set that $\widehat{D} = m+n$ and obtain the mistake bound as follows:

$$O\left(\frac{m+n}{\gamma^2} + 2\text{hloss}_{\mathbf{P}\mathbf{Q}^T \in \mathbb{R}^{m \times n}}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma)\right).$$

Nevertheless, the side information indeed matters in non-trivial cases. When the comparator matrix $\mathbf{U}$ contains some latter structure, specifically, when $\mathbf{U}$ is (k, l) -biclustered, the quasi-dimension estimator $\widehat{D} \in O(k+l)$, which is strictly smaller than $O(m+n)$, if the

side information M, N is chosen as a special matrix according to the structure of (P, Q) (the details are in Appendix A). According to this instance, the accuracy of the prediction can be effectively improved when the side information is selected correlating to the structure of the underlying matrix.

Note that our mistake bound performs better than previous bound especially in the realizable case. Compared with the bound $O\left(\frac{\hat{D}}{\gamma^2} \ln(m+n)\right)$ in [2], our mistake bound is $O\left(\frac{\hat{D}}{\gamma^2}\right)$, which removes the logarithmic factor $\ln(m+n)$. In addition, our bound recovers the known lower bound of Herbster et al. [1] up to a constant factor. If U contains a (k, l) -biclustered structure ($k \geq l$), by setting $\gamma = \frac{1}{\sqrt{l}}$, our mistake bound can become $O(kl)$. On the other hand, the lower bound of Herbster et al. is $\Omega(kl)$. Thus, the mistake bound of Theorem 4 is optimal.

In the next subsection, we show an example, online similarity prediction with side information, where the comparator matrix is (k, l) -biclustered. With the ideal side information, we can effectively improve the mistake bound.

4.5. Online Similarity Prediction with Side Information

In this subsection, we show the application of our reduction method and generalized log-determinant regularizer to online similarity prediction with side information.

Before we introduce the online similarity prediction problem, we need involve some notations and basic concepts. Denote $G = (V, E)$ be an undirected and connected graph, where $n = |V|$ and $m = |E|$. In graph G , if all the vertices are assigned into K different classes, i.e., there is n -dimensional vector $\mathbf{y} = \{y_1, \dots, y_n\}$, where $y_i \in \{1, \dots, K\}$, and the classification of each vertex v_i is represented by y_i in vector $\mathbf{y}$, we denote this graph with respect to the assignment $\mathbf{y}$ as $(G, \mathbf{y})$.

Next we define the *cut-edges* of $(G, \mathbf{y})$, $\Phi^G(\mathbf{y}) = \{(i, j) \in E : y_i \neq y_j\}$, and Φ^G in abbreviation. The cardinality of $|\Phi^G(\mathbf{y})|$ is the cut size. For each graph G , we denote the adjacency matrix of G as A if $A_{ij} = A_{ji} = 1$ if $(i, j) \in E(G)$ and $A_{ij} = 0$, otherwise. The degree matrix of G , $D \in \mathbb{R}^{n \times n}$ is defined as a diagonal matrix where D_{ii} is the degree of vertex i . The Laplacian is defined as $L = D - A$. We define the PD-Laplacian $\tilde{L} = L + \left(\frac{1}{n}\right) \left(\frac{1}{n}\right)^\top \alpha_L$, where $\mathbf{1}$ is a n -dimensional vector that all entries are 1. Given a graph $G = (V, E)$ and its Laplacian L , assume that each edge of G is a unit resistor, the effective resistance of any pair of vertices $(i, j) \in V \times V$, $R_{i,j}^G = (e_i - e_j)L^+(e_i - e_j)$, where e_i is the standard basis in $\mathbb{R}^n$, and L^+ is the pseudo-inverse matrix of L .

Now the online similarity prediction is formulated as follows. Given a K classified graph $G = (V, E)$. On each round $t \in [T]$, we have the following:

1. The environment confirms a pair of vertices $(i_t, j_t) \in V \times V$.
2. The algorithm predicts whether these two vertices belong to the same class. If they are classified in the same class, the algorithm predicts $\hat{y}_{i_t, j_t} = 1$ and -1 , otherwise.
3. The environment reveals the true answer y_{i_t, j_t} . If they are in the same class, then $y_{i_t, j_t} = 1$; $y_{i_t, j_t} = -1$ otherwise.

