



Continuous optimization

A distributionally robust chance-constrained kernel-free quadratic surface support vector machine

Fengming Lin^a, Shu-Cherng Fang^a, Xiaolei Fang^a, Zheming Gao^{b,d,*}, Jian Luo^c^a Edward P. Fitts Department of Industrial and Systems Engineering, North Carolina State University, Raleigh, North Carolina, 27695, USA^b College of Information Science and Engineering, Northeastern University, Shenyang, Liaoning, 110819, China^c International Business School, Hainan University, Haikou, Hainan, 570228, China^d Yunnan Key Laboratory of Service Computing, Kunming, Yunnan, 650221, China

ARTICLE INFO

Keywords:

Data science

Kernel-free support vector machine

Robust classification

Distributionally robust optimization

Chance-constrained optimization

ABSTRACT

This paper studies the problem of constructing a robust nonlinear classifier when the data set involves uncertainty and only the first- and second-order moments are known a priori. A distributionally robust chance-constrained kernel-free quadratic surface support vector machine (SVM) model is proposed using the moment information of the uncertain data. The proposed model is reformulated as a semidefinite programming problem and a second-order cone programming problem for efficient computations. A geometric interpretation of the proposed model is also provided. For commonly used data without prescribed uncertainty, a cluster-based data-driven approach is introduced to retrieve the hidden moment information that enables the proposed model for robust classification. Extensive computational experiments using synthetic and public benchmark data sets with or without uncertainty involved support the superior performance of the proposed model over other state-of-the-art SVM models, particularly when the data sets are massive and/or imbalanced.

1. Introduction

Support Vector Machines (SVMs) are often used for classification in supervised machine learning. Given a set of N data points $\{(x^i, y^i) | x^i \in \mathbb{R}^n, y^i \in \{-1, 1\}, i = 1, \dots, N\}$, a linear soft support vector machine (LSSVM) can be represented as the following linearly constrained convex quadratic programming problem (Cortes & Vapnik, 1995):

$$\begin{aligned} \min \quad & \frac{1}{2} \|w\|_2^2 + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & y^i (w^T x^i + b) \geq 1 - \xi_i, \quad i = 1, \dots, N, \\ & w \in \mathbb{R}^n, \quad b \in \mathbb{R}, \quad \xi \in \mathbb{R}_+^N, \end{aligned} \quad (\text{LSSVM})$$

where $C > 0$ is a given parameter. The variables of w and b determine a separation hyperplane $H(w, b) \triangleq \{x \in \mathbb{R}^n | w^T x + b = 0\}$, and the slack vector ξ introduces a “soft margin” to accommodate the data that are not linearly separable. For nonlinear classification, the data can be lifted to a higher dimensional space for linear separation using a feature map $\phi: \mathbb{R}^n \rightarrow \mathbb{R}^l$ with $l > n$. In this case, we consider the following optimization problem:

$$\begin{aligned} \min \quad & \frac{1}{2} \|v\|_2^2 + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & y^i (v^T \phi(x^i) + d) \geq 1 - \xi_i, \quad i = 1, \dots, N, \\ & v \in \mathbb{R}^l, \quad d \in \mathbb{R}, \quad \xi \in \mathbb{R}_+^N. \end{aligned} \quad (\text{KSSVM})$$

The kernel trick is used to solve the dual problem of (KSSVM) by introducing the kernel function $K(x^i, x^j) \triangleq \phi(x^i)^T \phi(x^j)$ (Zhou, 2021). Commonly used kernel functions $K(\cdot, \cdot)$ include the polynomial kernel and radial basis function kernel (i.e., Gaussian kernel). Considering the drawbacks of selecting a proper kernel function (Jiménez-Cordero, Morales, & Pineda, 2021) and adjusting its embedded parameters, some kernel-free nonlinear SVMs have recently been proposed (Dagher, 2008; Luo, Fang, Deng, & Guo, 2016; Luo, Yan, & Tian, 2020). One representative model is the following quadratic surface support vector machine (QSSVM):

$$\begin{aligned} \min \quad & \sum_{i=1}^N \|M x^i + w\|_2^2 + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & y^i \left(\frac{1}{2} (x^i)^T M x^i + w^T x^i + b \right) \geq 1 - \xi_i, \quad i = 1, \dots, N, \\ & M \in \mathbb{S}^n, \quad w \in \mathbb{R}^n, \quad b \in \mathbb{R}, \quad \xi \in \mathbb{R}_+^N, \end{aligned} \quad (\text{QSSVM})$$

where $\mathbb{S}^n$ is the set of n -dimensional symmetric matrices; M , w , and b determine a separation quadratic surface $Q(M, w, b) \triangleq \{x \in \mathbb{R}^n | \frac{1}{2} x^T M x + x^T w + b = 0\}$. The (QSSVM) model has been extended to a kernel-free quartic surface SVM model by utilizing the double well potential function of degree four (Gao, Fang, Luo, & Medhin, 2021).

While both the kernel-based and kernel-free nonlinear SVMs have achieved promising performance in some real-world applications, their

* Corresponding author at: College of Information Science and Engineering, Northeastern University, Shenyang, Liaoning, 110819, China.

E-mail address: gaozheming@ise.neu.edu.cn (Z. Gao).

classification performances still need to be investigated when uncertainty is involved in the training data. For instance, the classification task on benign and malignant tumors (Bertsimas, Dunn, Pawlowski, & Zhuo, 2019), whose training data includes features derived from digitized images, such as cell nuclei radius, texture, and symmetry. Even though these features are precisely measured, the existence of image noise and measurement inaccuracy may yield data uncertainty and affect the classification accuracy. In addition to medical applications, challenges raised by data uncertainties are urgent to be resolved in other fields, including battery failure detection (Luo, Fang, Deng, & Tian, 2022) and biological gene expression (Ben-Tal, Bhadra, Bhattacharya, & Nath, 2011). In datasets requiring imputation for missing data (Shivaswamy, Bhattacharya, & Smola, 2006), additional uncertainties are introduced. Given the presence of data uncertainty within real-world applications, neglecting to recognize the uncertainty might lead to a substantial decline in classification performance (Goldfarb & Iyengar, 2003).

Recent studies indicate that classifiers explicitly addressing uncertainty in the training data outperform those ignoring such information (Wang, Fan, & Pardalos, 2018). This paper introduces a novel maximum-margin nonlinear SVM resilient to data uncertainty. It handles the underlying data uncertainty by utilizing moment information instead of the distributional assumptions on data.

1.1. Relevant works

Optimization under uncertainty has been addressed by several complementary modeling paradigms that differ mainly in the representation of uncertainty. SVM models applying robust optimization techniques are developed for applications whose data points are fluctuating within an uncertain set, specified by the l_p -norm uncertainty (Trafalis & Gilbert, 2006), ellipsoidal uncertainty (Bhattacharyya, Grate, Jordan, Ghaoui, & Mian, 2004), and others (Singla, Ghosh, & Shukla, 2020; Wang & Pardalos, 2014). Applying robust optimization in a principled way of uncertain data, Bertsimas et al. (2019) investigate the SVMs, logistic regression, and decision trees, among which the robust SVM performed the best. Nonetheless, the robust models generally tend to be on the conservative side since they ignore the hidden distribution information embedded in the data sets.

Consider a set of data points with uncertain inputs following some underlying probability distributions F_i , i.e., $\tilde{\mathbf{x}}^i \sim F_i$, for $i = 1, \dots, N$. For a given tolerance level $0 < \epsilon < 1$, a chance constraint at the point $(\tilde{\mathbf{x}}^i, y^i)$,

$$\mathbb{P}_{F_i} \{y^i (\mathbf{w}^T \tilde{\mathbf{x}}^i + b) \leq 1 - \xi_i\} \leq \epsilon,$$

can be used to ensure that the probability of misclassifying $\tilde{\mathbf{x}}^i$ is no larger than ϵ . The chance-constrained optimization problems are non-convex and hard to solve in general. In the literature, Peng, Gianpiero, and Zhihua (2023) adopt the sample average approximation method to formulate a mixed integer programming problem for a chance-constrained conic-segmentation SVM with an empirical distribution. In fact, a true distribution is hard to estimate, and even a good estimation may still cause the “optimizer’s curse” (Kuhn, Esfahani, Nguyen, & Shafieezadeh-Abadeh, 2019) with discontent performance.

Instead of relying on a single estimate of F_i , a distribution family D_i could hedge against the uncertainty in data distribution. D_i , which is also known as the ambiguity set (Lin, Fang, & Gao, 2022), consists of probability distributions possessing certain properties of the true distribution F_i . The following distributionally robust chance constraints are developed to ensure that a linear SVM works best in the worst case over D_i :

$$\sup_{F_i \in D_i} \mathbb{P}_{F_i} \{y^i (\mathbf{w}^T \tilde{\mathbf{x}}^i + b) \leq 1 - \xi_i\} \leq \epsilon, \quad i = 1, \dots, N. \quad (1)$$

When the ambiguity set D_i is constructed based on the first- and second-order moments, the distributionally robust chance-constrained

linear soft SVM (DRC-LSSVM) model is developed with the chance constraints defined by (1). Shivaswamy et al. (2006) adopt the multivariate Chebyshev inequality to derive a second-order cone programming (SOCP) reformulation for the DRC-LSSVM model. Ben-Tal et al. (2011) employ the Bernstein bounds to include richer partial information for constructing a less conservative SOCP reformulation. Wang et al. (2018) derive both semidefinite programming (SDP) and SOCP reformulations for DRC-LSSVM, and they further design a stochastic gradient-based method for improving the computational efficiency in large-scale classification cases (Wang, Fan, & Pardalos, 2017). Considering dependency among the random input points, Khanjani-Shiraz, Babapour-Azar, Hosseini-Nodeh, and Pardalos (2023) propose a robust joint chance-constrained linear SVM. The DRC-LSSVM model has also been applied to different contexts with promising performance, such as data with missing values (Shivaswamy et al., 2006) and semi-supervised classifications (Huang, Song, Gupta, & Wu, 2013). A kernel-free DRC support vector regression (SVR) model is proposed to solve regression problems, which also shows superior performance over other well-established SVR models (Luo et al., 2022). Similar links to supervised training with uncertain data employing distributionally robust optimization under the Wasserstein metric have been investigated for linear SVMs (Ma & Wang, 2021), regression models (Chen & Paschalidis, 2020; Kuhn et al., 2019) and reinforcement learning models (Chen & Paschalidis, 2020).

In summary, the literature indicates that the application of distributionally robust optimization enhances the capability of conventional SVMs in addressing uncertain classification tasks. Previous studies on distributionally robust chance-constrained SVMs demonstrate higher accuracy compared to nominal methods in some cases, but most of them focus on linear SVMs, limiting the applicability to nonlinear classification. Our contribution builds on these efforts by proposing a distributionally robust chance-constrained kernel-free nonlinear SVM, utilizing a quadratic surface for increased flexibility in handling nonlinear data. We compare this approach to state-of-the-art SVMs, assessing the impact of adding robustness to different models and evaluating their performance through computational experiments with balanced, imbalanced, and massive datasets.

1.2. Contributions

To deal with the nonlinear binary classification cases with uncertain data, in this paper, we propose a kernel-free quadratic surface SVM model that considers the distributionally robust chance constraints (2).

$$\sup_{F_i \in D_i} \mathbb{P}_{F_i} \left\{ y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i \right\} \leq \epsilon, \quad i = 1, \dots, N. \quad (2)$$

Certain analytic properties of the proposed model are rigorously investigated. In addition, extensive computational experiments are conducted to validate the effectiveness and efficiency of the proposed model in solving binary classification problems with and without data uncertainty. The main contributions of this paper are summarized as follows.

- We propose a distributionally robust chance-constrained quadratic SVM (DRC-QSSVM) model utilizing the first- and second-order moments embedded in the data set, which characterizes the uncertainty of the classification problem. To the best of our knowledge, it is the first study of utilizing kernel-free nonlinear SVM models to deal with classification problems under data uncertainty. As the quadratic structure of the distributionally robust chance constraints in the proposed model complicates the analysis, we explicitly derive the SDP and the SOCP reformulations of the proposed model for computational efficiency. In addition, a geometric interpretation of the distributionally robust quadratic chance constraints is provided for a better understanding of the proposed model.

- We extend the proposed model to handle commonly used data without uncertainty. To retrieve the moment information needed by the proposed model, a cluster-based data-driven approach is designed. Surprisingly, the proposed model provides higher classification accuracy than the other tested state-of-the-art SVM models in the computational experiments. It strengthens the applicability of the proposed model to real-life applications.
- The computational results verify the classification effectiveness of the proposed DRC-QSSVM model. As a maximum-margin SVM model that can explicitly use moment information to handle input data with uncertainty, the proposed model outperforms the most related DRC-LSSVM model on some synthetic and public benchmark data sets. Also, the results from some extensive computational experiments on massive and imbalanced data sets verify the dominant classification accuracy of the proposed model over other state-of-the-art SVM models. It reveals the significance of the proposed model that reframing a specific problem as one characterized by uncertainty and subsequently addressing the resultant uncertain formulation have the potential to yield remarkably improved outcomes.

The rest of the paper is organized as follows. In Section 2, we propose a distributionally robust chance-constrained quadratic surface support vector machine model for nonlinear classification with uncertain data knowing the first- and second-order moments. The SDP and SOCP reformulations are derived for computational efficiency. A geometric interpretation is also provided to show how the proposed model works on uncertain data. Section 3 presents a data-driven approach for applying the proposed model to classify commonly used data sets without moment information. Synthetic data sets and public benchmark data sets are included in Section 4 for validating the effectiveness and efficiency of the proposed model and comparing the performance of the proposed model with other well-known SVM models. Section 5 concludes the paper.

Notations: In this paper, we use lower-case boldface letters to denote vectors and upper-case boldface letters to denote matrices. Random variables are represented by symbols with tildes, while their realizations are denoted by the same symbols without tildes. $\mathbb{S}^n$ denotes the set of symmetric matrices of dimension n . For any two matrices $\mathbf{A}, \mathbf{B} \in \mathbb{S}^n$, $\mathbf{A} \bullet \mathbf{B} = \text{Trace}(\mathbf{A}\mathbf{B})$ denotes the trace of the product of $\mathbf{A}$ and $\mathbf{B}$.

2. Distributionally robust chance-constrained quadratic SVM

This section considers a binary classification problem for uncertain data sets with known first- and second-order moments information and proposes the DRC-QSSVM model. Section 2.1 constructs an ambiguity set based on the first two moments to formalize the proposed model. An equivalent SDP model is derived in Section 2.2. Section 2.3 presents an explicit geometric interpretation of the conceptual chance constraints. An SOCP model for efficient computation is derived in Section 2.4.

2.1. DRC-QSSVM model

In this paper, each uncertain input $\tilde{\mathbf{x}}^i$ in a data set of $\{(\tilde{\mathbf{x}}^i, y^i) | \tilde{\mathbf{x}}^i \in \mathbb{R}^n, y^i \in \{-1, 1\}, i = 1, \dots, N\}$ is considered as a random vector, i.e., $\tilde{\mathbf{x}}^i : \Xi_i \rightarrow \mathbb{R}$ and $\tilde{\mathbf{x}}^i \sim F_i$, for an outcome space Ξ_i and its σ -algebra $\mathcal{F}_i \subseteq 2^{\Xi_i}$, and $F_i : \mathcal{F}_i \rightarrow \mathbb{R}$ is a probability measure on $(\Xi_i, \mathcal{F}_i)$, for $i = 1, \dots, N$. Let $\mathcal{M}(\Xi_i, \mathcal{F}_i)$ denote the space of all probability measures defined on $(\Xi_i, \mathcal{F}_i)$. In this way, $F_i \in \mathcal{M}(\Xi_i, \mathcal{F}_i)$ and they are assumed to be mutually independent for $i = 1, \dots, N$. In Section 2, the mean $\boldsymbol{\mu}_i \triangleq \mathbb{E}_{F_i}[\tilde{\mathbf{x}}_i] \in \mathbb{R}^n$ and covariance matrix $\boldsymbol{\Sigma}_i \triangleq \mathbb{E}_{F_i}[(\tilde{\mathbf{x}}_i - \mathbb{E}_{F_i}[\tilde{\mathbf{x}}_i])(\tilde{\mathbf{x}}_i - \mathbb{E}_{F_i}[\tilde{\mathbf{x}}_i])^T] \in \mathbb{S}_+^n$ are particularly assumed to be known. Without loss of generality, $\boldsymbol{\Sigma}_i$ is

considered to be positive definite. A moment-based ambiguity set $\mathcal{D}$ is then defined by $\mathcal{D} \triangleq \bigcup_{i=1}^N \mathcal{D}_i(\tilde{\mathbf{x}}^i; \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$ with

$$\mathcal{D}_i(\tilde{\mathbf{x}}^i; \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i) \triangleq \left\{ F_i \in \mathcal{M}(\Xi_i, \mathcal{F}_i) \mid \begin{cases} \mathbb{P}(\tilde{\mathbf{x}}^i \in \Xi_i) = 1, \\ \mathbb{E}_{F_i}[\tilde{\mathbf{x}}^i] = \boldsymbol{\mu}_i, \\ \mathbb{E}_{F_i}[(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)^T] = \boldsymbol{\Sigma}_i \end{cases} \right\}, \quad (3)$$

and $\mathcal{D}_i$ abbreviates $\mathcal{D}_i(\tilde{\mathbf{x}}^i; \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$ for $i = 1, \dots, N$. In our model, we set $\mathbb{R}^n$ as the support set Ξ_i for $i = 1, \dots, N$.

