



Multi-class Support Vector Machine with Uncertain Inputs by Probabilistic Constraints

Mohammadi, T.¹ , Jabbari, H.¹ , Effati, S.² 

¹Department of Statistics, Ferdowsi University of Mashhad, Mashhad, Iran.

²Department of Applied Mathematics, Ferdowsi University of Mashhad, Mashhad, Iran.

Corresponding author: Hadi Jabbari, jabbarinh@um.ac.ir

Received: 12/11/2024 Revised: 24/1/2025 Accepted and Published Online: 27/11/2025.

Introduction

Support Vector Machine (SVM) is a robust supervised learning algorithm primarily used for classification and regression. It is initially invented for binary classification. However, many real-world problems involve multi-class classification, where the goal is to distinguish between more than two classes. Various strategies have been proposed to address this challenge, including decomposition methods that reduce the multi-class problem into multiple binary subproblems, as well as direct approaches that optimize for all classes simultaneously. Two of the methods that distinct all classes simultaneously are Cramer-Singer and simplified multi-class SVM (SimMSVM). The challenge in multi-class SVM lies in accurate inputs while the data are uncertain in the real world due to measurement errors. Some approaches include Crammer and Singer's method or SimMSVM. Despite having nice strengths and performance, they consider inputs in an accurate state. This work has explored the theoretical foundations and practical implementations of SimMSVM in which the inputs are random variables and the constraints are probabilistic. The constraints have been released from the probabilistic state by leveraging statistical theories, and the expectation value has been entered into the model. Hence, the model has been converted into ordinary quadratic programming. This paper will show that the model with uncertain inputs has a better procedure than the traditional one.

Material and Methods

This paper applies the expected value to release the constraints from the probabilistic state. Then, the sample mean is used to estimate the expected value in the unknown population state. Also, the Monte Carlo method is implemented to simulate the data, and then the bootstrap resampling method is used to generate random samples to estimate the expected value. Moreover, the confusion matrix statistical indices and accuracy criterion are used to evaluate the proposed and competing models.

Results and Discussion

Using three kernel functions, linear, polynomial, and Gaussian, we investigate the accuracy values of the three models. In the simulation subsection, we examine the performance of the three proposed models, basic and Cramer-Singer, using the statistical indices of mean and standard deviation, and we compare the three models using the Friedman test statistic. Also, in the real example subsection, we evaluate these models in different parameter ranges for three real data sets of seeds, ecoli, and leukemia.

Conclusion

Given that most of the data we encounter in the real world is not precise and certain and has inherent flaws due to measurement errors, this paper proposes a model that addresses this issue. The results show that the model based on random inputs is far better than the models based on definite and precise inputs.

Keywords: Multi-class SVM, Stochastic input, Probabilistic constraint, Expectation.

Mathematics Subject Classification (2010): 62H30, 62G35.



©The Author(s). The Publisher is Iranian Statistical Society.

This is an open access article distributed under the terms and conditions of [CC BY-NC 4.0](https://creativecommons.org/licenses/by-nc/4.0/)

ماشین بردار پشتیبان چند-کلاسه با ورودی‌های تصادفی با محدودیت‌های احتمالی

تارا محمدی^۱، هادی جباری نوقابی^۱ و سهراب عفتی^۲

^۱گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

^۲گروه ریاضی کاربردی، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

نویسنده مسئول: هادی جباری نوقابی، jabbarinh@um.ac.ir

تاریخ دریافت: ۱۴۰۳/۸/۲۲ تاریخ بازنگری: ۱۴۰۳/۱۱/۵ تاریخ پذیرش و انتشار: ۱۴۰۳/۱۱/۸

چکیده: ماشین بردار پشتیبان به عنوان یک الگوریتم با نظارت در ابتدا برای حالت دودویی ابداع شد، سپس به علت کاربردهای آن، الگوریتم‌های چند-کلاسه نیز طراحی شدند و همچنان به عنوان یک پژوهش در حال بررسی است. اخیراً مدل‌هایی برای بهبود روش‌های چند-کلاسه ارائه گردیده است. اغلب آن‌ها حالتی را بررسی می‌کنند که در آن ورودی‌ها غیرتصادفی هستند در حالی که در دنیای واقعی با داده‌های غیرقطعی و نادقیق مواجه هستیم. لذا در این مقاله به بررسی مدلی پرداخته شده است که در آن ورودی‌ها تصادفی و محدودیت‌های مسئله نیز احتمالی هستند. با استفاده از قضایای آماری و با استفاده از امیدریاضی، محدودیت‌های مسئله از حالت احتمالی خارج شده است. سپس از روش برآورد گشتاوری، برای برآورد امیدریاضی استفاده شده است. با استفاده از شبیه‌سازی مونت کارلو به تولید داده‌های مصنوعی پرداخته و از روش بازنمونه‌گیری بوت استرپ نمونه‌ها را به عنوان ورودی به مدل داده و دقت مدل مورد بررسی قرار گرفته است. در نهایت مدل پیشنهادی را با داده‌های واقعی آموزش داده و دقت آن با شاخص‌های آماری ارزیابی شده است. نتایج حاصل از شبیه‌سازی و مثال واقعی برتری کارایی مدل پیشنهادی را بر مدل‌های مبتنی بر ورودی‌های قطعی نشان می‌دهد.

واژه‌های کلیدی: ماشین بردار پشتیبان چند-کلاسه، ورودی تصادفی، محدودیت احتمالی، امیدریاضی. کد موضوع بندی ریاضی (۲۰۱۰): 62H30، 62G35.