The target of the algorithm is to minimize the number of the prediction mistakes $M = \sum_{t=1}^T \mathbb{I}_{\hat{y}_{i_t, j_t} \neq y_{i_t, j_t}}$. Due to this formulation of online similarity prediction, this problem is a special case of an OBMC problem. We denote also a sequence $\mathcal{S} = \{((i_t, j_t), y_t)\}_{t=1}^T \subseteq ([n] \times [n], \{-1, 1\})^T$ for our online similarity prediction problem.

The side information we defined for online similarity prediction is the PD-Laplacian $\tilde{L}$ of graph G in our paper. Note that the side information is restricted about only the graph G itself, and irrelevant to the classification vector $\mathbf{y}$. Gentile et al. [15] explored this problem by involving the graph G as the prior information to the algorithm. It is equivalent to our problem setting. The mistake bound from [15] is described in the following proposition:

Proposition 3. Let $(G, \mathbf{y})$ be a labeled graph. If we run Matrix Winnow with G as an input graph, we have the following mistake bound:

$$M^W = O\left(|\Phi^G| \max_{(i,j) \in V^2} R_{ij}^G \ln n\right). \quad (35)$$

From the definitions of the online similarity prediction with side information and OBMC with side information, online similarity prediction is a special case of OBMC, by considering the comparator matrix $\mathbf{U} \in \{1, -1\}^{n \times n}$, where $\mathbf{U}_{ij}$ indicates the classifications of vertices i and j . More specifically, if vertices i, j are in the same class, then $\mathbf{U}_{ij} = 1$, and $\mathbf{U}_{ij} = -1$, otherwise. Meanwhile, the side information unifies the symmetric positive definite matrices $\mathbf{M}, \mathbf{N}$ into $\tilde{\mathbf{L}}$.

Moreover, due to [2,15] (see details in Appendix A.2), the comparator matrix $\mathbf{U}$ is actually a (K, K) -biclustered $n \times n$ binary matrix. According to the aforementioned result from [2], there exists a matrix $\mathbf{R} \in \mathcal{B}^{n, k}$ such that $\mathbf{U} = \mathbf{R}\mathbf{U}^*\mathbf{R}^\top$, where $\mathbf{U}^* = 2\mathbf{I}_K - \mathbf{1}\mathbf{1}^\top$, denoting $\mathbf{1}$ as a K -dimensional vector for which all entries are 1.

As same as the previous reductions, the reduced Γ specified OSDP problem corresponding to the online similarity prediction with side information is as follows. Firstly, we define the side information $\tilde{\mathbf{L}}$ parameterized Γ in the following matrix.

$$\Gamma = \begin{bmatrix} \sqrt{\alpha_{\tilde{\mathbf{L}}}} \tilde{\mathbf{L}} & 0 \\ 0 & \sqrt{\alpha_{\tilde{\mathbf{L}}}} \tilde{\mathbf{L}} \end{bmatrix}, \quad (36)$$

where $\alpha_{\tilde{\mathbf{L}}} = \max_i (\tilde{\mathbf{L}}^{-1})_{ii}$.

Next the decision space $\mathcal{K}$ and the loss space $\mathcal{L}$ are defined as previously as in Equations (30) and (31), respectively.

Thus, the following proposition demonstrates that the online similarity prediction with side information $\tilde{\mathbf{L}}$ can be reduced to an OSDP problem $(\mathcal{K}, \mathcal{L})$ parameterized with Γ .

Proposition 4. Given an online similarity prediction problem with graph $(G, \mathbf{y})$, set the side information as $\tilde{\mathbf{L}}$. Then we can reduce this problem to a generalized OSDP problem $(\mathcal{K}, \mathcal{L})$ with bounded Γ -trace norm such that

$$\begin{aligned} \mathcal{K} &= \left\{ \mathbf{X} \in \mathbb{S}_{++}^{n \times n} : \|\text{vec}(\mathbf{X})\|_\infty \leq 1 \wedge \text{Tr}(\Gamma \mathbf{X} \Gamma) \leq \hat{\mathcal{D}} \right\} \\ \mathcal{L} &= \left\{ \frac{1}{\gamma} \mathbf{Z}\langle i, j \rangle, \frac{-1}{\gamma} \mathbf{Z}\langle i, j \rangle : \forall i \in [n], j \in [n] \right\}, \end{aligned}$$

where Γ is defined as above, and

$$\mathbf{Z}\langle i, j \rangle = \frac{1}{2}(\mathbf{e}_i \mathbf{e}_{n+j}^\top + \mathbf{e}_{n+j} \mathbf{e}_i^\top). \quad (37)$$

Therefore, the mistake bound of online similarity prediction is bounded as follows:

$$M = \sum_{t=1}^T \mathbb{I}_{\hat{y}_{i_t, j_t} \neq y_{i_t, j_t}} \leq \text{Regret}_{\text{OSDP}}(M, \mathcal{K}, \mathcal{L}) + \text{hloss}(\mathcal{S}, \gamma), \quad (38)$$

where γ is the margin parameter.