We aim to determine a quadratic surface $Q(\mathbf{M}, \mathbf{w}, b) = \{\mathbf{x} \in \mathbb{R}^n | \frac{1}{2} \mathbf{x}^T \mathbf{M} \mathbf{x} + \mathbf{x}^T \mathbf{w} + b = 0\}$, where $\mathbf{M} \in \mathbb{S}^n$, $\mathbf{w} \in \mathbb{R}^n$ and $b \in \mathbb{R}$, which separates the two classes of points with the maximum margin and bounded misclassification probability with respect to all distributions in $\mathcal{D}$. Adopting the concept of “total approximated relative geometric margins” used in (QSSVM) (Luo et al., 2016), the min-max approach used in robust optimization helps form the following objective function:

$$\min_{\mathbf{M}, \mathbf{w}, b, \xi} \sup_{F_i \in \mathcal{D}_i} \left\{ \sum_{i=1}^N \mathbb{E}_{F_i} \|\mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}\|_2^2 \right\} + C \sum_{i=1}^N \xi_i. \quad (4)$$

Note that for any $F_i \in \mathcal{D}_i$,

$$\begin{aligned} \mathbb{E}_{F_i} \|\mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}\|_2^2 &= \mathbb{E}_{F_i} [(\tilde{\mathbf{x}}^i)^T \mathbf{M}^T \mathbf{M} \tilde{\mathbf{x}}^i] + 2\mathbf{w}^T \mathbf{M} \mathbb{E}_{F_i}[\tilde{\mathbf{x}}^i] + \mathbf{w}^T \mathbf{w} \\ &= \boldsymbol{\mu}_i^T \mathbf{M}^T \mathbf{M} \boldsymbol{\mu}_i + (\mathbf{M}^T \mathbf{M}) \bullet \boldsymbol{\Sigma}_i + 2\mathbf{w}^T \mathbf{M} \boldsymbol{\mu}_i + \mathbf{w}^T \mathbf{w} \\ &= \|\mathbf{M} \boldsymbol{\mu}_i + \mathbf{w}\|_2^2 + \|\boldsymbol{\Sigma}_i^{\frac{1}{2}} \mathbf{M}\|_F^2, \end{aligned} \quad (5)$$

where $\|\cdot\|_F$ is the Frobenius norm. When $\{F_i\}_i$ are mutually independent, the supremum and summation operations are exchangeable, and consequently, $\sup_{F_i \in \mathcal{D}_i} \{\sum_{i=1}^N \mathbb{E}_{F_i} \|\mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}\|_2^2\} = \sum_{i=1}^N \sup_{F_i \in \mathcal{D}_i} \mathbb{E}_{F_i} \|\mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}\|_2^2 = \sum_{i=1}^N \{\|\mathbf{M} \boldsymbol{\mu}_i + \mathbf{w}\|_2^2 + \|\boldsymbol{\Sigma}_i^{\frac{1}{2}} \mathbf{M}\|_F^2\}$. The first term of the results in (5) is the approximated relative geometrical margin at the mean vector $\boldsymbol{\mu}_i$ (Luo et al., 2016), and the second term is similar to the “G-margin” defined in Gao et al. (2021). In general, the total relative geometrical margin dominates the G-margin. And the second term may serve as a regularization term that shapes the target quadratic surface. Extensive computational experiments indicate such a regularization term can be neglected without changing much of the final classifier (similar results showed by Luo et al. (2016)). Moreover, to avoid the computational difficulty induced by $\|\boldsymbol{\Sigma}_i^{\frac{1}{2}} \mathbf{M}\|_F^2$, we omit this term in the objective function. A robust classifier $Q(\mathbf{M}, \mathbf{w}, b)$ could bound the misclassification probability by ϵ ($0 < \epsilon < 1$) employing the distributionally robust chance constraints (2). Consequently, we propose the following distributionally robust chance-constrained quadratic SVM model:

$$\begin{aligned} \min \quad & \sum_{i=1}^N \|\mathbf{M} \boldsymbol{\mu}_i + \mathbf{w}\|_2^2 + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & \sup_{F_i \in \mathcal{D}_i} \mathbb{P}_{F_i} \left\{ y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i \right\} \leq \epsilon, \quad i = 1, \dots, N, \\ & \mathbf{M} \in \mathbb{S}^n, \quad \mathbf{w} \in \mathbb{R}^n, \quad b \in \mathbb{R}, \quad \xi \in \mathbb{R}_+^N. \end{aligned} \quad (\text{DRC-QSSVM})$$

It is often the case that the ambiguous chance constraints are hard to solve directly, not to mention the nonlinear functions used in (DRC-QSSVM). For $i = 1, \dots, N$, let

$$\begin{aligned} \mathcal{V}_i \triangleq \left\{ (\mathbf{M}, \mathbf{w}, b) \in \mathbb{S}^n \times \mathbb{R}^n \times \mathbb{R} \mid \sup_{F_i \in \mathcal{D}_i} \mathbb{P}_{F_i} \left\{ y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i \right\} \leq \epsilon \right\} \end{aligned} \quad (6)$$

denote the feasible sets of (DRC-QSSVM). We shall demonstrate that for any i , $\mathcal{V}_i$ has tractable SDP and SOCP representations for efficient computations in later subsections.

2.2. SDP reformulation of (DRC-QSSVM)

We first show an equivalent SDP reformulation of (DRC-QSSVM) in the next theorem.

Theorem 2.1. For the ambiguity set D defined by (3), (DRC-QSSVM) can be equivalently reformulated as the following SDP problem:

$$\begin{aligned} \min \quad & \sum_{i=1}^N \|\mathbf{M}\boldsymbol{\mu}_i + \mathbf{w}\|_2^2 + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & \beta_i - \frac{1}{\epsilon} \boldsymbol{\Gamma}_i \bullet \mathbf{R}_i \geq 0, \quad i = 1, \dots, N, \\ & \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^i \mathbf{M} & \frac{1}{2} y^i \mathbf{w} \\ \frac{1}{2} y^i \mathbf{w}^T & y^i b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0, \quad i = 1, \dots, N, \\ & \mathbf{R}_i \geq 0, \quad i = 1, \dots, N, \\ & \mathbf{M} \in \mathbb{S}^n, \mathbf{w} \in \mathbb{R}^n, b \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}_+^N, \boldsymbol{\beta} \in \mathbb{R}^N, \mathbf{R}_i \in \mathbb{S}^{n+1}, \quad i = 1, \dots, N, \end{aligned} \quad (7)$$

where $\boldsymbol{\Gamma}_i = \begin{bmatrix} \boldsymbol{\Sigma}_i + \boldsymbol{\mu}_i \boldsymbol{\mu}_i^T & \boldsymbol{\mu}_i \\ \boldsymbol{\mu}_i^T & 1 \end{bmatrix}$ denotes the second-order moment matrix.

Proof. For $i = 1, \dots, N$, we denote the indicator functions by

$$\mathbb{1}_{A_i}(\tilde{\mathbf{x}}^i) = \begin{cases} 1, & \text{if } \tilde{\mathbf{x}}^i \in A_i, \\ 0, & \text{if } \tilde{\mathbf{x}}^i \notin A_i, \end{cases}$$

where $A_i \triangleq \{\tilde{\mathbf{x}}^i \in \Xi_i | y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i\}$. Then $\varphi_i \triangleq \sup_{F_i \in \mathcal{D}_i} \mathbb{P}_{F_i} \{y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i\} = \sup_{F_i \in \mathcal{D}_i} \mathbb{E}_{F_i} [\mathbb{1}_{A_i}(\tilde{\mathbf{x}}^i)]$, which can be obtained by solving the following problem:

$$\sup \int_{\Xi_i} \mathbb{1}_{A_i}(\tilde{\mathbf{x}}^i) dF_i(\tilde{\mathbf{x}}^i) \quad (8a)$$

$$\text{s.t.} \quad \int_{\Xi_i} dF_i(\tilde{\mathbf{x}}^i) = 1, \quad (8b)$$

$$\int_{\Xi_i} \tilde{\mathbf{x}}^i dF_i(\tilde{\mathbf{x}}^i) = \boldsymbol{\mu}_i, \quad (8c)$$

$$\int_{\Xi_i} (\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)^T dF_i(\tilde{\mathbf{x}}^i) = \boldsymbol{\Sigma}_i, \quad (8d)$$

$$F_i \in \mathcal{M}(\Xi_i, \mathcal{F}_i).$$

Notice that constraint (8d) is equivalent to $\int_{\Xi_i} \tilde{\mathbf{x}}^i (\tilde{\mathbf{x}}^i)^T dF_i(\tilde{\mathbf{x}}^i) = \boldsymbol{\Sigma}_i + \boldsymbol{\mu}_i \boldsymbol{\mu}_i^T$, and the difficulty of this problem can be circumvented by using the duality theory involving moment information. Let $r_i \in \mathbb{R}$, $\mathbf{p}_i \in \mathbb{R}^n$ and $\mathbf{Q}_i \in \mathbb{S}^n$ be the dual variables corresponding to (8b), (8c) and (8d), respectively. Then the dual problem of (8) becomes

$$\begin{aligned} \inf \quad & (\boldsymbol{\Sigma}_i + \boldsymbol{\mu}_i \boldsymbol{\mu}_i^T) \bullet \mathbf{Q}_i + \mathbf{p}_i^T \boldsymbol{\mu}_i + r_i \\ \text{s.t.} \quad & (\tilde{\mathbf{x}}^i)^T \mathbf{Q}_i \tilde{\mathbf{x}}^i + (\tilde{\mathbf{x}}^i)^T \mathbf{p}_i + r_i \geq \mathbb{1}(\tilde{\mathbf{x}}^i), \quad \forall \tilde{\mathbf{x}}^i \in \Xi_i, \\ & \mathbf{Q}_i \in \mathbb{S}^n, \mathbf{p}_i \in \mathbb{R}^n, r_i \in \mathbb{R}. \end{aligned} \quad (9)$$

Let $\mathbf{N}_i = \begin{bmatrix} \mathbf{Q}_i & \frac{1}{2} \mathbf{p}_i \\ \frac{1}{2} \mathbf{p}_i^T & r_i \end{bmatrix}$, then the objective function of (9) becomes $\mathbf{N}_i \bullet \boldsymbol{\Gamma}_i$. Restoring the indicator function, the constraint of (9) becomes

$$(\tilde{\mathbf{x}}^i)^T \mathbf{Q}_i \tilde{\mathbf{x}}^i + (\tilde{\mathbf{x}}^i)^T \mathbf{p}_i + r_i \geq 1, \quad \text{if } y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i, \quad (10a)$$

$$(\tilde{\mathbf{x}}^i)^T \mathbf{Q}_i \tilde{\mathbf{x}}^i + (\tilde{\mathbf{x}}^i)^T \mathbf{p}_i + r_i \geq 0, \quad \forall \tilde{\mathbf{x}}^i \in \Xi_i. \quad (10b)$$

Constraint (10b) implies that $[(\tilde{\mathbf{x}}^i)^T \ 1] \mathbf{N}_i \begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix} \geq 0, \forall \tilde{\mathbf{x}}^i \in \Xi_i \Leftrightarrow \mathbf{N}_i \geq 0$. Constraint (10a) can be further transformed as below using the S-Lemma:

$$\begin{aligned} & (\tilde{\mathbf{x}}^i)^T \mathbf{Q}_i \tilde{\mathbf{x}}^i + (\tilde{\mathbf{x}}^i)^T \mathbf{p}_i + r_i - 1 + \alpha_i \left(y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) - 1 + \xi_i \right) \\ & \geq 0, \quad \alpha_i \geq 0, \end{aligned}$$

which means

$$[(\tilde{\mathbf{x}}^i)^T \ 1] \mathbf{N}_i \begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix} - 1 + \alpha_i \left(y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) - 1 + \xi_i \right) \geq 0, \quad \alpha_i \geq 0. \quad (11)$$

If $\alpha_i = 0$, constraint (11) implies that $[(\tilde{\mathbf{x}}^i)^T \ 1] \mathbf{N}_i \begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix} \geq 1$. This contradicts the fact of $\mathbf{N}_i \bullet \boldsymbol{\Gamma}_i \leq \epsilon, 0 < \epsilon < 1$, because $\boldsymbol{\Gamma}_i = \mathbb{E}_{F_i} \left[\begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix} \begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix}^T \right]$. Thus, we have $\alpha_i > 0$. Let $\mathbf{R}_i = \frac{1}{\alpha_i} \mathbf{N}_i$ and $\beta_i = \frac{1}{\alpha_i}$, we see that

$$\begin{aligned} (11) \Leftrightarrow & [(\tilde{\mathbf{x}}^i)^T \ 1] \frac{\mathbf{N}_i}{\alpha_i} \begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix} + \left(-\frac{1}{\alpha_i} + y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) - 1 + \xi_i \right) \geq 0 \\ \Leftrightarrow & [(\tilde{\mathbf{x}}^i)^T \ 1] \left(\mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^i \mathbf{M} & \frac{1}{2} y^i \mathbf{w} \\ \frac{1}{2} y^i \mathbf{w}^T & y^i b + \xi_i - 1 - \frac{1}{\alpha_i} \end{bmatrix} \right) \begin{bmatrix} \tilde{\mathbf{x}}^i \\ 1 \end{bmatrix} \geq 0 \\ \Leftrightarrow & \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^i \mathbf{M} & \frac{1}{2} y^i \mathbf{w} \\ \frac{1}{2} y^i \mathbf{w}^T & y^i b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0. \end{aligned}$$

Therefore, the dual problem (9) becomes

$$\begin{aligned} \inf \quad & \frac{1}{\beta_i} \boldsymbol{\Gamma}_i \bullet \mathbf{R}_i \\ \text{s.t.} \quad & \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^i \mathbf{M} & \frac{1}{2} y^i \mathbf{w} \\ \frac{1}{2} y^i \mathbf{w}^T & y^i b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0, \\ & \mathbf{R}_i \geq 0, \\ & \mathbf{R}_i \in \mathbb{S}^n, \beta_i \in \mathbb{R}_+. \end{aligned}$$

Since the strong duality holds for the pair of (8) and (9) (Similar to Delage and Ye (2010)), requiring $\varphi_i \leq \epsilon$ yields $\frac{1}{\beta_i} \boldsymbol{\Gamma}_i \bullet \mathbf{R}_i \leq \epsilon \Leftrightarrow \beta_i - \frac{1}{\epsilon} \boldsymbol{\Gamma}_i \bullet \mathbf{R}_i \geq 0$. This completes the proof. $\square$

Notice that one can reformulate (7) as a standard SDP by rewriting the summation term in the objective function in the matrix form, i.e., $\|\mathbf{M}\boldsymbol{\mu}_i + \mathbf{w}\|_2^2 + \frac{C}{N} \xi_i \leq \eta_i \Leftrightarrow \begin{bmatrix} \mathbf{I}_n & \mathbf{M}\boldsymbol{\mu}_i + \mathbf{w} \\ (\mathbf{M}\boldsymbol{\mu}_i + \mathbf{w})^T & -\frac{C}{N} \xi_i + \eta_i \end{bmatrix} \geq 0$, where $\mathbf{I}_n$ denotes the n -dimensional identical matrix. This leads to

$$\begin{aligned} \min \quad & \sum_{i=1}^N \eta_i \\ \text{s.t.} \quad & \begin{bmatrix} \mathbf{I}_n & \mathbf{M}\boldsymbol{\mu}_i + \mathbf{w} \\ (\mathbf{M}\boldsymbol{\mu}_i + \mathbf{w})^T & -\frac{C}{N} \xi_i + \eta_i \end{bmatrix} \geq 0, \quad i = 1, \dots, N, \\ & \boldsymbol{\Gamma}_i \bullet \mathbf{R}_i \leq \epsilon \beta_i, \quad i = 1, \dots, N, \\ & \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^i \mathbf{M} & \frac{1}{2} y^i \mathbf{w} \\ \frac{1}{2} y^i \mathbf{w}^T & y^i b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0, \quad i = 1, \dots, N, \\ & \mathbf{R}_i \geq 0, \quad i = 1, \dots, N, \\ & \mathbf{M} \in \mathbb{S}^n, \mathbf{w} \in \mathbb{R}^n, b \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}_+^N, \boldsymbol{\eta}, \boldsymbol{\beta} \in \mathbb{R}^N, \mathbf{R}_i \in \mathbb{S}^{n+1}, \quad i = 1, \dots, N. \end{aligned} \quad (12)$$

In this way, (12) provides an SDP reformulation of (DRC-QSSVM) for using off-the-shelf solvers.