According to [2], there exists $\mathbf{P} = \mathbf{R}\mathbf{P}^*, \mathbf{Q} = \mathbf{R}\mathbf{Q}^*$ where

$$\mathbf{P}^* = (\mathbf{P}_{ij}^* = \sqrt{2}\mathbb{I}_{i=j} + \mathbb{I}_{j=k+1})_{i \in [k], j \in [k+1]} \quad \text{and} \quad \mathbf{Q}^* = (\mathbf{Q}_{ij}^* = \sqrt{2}\mathbb{I}_{i=j} - \mathbb{I}_{j=k+1})_{i \in [k], j \in [k+1]}$$

such that $\mathbf{U} = \mathbf{P}\mathbf{Q}^\top$. It implies that the hinge loss $\text{hloss}(\mathcal{S}, \gamma) = 0$, when $\gamma \in O(1)$, more specifically $\gamma = \frac{1}{3}$.

Remark 3. According to Theorem 3 and Section 4.2 in [2], and the characterization of online similarity prediction with side information, the quasi-dimension estimator $\widehat{D} \leq O(\text{Tr}(\mathbf{R}^\top \mathbf{L} \mathbf{R}) \alpha_L)$, where $\mathbf{L}$ is the Laplacian of the corresponding graph G , if we set the side information is $\widetilde{\mathbf{L}}$.

Due to our Theorem 4 and running Algorithm 1 in main part, finally, we have the mistake bound as follows:

$$M \leq O(\text{Tr}(\mathbf{R}^\top \mathbf{L} \mathbf{R}) \alpha_L). \quad (39)$$

Remark 4. In the work of Herbster et al. [2], the resulting mistake bound is $O(\text{Tr}(\mathbf{R}^\top \mathbf{L} \mathbf{R}) \alpha_L \ln n)$, which can recover the bound of $O(|\Phi^G| \max_{(i,j) \in V^2} R_{i,j}^G \ln n)$ in [15] up to a constant factor. Moreover, our bound improves the logarithmic factor $\ln n$ compared with [2].

5. Connection to a Batch Setting

In this section, we employ the well-known online-to-batch conversion technique (see, for example, [16]) and obtain a batch learning algorithm with generalization error bounds. The results imply that the algorithm performs nearly as well as the support vector machine (SVM) running over the optimal feature space, although the side information is vacuous. Moreover, with the assistance of the side information, a more refined mistake bound for batch setting follows from our online version analysis.

First, we describe our setting formally. We consider the problem in the standard probably approximately correct (PAC) learning framework [16,19,20]. The algorithm is given the side information matrix $\mathbf{M} \in \mathbb{S}_{++}^{m \times m}$ and $\mathbf{N} \in \mathbb{S}_{++}^{n \times n}$ and a sample sequence $\mathcal{S}$:

$$\mathcal{S} = ((i_1, j_1, y_1), (i_2, j_2, y_2), \dots, (i_T, j_T, y_T))$$

where each triple (i_t, j_t, y_t) is randomly and independently generated according to some unknown probability distribution $\mathcal{V}$ over $[m] \times [n] \times \{-1, 1\}$. Then the algorithm outputs a hypothesis $f : [m] \times [n] \rightarrow [-1, 1]$. The goal is to find, with high probability, a hypothesis f that has small generalization error

$$R(f) = \Pr_{(i,j,y) \sim \mathcal{V}} (\text{sgn}(f(i,j)) \neq y).$$

In particular, we consider a hypothesis of the form

$$f_{\mathbf{W}} : (i, j) \mapsto \mathbf{W} \bullet \mathbf{Z} \langle i, j \rangle$$

where $\mathbf{W} \in \mathcal{K} = \{\mathbf{W} \in \mathbb{S}_{++}^{N \times N} : \forall i \in [N], \mathbf{W}_{i,i} \leq 1 \wedge \text{Tr}(\mathbf{\Gamma} \mathbf{W} \mathbf{\Gamma}) \leq \widehat{D}\}$, where $\mathbf{\Gamma}$ is defined as in Equation (28).

In Algorithm 2, we give the algorithm obtained by the online-to-batch conversion.

Algorithm 2 Binary matrix completion in the batch setting.

- 1: Parameter: $\gamma > 0$
 - 2: Input: a sample $\mathcal{S}$ of size T .
 - 3: Run Algorithm 1 over $\mathcal{S}$ and get its predictions $\mathbf{W}_1, \mathbf{W}_2, \dots, \mathbf{W}_T$.
 - 4: Choose $\mathbf{W}$ from $\{\mathbf{W}_1, \mathbf{W}_2, \dots, \mathbf{W}_T\}$ uniformly at random.
 - 5: Output $f_{\mathbf{W}}$.
-

To bound the generalization error, we use the following lemma, which is straightforward from Lemma 7.1 of [16].

Lemma 8. Let $L : [-1, 1] \times \{-1, 1\} \rightarrow [-B, B]$ be a function and $\mathbf{W}_1, \dots, \mathbf{W}_T$ and $\mathbf{W}$ be the matrices obtained in Algorithm 2. Then, for any $\delta > 0$, with probability at least $1 - \delta$, the following holds:

$$\begin{aligned} \mathbb{E}_{(i,j,y) \sim \mathcal{V}, \mathbf{W}}[L(f_{\mathbf{W}}(i, j), y)] &= \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{(i,j,y) \sim \mathcal{V}}[L(f_{\mathbf{W}}(i, j), y)] \\ &\leq \frac{1}{T} \sum_{t=1}^T L(f_{\mathbf{W}_t}(i_t, j_t), y_t) + B \sqrt{\frac{2 \ln 1/\delta}{T}}. \end{aligned}$$

Applying the lemma with the zero-one loss $L(r, y) = \mathbf{1}(\text{sgn}(r) \neq y)$ combined with the mistake bound (34) of Theorem 4, we have the following generalization bound.

Theorem 5. For any $\delta > 0$, with probability of at least $1 - \delta$, Algorithm 2 produces $f_{\mathbf{W}}$ with the following property:

$$\mathbb{E}_{\mathbf{W}}[R(f_{\mathbf{W}})] \leq \frac{O\left(\frac{\hat{D}}{\gamma^2} + \text{hloss}(\mathcal{S}, \gamma)\right)}{T} + \sqrt{\frac{2 \ln 1/\delta}{T}}. \quad (40)$$

On the other hand, when applying the lemma with the hinge loss $L(r, y) = h_{\gamma}(ry)$ combined with an $O(\sqrt{\hat{D}T/\gamma^2})$ regret bound of (32) with the minimizer $\eta = \sqrt{\gamma^2 \hat{D}/T}$, then we have

$$\begin{aligned} \mathbb{E}_{\mathbf{W}}[R(f_{\mathbf{W}})] &\leq \mathbb{E}_{(i,j,y) \sim \mathcal{V}, \mathbf{W}}[h_{\gamma}(y f_{\mathbf{W}}(i, j))] \\ &\leq O\left(\sqrt{\frac{\hat{D}}{\gamma^2 T}} + \frac{\text{hloss}(\mathcal{S}, \gamma)}{T}\right) + (1 + \gamma) \sqrt{\frac{2 \ln 1/\delta}{T}}, \end{aligned} \quad (41)$$

which is slightly worse than (40). Note that $\hat{D} = O(m)$, if the side information is vacuous.

Now we examine some implications of our generalization bounds. First, we assume without loss of generality that $m \geq n$, because otherwise, we can make everything transposed.