Remark 2.1. Let $\mathcal{V}_i^{SDP}$ denote the feasible region described by the constraints associated with the i th random input in the SDP model, i.e.,

$$\mathcal{V}_i^{SDP} \triangleq \left\{ (\mathbf{M}, \mathbf{w}, b) \in \mathbb{S}^n \times \mathbb{R}^n \times \mathbb{R} \mid \begin{array}{l} \exists \beta_i, \mathbf{R}_i \geq 0, \beta_i - \frac{1}{\epsilon} \boldsymbol{\Gamma}_i \bullet \mathbf{R}_i \geq 0, \\ \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^i \mathbf{M} & \frac{1}{2} y^i \mathbf{w} \\ \frac{1}{2} y^i \mathbf{w}^T & y^i b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0 \end{array} \right\}. \quad (13)$$

Then Theorem 2.1 implies that $\mathcal{V}_i = \mathcal{V}_i^{SDP}$, for $i = 1, \dots, N$.

2.3. Geometric interpretation

To study the geometric interpretation of the distributionally robust chance constraints in (DRC-QSSVM), we first provide the following lemma:

Lemma 2.2. For any i , $\mathcal{V}_i^{SDP}$ defined in (13) is equivalent to

$$\mathcal{V}_i^{SDP} = \left\{ (\mathbf{M}, \mathbf{w}, b) \in \mathbb{S}^n \times \mathbb{R}^n \times \mathbb{R} \mid \begin{array}{l} y^i \left(\frac{1}{2} \mathbf{d}_i^T \mathbf{M} \mathbf{d}_i + \mathbf{w}^T \mathbf{d}_i + b \right) \geq 1 - \xi_i - \frac{1}{2} y^i \mathbf{M} \cdot \mathbf{D}_i^0, \\ \left[\begin{array}{cc} \Sigma_i - \epsilon \mathbf{D}_i^0 & \mu_i - \mathbf{d}_i \\ \mu_i^T - \mathbf{d}_i^T & \frac{1-\epsilon}{\epsilon} \end{array} \right] \geq 0, \mathbf{D}_i^0 \in \mathbb{S}_+^n, \mathbf{d}_i \in \mathbb{R}^n \end{array} \right\}.$$

Proof. Please see Appendix A.1. $\square$

Now, let $\mathcal{E}(\mu_i, \Sigma_i, \frac{1-\epsilon}{\epsilon}) \triangleq \{ \mathbf{x} \in \mathbb{R}^n \mid (\mathbf{x} - \mu_i)^T \Sigma_i^{-1} (\mathbf{x} - \mu_i) \leq \frac{1-\epsilon}{\epsilon} \}$ represent an ellipsoid centered at μ_i , whose shape and size are determined by Σ_i and ϵ , respectively. Considering the case of correctly classifying $\mathbf{x}^i$ for all $\mathbf{x}^i \in \mathcal{E}(\mu_i, \Sigma_i, \frac{1-\epsilon}{\epsilon})$, we define

$$\mathcal{V}_i^E \triangleq \left\{ (\mathbf{M}, \mathbf{w}, b) \in \mathbb{S}^n \times \mathbb{R}^n \times \mathbb{R} \mid y^i \left(\frac{1}{2} (\mathbf{x}^i)^T \mathbf{M} \mathbf{x}^i + \mathbf{w}^T \mathbf{x}^i + b \right) \geq 1 - \xi_i, \forall \mathbf{x}^i \in \mathcal{E} \left(\mu_i, \Sigma_i, \frac{1-\epsilon}{\epsilon} \right) \right\}.$$

In the following lemma, a geometrical interpretation of $\mathcal{V}_i^{SDP}$ is presented by discussing its relation with $\mathcal{V}_i^E$.

Lemma 2.3. For any given $(\mu_i, \Sigma_i, \epsilon)$, $\mathcal{V}_i^E \subseteq \mathcal{V}_i^{SDP}$ for $i = 1, \dots, N$. If $y^i \mathbf{M} \geq 0$, then $\mathcal{V}_i^E = \mathcal{V}_i^{SDP}$. Moreover, for the linear case with $\mathbf{M} = \mathbf{0}$, $\mathcal{V}_i^E = \mathcal{V}_i^{SDP}$.

Proof. From Lemma 2.2, the constraints in $\mathcal{V}_i^{SDP}$ can be rewritten as

$$\begin{aligned} y^i \left(\frac{1}{2} \mathbf{d}_i^T \mathbf{M} \mathbf{d}_i + \mathbf{w}^T \mathbf{d}_i + b \right) &\geq 1 - \xi_i - \frac{1}{2} y^i \mathbf{M} \cdot \mathbf{D}_i^0, \\ \forall \mathbf{d}_i \in \mathcal{E} \left(\mu_i, \Sigma_i - \epsilon \mathbf{D}_i^0, \frac{1-\epsilon}{\epsilon} \right), \end{aligned} \quad (14)$$

with $\Sigma_i - \epsilon \mathbf{D}_i^0 > 0$ and $\mathbf{D}_i^0 \geq 0$ being assumed without loss of generality. For each feasible $\mathbf{D}_i^0$, this means that the quadratic surface defined by $(\mathbf{M}, \mathbf{w}, b)$ separates all points in $\mathcal{E} \left(\mu_i, \Sigma_i - \epsilon \mathbf{D}_i^0, \frac{1-\epsilon}{\epsilon} \right)$ softly. Obviously $\mathcal{V}_i^E$ is a special case of $\mathcal{V}_i^{SDP}$ when $\mathbf{D}_i^0 = \mathbf{0}$. Therefore, $\mathcal{V}_i^E \subseteq \mathcal{V}_i^{SDP}$. If $y^i \mathbf{M} \geq 0$, we have $y^i \mathbf{M} \cdot \mathbf{D}_i^0 \geq 0$ since $\mathbf{D}_i^0 \geq 0$. The upper bound of the right-hand side in inequality (14) is obtained when $\mathbf{D}_i^0 = \mathbf{0}$. In this case, we have $\mathcal{V}_i^E = \mathcal{V}_i^{SDP}$. For the linear case with $\mathbf{M} = \mathbf{0}$, $\mathbf{D}_i^0 = \mathbf{0}$ can be derived from (A.3) in the proof of Lemma 2.2. Thus, $\mathcal{V}_i^E = \mathcal{V}_i^{SDP}$. $\square$

Remark 2.2. Lemma 2.3 implies that the proposed (DRC-QSSVM) model views each uncertain input as a set $\mathcal{E} \left(\mu_i, \Sigma_i - \epsilon \mathbf{D}_i^0, \frac{1-\epsilon}{\epsilon} \right)$ and seeks a maximum-margin classification using these ellipsoids (See Fig. 1). When $\mathbf{M} = \mathbf{0}$, this result is reduced to the linear case, which leads to a maximum-margin classification using $\mathcal{E} \left(\mu_i, \Sigma_i, \frac{1-\epsilon}{\epsilon} \right)$ as discussed by Shivaswamy et al. (2006) and Wang et al. (2018).

Remark 2.3. Note that for each i , the size of the set $\mathcal{E}(\mu_i, \Sigma_i - \epsilon \mathbf{D}_i^0, \frac{1-\epsilon}{\epsilon})$ depends on ϵ . As ϵ decreases, the size increases (See the trends shown by Fig. 1). Consider the following two extreme cases:

- $\epsilon = 0$. In this case, $\mathbb{P}_{F_i} \left\{ y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i \right\} \leq \epsilon = 0, \forall F_i \in \mathcal{D}_i$, hence, each chance constraint becomes a deterministic constraint, $y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) > 1 - \xi_i, \forall F_i \in \mathcal{D}_i$. The quadratic surface $Q(\mathbf{M}, \mathbf{w}, b)$ is required to separate data points generated from any potential distribution with fixed mean and covariance. However, there may be numerous possible distributions with the same given mean and covariance, for data

points to spread everywhere such that it becomes impossible to find such a separating surface. Also, note that when $\epsilon \rightarrow 0$, $\frac{1-\epsilon}{\epsilon} \rightarrow \infty$, which means the radius of the ellipsoid implied by $\mathcal{V}_i^{SDP}$ approaches to infinity. It is impossible to find a classifier to separate infinitely large ellipsoids.

- $\epsilon = 1$. In this case, $\mathbb{P}_{F_i} \left\{ y^i \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i \right\} \leq \epsilon = 1, \forall F_i \in \mathcal{D}_i$, becomes a trivial requirement. This means that for any data point, it is totally random to be correctly classified or not. Then, statistically speaking, the separation surface may be obtained by separating the mean vectors of random data points (See Fig. 1's last column). Note that, $\frac{1-\epsilon}{\epsilon} = 0$ as $\epsilon = 1$. This result is consistent with Lemma 2.3 when the ellipsoid reduces to the center point μ_i with a zero radius.

2.4. SOCP reformulation of (DRC-QSSVM)

SDP often requires heavy computational efforts even if it is a tractable convex program, while an equivalent SOCP model may have fewer variables for more efficient computations. This subsection presents SOCP constraints derived from the SDP reformulation of (DRC-QSSVM).

Theorem 2.4. For $i = 1, \dots, N$, $\mathcal{V}_i^{SDP}$ is equivalent to

$$\mathcal{V}_i^{SOC} \triangleq \left\{ (\mathbf{M}, \mathbf{w}, b) \in \mathbb{S}^n \times \mathbb{R}^n \times \mathbb{R} \mid \begin{array}{l} y^i \left(\frac{1}{2} \mu_i^T \mathbf{M} \mu_i + \mu_i^T \mathbf{w} + b \right) \geq 1 - \xi_i \\ + \sqrt{\frac{1-\epsilon}{\epsilon}} \|\Sigma_i^{\frac{1}{2}} (\mathbf{M} \mu_i + \mathbf{w})\|_2 - \frac{1-\epsilon}{2\epsilon} \Sigma_i \cdot y^i \mathbf{M} \end{array} \right\}. \quad (15)$$

Proof. Please see Appendix A.2. $\square$

Similar to Lemma 2.2, a geometric interpretation of $\mathcal{V}_i^{SOC}$ is given in the next lemma.

Lemma 2.5. For any given $(\mu_i, \Sigma_i, \epsilon)$, $\mathcal{V}_i^E \subseteq \mathcal{V}_i^{SOC}$ for $i = 1, \dots, N$. Moreover, for the linear case when $\mathbf{M} = \mathbf{0}$, $\mathcal{V}_i^E = \mathcal{V}_i^{SOC}$.

Proof. Please see Appendix A.3. $\square$

For more efficient computations, we adopt the vectorization technique for $\mathcal{V}_i^{SOC}$. Define the vectorizations of $\mathbf{M} \in \mathbb{S}^n$ as $\text{vec}(\mathbf{M}) \triangleq [M_{11}, \dots, M_{1n}, M_{21}, \dots, M_{2n}, M_{n1}, \dots, M_{nn}]^T \in \mathbb{R}^{n^2}$, and $\text{hvec}(\mathbf{M}) \triangleq [M_{11}, \dots, M_{1n}, M_{22}, \dots, M_{2n}, M_{n-1,n-1}, M_{n-1,n}, M_{nn}]^T \in \mathbb{R}^{\frac{n(n+1)}{2}}$. Let $\mathbf{D}_n \in \mathbb{R}^{n^2 \times \frac{n(n+1)}{2}}$ be the matrix that satisfies $\mathbf{D}_n \text{hvec}(\mathbf{M}) = \text{vec}(\mathbf{M})$. For any i , define

$$\mathbf{H}^i \triangleq [\mathbf{I}_n \otimes (\mu_i)^T \mathbf{D}_n \quad \mathbf{I}_n] \in \mathbb{R}^{n \times (\frac{n(n+1)}{2} + n)}, \quad (16)$$

where $\otimes$ denotes the Kronecker product.

Let $\mathbf{z} = [\text{hvec}(\mathbf{M})^T \mathbf{w}^T]^T \in \mathbb{R}^{\frac{n(n+1)}{2} + n}$ be the reorganized variable of $\mathbf{M}$ and $\mathbf{w}$. The objective function of (7) can be rewritten as $\sum_{i=1}^N \|\mathbf{M} \mu_i + \mathbf{w}\|_2^2 = \sum_{i=1}^N (\mathbf{H}^i \mathbf{z})^T (\mathbf{H}^i \mathbf{z}) = \mathbf{z}^T (\sum_{i=1}^N (\mathbf{H}^i)^T \mathbf{H}^i) \mathbf{z}$. Letting $\mathbf{W} = \sum_{i=1}^N (\mathbf{H}^i)^T \mathbf{H}^i$, we further have $\sum_{i=1}^N \|\mathbf{M} \mu_i + \mathbf{w}\|_2^2 = \mathbf{z}^T \mathbf{W} \mathbf{z}$. Let $\mathbf{V}^{\frac{n(n+1)}{2}} \triangleq 2\mathbf{I}_{\frac{n(n+1)}{2}} - \text{Diag}(\text{hvec}(\mathbf{I}_n))$. Then we define

$$\mathbf{r}^i \triangleq \left[\frac{1}{2} (\mathbf{V}^{\frac{n(n+1)}{2}} \text{hvec}(\frac{1-\epsilon}{\epsilon} \Sigma_i + \mu_i \mu_i^T))^T \mu_i^T \right]^T \in \mathbb{R}^{\frac{n(n+1)}{2} + n}. \quad (17)$$

We have $\frac{1}{2} \mathbf{M} \cdot (\frac{1-\epsilon}{\epsilon} \Sigma_i + \mu_i \mu_i^T) + \mu_i^T \mathbf{w} = \mathbf{z}^T \mathbf{r}^i$. Hence, (DRC-QSSVM) with a feasible set in the form of $\mathcal{V}_i^{SOC}$ has the following SOCP reformulation:

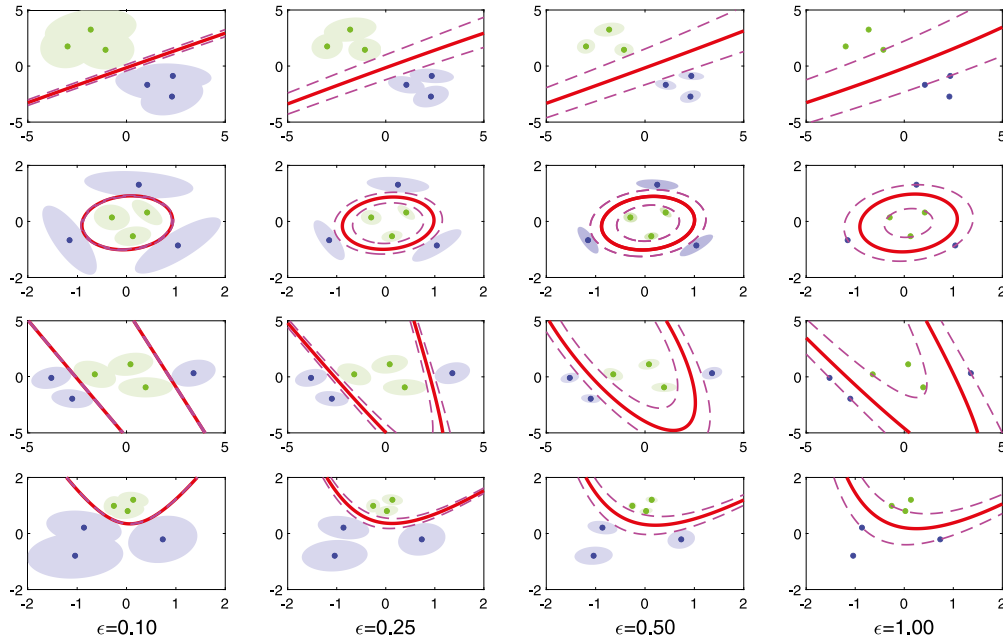


Fig. 1. Geometric interpretation of (DRC-QSSVM) on 2D synthetic data. The green and blue solid points represent $\{\mu_i\}_{i=1}^N$ in two classes, respectively. Shaded areas depict the corresponding ellipsoids $E(\mu_i, \Sigma_i, \frac{1-\epsilon}{\epsilon})$. The learned quadratic classifier, $\frac{1}{2}\mathbf{x}^T\mathbf{M}\mathbf{x} + \mathbf{x}^T\mathbf{w} + b = 0$, is represented by the red solid line, and the pink dashed lines represent quadratic curves defined by $\frac{1}{2}\mathbf{x}^T\mathbf{M}\mathbf{x} + \mathbf{x}^T\mathbf{w} + b = \pm 1$.

$$\begin{aligned} \min \quad & \mathbf{z}^T \mathbf{W} \mathbf{z} + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & y^i (\mathbf{z}^T \mathbf{r}^i + b) \geq 1 - \xi_i + \sqrt{\frac{1-\epsilon}{\epsilon}} \|\Sigma_i^{\frac{1}{2}} \mathbf{H}^i \mathbf{z}\|_2, \quad i = 1, \dots, N, \\ & \mathbf{z} \in \mathbb{R}^{\frac{n(n+1)}{2} + n}, \quad b \in \mathbb{R}, \quad \xi \in \mathbb{R}_+^N. \end{aligned} \quad (18)$$

Since matrix $\mathbf{W}$ is real, symmetric, and positive semi-definite, its Cholesky factorization leads to $\mathbf{W} = \mathbf{Q}^T \mathbf{Q}$ with $\mathbf{Q} \in \mathbb{S}_+^{\frac{n(n+1)}{2} + n}$. Consequently, we have the following standard SOCP formulation:

$$\begin{aligned} \min \quad & \theta + C \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & \|\mathbf{Q}\mathbf{z}\|_2 \leq \theta, \\ & \|\Sigma_i^{\frac{1}{2}} \mathbf{H}^i \mathbf{z}\|_2 \leq \sqrt{\frac{\epsilon}{1-\epsilon}} (y^i (\mathbf{z}^T \mathbf{r}^i + b) - 1 + \xi_i), \quad i = 1, \dots, N, \\ & \mathbf{z} \in \mathbb{R}^{\frac{n(n+1)}{2} + n}, \quad b \in \mathbb{R}, \quad \theta \in \mathbb{R}, \quad \xi \in \mathbb{R}_+^N. \end{aligned} \quad (19)$$

Some observations on the SDP model (12) and SOCP model (19) can be made here.

- Excluding the common variables $(b, \xi) \in \mathbb{R} \times \mathbb{R}_+^N$, the SDP model has N matrix variables in $\mathbb{S}^{n+1}$, one matrix variable in $\mathbb{S}^n$, two vector variables in $\mathbb{R}^n$ and one vector variable in $\mathbb{R}^{\frac{n(n+1)}{2} + n}$, while the SOCP model only has one vector variable in $\mathbb{R}^{\frac{n(n+1)}{2} + n}$ and one scalar variable.
- The SDP model has $3N$ semidefinite constraints involving $(n+1)$ -dimensional positive semi-definite cones, while the SOCP model only has one constraint involving $(n(n+1)/2 + n + 1)$ -dimensional second-order cones and N constraints involving $(n+1)$ -dimensional second-order cones.

For general practice, the number of features n is much smaller than that of input points N . Therefore, the SOCP model is more computationally friendly than the SDP model.

3. Data-driven approach for data sets without moment information

The results presented in the previous section count on the first- and second-order moments of the data set with probability uncertainty. This section utilizes the proposed (DRC-QSSVM) for classifying data sets without moment information. Fig. 2 illustrates the basic workflow of our data-driven approach. For a finite set of binary samples $\{(\mathbf{x}^i, y^i) \in \mathbb{R}^n \times \{-1, 1\}, i = 1, \dots, N\}$ without any moment information (Fig. 2a), we intend to use the proposed (DRC-QSSVM) model for building a robust classifier. The key idea is to group similar data points together and represent them by the mean and covariance of the group. Clustering is widely used to identify the inherent structure of data, while it can also serve as a pre-processing technique for classification (Zhou, 2021). Clustering aims to partition a data set into disjoint subsets, called clusters, where data points within the same cluster have high similarities. As shown in Fig. 2(b), we first use clustering algorithms to partition the given data set into N_K clusters, $C_k, k = 1, \dots, N_K$, with $N_K < N$. The K-means clustering algorithm could partition data into a finite number of homogeneous and separate clusters without using any prior knowledge. Hence we adopt the K-means++ algorithm proposed by Vassilvitskii and Arthur (2006) that enhances the performance of ordinary K-means algorithms. The well-known ELBOW validation method could help determine an optimal value of N_K . Once the clusters were obtained, the sample mean and covariance matrix can be easily calculated to estimate $(\mu_k, \Sigma_k), k = 1, \dots, N_K$ (Fig. 2(c)). Moreover, the sample mean is denoted as $\bar{\mathbf{x}}^k \triangleq \frac{1}{|C_k|} \sum_{i=1}^N \mathbf{x}^i \mathbb{1}_{C_k}(\mathbf{x}^i)$, and the sample covariance as $\mathbf{S}^k \triangleq \frac{1}{|C_k|-1} \sum_{i=1}^N (\mathbf{x}^i \mathbb{1}_{C_k}(\mathbf{x}^i) - \bar{\mathbf{x}}^k)(\mathbf{x}^i \mathbb{1}_{C_k}(\mathbf{x}^i) - \bar{\mathbf{x}}^k)^T$, where $|C_k|$ is the cardinality of the k th cluster C_k . Then we could construct the ambiguity set (3) to employ the proposed (DRC-QSSVM) model for a robust classifier described in Fig. 2(d).

The proposed process illustrated by Fig. 2 leads to Algorithm 1, where Step 1 partitions the given binary data shown in 2(a) into two classes first. The clustering shown in Fig. 2(b) is accomplished by Steps 2–8. Step 9 computes the mean and covariance shown in Fig. 2(c). Steps 10–12 employ the proposed model to obtain a robust quadratic

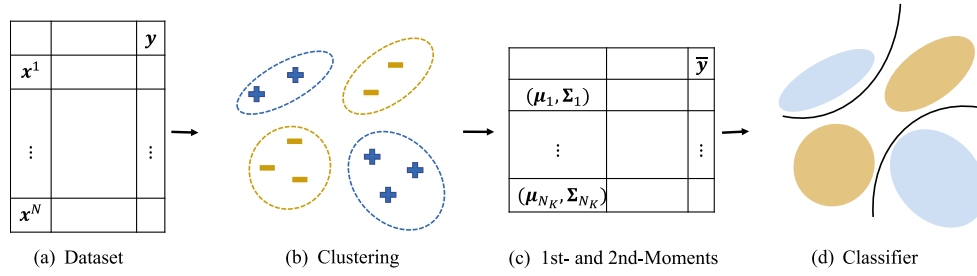


Fig. 2. Data-driven approach for data sets without moment information.

classifier (Fig. 2(d)) determined by (M^*, w^*, b^*) by solving the SOCP reformulation (18).

Algorithm 1 Data-driven distributionally robust classification algorithm.

Input: Data set $C = \{(x^i, y^i) \in \mathbb{R}^n \times \{-1, 1\}, i = 1, \dots, N\}$, $\epsilon \in (0, 1)$.

Output: A quadratic surface defined by (M^*, w^*, b^*) .

- 1: Extract two classes from the original data set C :
 $C_+ = \{x^i \in \mathbb{R}^n \mid y^i = 1, i = 1, \dots, N\}$, $C_- = \{x^i \in \mathbb{R}^n \mid y^i = -1, i = 1, \dots, N\}$.
- 2: Compute the cardinality: $N_+ = |C_+|$, $N_- = |C_-|$.
- 3: Let N_+^U be the upper bound of possible cluster numbers for C_+ , and N_-^U for C_- .
Set $N_+^U = \min\{25, 10\%N_+\}$, and $N_-^U = \min\{25, 10\%N_-\}$.
**The values 25 and 10% are user-defined based on the size of the data set.*
- 4: **for** $I \in \{+, -\}$ **do**
- 5: **for** $K = 1, 2, \dots, N_I^U$ **do**
- 6: Apply the K-means++ algorithm to divide the set C_I into K clusters.
- 7: **end for**
- 8: Use the Elbow method to find the optimal value of K denoted as K_I^* .
- 9: **end for**
- 10: Set $N_K = K_+^* + K_-^*$.
- 11: Calculate (μ_k, Σ_k) for C_+^k and number them by $k = 1, \dots, K_+^*$. Calculate (μ_k, Σ_k) for C_-^k and number them by $k = K_+^* + 1, \dots, N_K$. Set $\bar{y} = [e_{K_+^*}; -e_{K_-^*}] \in \mathbb{R}^{N_K}$ as the new label vector, where $e_{K_+^*}$ and $e_{K_-^*}$ are K_+^* -dimensional and K_-^* -dimensional vectors of all ones, respectively.
- 12: Use (μ_k, Σ_k) to compute H^k by (16), r^k by (17), $k = 1, \dots, N_K$. Set $W = \sum_{k=1}^{N_K} (H^k)^T H^k$.
- 13: With W , H^k , r^k , $k = 1, \dots, N_K$, and $\bar{y}$, solve the SOCP problem (18) of (DRC-QSSVM) to find an optimal solution (z^*, b^*) .
- 14: Compute (M^*, w^*) from z^* using the vectorization technique discussed in Section 2.4.
- 15: **Return** (M^*, w^*, b^*) .

Algorithm 1 implies that the robust quadratic classifier is learned by separating N_K ellipsoids instead of separating N data points with $N_K \ll N$ in general. This observation indicates the potential benefit of our proposed in dealing with massive data that shall be explored by computational experiments in Section 4.2.2. In addition, we notice that the classification objects become ellipsoids that cover data points with similarity, which might help avoid outliers and balance the magnitude of the two classes. It implies the potential of Algorithm 1 for treating imbalanced data. This works especially well for applications on rare case detection. SVMs with non-robust counterparts may perform poorly on the minority class since they may focus on the majority class while maximizing the overall accuracy (Thabtah, Hammoud, Kamalov, & Gonsalves, 2020). To illustrate this idea, for a highly imbalanced data set with a ratio of 5/21, Fig. 3(a) shows that the quadratic classifier obtained by (QSSVM) sacrifices 4 minority points in brown by treating them as outliers. However, instead of classifying the 26 imbalanced

points, the proposed approach finds a quadratic classifier by separating 4 ellipsoids with a more balanced ratio of 2/2 (See Fig. 3(b)). Further experiments on classifying imbalanced data will be conducted in Section 4.2.3.

Remark 3.1. In practice, decision-makers can rarely be completely confident in the sample mean and covariance matrix for estimated moments. The challenge here is that the sample means and covariance matrices themselves are uncertain. Hence their uncertainty is factored into chance constraints by considering a confidence region for the sample mean and covariance matrix. Appendix B.1 further extends the work.

Remark 3.2. As shown in Fig. 2 and Algorithm 1, the data-driven approach first takes the clustering process and then extracts the moment information. For real-world applications, historical data might not be sufficient to reflect the whole data structure, and we may need to collect new data over time to form a dynamic approach. When new instances are added to the data set, we need to update the clustering and classification processes as well.

4. Computational experiments

This section studies the proposed (DRC-QSSVM) model by computational experiments. For uncertain data with given first- and second-order moments, in Section 4.1, we validate the effectiveness of the proposed model and analyze its performance in terms of the parameter ϵ . Then we compare the proposed model with the DRC linear soft SVM (DRC-LSSVM) model which is the only maximum-margin SVM model using the first- and second-order moments information in the literature. For data without moment information, in Section 4.2, we compare the data-driven approach proposed in Section 3 with some state-of-the-art SVM models. In particular, we explore the potential benefits of using the proposed model for problems with massive and/or imbalanced data. In this section, all computational experiments were conducted using MATLAB (R2021a) software on a desktop equipped with Intel(R) Core(TM) i3-9100 CPU @ 3.60 GHz CPUs and 32 GB RAM. The commercial solver SDPT3 (Toh, Todd, & Tütüncü, 1999) is employed to solve SDP and SOCP problems.

4.1. Data sets with first- and second-order moments

As discussed in Section 2, for uncertain data with first- and second-order moments, the proposed (DRC-QSSVM) model has computable SOCP and SDP reformulations. In Section 4.1.1, we test these two formulations on synthetic data sets regarding classification accuracy and computational efficiency. For the proposed model, the parameter C controls the trade-off between maximizing the margin and minimizing the misclassification loss, as commonly adopted in most SVM models. While the parameter ϵ determines the upper bound of misclassification probability that affects the quality of robust classification. Here we skip the detailed analysis on the parameter C , but focus on the parameter

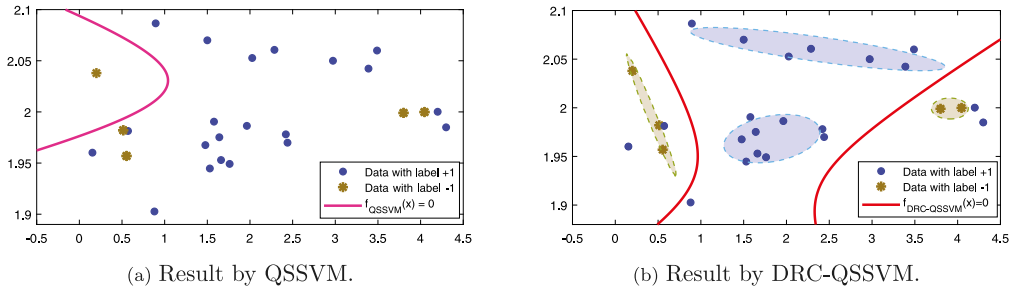
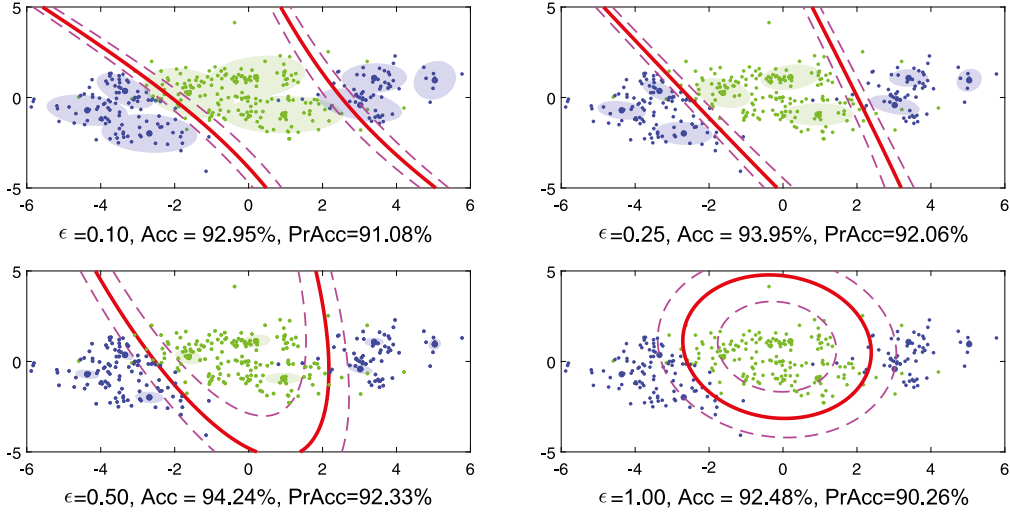


Fig. 3. Comparison on an imbalanced data set.

Fig. 4. Results of (DRC-QSSVM) on Syn-Hype-2d-9 data. Shaded areas depict the given ellipsoids $E(\mu_i, \Sigma_i, \frac{1}{\epsilon})$. The learned quadratic classifier, $\frac{1}{2} \mathbf{x}^T \mathbf{M} \mathbf{x} + \mathbf{x}^T \mathbf{w} + b = 0$, is represented by the red solid line, and the pink dashed lines represent $\frac{1}{2} \mathbf{x}^T \mathbf{M} \mathbf{x} + \mathbf{x}^T \mathbf{w} + b = \pm 1$. Solid points represent random testing points.

ϵ . For all computational experiments, we take the grid method to set $C \in \{2^{-1}, 2^1, \dots, 2^{14}\}$ and $\epsilon \in \{0.10, 0.25, 0.50, 1.00\}$.