As explained in the aforementioned sections, we can think of each $\bar{\mathbf{Q}}_j$ as a feature vector of item j . Assume all feature vectors $\bar{\mathbf{Q}} \in \mathbb{R}^{n \times k}$ are given and consider the problem of finding a good linear classifier $\bar{\mathbf{P}}_i$ for each user i independently. A natural way is to use the SVM, which solves

$$\inf_{\gamma > 0, \bar{\mathbf{P}}_i \in \mathbb{R}^k} \left(1/\gamma^2 + C \sum_{t: i_t=i} h_{\gamma}(y_t \bar{\mathbf{P}}_i \bar{\mathbf{Q}}_{j_t}^{\top}) \right)$$

for every $i \in [m]$ for some constant $C > 0$. Now if we fix γ to be a constant for all i , then the optimization problems above are summarized as

$$\begin{aligned} \inf_{\mathbf{P} \in \mathbb{R}^{m \times k}} \sum_{i=1}^m \left(1/\gamma^2 + C \sum_{t: i_t=i} h_{\gamma}(y_t \bar{\mathbf{P}}_i \bar{\mathbf{Q}}_{j_t}^{\top}) \right) &= \inf_{\mathbf{P} \in \mathbb{R}^{m \times k}} \left(m/\gamma^2 + C \sum_t h_{\gamma}(y_t \bar{\mathbf{P}}_{i_t} \bar{\mathbf{Q}}_{j_t}^{\top}) \right) \\ &= \frac{m}{\gamma^2} + C \inf_{\mathbf{P}} \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma). \end{aligned}$$

So, if we further optimize feature vectors, we obtain

$$\frac{m}{\gamma^2} + C \inf_{\mathbf{P}, \mathbf{Q}^{\top} \in \mathbb{R}^{m \times n}} \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma) = \frac{m}{\gamma^2} + C \text{hloss}(\mathcal{S}, \gamma) \quad (42)$$

which roughly is proportional to our generalization bound (40), when the side information is vacuous (i.e., $\mathbf{M}$ and $\mathbf{N}$ are identity matrices). This result implies that our generalization

bound is upper bounded by the objective function value of the SVM running over the optimal choice of feature vectors. Meanwhile, we expect a more refined bound when the side information is not vacuous for the batch setting.

Moreover, a well known generalization bound for linear classifiers (see, for example, [16]) gives

$$\Pr_{(j,y) \sim \mathcal{V}_i} (\text{sgn}(\bar{\mathbf{P}}_i \bar{\mathbf{Q}}_j^\top) \neq y) \leq \frac{1}{T_i} \sum_{t: i_t=i} h_\gamma(y_t \bar{\mathbf{P}}_{i_t} \bar{\mathbf{Q}}_{j_t}^\top) + 2\sqrt{\frac{1}{\gamma^2 T_i}} + \sqrt{\frac{\ln(1/\delta)}{2T_i}}$$

for every $i \in [m]$, where $\mathcal{V}_i$ is the conditional distribution of $\mathcal{V}$ given that the first component is i , and T_i is the number of $t \in [T]$ that satisfies $i_t = i$. Assume for simplicity that $\mathcal{V}(i) = \sum_{j,y} \mathcal{V}(i,j,y) = 1/m$ and $T_i = T/m$ for every i . Then,

$$\begin{aligned} R(f_{\mathbf{W}_{\mathbf{P},\mathbf{Q}}}) &= \sum_{i=1}^m \mathcal{V}(i) \Pr_{(j,y) \sim \mathcal{V}_i} (\text{sgn}(\bar{\mathbf{P}}_i \bar{\mathbf{Q}}_j^\top) \neq y) \\ &\leq \sum_{i=1}^m \mathcal{V}(i) \left(\frac{1}{T_i} \sum_{t: i_t=i} h_\gamma(y_t \bar{\mathbf{P}}_{i_t} \bar{\mathbf{Q}}_{j_t}^\top) + 2\sqrt{\frac{1}{\gamma^2 T_i}} + \sqrt{\frac{\ln(1/\delta)}{2T_i}} \right) \\ &= \frac{1}{T} \text{hloss}(\mathcal{S}, (\mathbf{P}, \mathbf{Q}), \gamma) + 2\sqrt{\frac{m}{\gamma^2 T}} + \sqrt{\frac{m \ln(1/\delta)}{2T}}. \end{aligned}$$

Minimizing $R(f_{\mathbf{W}_{\mathbf{P},\mathbf{Q}}})$ over all $\mathbf{P}$ and $\mathbf{Q}$ such that $\mathbf{P}\mathbf{Q}^\top \in \mathbb{R}^{m \times n}$, the bound obtained is very similar to our bound (41). This observation implies that our hypothesis has generalization ability competitive with the optimal linear classifiers $\bar{\mathbf{P}}$ over the optimal feature vectors $\bar{\mathbf{Q}}$.