We first introduce some error measures to be used. For a quadratic surface obtained by solving (DRC-QSSVM) with an output $(\mathbf{M}^*, \mathbf{w}^*, b^*) \in \mathbb{S}^n \times \mathbb{R}^n \times \mathbb{R}$, the predicted label $\hat{y} = \text{sign}(\frac{1}{2}(\mathbf{x})^T \mathbf{M}^* \mathbf{x} + (\mathbf{w}^*)^T \mathbf{x} + b^*)$ can be determined for a test data point $(\mathbf{x}, y)$. A commonly used measure is the accuracy score (Acc) computed by $\sum_{i=1}^N \mathbb{1}(\hat{y}^i = y^i) / N \times 100\%$ where N is the total number of tested points. However, for the uncertain classification with data points from a distribution, we further consider a probabilistic accuracy score (PrAcc). Note that Ben-Tal et al. (2011) and Wang et al. (2018) adopted an “optimal error” to quantify the probabilistic error. For DRC-LSSVM, we have $\text{PrAcc} = 1 - \text{“optimal error”}$. Here we further extend the measure for quadratic classifiers. At each uncertain data point $\tilde{\mathbf{x}}^i$ associated with the ambiguity set $\mathcal{D}_i$ in (3), the robust chance constraint, $\sup_{F_i \in \mathcal{D}_i} \mathbb{P}_{F_i} \left\{ y^i \left(\frac{1}{2}(\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 0 \right\} \leq \epsilon$, ensures an upper bound ϵ of the misclassification probability for the quadratic classifier. Using Theorem 2.4 and the SOCP reformulation (19), the true probability of misclassification at $\tilde{\mathbf{x}}^i$ should be no more than ϵ , if $\|\Sigma_i^{-\frac{1}{2}} T^i \mathbf{z}\|_2 \leq \sqrt{\epsilon / (1 - \epsilon)} y^i (\mathbf{z}^T T^i \mathbf{r}^i + b)$. And this could imply that the least value of ϵ for $\tilde{\mathbf{x}}^i$ is $\epsilon_i^* = \frac{\mathbf{z}^T (T^i)^T \Sigma_i T^i \mathbf{z}}{(\mathbf{z}^T \mathbf{r}^i + b)^2 + \mathbf{z}^T (T^i)^T \Sigma_i T^i \mathbf{z}}$. Consequently, we define

$$\text{PrAcc} = (1 - \sum_{i=1}^N \text{err}_i / N) \times 100\%, \quad (20)$$

$$\text{where } \text{err}_i = \begin{cases} 1, & \text{if } \hat{y}^i \neq y^i \\ \epsilon_i^*, & \text{if } \hat{y}^i = y^i \end{cases}, i = 1, \dots, N.$$

4.1.1. Validation and analysis of the proposed model

Given that we have uncertain data set with known $\{(\mu_i, \Sigma_i) \in \mathbb{R}^n \times \mathbb{S}^n, i = 1, \dots, N\}$, to validate the proposed model and its effectiveness, we first generate some synthetic data in different quadratic patterns including hyperbolic, elliptic, and parabolic structures (See Fig. 1 for illustration). The corresponding synthetic data sets are named in the format of “Syn-Pattern-nd-N”, for example, the data set “Syn-Hype-4d-50” denotes a set of $\{(\mu_i, \Sigma_i) \in \mathbb{R}^4 \times \mathbb{S}^4, i = 1, \dots, 50\}$ where μ_i are generated along with a hyperbolic surface, and Σ_i are random positive definite matrices with eigenvalues in $[0, 1]$. For each i , we generate 50 random points following the normal distribution with mean μ_i and covariance Σ_i as the testing data points. Note that the proposed model can handle distribution-free data and we choose the normal distribution for simplicity in this section. We generate 8 data sets by selecting $n \in \{2, 4, 8, 16\}$ and $N \in \{50, 100\}$. We also record the testing accuracy by the average Acc and PrAcc. The average training CPU time is recorded for both results solved by the SDP model (12) and SOCP model (19).

For a simple illustration, first, we show a 2-dimensional example tested on the “Syn-Hype-2d-9” data set. Fig. 4 shows that the classifiers learned based on the first- and second-order moments depend on the value of ϵ . For example, the proposed model provides hyperbolic curves when $\epsilon = 0.10, 0.25$, a parabolic curve when $\epsilon = 0.50$, and an ellipsoidal curve when $\epsilon = 1.00$. Figs. 4 and 5(a) show that ϵ indeed affects the classification accuracy. The Acc and PrAcc shown in Fig. 5(a) depict how ϵ affects the performance.

Synthetic data sets with bigger sample sizes in higher dimensions have been tested to investigate the proposed model further. Table 1 shows one group of results on “hyperbolic” synthetic data sets. More results on the “elliptic” and “parabolic” data sets are shown in Appendix B.2. For all testing problems, we ensure that the SOCP model

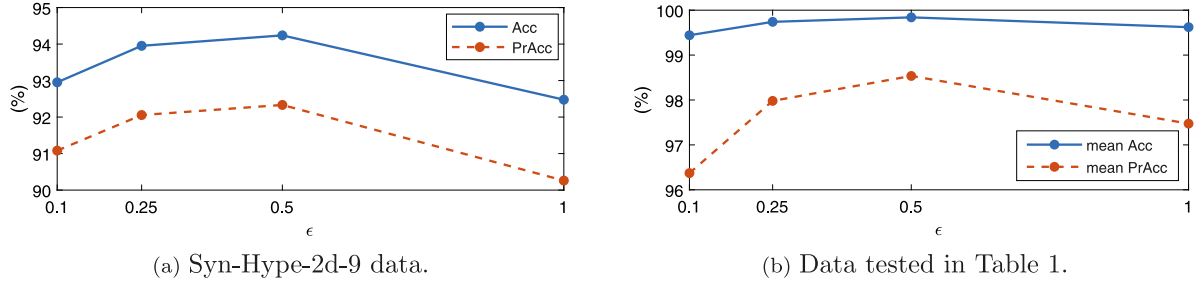


Fig. 5. Performance of (DRC-QSSVM) on “hyperbolic” synthetic data sets in terms of ϵ .

Table 1
Testing results on “hyperbolic” synthetic data sets by DRC-QSSVM.

Data	Syn-Hype-4d-50				Syn-Hype-8d-50				Syn-Hype-16d-50			
ϵ	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
Acc(%)	99.96	100.00	99.96	99.24	97.92	98.64	99.28	99.04	98.88	99.84	99.92	99.68
PrAcc(%)	98.54	99.18	99.39	98.69	92.57	95.49	96.46	93.61	88.67	94.76	96.94	94.11
CPU time (s)	SOCP				SOCP				SOCP			
	3.01	2.99	2.99	2.22	3.14	3.13	3.08	2.38	5.72	5.71	5.98	3.85
CPU time (s)	SDP				SDP				SDP			
	4.55	4.50	4.45	3.35	6.98	6.99	6.89	5.10	34.96	30.90	29.54	21.77

Data	Syn-Hype-4d-100				Syn-Hype-8d-100				Syn-Hype-16d-100			
ϵ	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
Acc(%)	100.00	99.96	99.88	99.88	99.89	100.00	100.00	100.00	100.00	100.00	100.00	99.89
PrAcc(%)	99.35	99.36	99.27	99.26	99.59	99.70	99.80	99.83	99.51	99.39	99.35	99.34
CPU time (s)	SOCP				SOCP				SOCP			
	5.14	5.25	5.22	2.97	5.50	5.56	5.62	3.89	12.43	12.41	12.36	7.70
CPU time (s)	SDP				SDP				SDP			
	8.05	8.03	7.95	4.74	19.40	18.20	17.81	11.74	187.10	176.16	142.67	136.87

achieves the same results as the SDP model, but in a much more efficient way. From Table 1, we see that (i) the proposed (DRC-QSSVM) model performs well in terms of Acc and PrAcc measures; (ii) it is not surprising that PrAcc is always smaller than the corresponding Acc since the former considers the potential misclassification probability when the predicted label is correct; (iii) For fixed N and n , both Acc and PrAcc change depending on ϵ , but in the same trend, as illustrated in Fig. 5(b).

4.1.2. Comparison with the DRC-LSSVM model

In the literature, DRC-LSSVM (Wang et al., 2018) is the only maximum-margin SVM model using the means and covariance matrices for distributionally robust classification. Hence we compare the proposed (DRC-QSSVM) with DRC-LSSVM using the well-known data sets Wisconsin breast cancer (WIBC) and the Ionosphere from the UCI dataset. For fair comparisons, we adopt the same settings for data preprocessing as in Wang et al. (2018). WIBC data contains 683 samples with 9 features, i.e. $N = 683$, $n = 9$, and extracted Ionosphere data has $N = 351$, $n = 15$ (extracted from $n = 34$ as Wang et al. (2018)). Moreover, for computational efficiency, SOCP reformulations are used for both (DRC-QSSVM) and DRC-LSSVM. Similarly to Wang et al. (2018), μ_i is set to be the value of each training point, and Σ^i is calculated based on the covariance matrix of all training points in the same class. Table 2 shows the results where (i) 20% of the data are used for training and the remaining 80% for testing; (ii) 80% for training and 20% for testing.

The proposed DRC-QSSVM model generates a quadratic surface, which increases the flexibility when handling nonlinear data. However, it increases the model complexity as well compared with the linear DRC-LSSVM model. Table 2 clearly shows that the performance of (DRC-QSSVM) dominates that of DRC-LSSVM in all cases, taking a reasonably longer running time. The superiority of the proposed (DRC-QSSVM) model becomes particularly evident when applied to the

Extracted Ionosphere data, which is more nonlinearly complex than the WIBC data.

4.2. Data sets without moment information

Section 3 provides a data-driven approach to apply (DRC-QSSVM) for robustly classifying exact data points $\{(\mathbf{x}^i, y^i) \in \mathbb{R}^n \times \{-1, +1\}, i = 1, \dots, N\}$. We conduct computational experiments to compare the data-driven approach with well-known state-of-the-art SVMs using some commonly used public benchmark data sets. Table 3 lists the tested models, including their abbreviations, solvers, and parameters. Kernelized SVMs are solved by utilizing LIBSVM (Chang & Lin, 2011), and other SVMs are solved by SDPT3. Note that (DRC-QSSVM) is realized by the data-driven Algorithm 1 in which the SOCP problem (18) is solved by SDPT3.

For all tests, the 10-fold cross-validation and grid methods are adopted to select the best parameters of C , ϵ , and σ from the ranges of $C \in \{2^{-1}, 2^1, \dots, 2^{14}\}$, $\epsilon \in \{0.1, 0.2, \dots, 1\}$, and $\sigma \in \{2^{-5}, 2^{-4}, \dots, 2^5\}$, respectively. All test results are based on the best-selected parameters. Some public benchmark data sets from UCI databases (See Table 4) are chosen. Section 4.2.1 presents the results of “balanced” data sets, including the commonly used Scale, Pima Indians Diabetes (Pima), WIBC, and Ionosphere. Section 4.2.2 reports the performance of “massive” data sets including Skin and Cod-RNA with large sample sizes. Section 4.2.3 explores the results of “imbalanced” data sets, including Car Evaluation (Careval) and Heart Disease (Heart) with skewed class proportions. All the classical SVMs are tested on the original data, and the two robust models including DRC-LSSVM and the proposed DRC-QSSVM utilize the moment information of the data retrieved by the data-driven approach described in Algorithm 1.

4.2.1. Benchmark tests on balanced data

Four popular balanced benchmark data sets: Scale, Pima, WIBC, and Ionosphere are tested. Since they are exact data sets, we report the Acc

Table 2
Testing results on WIBC and extracted Ionosphere data sets.

Data (20% training)		WIBC				Extracted Ionosphere			
ϵ		0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
DRC-LSSVM	Acc(%)	96.60	96.38	96.30	96.12	84.26	84.54	84.40	84.26
	PrAcc(%)	93.98	93.70	93.61	93.00	82.61	82.91	82.09	81.92
	CPU time (s)	5.77	5.48	5.05	5.01	3.36	3.37	3.39	3.40
DRC-QSSVM	Acc(%)	96.63	96.63	96.52	96.56	92.81	91.95	92.09	92.09
	PrAcc(%)	94.05	94.39	94.69	94.82	91.64	91.21	91.38	91.30
	CPU time (s)	6.66	6.13	5.71	5.77	5.56	5.53	5.56	5.61
Data (80% training)		WIBC				Extracted Ionosphere			
ϵ		0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
DRC-LSSVM	Acc(%)	96.63	96.63	96.49	96.49	87.46	88.31	88.02	87.74
	PrAcc(%)	95.45	95.57	95.35	95.15	83.36	83.70	84.19	83.67
	CPU time (s)	16.97	17.56	16.37	18.14	10.51	10.54	12.15	11.37
DRC-QSSVM	Acc(%)	97.22	97.37	97.22	96.93	96.58	96.02	95.73	95.16
	PrAcc(%)	95.86	95.91	95.76	95.61	92.87	93.48	92.94	92.18
	CPU time (s)	32.25	30.82	34.77	30.51	33.05	32.77	30.47	31.01

Table 3
Models and solvers of the tested models.

Model	Abbreviation	Solver/Package	Parameter
Linear soft SVM	LSSVM	LibSVM	C
Quadratic soft surface SVM	QSSVM	SDPT3	C
SVM with quadratic kernel	KQSSVM	LibSVM	(C, σ)
SVM with Gaussian kernel	KGSSVM	LibSVM	(C, σ)
DRC linear SVM	DRC-LSSVM	SDPT3	(C, ϵ)
The proposed model	DRC-QSSVM	SDPT3	(C, ϵ)

measure only. The mean and standard deviation of Acc are shown in Table 5. Same as in Section 4.1.2, we select 20% and 80% of data sets for training, respectively. The average training CPU time of each model is also reported.

The following observations can be made:

- The proposed (DRC-QSSVM) model produces much more accurate classifications than other tested SVM models on all tested balanced data sets. It shows the special value of the data-driven based robust (DRC-QSSVM) for commonly used data sets without prescribed uncertainty.
- For most data sets, the classification accuracy obtained by (DRC-QSSVM) changes very little in terms of different training rates of 20% vs 80%. It indicates the potential advantage of the proposed model when we have limited data points for training.
- The CPU time consumed by the proposed (DRC-QSSVM) is acceptable overall considering its classification accuracy. Also note that (DRC-QSSVM), (QSSVM), and DRC-LSSVM are solved using the solver SDPT 3.0, while others are solved using an integrated software LIBSVM.

4.2.2. Benchmark tests on massive data

Two massive benchmark data sets, Skin and Cod-RNA, are used for testing. We also use 20% and 80% of data points for training, respectively. The mean, standard deviation of accuracy scores, and the average training CPU time are reported in Table 6.

The following observations can be made:

- The proposed (DRC-QSSVM) model significantly outperforms others in accuracy for the Skin data. For the Cod-RNA data, (DRC-QSSVM) also outperforms other models considering both the accuracy and CPU time.

- (DRC-QSSVM) can achieve high accuracy using only 20% data points for training. This supports the stability of the proposed model and the promising potential of practical use for massive data classification.

4.2.3. Benchmark tests on imbalanced data

Imbalanced data sets, where one class greatly outnumbers the other, are a common issue in many real-life applications. Classifying imbalanced data presents a challenge for standard classification algorithms. In this subsection, two imbalanced data sets, Careval and Heart, are tested. For the classification of imbalanced data, a good SVM model should (i) enhance recognition success specifically for the minority class, and/or (ii) balance recognition capabilities between both classes (Sun, Wong, & Kamel, 2009). Additional error measurements are often used to evaluate such performance. The Area Under Curve (AUC) score could help evaluate the first performance, while the G-mean score could help the second one (Details refer to Sun et al. (2009)). In this subsection, we elect 80% as the training rate due to the potential inadequacy of the minority class sample size to facilitate training at the 20% rate. Table 7 displays the average CPU time, and the mean and standard deviation of the Acc, AUC, and G-mean scores.

The following observations can be made:

- The proposed (DRC-QSSVM) model outperforms other models in all three measures. The dominance is particularly significant in AUC and G-mean scores.
- Note that the imbalance ratio of the Heart data is higher than that of the Careval data. Most models have AUC and G-mean around 50%, which means these classifiers cannot distinguish two classes clearly. However, the proposed model still shows good performance. This means that the proposed model may have a better capability of handling highly imbalanced data.

In summary, applying the DRO approach, the proposed model extends the QSSVM framework, enabling a kernel-free nonlinear SVM with a quadratic classifier to handle uncertainties in data. It outperforms the QSSVM as well as other state-of-the-art SVMs in general classification tasks without uncertainty. While the QSSVM is tested on original data, solving a certain problem, the proposed model leverages hidden moment information to address uncertain problems. This highlights the significant finding that transforming a certain problem into an uncertain one and then solving it may lead to surprisingly better outcomes.

Table 4
Summary of the benchmark data sets.

Data set		Balanced				Massive		Imbalanced	
		Scale	Pima	WIBC	Ionosphere	Skin	Cod-RNA	Careval	Heart
Dimension	n	4	8	9	34	3	8	6	9
Sample size ^a N	N_+	288	268	239	225	50,859	19,845	1,210	3,101
	N_-	288	500	444	126	194,198	39,690	384	557

^a $N_{\pm}$ = sample size of points in the class labeled ‘ ± 1 ’ and $N = N_+ + N_-$.

Table 5
Testing results on the balanced benchmark data sets.

Model/ Training rate		Data set			Scale			Pima			WIBC			Ionosphere		
		Acc(%)		CPU(s)	Acc(%)		CPU(s)	Acc(%)		CPU(s)	Acc(%)		CPU(s)	Acc(%)		CPU(s)
		mean	std		mean	std		mean	std		mean	std		mean	std	
LSSVM	20%	70.24	6.99	0.27	74.97	3.69	0.01	92.65	3.16	0.01	82.48	1.44	0.01			
	80%	73.02	5.63	1.50	76.14	2.51	0.08	96.03	1.26	3.90	88.76	3.19	0.28			
QSSVM	20%	97.35	0.89	0.71	74.31	3.87	0.78	94.12	2.38	2.44	87.81	2.64	3.18			
	80%	97.53	0.91	8.50	76.67	2.43	9.87	95.66	0.95	11.58	93.52	2.97	10.29			
KQSSVM	20%	97.65	0.78	0.16	76.99	2.68	0.23	95.44	1.66	0.15	87.43	3.12	0.05			
	80%	97.65	0.88	1.11	78.24	1.97	1.01	96.10	0.98	0.54	91.62	4.06	0.13			
KGSSVM	20%	97.82	1.00	0.01	76.34	2.09	0.01	92.18	1.42	0.01	88.57	1.56	0.01			
	80%	98.24	0.83	0.02	77.58	2.09	0.11	95.96	0.93	0.02	93.74	1.17	0.02			
DRC-LSSVM	20%	68.65	5.03	0.40	76.99	2.80	0.72	94.49	2.00	2.10	88.19	3.73	0.50			
	80%	74.18	4.59	0.48	76.93	3.28	0.77	94.71	1.54	2.03	89.05	5.33	0.74			
DRC-QSSVM	20%	98.18	1.01	0.44	80.77	4.15	0.77	96.55	2.03	2.20	93.52	1.17	3.04			
	80%	98.30	0.79	0.51	81.36	3.42	0.82	96.66	1.29	2.14	94.00	2.97	3.57			

Table 6
Testing results on the massive benchmark data sets.