6. Conclusions

In this paper, on the one hand, we consider a variant of the OSDP problem, whose constraint to the decision space is the bounded with Γ -trace norm, a generalization of the trace norm in the standard OSDP problem. By involving an FTRL-based algorithm with a Γ -calibrated log-determinant regularizer, we achieve a regret bound irrelevant to the size of the matrix. On the other hand, we utilize our result to the online binary matrix completion problem with side information, particularly. We reduce OBMC with side information to our new OSDP framework and parameterize the side information into Γ . With our proposed algorithm, combining the result to the generalized OSDP framework, we obtain a tighter mistake bound than the previous work by removing the logarithmic factor. Furthermore, our result in the off-line version is better than the traditional margin-based SVMs with the best kernel, when the side information is not vacuous.

Author Contributions: Conceptualization, K.-i.M., Y.L., K.H. and E.T.; methodology, K.-i.M. and Y. L.; validation, K.H. and E.T.; formal analysis, Y.L., K.H. and E.T.; writing—original draft preparation, Y.L.; writing—review and editing, Y.L., K.H. and E.T. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by JSPS KAKENHI Grant Numbers JP19H04174, JP19H04067 and JP20H05967, respectively.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Not applicable.

Acknowledgments: We thank the reviewers for their suggestions and comments.

Conflicts of Interest: The authors declare no conflict of interest.

Appendix A

Appendix A.1. Proof of Main Proposition and Main Theorem

Lemma A1 ([7]). Let P and Q be probability distributions over $\mathbb{R}^N$ and $\phi_P(u)$ and $\phi_Q(u)$ be their characteristic functions, respectively. Then

$$\max_{u \in \mathbb{R}^N} |\phi_P(u) - \phi_Q(u)| \leq \int_{\mathbb{R}^N} |P(x) - Q(x)| dx, \quad (\text{A1})$$

the right hand side is the total variation distance between any distribution Q and P .

Lemma A2 ([18]). Let P and Q be probability distributions over $\mathbb{R}^N$ with total variation distance δ . Then

$$H(\alpha P + (1 - \alpha)Q) \leq \alpha H(P) + (1 - \alpha)H(Q) - \alpha(1 - \alpha)\delta^2, \quad (\text{A2})$$

where $H(P) = \mathbb{E}_{x \sim P}[\ln P(x)]$.

Lemma A3 ([7]). For any probability distribution P over $\mathbb{R}^N$ with zero mean and covariance matrix Σ , its entropy is bounded by the log-determinant of covariance matrix. That is,

$$-H(P) \leq \frac{1}{2} \ln(\det(\Sigma)(2\pi e)^N). \quad (\text{A3})$$

Note that the equality holds if and only if P is a Gaussian distribution.

Lemma A4 ([7]).

$$e^{-\frac{x}{2}} - e^{-\frac{1-x}{2}} \geq \frac{e^{-1/4}}{2}(1 - 2x), \quad (\text{A4})$$

for $0 \leq x \leq 1/2$.

Appendix A.2. Definition of Biclustered Structure and Ideal Quasi Dimension

As in [2], we define the class of (k, l) -biclustered structure matrices as follows.

Definition A1. For $m \geq k$ and $n \geq l$, the class of (k, l) -binary biclustered matrices is defined as

$$\mathbb{B}_{k,l}^{m \times n} = \{\mathbf{U} \in \{-1, +1\}^{m \times n} : \mathbf{r} \in [k]^m, \mathbf{c} \in [l]^n, \mathbf{V} \in \{1, -1\}^{k \times l}, \mathbf{U}_{i,j} = \mathbf{V}_{r_i, c_j}, i \in [m], j \in [n]\}.$$

Visually, if $\mathbf{U} \in \mathbb{B}_{k,l}^{m \times n}$, there exists a permutation of rows and columns over $\mathbf{U}$ such that after this permutation, $\mathbf{U}$ becomes a $k \times l$ block matrix where in each block, all the entries are uniformly labeled -1 or $+1$. Formally, for any matrix $\mathbf{U} \in \mathbb{B}_{k,l}^{m \times n}$ we can decompose $\mathbf{U} = \mathbf{R}\mathbf{U}^*\mathbf{C}^\top$ for some $\mathbf{U}^* \in \{-1, +1\}^{k \times l}$, $\mathbf{R} \in \mathcal{B}^{m,k}$ and $\mathbf{C} \in \mathcal{B}^{n,l}$, where $\mathcal{B}^{m,d} = \{\mathbf{R} \in \{0, 1\}^{m \times d} : \|\mathbf{R}_i\|_2 = 1, i \in [m], \text{rank}(\mathbf{R}) = d\}$.