Training rate	Skin						Cod-RNA					
	20%			80%			20%			80%		
	Acc (%)		CPU(s)	Acc (%)		CPU(s)	Acc (%)		CPU(s)	Acc (%)		CPU(s)
	mean	std		mean	std		mean	std		mean	std	
LSSVM	75.01	8.31	1.52	80.01	5.29	6.08	78.51	10.32	10.88	84.67	5.21	44.90
QSSVM	85.23	1.21	25.01	91.24	3.20	1656.43	91.17	0.78	138.83	92.04	0.46	2942.85
KQSSVM	80.53	21.95	0.97	92.29	5.42	4.12	91.61	0.98	4.13	92.55	0.59	30.52
KGSSVM	79.56	0.11	0.01	81.07	0.11	0.08	88.59	0.49	0.07	90.86	0.24	1.66
DRC-LSSVM	87.73	7.84	0.62	88.82	5.49	0.81	80.45	9.79	0.94	89.47	1.64	0.98
DRC-QSSVM	96.51	0.51	0.66	97.85	0.75	1.36	92.54	0.64	0.98	92.21	0.60	1.26

Table 7
Testing results on the imbalanced benchmark data sets.

Model	Careval							Heart						
	Acc(%)		AUC(%)		G-mean(%)		CPU(s)	Acc(%)		AUC(%)		G-mean(%)		CPU(s)
	mean	std	mean	std	mean	std		mean	std	mean	std	mean	std	
LSSVM	79.88	10.68	84.64	10.63	74.51	8.06	0.70	53.17	1.11	51.30	0.40	50.68	0.28	0.02
QSSVM	96.22	0.96	98.59	0.44	92.37	1.93	2.98	76.88	3.06	56.07	2.75	48.33	2.32	35.88
KQSSVM	94.72	0.55	99.11	0.28	93.45	0.91	0.70	67.09	0.60	50.11	1.01	55.10	1.53	3.10
KGSSVM	96.00	0.46	98.74	0.20	94.34	1.27	0.01	84.79	^a 0.00	57.85	^a 0.00	^a 0.00	^a 0.00	0.04
DRC-LSSVM	95.66	0.79	96.51	0.59	83.02	1.18	0.49	78.61	6.33	54.65	1.34	53.68	0.85	0.58
DRC-QSSVM	99.65	0.70	99.82	0.78	95.97	1.32	0.50	85.77	0.02	71.32	0.90	66.24	1.47	1.21

^a A G-mean score of value 0 and std of 0 indicate the classifier simply assigns all instances to the majority class.

5. Conclusion

In this paper, we have established a novel distributionally robust chance-constrained kernel-free quadratic surface support vector machine model that can robustly conduct nonlinear classification for data sets involving stochastic uncertainties, in which only the first- and second-order moments are known a priori. SDP and SOCP reformulations of the proposed model have been derived for computational

efficiency. Additionally, an explicit geometric interpretation of the conceptual distributionally robust chance constraints has been presented to show how the proposed model handle uncertain data. Our computational experiments show that the proposed model clearly outperforms the DRC-LSSVM model, the only maximum-margin SVM model explicitly using moment information for classifying uncertain data, on synthetic and public benchmark data sets.

For commonly used data sets without uncertainty involved, we design a cluster-based data-driven approach that retrieves the hidden

moment information first and enables the proposed model to leverage the moments for robust classification. This approach aids in further exploring the applicability of the proposed model. Extensive computational experiments using public benchmark data sets exhibit the surprisingly dominant performance of the proposed model over other state-of-the-art SVM models, especially for massive and/or imbalanced data sets.

Our investigation of the proposed model leads to some potential research works. First, in real-world applications, historical data may not be sufficient to capture the whole structure of the data set (Hsu, Xu, Lin, & Bell, 2022; Mi, Quan, Shi, & Wang, 2022). Collecting new data over time is necessary to adapt to changes in the data and the evolving moment information. We are interested in developing a dynamic approach to update the clustering and classification processes to extend the proposed model further. Besides, we are interested in how the proposed model performs in healthcare applications (Jiang, Han, Yu, & Ding, 2023; Naumzik, Feuerriegel, & Nielsen, 2023) with large-scale uncertain data.

Acknowledgments

This work has been sponsored by the National Science Foundation CNS-2229245, the National Natural Science Foundation of China Grants #72201052 and #72261008, Hainan Provincial Natural Science Foundation of China Grant 724RC488 and the Foundation of Yunnan Key Laboratory of Service Computing Grant #YNSC23115.

Appendix A. Proofs

A.1. Proof of Lemma 2.2

Proof. For any i , the constraints in $\mathcal{V}_i^{SDP}$ require finding β_i and $\mathbf{R}_i \geq 0$ to satisfy $\frac{1}{\epsilon} \mathbf{G}_i \bullet \mathbf{R}_i - \beta_i \leq 0$ and the matrix inequality. Notice that finding β_i and $\mathbf{R}_i \geq 0$ with $\mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0$ is not difficult. However, it is hard to guarantee such β_i and $\mathbf{R}_i$ satisfies $\frac{1}{\epsilon} \mathbf{G}_i \bullet \mathbf{R}_i - \beta_i \leq 0$. Requiring $\mathcal{V}_i^{SDP} \neq \emptyset$ is equivalent to requiring that the optimal value of the following optimization problem is less than or equal to 0:

$$\inf \frac{1}{\epsilon} \mathbf{G}_i \bullet \mathbf{R}_i - \beta_i \quad (\text{A.1a})$$

$$\text{s.t. } \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0, \quad (\text{A.1b})$$

$$\mathbf{R}_i \geq 0. \quad (\text{A.1c})$$

Since one can easily find a proper $\beta_i \in \mathbb{R}$ and a matrix $\mathbf{R}_i > 0$ with eigenvalues large enough to satisfy $\mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \beta_i \end{bmatrix} > 0$, the Slater's condition of problem (A.1) is satisfied and the strong duality holds for (A.1) and its dual problem. Let $\bar{\mathbf{D}}_i = \begin{bmatrix} \mathbf{D}_i & \mathbf{d}_i \\ \mathbf{d}_i^T & d_{0i} \end{bmatrix} \in \mathbb{S}_+^{n+1}$ and $\mathbf{C}_i \in \mathbb{S}_+^{n+1}$ be the dual variables corresponding to (A.1b) and (A.1c), respectively. Then the Lagrangian becomes

$$\begin{aligned} & \sup_{\bar{\mathbf{D}}_i \geq 0, \mathbf{C}_i \geq 0} \inf_{\mathbf{R}_i, \beta_i} \mathcal{L}(\mathbf{R}_i, \beta_i, \bar{\mathbf{D}}_i, \mathbf{C}_i) \\ &= \sup_{\bar{\mathbf{D}}_i \geq 0, \mathbf{C}_i \geq 0} \inf_{\mathbf{R}_i, \beta_i} \left\{ \frac{1}{\epsilon} \mathbf{G}_i \bullet \mathbf{R}_i - \beta_i - \bar{\mathbf{D}}_i \bullet \left(\mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \beta_i \end{bmatrix} \right) \right. \\ & \quad \left. - \mathbf{C}_i \bullet \mathbf{R}_i \right\} \\ &= \sup_{\bar{\mathbf{D}}_i \geq 0, \mathbf{C}_i \geq 0} \inf_{\mathbf{R}_i, \beta_i} \left\{ \mathbf{R}_i \bullet \left(\frac{1}{\epsilon} \mathbf{G}_i - \mathbf{C}_i - \bar{\mathbf{D}}_i \right) + (d_{0i} - 1) \beta_i - d_{0i} (y^j b + \xi_i - 1) \right. \\ & \quad \left. - \frac{1}{2} y^j \mathbf{M} \bullet \mathbf{D}_i - y^j \mathbf{w}^T \mathbf{d}_i \right\} \\ &= \sup_{\bar{\mathbf{D}}_i \geq 0} \begin{cases} -\frac{1}{2} y^j \mathbf{M} \bullet \mathbf{D}_i - y^j \mathbf{w}^T \mathbf{d}_i - (y^j b + \xi_i - 1), & d_{0i} - 1 = 0, \frac{1}{\epsilon} \mathbf{G}_i - \bar{\mathbf{D}}_i \geq 0, \\ -\infty, & \text{otherwise.} \end{cases} \end{aligned}$$

Consequently, we have dual problem of (A.1):

$$\begin{aligned} & \sup \quad -\frac{1}{2} y^j \mathbf{M} \bullet \mathbf{D}_i - y^j \mathbf{w}^T \mathbf{d}_i - (y^j b + \xi_i - 1) \\ & \text{s.t.} \quad \begin{bmatrix} \Sigma_i + \mu_i \mu_i^T - \epsilon \mathbf{D}_i & \mu_i - \epsilon \mathbf{d}_i \\ \mu_i^T - \epsilon \mathbf{d}_i^T & 1 - \epsilon \end{bmatrix} \geq 0, \\ & \quad \begin{bmatrix} \mathbf{D}_i & \mathbf{d}_i \\ \mathbf{d}_i^T & 1 \end{bmatrix} \geq 0, \\ & \quad \mathbf{D}_i \in \mathbb{S}^n, \mathbf{d}_i \in \mathbb{R}^n. \end{aligned} \quad (\text{A.2})$$

The Schur Complement Lemma implies that $\begin{bmatrix} \mathbf{D}_i & \mathbf{d}_i \\ \mathbf{d}_i^T & 1 \end{bmatrix} \geq 0 \Leftrightarrow \mathbf{D}_i - \mathbf{d}_i \mathbf{d}_i^T \geq 0$. Let $\mathbf{D}_i^0 = \mathbf{D}_i - \mathbf{d}_i \mathbf{d}_i^T$. Substituting it into the first constraint, we have

$$\begin{aligned} & \begin{bmatrix} \Sigma_i + \mu_i \mu_i^T - \epsilon \mathbf{D}_i & \mu_i - \epsilon \mathbf{d}_i \\ \mu_i^T - \epsilon \mathbf{d}_i^T & 1 - \epsilon \end{bmatrix} \geq 0 \\ & \Leftrightarrow \Sigma_i + \mu_i \mu_i^T - \epsilon \mathbf{D}_i - \frac{1}{1-\epsilon} (\mu_i - \epsilon \mathbf{d}_i)(\mu_i - \epsilon \mathbf{d}_i)^T \geq 0 \\ & \Leftrightarrow \begin{bmatrix} \Sigma_i - \epsilon \mathbf{D}_i^0 & \mu_i - \mathbf{d}_i \\ \mu_i^T - \mathbf{d}_i^T & \frac{1-\epsilon}{\epsilon} \end{bmatrix} \geq 0. \end{aligned}$$

Therefore, we can rewrite (A.2) as

$$\begin{aligned} & \sup \quad -y^j \left(\frac{1}{2} \mathbf{d}_i^T \mathbf{M} \mathbf{d}_i + \mathbf{w}^T \mathbf{d}_i + b \right) - \xi_i + 1 - \frac{1}{2} y^j \mathbf{M} \bullet \mathbf{D}_i^0 \\ & \text{s.t.} \quad \begin{bmatrix} \Sigma_i - \epsilon \mathbf{D}_i^0 & \mu_i - \mathbf{d}_i \\ \mu_i^T - \mathbf{d}_i^T & \frac{1-\epsilon}{\epsilon} \end{bmatrix} \geq 0, \\ & \quad \mathbf{D}_i^0 \geq 0, \\ & \quad \mathbf{D}_i^0 \in \mathbb{S}^n, \mathbf{d}_i \in \mathbb{R}^n. \end{aligned} \quad (\text{A.3})$$

Remember that the strong duality holds for (A.1) and (A.3). By satisfying the requirement that the optimal value is no larger than 0, the constraints in $\mathcal{V}_i^{SDP'}$ are obtained and thus the claim follows. $\square$

A.2. Proof of Theorem 2.4

Proof. When deriving Lemma 2.2, we notice that the SDP constraints in (7) could be obtained by requiring the optimal value of the following convex SDP less than or equal to 0:

$$\begin{aligned} & \sup \quad -\frac{1}{2} y^j \mathbf{M} \bullet \mathbf{D}_i - y^j \mathbf{w}^T \mathbf{d}_i - (y^j b + \xi_i - 1) \\ & \text{s.t.} \quad \Sigma_i + \mu_i \mu_i^T - \epsilon \mathbf{D}_i - \frac{1}{1-\epsilon} (\mu_i - \epsilon \mathbf{d}_i)(\mu_i - \epsilon \mathbf{d}_i)^T \geq 0, \\ & \quad \epsilon \mathbf{D}_i - \epsilon \mathbf{d}_i \mathbf{d}_i^T \geq 0, \\ & \quad \mathbf{D}_i \in \mathbb{S}^n, \mathbf{d}_i \in \mathbb{R}^n. \end{aligned} \quad (\text{A.4})$$

Let $\mathbf{D}_i^* = \Sigma_i + \mu_i \mu_i^T > 0$ and $\mathbf{d}_i^* = \mu_i$, the constraints are satisfied strictly with positive definite matrices. Slater's condition is satisfied, and the strong duality holds. It is hard to give an explicit optimal solution and optimal objective value directly by solving (A.4). We consider its dual problem. Let $\mathbf{H}_i, \mathbf{G}_i \in \mathbb{S}_+^n$ be the dual variables of (A.4). The corresponding Lagrangian dual is

$$\begin{aligned} & \inf_{\mathbf{H}_i \geq 0, \mathbf{G}_i \geq 0} \sup_{\mathbf{D}_i, \mathbf{d}_i} \mathcal{L}(\mathbf{D}_i, \mathbf{d}_i, \mathbf{H}_i, \mathbf{G}_i) \\ &= \inf_{\mathbf{H}_i \geq 0, \mathbf{G}_i \geq 0} \sup_{\mathbf{D}_i, \mathbf{d}_i} \left\{ -\frac{1}{2} y^j \mathbf{M} \bullet \mathbf{D}_i - y^j \mathbf{w}^T \mathbf{d}_i - (y^j b + \xi_i - 1) \right. \\ & \quad \left. + \mathbf{H}_i \bullet \left(\Sigma_i + \mu_i \mu_i^T - \epsilon \mathbf{D}_i - \frac{1}{1-\epsilon} (\mu_i - \epsilon \mathbf{d}_i)(\mu_i - \epsilon \mathbf{d}_i)^T \right) + \mathbf{G}_i \bullet (\epsilon \mathbf{D}_i - \epsilon \mathbf{d}_i \mathbf{d}_i^T) \right\} \\ &= \inf_{\mathbf{H}_i \geq 0, \mathbf{G}_i \geq 0} \sup_{\mathbf{D}_i, \mathbf{d}_i} \left\{ -(y^j b + \xi_i - 1) + \mathbf{H}_i \bullet \Sigma_i - \frac{\epsilon}{1-\epsilon} \mu_i^T \mathbf{H}_i \mu_i \right. \\ & \quad \left. + \mathbf{D}_i \bullet \left(-\frac{1}{2} y^j \mathbf{M} - \epsilon \mathbf{H}_i + \epsilon \mathbf{G}_i \right) - \mathbf{d}_i^T \left(\frac{\epsilon}{1-\epsilon} \mathbf{H}_i + \epsilon \mathbf{G}_i \right) \mathbf{d}_i + (-y^j \mathbf{w} + \frac{2\epsilon}{1-\epsilon} \mathbf{H}_i \mu_i)^T \mathbf{d}_i \right\}. \end{aligned}$$