In the following theorem, we demonstrate that in the OBMC problem with side information, the quasi-dimension can be effectively bounded with an appropriate choice of side information, if the underlying matrix $\mathbf{U}$ is biclustered.

Theorem A1 ([2]). If $\mathbf{U} \in \mathbb{B}_{k,l}^{m \times n}$ define $\mathcal{D}_{M,N}^o(\mathbf{U})$ as

$$\mathcal{D}_{M,N}^o(\mathbf{U}) = 2\text{Tr}(\mathbf{R}^\top \mathbf{M} \mathbf{R})\alpha_M + 2\text{Tr}(\mathbf{C}^\top \mathbf{N} \mathbf{C})\alpha_N + 2k + 2l, \quad (\text{A5})$$

where $\mathbf{M}, \mathbf{N}$ are PD-Laplacian, as the minimum over all decompositions of $\mathbf{U} = \mathbf{R}\mathbf{U}^*\mathbf{C}^\top$ for some $\mathbf{U}^* \in \{-1, +1\}^{k \times l}$, $\mathbf{R} \in \mathcal{B}^{m,k}$ and $\mathbf{C} \in \mathcal{B}^{n,l}$. Thus, for $\mathbf{U} \in \mathbb{B}_{k,l}^{m \times n}$,

$$\mathcal{D}_{M,N}^\gamma(\mathbf{U}) \leq \mathcal{D}_{M,N}^o(\mathbf{U}), \quad (\text{A6})$$

if $\|\mathbf{U}\|_{\max} \leq \frac{1}{\gamma}$.

Moreover, we define the max-norm of a matrix $\mathbf{U} \in \mathbb{R}^{m \times n}$ as follows:

$$\|\mathbf{U}\|_{\max} = \min_{\mathbf{P}\mathbf{Q}^T = \mathbf{U}} \left\{ \max_{1 \leq i \leq m} \|\mathbf{P}_i\| \max_{1 \leq j \leq n} \|\mathbf{Q}_j\| \right\}. \quad (\text{A7})$$

Furthermore we define the quasi-dimension of a matrix $\mathbf{U}$ with $\mathbf{M} \in \mathbb{S}_{++}^{m \times m}$ and $\mathbf{N} \in \mathbb{S}_{++}^{n \times n}$ at margin γ as

$$\mathcal{D}_{\mathbf{M}, \mathbf{N}}^{\gamma}(\mathbf{U}) = \min_{\bar{\mathbf{P}}\bar{\mathbf{Q}}^T = \gamma\mathbf{U}} \alpha_{\mathbf{M}} \text{Tr}(\bar{\mathbf{P}}^T \mathbf{M} \bar{\mathbf{Q}}) + \alpha_{\mathbf{N}} \text{Tr}(\bar{\mathbf{Q}}^T \mathbf{N} \bar{\mathbf{Q}}). \quad (\text{A8})$$

See Section 4.1 from [2]; if $\mathbf{U}$ is a (k, l) -biclustered structured matrix, they show an example where $\mathcal{D}_{\mathbf{M}, \mathbf{N}}^{\gamma}(\mathbf{U}) \in O(k + l)$ with ideal side information. It implies that in a realizable case, such as that for the sequence $\mathcal{S} = ((i_t, j_t), y_t)_{t=1}^T$ where $y_t = (\bar{\mathbf{P}}\bar{\mathbf{Q}}^T)_{i_t, j_t} = \mathbf{U}_{i_t, j_t}$ for some $\mathbf{U}$ satisfying the conditions in [2], and $(\bar{\mathbf{P}}, \bar{\mathbf{Q}}) = \arg \min_{\mathbf{P}, \mathbf{Q}} \mathcal{D}_{\mathbf{M}, \mathbf{N}}^{\gamma}(\mathbf{U})$, with appropriate side information, the quasi-dimension estimator $\hat{D} \in O(k + l)$. It follows that with ideal side information, our proposed algorithm in the main part (Algorithm 1) is effective when the binary comparator matrix $\mathbf{U}$ is (k, l) -biclustered.