The dual problem is finite if and only if $-\frac{1}{2} y^j \mathbf{M} - \epsilon \mathbf{H}_i - \epsilon \mathbf{G}_i = 0$, which gives $\epsilon \mathbf{G}_i = \frac{1}{2} y^j \mathbf{M} + \epsilon \mathbf{H}_i \geq 0$. Substitute this into the above, we get

$$\begin{aligned} & \inf_{\mathbf{H}_i \geq 0, \mathbf{G}_i \geq 0} \sup_{\mathbf{D}_i, \mathbf{d}_i} \mathcal{L}(\mathbf{D}_i, \mathbf{d}_i, \mathbf{H}_i, \mathbf{G}_i) \\ &= \inf_{\mathbf{H}_i \geq 0} \sup_{\mathbf{d}_i} \begin{cases} -(y^j b + \xi_i - 1) + \mathbf{H}_i \bullet \Sigma_i - \frac{\epsilon}{1-\epsilon} \mu_i^T \mathbf{H}_i \mu_i + q(\mathbf{d}_i), & \text{if } \frac{1}{2} y^j \mathbf{M} + \epsilon \mathbf{H}_i \geq 0, \\ +\infty, & \text{otherwise,} \end{cases} \end{aligned}$$

where $q(d_i) = -d_i^T(\frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM)d_i + (-y^i w + \frac{2\epsilon}{1-\epsilon}H_i\mu_i)^T d_i$. Since the dual variable $H_i \geq 0$, we have $\frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM \geq \epsilon H_i + \frac{1}{2}y^iM \geq 0$ by $0 < \epsilon < 1$. Thus, $q(d_i)$ is a concave function of d_i since $\frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM \geq 0$. To make it strictly concave, we can add ηI_n to $\frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM$ such that $A_i(\eta) \triangleq \frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM + \eta I_n > 0$ and $\lim_{\eta \rightarrow 0} A_i(\eta) = \frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM$. Then solving $\nabla_{d_i} q(d) = -2A_i(\eta)d_i - (y^i w - \frac{2\epsilon}{1-\epsilon}H_i\mu_i) = 0$, we have $d = -\frac{1}{2}A_i(\eta)^{-1}(y^i w - \frac{2\epsilon}{1-\epsilon}H_i\mu_i)$. Then we get the dual function as follows:

$$\begin{aligned} g(H_i) &= -(y^i b + \xi_i - 1) + H_i \bullet \Sigma_i - \frac{\epsilon}{1-\epsilon}\mu_i^T H_i \mu_i \\ &\quad + \frac{1}{4}(y^i w - \frac{2\epsilon}{1-\epsilon}H_i\mu_i)^T A_i(\eta)^{-1}(y^i w - \frac{2\epsilon}{1-\epsilon}H_i\mu_i) \\ &= -(y^i b + \xi_i - 1) + H_i \bullet \Sigma_i - \frac{\epsilon}{1-\epsilon}\mu_i^T H_i \mu_i + \frac{1}{4}w^T A_i(\eta)^{-1}w \\ &\quad - y^i w^T A_i(\eta)^{-1}(\frac{\epsilon}{1-\epsilon}H_i\mu_i) + (\frac{\epsilon}{1-\epsilon}H_i\mu_i)^T A_i(\eta)^{-1}(\frac{\epsilon}{1-\epsilon}H_i\mu_i). \end{aligned}$$

For the last two terms, We have $y^i w^T A_i(\eta)^{-1}(\frac{\epsilon}{1-\epsilon}H_i\mu_i) = y^i w^T \mu_i - y^i w^T A_i(\eta)^{-1}(\frac{1}{2}y^i M \mu_i + \eta \mu_i)$, and $(\frac{\epsilon}{1-\epsilon}H_i\mu_i)^T A_i(\eta)^{-1}(\frac{\epsilon}{1-\epsilon}H_i\mu_i) = \frac{\epsilon}{1-\epsilon}\mu_i^T H_i \mu_i - \frac{1}{2}y^i \mu_i^T M \mu_i - \eta \mu_i^T \mu_i + (\frac{1}{2}y^i M \mu_i + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i M \mu_i + \eta \mu_i)$. Hence, we can derive

$$\begin{aligned} g(H_i) &= \frac{1}{4}w^T A_i(\eta)^{-1}w + (\frac{1}{2}y^i M \mu_i + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i M \mu_i + \eta \mu_i) \\ &\quad + y^i w^T A_i(\eta)^{-1}(\frac{1}{2}y^i M \mu_i + \eta \mu_i) \\ &\quad + H_i \bullet \Sigma_i - \frac{1}{2}y^i \mu_i^T M \mu_i - y^i w^T \mu_i - \eta \mu_i^T \mu_i - (y^i b + \xi_i - 1) \\ &= (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \\ &\quad + H_i \bullet \Sigma_i - \frac{1}{2}y^i \mu_i^T M \mu_i - y^i w^T \mu_i - \eta \mu_i^T \mu_i - (y^i b + \xi_i - 1). \end{aligned}$$

The dual problem of (A.4) is followed by

$$\begin{aligned} \inf \quad & g(H_i) = (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \\ & + H_i \bullet \Sigma_i - \frac{1}{2}y^i \mu_i^T M \mu_i - y^i w^T \mu_i - \eta \mu_i^T \mu_i - (y^i b + \xi_i - 1) \\ \text{s.t.} \quad & \epsilon H_i + \frac{1}{2}y^i M \geq 0, \\ & H_i \geq 0. \end{aligned} \quad (\text{A.5})$$

A bounded unconstrained convex program (A.5) is obtained. And $\nabla g(H_i) = 0$ implies that

$$A_i(\eta)^{-1}(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T A_i(\eta)^{-1} = \frac{1-\epsilon}{\epsilon} \Sigma_i. \quad (\text{A.6})$$

Multiply $(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T$ on the left hand side and $(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)$ on the right hand side of (A.6), we have

$$\begin{aligned} & \left((\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \right)^2 \\ &= \frac{1-\epsilon}{\epsilon} (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T \Sigma_i (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \\ &\Rightarrow (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \\ &= \sqrt{\frac{1-\epsilon}{\epsilon}} \| \Sigma_i^{\frac{1}{2}} (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \|_2. \end{aligned} \quad (\text{A.7})$$

Multiply $A_i(\eta)$ on the left hand side of (A.6) and take the trace, we have

$$\begin{aligned} & (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T \bullet A_i(\eta)^{-1} \\ &= (\frac{\epsilon}{1-\epsilon}H_i + \frac{1}{2}y^iM + \eta I_n) \bullet \frac{1-\epsilon}{\epsilon} \Sigma_i \\ &\Rightarrow (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i)^T A_i(\eta)^{-1}(\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \\ &= H_i \bullet \Sigma_i + \frac{1-\epsilon}{2\epsilon} y^i M \bullet \Sigma_i + \eta \frac{1-\epsilon}{\epsilon} \text{Trace}(\Sigma_i). \end{aligned} \quad (\text{A.8})$$

By (A.7) and (A.8), we have $\inf g(H_i) = -\frac{1}{2}y^i \mu_i^T M \mu_i - y^i w^T \mu_i + \sqrt{\frac{1-\epsilon}{\epsilon}} \| \Sigma_i^{\frac{1}{2}} (\frac{1}{2}y^i(M \mu_i + w) + \eta \mu_i) \|_2 - \frac{1-\epsilon}{2\epsilon} y^i M \bullet \Sigma_i - \eta \frac{1-\epsilon}{\epsilon} \text{Trace}(\Sigma_i) - \eta \mu_i^T \mu_i - (y^i b + \xi_i - 1)$, and $\lim_{\eta \rightarrow 0} \inf g(H_i) = -\frac{1}{2}y^i \mu_i^T M \mu_i - y^i w^T \mu_i + \sqrt{\frac{1-\epsilon}{\epsilon}} \| \Sigma_i^{\frac{1}{2}} (M \mu_i + w) \|_2 - \frac{1-\epsilon}{2\epsilon} y^i M \bullet \Sigma_i - (y^i b + \xi_i - 1)$. Since the strong duality

holds, we require that $\lim_{\eta \rightarrow 0} \inf g(H_i) \leq 0$ which yields SOC constraints $y^i(\frac{1}{2}\mu_i^T M \mu_i + \mu_i^T w + b) \geq 1 - \xi_i + \sqrt{\frac{1-\epsilon}{\epsilon}} \| \Sigma_i^{\frac{1}{2}} (M \mu_i + w) \|_2 - \frac{1-\epsilon}{2\epsilon} y^i M \bullet \Sigma_i$. This completes the proof. $\square$

A.3. Proof of Lemma 2.5

Proof. In $\mathcal{V}_i^E$, for any $x^i \in \mathcal{E}(\mu_i, \Sigma_i, \frac{1-\epsilon}{\epsilon})$, we require that $y^i \left(\frac{1}{2}(x^i)^T M x^i + w^T x^i + b \right) \geq 1 - \xi_i$ which is equivalent to

$$\begin{aligned} y^i \left(\frac{1}{2}\mu_i^T M \mu_i + \mu_i^T w + b \right) &\geq 1 - \xi_i - \sqrt{\frac{1-\epsilon}{\epsilon}} y^i \left(\Sigma_i^{\frac{1}{2}} (M \mu_i + w) \right)^T v \\ &\quad - \frac{1-\epsilon}{2\epsilon} y^i v^T \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} v, \end{aligned} \quad (\text{A.9})$$

for any $v \in \mathbb{R}^n$ with $\|v\|_2 \leq 1$. To eliminate v , we need to know the maximum of the RHS of the above inequality. We need to solve the following subproblem:

$$\begin{aligned} \inf_{v \in \mathbb{R}^n} \quad & \sqrt{\frac{1-\epsilon}{\epsilon}} y^i (\Sigma_i^{\frac{1}{2}} (M \mu_i + w))^T v + \frac{1-\epsilon}{2\epsilon} y^i v^T \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} v \\ \text{s.t.} \quad & \|v\|_2 \leq 1. \end{aligned} \quad (\text{A.10})$$

Problem (A.10) is a typical trust-region subproblem. The dual problem of (A.10) can be derived as follows:

$$\begin{aligned} \sup_{\lambda \in \mathbb{R}^{++}} \quad & \phi(\lambda) \\ \text{s.t.} \quad & \frac{1-\epsilon}{\epsilon} \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} + \lambda I > 0, \end{aligned} \quad (\text{A.11})$$

where $\phi(\lambda) = -\frac{1}{2} \sqrt{\frac{1-\epsilon}{\epsilon}} (\Sigma_i^{\frac{1}{2}} (M \mu_i + w))^T \sqrt{\frac{1-\epsilon}{\epsilon}} (\frac{1-\epsilon}{\epsilon} \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} + \lambda I)^{-1} (\Sigma_i^{\frac{1}{2}} (M \mu_i + w)) + \lambda$. A pair (v^*, λ^*) provides primal and dual optimal solutions if and only if

$$\begin{cases} \left(\frac{1-\epsilon}{\epsilon} \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} + \lambda I \right) v^* = -\sqrt{\frac{1-\epsilon}{\epsilon}} \left(\Sigma_i^{\frac{1}{2}} (M \mu_i + w) \right) \\ \lambda^* (\|v^*\|_2 - 1) = 0, \|v^*\|_2 \leq 1 \\ \frac{1-\epsilon}{\epsilon} \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} + \lambda^* I > 0, \lambda^* > 0. \end{cases} \quad (\text{A.12})$$

Let $\hat{v} = -y^i \frac{\Sigma_i^{\frac{1}{2}} (M \mu_i + w)}{\| \Sigma_i^{\frac{1}{2}} (M \mu_i + w) \|_2}$. We have $\|\hat{v}\|_2 = 1$, thus $\hat{v}$ is a feasible solution of the primal problem (A.10). The objective value respect to $\hat{v}$ is $-\sqrt{\frac{1-\epsilon}{\epsilon}} \| \Sigma_i^{\frac{1}{2}} (M \mu_i + w) \|_2 + y^i \frac{1-\epsilon}{2\epsilon} \hat{v}^T \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} \hat{v}$. For the second term of the objective value, we have

$$\begin{aligned} \hat{v}^T \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} \hat{v} &= \frac{(M \mu_i + w)^T \Sigma_i M \Sigma_i (M \mu_i + w)}{(M \mu_i + w)^T \Sigma_i (M \mu_i + w)} \\ &= \frac{M \Sigma_i \bullet (M \mu_i + w)(M \mu_i + w)^T \Sigma_i}{(M \mu_i + w)^T \Sigma_i (M \mu_i + w)}. \end{aligned}$$

According to Von Neumann's trace inequality, we have that $\theta_{\min}(M \Sigma_i) \text{Trace}((M \mu_i + w)(M \mu_i + w)^T \Sigma_i) \leq M \Sigma_i \bullet (M \mu_i + w)(M \mu_i + w)^T \Sigma_i \leq \theta_{\max}(M \Sigma_i) \text{Trace}((M \mu_i + w)(M \mu_i + w)^T \Sigma_i)$, where θ denotes the eigenvalue of $M \Sigma_i$. Since $\text{Trace}((M \mu_i + w)(M \mu_i + w)^T \Sigma_i) = (M \mu_i + w)^T \Sigma_i (M \mu_i + w)$, then $\theta_{\min}(M \Sigma_i) \leq \hat{v}^T \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} \hat{v} \leq \theta_{\max}(M \Sigma_i)$. If M is positive semi-definite, we have $\hat{v}^T \Sigma_i^{\frac{1}{2}} M \Sigma_i^{\frac{1}{2}} \hat{v} \leq \text{Trace}(M \Sigma_i)$. In this case, substituting the corresponding objective value make the inequality (A.9) become the constraint of $\mathcal{V}_i^{\text{SOC}}$ in (15) for each i . However, we can verify that $\hat{v}$ does not satisfy the optimal conditions (A.12), which means that $\mathcal{V}_i^E$ is a conservative cut of $\mathcal{V}_i^{\text{SOC}}$. Hence, we have $\mathcal{V}_i^E \subset \mathcal{V}_i^{\text{SOC}}$.

For the linear case when $M = 0$, the problem (A.10) becomes $\inf_{\|v\|_2 \leq 1} \sqrt{\frac{1-\epsilon}{\epsilon}} y^i (\Sigma_i^{\frac{1}{2}} w)^T v$. The optimal solution is that $v^* = -y^i \Sigma_i^{\frac{1}{2}} w / \| \Sigma_i^{\frac{1}{2}} w \|_2$. Consequently, the constraint (A.9) becomes $y^i (\mu_i^T w + b) \geq 1 - \xi_i + \sqrt{\frac{1-\epsilon}{\epsilon}} \| \Sigma_i^{\frac{1}{2}} w \|_2$, which is equivalent to $y^i (x^T w + b) \geq 1 - \xi_i, \forall x \in \mathcal{V}_i^E$ by utilizing the multivariate Chebyshev inequality (Bertsimas &

Table B.8

Testing results on “elliptic” synthetic data sets by DRC-QSSVM.

Data	Syn-Ellip-4d-50				Syn-Ellip-8d-50				Syn-Ellip-16d-50			
ϵ	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
Acc(%)	99.52	99.60	99.20	98.64	94.04	99.88	99.52	99.08	98.88	99.84	99.92	99.68
PrAcc(%)	97.17	97.97	96.46	95.32	87.12	88.74	88.53	88.31	78.67	84.76	86.94	87.11
CPU time (s)	SOCP	4.66	4.69	4.71	3.78	4.78	4.76	4.72	3.85	5.72	5.71	5.98
	SDP	6.87	7.14	7.13	5.71	9.97	9.59	9.65	8.28	34.96	30.90	29.54
Data	Syn-Ellip-4d-100				Syn-Ellip-8d-100				Syn-Ellip-16d-100			
ϵ	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
Acc(%)	99.40	99.64	99.56	99.60	100.00	100.00	100.00	98.80	98.60	99.64	99.64	99.64
PrAcc(%)	97.99	98.43	98.21	98.27	92.38	92.46	92.47	89.49	84.49	89.22	89.83	90.72
CPU time (s)	SOCP	7.19	7.37	7.38	5.46	5.98	5.90	6.14	3.67	13.49	13.90	13.92
	SDP	11.78	12.99	12.93	8.68	18.04	16.10	16.16	11.46	34.96	30.90	29.54

Popescu, 2005): for an arbitrary closed convex set S , $\mathbb{P}(\mathbf{x} \in S) \leq \frac{1}{1+c^2}$, where $c^2 = \inf_{\mathbf{x} \in S} (\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu})$. Notice that the set $\{\mathbf{x} | y^j (\mathbf{w}^T \mathbf{x} + b) \geq 1 - \xi_i\}$ is a convex set for each i . Hence, we have $\mathcal{V}_i^{SOC} = \mathcal{V}_i^E$. This completes the proof. $\square$

Appendix B. Auxiliary information

B.1. Moments uncertainty

The uncertainty of moments mentioned in Remark 3.1 adopts a general bounded set (Delage & Ye, 2010):

$$\left(\mathbb{E}_{F_i}[\tilde{\mathbf{x}}^i] - \boldsymbol{\mu}_i \right)^T \boldsymbol{\Sigma}_i^{-1} \left(\mathbb{E}_{F_i}[\tilde{\mathbf{x}}^i] - \boldsymbol{\mu}_i \right) \leq \omega_1, \quad (\text{B.1a})$$

$$\mathbb{E}_{F_i}[(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)^T] \leq \omega_2 \boldsymbol{\Sigma}_i, \quad (\text{B.1b})$$

where the parameters $\omega_1 \geq 0$ and $\omega_2 \geq 1$ provide natural means of quantifying one's confidence in $\boldsymbol{\mu}_i$ and $\boldsymbol{\Sigma}_i$. Constraint (B.1a) provides an ellipsoidal uncertainty of $\boldsymbol{\mu}_i$, and constraint (B.1b) assumes that $\boldsymbol{\Sigma}_i$ lies in a positive semidefinite cone. In what follows, to overcome the possible estimation errors of moments, we consider the ambiguity set for each i ,

$$D_i^{DY}(\tilde{\mathbf{x}}^i; \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i, \omega_1, \omega_2) \triangleq \left\{ F_i \in \mathcal{M}(\Xi_i, \mathcal{F}_i) \mid \begin{array}{l} \mathbb{P}(\tilde{\mathbf{x}}^i \in \Xi_i) = 1, \\ \left(\mathbb{E}_{F_i}[\tilde{\mathbf{x}}^i] - \boldsymbol{\mu}_i \right)^T \boldsymbol{\Sigma}_i^{-1} \left(\mathbb{E}_{F_i}[\tilde{\mathbf{x}}^i] - \boldsymbol{\mu}_i \right) \leq \omega_1, \\ \mathbb{E}_{F_i}[(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)(\tilde{\mathbf{x}}^i - \boldsymbol{\mu}_i)^T] \leq \omega_2 \boldsymbol{\Sigma}_i \end{array} \right\}. \quad (\text{B.2})$$

Let $D^{DY} \triangleq \bigcup_i D_i^{DY}$. Similarly, we can circumvent the difficulty of solving distributionally robust chance-constrained problems by duality theory. An SDP model can be obtained accordingly.