References

- Herbster, M.; Pasteris, S.; Pontil, M. Mistake bounds for binary matrix completion. In Proceedings of the Advances in Neural Information Processing Systems, Barcelona, Spain, 5–10 December 2016; pp. 3954–3962.
- Herbster, M.; Pasteris, S.; Tse, L. Online Matrix Completion with Side Information. *Adv. Neural Inf. Process. Syst.* **2020**, *33*, 20402–20414.
- Zhang, X.; Wang, L.; Yu, Y.; Gu, Q. A primal-dual analysis of global optimality in nonconvex low-rank matrix recovery. In Proceedings of the International Conference on Machine Learning, PMLR, Stockholm, Sweden, 10–15 July 2018; pp. 5862–5871.
- Beckerleg, M.; Thompson, A. A divide-and-conquer algorithm for binary matrix completion. *Linear Algebra Its Appl.* **2020**, *601*, 113–133. [\[CrossRef\]](#)
- Bennett, J.; Lanning, S. The netflix prize. In Proceedings of the KDD Cup and Workshop, Halifax, NS, Canada, 13–17 August 2007; Volume 2007, p. 35.
- Hazan, E.; Kale, S.; Shalev-Shwartz, S. Near-optimal algorithms for online matrix prediction. *SIAM J. Comput.* **2016**, *46*, 744–773. [\[CrossRef\]](#)
- Moridomi, K.I.; Hatano, K.; Takimoto, E. Online linear optimization with the log-determinant regularizer. *IEICE Trans. Inf. Syst.* **2018**, *101*, 1511–1520. [\[CrossRef\]](#)
- Cesa-Bianchi, N.; Lugosi, G. *Prediction, Learning, and Games*; Cambridge University Press: Cambridge, UK, 2006.
- Shalev-Shwartz, S. Online learning and online convex optimization. *Found. Trends Mach. Learn.* **2012**, *4*, 107–194. [\[CrossRef\]](#)
- Hazan, E. Introduction to Online Convex Optimization. *Found. Trends Optim.* **2016**, *2*, 157–325. [\[CrossRef\]](#)
- Abernethy, J.D. Can We Learn to Gamble Efficiently? In Proceedings of the 23rd Annual Conference on Learning Theory (COLT), Haifa, Israel, 27–29 June 2010; pp. 318–319.
- Shamir, O.; Shalev-Shwartz, S. Collaborative filtering with the trace norm: Learning, bounding, and transducing. In Proceedings of the 24th Annual Conference on Learning Theory, Budapest, Hungary, 9–11 July 2011; pp. 661–678.
- Cesa-Bianchi, N.; Shamir, O. Efficient online learning via randomized rounding. In Proceedings of the Advances in Neural Information Processing Systems, Granada, Spain, 12–15 December 2011; pp. 343–351.
- Koltchinskii, V.; Lounici, K.; Tsybakov, A.B. Nuclear-norm penalization and optimal rates for noisy low-rank matrix completion. *Ann. Stat.* **2011**, *39*, 2302–2329. [\[CrossRef\]](#)
- Gentile, C.; Herbster, M.; Pasteris, S. Online Similarity Prediction of Networked Data from Known and Unknown Graphs. In Proceedings of the Conference on Learning Theory, Princeton, NJ, USA, 12–14 June 2013; pp. 662–695.
- Mohri, M.; Rostamizadeh, A.; Talwalkar, A. *Foundations of Machine Learning*; Adaptive Computation and Machine Learning Series; MIT Press: Cambridge, MA, USA, 2012.
- Liu, Y.; Moridomi, K.I.; Hatano, K.; Takimoto, E. An online semi-definite programming with a generalised log-determinant regularizer and its applications. In Proceedings of the Asian Conference on Machine Learning, PMLR, Virtual, 17–19 November 2021; pp. 1113–1128.
- Christiano, P. Online local learning via semidefinite programming. In Proceedings of the Forty-Sixth Annual ACM Symposium on Theory of Computing, New York, NY, USA, 31 May–3 June 2014; pp. 468–474.
- Valiant, L.G. A theory of the learnable. *Commun. ACM* **1984**, *27*, 1134–1142. [\[CrossRef\]](#)
- Shalev-Shwartz, S.; Ben-David, S. *Understanding Machine Learning: From Theory to Algorithms*; Cambridge University Press: Cambridge, UK, 2014.