Lemma B.1. Suppose that $\omega_1 \geq 0$, $\omega_2 \geq 1$ and $\boldsymbol{\Sigma}_i > 0$ for any i . Then, the DRC-QSSVM model under the ambiguity set D^{DY} can be equivalently reformulated as the following SDP model:

$$\begin{aligned} \min \quad & \sum_{i=1}^N \eta_i \\ \text{s.t.} \quad & \begin{bmatrix} \mathbf{I}_n & \mathbf{M}\boldsymbol{\mu}_i + \mathbf{w} \\ (\mathbf{M}\boldsymbol{\mu}_i + \mathbf{w})^T & -\frac{c}{N}\xi_i + \eta_i \end{bmatrix} \geq 0, & i = 1, \dots, N, \\ & \beta_i - \frac{1}{\epsilon} \boldsymbol{\Gamma}_i^0 \bullet \mathbf{R}_i \geq \frac{1}{\epsilon} \sqrt{\omega_1} \|\boldsymbol{\Sigma}_i^0 \mathbf{R}_i \boldsymbol{\mu}_i^0\|, & i = 1, \dots, N, \\ & \mathbf{R}_i + \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \beta_i \end{bmatrix} \geq 0, & i = 1, \dots, N, \\ & \mathbf{R}_i \geq 0, & i = 1, \dots, N, \\ & \mathbf{M} \in \mathbb{S}^n, \mathbf{w} \in \mathbb{R}^n, b \in \mathbb{R}, \xi \in \mathbb{R}_+^N, \beta, \eta \in \mathbb{R}^N, \mathbf{R}_i \in \mathbb{S}^{n+1}, & i = 1, \dots, N, \end{aligned} \quad (\text{B.3})$$

where $\boldsymbol{\Gamma}_i^0 = \begin{bmatrix} \omega_2 \boldsymbol{\Sigma}_i + \boldsymbol{\mu}_i \boldsymbol{\mu}_i^T & \boldsymbol{\mu}_i \\ \boldsymbol{\mu}_i^T & 1 \end{bmatrix}$, $\boldsymbol{\Sigma}_i^0 = [\boldsymbol{\Sigma}_i^{\frac{1}{2}} \quad \mathbf{1}_n] \in \mathbb{R}^{n \times (n+1)}$, and $\boldsymbol{\mu}_i^0 = [\boldsymbol{\mu}_i \quad 0]^T \in \mathbb{R}^{n+1}$. The above result holds for the linear case when $\mathbf{M} = 0$.

Proof. Let $\varphi_i^{DY} \triangleq \sup_{F \in D_i^{DY}} \mathbb{P}\left\{ y^j \left(\frac{1}{2} (\tilde{\mathbf{x}}^i)^T \mathbf{M} \tilde{\mathbf{x}}^i + \mathbf{w}^T \tilde{\mathbf{x}}^i + b \right) \leq 1 - \xi_i \right\} = \sup_{F \in D_i^{DY}} \mathbb{E}_F[\mathbb{1}(\tilde{\mathbf{x}}^i)]$, for $i = 1, \dots, N$. From Lemma 1 in Delage and Ye (2010), φ_i^{DY} must be equal to the optimal value of the problem:

$$\inf r_i + t_i \quad (\text{B.4a})$$

$$\text{s.t. } r_i \geq \mathbb{1}(\tilde{\mathbf{x}}^i) - (\tilde{\mathbf{x}}^i)^T \mathbf{Q}_i \tilde{\mathbf{x}}^i - (\tilde{\mathbf{x}}^i)^T \mathbf{q}_i \quad \forall \tilde{\mathbf{x}}^i \in \Xi_i, \quad (\text{B.4b})$$

$$t_i \geq (\omega_2 \boldsymbol{\Sigma}_i + \boldsymbol{\mu}_i \boldsymbol{\mu}_i^T) \bullet \mathbf{Q}_i + \boldsymbol{\mu}_i^T \mathbf{q}_i + \sqrt{\omega_1} \|\boldsymbol{\Sigma}_i^{\frac{1}{2}} (\mathbf{q}_i + 2\mathbf{Q}_i \boldsymbol{\mu}_i)\|, \quad (\text{B.4c})$$

$$\mathbf{Q}_i \geq 0, \quad (\text{B.4d})$$

$$r_i, t_i \in \mathbb{R}, \mathbf{Q}_i \in \mathbb{S}^n, \mathbf{q}_i \in \mathbb{R}^n. \quad (\text{B.4e})$$

Let $\mathbf{N}_i = \begin{bmatrix} \mathbf{Q}_i & \frac{1}{2} \mathbf{q}_i \\ \frac{1}{2} \mathbf{q}_i^T & r_i \end{bmatrix}$ and $\boldsymbol{\Gamma}_i^0 = \begin{bmatrix} \omega_2 \boldsymbol{\Sigma}_i + \boldsymbol{\mu}_i \boldsymbol{\mu}_i^T & \boldsymbol{\mu}_i \\ \boldsymbol{\mu}_i^T & 1 \end{bmatrix}$. Using the same approach in the proof of Theorem 2.1, the constraint (B.4b) is equivalent to $\mathbf{N}_i + \alpha_i \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \frac{1}{\alpha_i} \end{bmatrix} \geq 0$, $\alpha_i > 0$. Denote $\boldsymbol{\Sigma}_i^0 = [\boldsymbol{\Sigma}_i^{\frac{1}{2}} \quad \mathbf{1}_n] \in \mathbb{R}^{n \times (n+1)}$, and $\boldsymbol{\mu}_i^0 = [\boldsymbol{\mu}_i \quad 0]^T \in \mathbb{R}^{n+1}$. Reformulate (B.4a) as $r_i + t_i \geq \boldsymbol{\Gamma}_i^0 \bullet \mathbf{N}_i + \sqrt{\omega_1} \|\boldsymbol{\Sigma}_i^0 \mathbf{N}_i \boldsymbol{\mu}_i^0\|$. Hence, we can rewrite (B.4) as follows:

$$\inf \boldsymbol{\Gamma}_i^0 \bullet \mathbf{N}_i + \sqrt{\omega_1} \|\boldsymbol{\Sigma}_i^0 \mathbf{N}_i \boldsymbol{\mu}_i^0\|$$

$$\text{s.t. } \mathbf{N}_i + \alpha_i \begin{bmatrix} \frac{1}{2} y^j \mathbf{M} & \frac{1}{2} y^j \mathbf{w} \\ \frac{1}{2} y^j \mathbf{w}^T & y^j b + \xi_i - 1 - \frac{1}{\alpha_i} \end{bmatrix} \geq 0,$$

$$\mathbf{N}_i \geq 0,$$

$$\alpha_i \in \mathbb{R}^{++}, \mathbf{N}_i \in \mathbb{S}^{n+1}.$$

Since $\varphi_i^{DY} \leq \epsilon$, by strong duality, we have $\boldsymbol{\Gamma}_i^0 \bullet \mathbf{N}_i + \sqrt{\omega_1} \|\boldsymbol{\Sigma}_i^0 \mathbf{N}_i \boldsymbol{\mu}_i^0\| \leq \epsilon$. The rest of the proof is similar to the proof of Theorem 2.1 and can be easily followed to complete the claim. $\square$

Notice that $D_i^{DY}(\tilde{\mathbf{x}}^i; \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i, 0, 1)$ relates closely to $D_i(\tilde{\mathbf{x}}^i; \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$, and in this event, we find that the model (B.3) reduces to the model (12).

B.2. Experimental results

See Tables B.8 and B.9.

Table B.9
Testing results on “parabolic” synthetic data sets by DRC-QSSVM.

Data		Syn-Para-4d-50				Syn-Para-8d-50				Syn-Para-16d-50			
ϵ		0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
Acc(%)		97.36	98.88	98.68	98.12	98.80	98.52	98.16	98.12	97.90	98.42	96.70	98.18
PrAcc(%)		93.85	98.21	98.01	97.31	97.25	97.24	96.97	96.96	86.90	90.03	88.03	88.30
CPU time (s)	SOCP	4.06	3.79	3.87	3.69	3.83	4.01	3.87	3.93	5.72	5.73	5.75	5.46
	SDP	6.28	6.06	5.96	6.57	8.76	8.22	8.32	8.49	33.49	31.97	30.44	30.31

Data		Syn-Para-4d-100				Syn-Para-8d-100				Syn-Para-16d-100			
ϵ		0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00	0.10	0.25	0.50	1.00
Acc(%)		99.40	99.32	99.20	99.04	98.88	98.76	98.48	98.32	98.44	97.00	96.06	96.66
PrAcc(%)		99.06	98.97	98.75	98.54	97.85	97.69	97.39	97.22	95.20	91.78	92.59	92.95
CPU time (s)	SOCP	6.97	6.89	6.80	6.60	5.62	23.74	22.63	5.58	10.21	10.42	10.27	10.26
	SDP	9.98	9.99	9.68	9.66	16.93	15.08	15.38	15.88	153.89	140.78	132.56	140.67

References

Ben-Tal, A., Bhadra, S., Bhattacharyya, C., & Nath, J. S. (2011). Chance constrained uncertain classification via robust optimization. *Mathematical Programming*, 127(1), 145–173.

Bertsimas, D., Dunn, J., Pawlowski, C., & Zhuo, Y. D. (2019). Robust classification. *INFORMS Journal on Optimization*, 1(1), 2–34.

Bertsimas, D., & Popescu, I. (2005). Optimal inequalities in probability theory: A convex optimization approach. *SIAM Journal on Optimization*, 15(3), 780–804.

Bhattacharyya, C., Grate, L., Jordan, M. I., Ghaoui, L. E., & Mian, I. S. (2004). Robust sparse hyperplane classifiers: Application to uncertain molecular profiling data. *Journal of Computational Biology*, 11(6), 1073–1089.

Chang, C.-C., & Lin, C.-J. (2011). LIBSVM: A library for support vector machines. *ACM Transactions on Intelligent Systems and Technology*, 2(3), 1–27.

Chen, R., & Paschalidis, I. C. (2020). Distributionally robust learning. *Foundations and Trends® in Optimization*, 4(1–2), 1–243.

Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297.

Dagher, I. (2008). Quadratic kernel-free non-linear support vector machine. *Journal of Global Optimization*, 41(1), 15–30.

Delage, E., & Ye, Y. (2010). Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Operations Research*, 58(3), 595–612.

Gao, Z., Fang, S.-C., Luo, J., & Medhin, N. (2021). A kernel-free double well potential support vector machine with applications. *European Journal of Operational Research*, 290(1), 248–262.

Goldfarb, D., & Iyengar, G. (2003). Robust convex quadratically constrained programs. *Mathematical Programming*, 97(3), 495–515.

Hsu, W.-K., Xu, J., Lin, X., & Bell, M. R. (2022). Integrated online learning and adaptive control in queueing systems with uncertain payoffs. *Operations Research*, 70(2), 1166–1181.

Huang, G., Song, S., Gupta, J. N., & Wu, C. (2013). A second order cone programming approach for semi-supervised learning. *Pattern Recognition*, 46(12), 3548–3558.

Jiang, R., Han, S., Yu, Y., & Ding, W. (2023). An access control model for medical big data based on clustering and risk. *Information Sciences*, 621, 691–707.

Jiménez-Cordero, A., Morales, J. M., & Pineda, S. (2021). A novel embedded min-max approach for feature selection in nonlinear support vector machine classification. *European Journal of Operational Research*, 293(1), 24–35.

Khanjani-Shiraz, R., Babapour-Azar, A., Hosseini-Nodeh, Z., & Pardalos, P. M. (2023). Distributionally robust joint chance-constrained support vector machines. *Optimization Letters*, 17(2), 299–332.

Kuhn, D., Esfahani, P. M., Nguyen, V. A., & Shafieezadeh-Abadeh, S. (2019). Wasserstein distributionally robust optimization: Theory and applications in machine learning. In N. Serguei, S. Douglas, & H. J. Greenberg (Eds.), *Operations research & management science in the age of analytics* (pp. 130–166). INFORMS.

Lin, F., Fang, X., & Gao, Z. (2022). Distributionally robust optimization: A review on theory and applications. *Numerical Algebra, Control & Optimization*, 12(1), 159–212.

Luo, J., Fang, S.-C., Deng, Z., & Guo, X. (2016). Soft quadratic surface support vector machine for binary classification. *Asia-Pacific Journal of Operational Research*, 33(6), Article 1650046.

Luo, J., Fang, S.-C., Deng, Z., & Tian, Y. (2022). Robust kernel-free support vector regression based on optimal margin distribution. *Knowledge-Based Systems*, 253, Article 109477.

Luo, J., Yan, X., & Tian, Y. (2020). Unsupervised quadratic surface support vector machine with application to credit risk assessment. *European Journal of Operational Research*, 280(3), 1008–1017.

Ma, Q., & Wang, Y. (2021). Distributionally robust chance constrained svm model with l_2 -wasserstein distance. *Journal of Industrial and Management Optimization*.

Mi, Y., Quan, P., Shi, Y., & Wang, Z. (2022). Concept-cognitive computing system for dynamic classification. *European Journal of Operational Research*, 301(1), 287–299.

Naumzik, C., Feuerriegel, S., & Nielsen, A. M. (2023). Data-driven dynamic treatment planning for chronic diseases. *European Journal of Operational Research*, 305(2), 853–867.

Peng, S., Gianpiero, C., & Zhihua, A.-Z. (2023). Chance constrained conic-segmentation support vector machine with uncertain data. *Annals of Mathematics and Artificial Intelligence*, 1–23.

Shivaswamy, P. K., Bhattacharyya, C., & Smola, A. J. (2006). Second order cone programming approaches for handling missing and uncertain data. *Journal of Machine Learning Research*, 7(47), 1283–1314.

Singla, M., Ghosh, D., & Shukla, K. (2020). A survey of robust optimization based machine learning with special reference to support vector machines. *International Journal of Machine Learning and Cybernetics*, 11(7), 1359–1385.

Sun, Y., Wong, A. K., & Kamel, M. S. (2009). Classification of imbalanced data: A review. *International Journal of Pattern Recognition and Artificial Intelligence*, 23(04), 687–719.

Thabtah, F., Hammoud, S., Kamalov, F., & Gonsalves, A. (2020). Data imbalance in classification: Experimental evaluation. *Information Sciences*, 513, 429–441.

Toh, K.-C., Todd, M. J., & Tütüncü, R. H. (1999). SDPT3 – a Matlab software package for semidefinite programming, Version 1.3. *Optimization Methods & Software*, 11(1–4), 545–581.

Trafalis, T. B., & Gilbert, R. C. (2006). Robust classification and regression using support vector machines. *European Journal of Operational Research*, 173(3), 893–909.

Vassilvitskii, S., & Arthur, D. (2006). K-means++: The advantages of careful seeding. In *Proceedings of the Eighteenth Annual ACM-SIAM Symposium on Discrete Algorithms* (pp. 1027–1035).

Wang, X., Fan, N., & Pardalos, P. M. (2017). Stochastic subgradient descent method for large-scale robust chance-constrained support vector machines. *Optimization Letters*, 11, 1013–1024.

Wang, X., Fan, N., & Pardalos, P. M. (2018). Robust chance-constrained support vector machines with second-order moment information. *Annals of Operations Research*, 263(1), 45–68.

Wang, X., & Pardalos, P. M. (2014). A survey of support vector machines with uncertainties. *Annals of Data Science*, 1(3–4), 293–309.

Zhou, Z. (2021). *Machine learning*. Springer Nature (Chapter 6).