۱ مقدمه

ماشین بردار پشتیبان^۱ یک الگوریتم با نظارت و از شاخه‌های یادگیری ماشین به شمار می‌رود. این الگوریتم در ابتدا توسط **وپنیک (۱۹۹۵)** برای کلاس‌های دودویی طراحی شد. او به همراه **کورتس (۱۹۹۵)** مدل خطی را به حالت غیرخطی ارتقا دادند. مدل ارائه شده توسط وپنیک فقط به کلاس‌های دودویی محدود می‌شد، در حالی که اغلب مواقع با داده‌های بیش از دو کلاس سروکار داریم. در میان مدل‌های چند-کلاسه، برخی کلاس‌ها را به صورت دودویی و برخی به صورت همزمان تفکیک می‌کنند. مدل‌های ارائه شده توسط **شیلکوپ و همکاران (۱۹۹۵)**، **وپنیک (۱۹۹۹)**، **نر و همکاران (۱۹۹۰)**، **پلات و همکاران (۱۹۹۹)** و **هستی و تیشیرانی (۱۹۹۷)** از جمله مدل‌های چند-کلاسه‌ای هستند که کلاس‌ها را به صورت دودویی تفکیک می‌کنند. همچنین مدل‌های ارائه شده توسط **وپنیک (۱۹۹۹)** و **وستون و واتکینز (۱۹۹۹)** از جمله مدل‌هایی هستند که کلاس‌ها را به صورت همزمان در نظر می‌گیرند. **گونن و همکاران (۲۰۰۸)** مدلی به نام ماشین بردار پشتیبان احتمال پسین چند-کلاسه ارائه کرد. آن‌ها مقادیر احتمال پسینی را برای هر کلاس در ساختار محدودیت مدل **وستون و واتکینز (۱۹۹۹)** در نظر گرفتند. در این مدل، اثر نقاط پرت کاهش می‌یابد. همچنین در روش آن‌ها، با کاهش تعداد ماشین‌های بردار پشتیبان، پیچیدگی مکانی و زمانی کاهش می‌یابد، زیرا محاسبات کمتری مورد نیاز است.

مدلی که توسط **هیی و همکاران (۲۰۱۲)** منتشر شده است یک ماشین بردار پشتیبان چند-کلاسه ساده شده را ارائه می‌کند که مدل **کرامر و سینگر (۲۰۰۱)** را با معرفی یک کران خطای جدید ارتقا می‌دهد و سبب افزایش دقت و عملکرد بهتر می‌گردد. همچنین مدل دیگری به نام ماشین بردار پشتیبان با محدودیت‌های احتمالی برای ورودی‌های نامعین توسط **یزدی و همکاران (۲۰۰۷)** پیشنهاد گردیده است که بر اساس توابع چگالی احتمال عمل می‌کند. در مدل آن‌ها ابتدا با لحاظ کردن یک خطای تصادفی در محدودیت مسئله اولیه، از حالت قطعی به حالت غیرقطعی تبدیل شد سپس با استفاده از روابط آماری و استفاده از تعریف تابع توزیع محدودیت را از حالت احتمالی خارج کرد. در نهایت نتایج نشان داد کارایی مدل آن‌ها نسبت به مدل استاندارد با ورودی‌های قطعی بالاتر است. علاوه بر این مدل‌ها، یک روش چند-کلاسه با محدودیت‌های احتمالی و داده‌های غیرقطعی توسط **بوش و همکاران (۲۰۱۳)** ارائه شد که در آن از طبقه‌بندی کننده‌های تصادفی دودویی و تصادفی چند-کلاسه یک در برابر همه استفاده شد. کارایی مدل پیشنهادی با عملکرد بهتری نسبت به مدل قطعی به دست آمد. همچنین یک مدل احتمالی توسط **لی و همکاران (۲۰۱۳)** برای داده‌های متأثر از اغتشاش منتشر شد که در آن مسئله اولیه بر اساس تابع چگالی احتمال وزن‌ها است. آن‌ها نشان دادند که در حضور اغتشاش می‌تواند کارایی بیشتری نسبت به حالت معمولی داشته باشد. مدل چند-کلاسه جدید دیگری توسط **چن و همکاران (۲۰۱۶)** پیشنهاد شد. در این مدل با به کارگیری یک الگوریتم متوالی برای کران خطا، سبب افزایش سرعت روند آموزش گردید. **عباس زاده و عفتی (۲۰۱۸)** یک رگرسیون بردار پشتیبان با ورودی و خروجی تصادفی پیشنهاد کردند که محدودیت‌های آن به صورت احتمالی است. محدودیت‌های احتمالی سبب استحکام مدل و کارایی بهتری نسبت به مدل استاندارد گردید. مدل ارائه شده توسط

¹Support Vector Machine

عباس زاده و عفتی (۲۰۱۹) یک مدل دودویی بر اساس محدودیت‌های احتمالی است. در آن نمونه‌های آموزشی به عنوان متغیرهای تصادفی و محدودیت مسئله نیز به صورت احتمالی در نظر گرفته شد. در مدل آن‌ها با استفاده از روش برآورد گشتاوری محدودیت مسئله از حالت احتمالی خارج گردید. نتایج نشان داد که کارایی آن به مراتب بهتر از مدل استاندارد است.

مدل دیگری توسط **دُن و یاکوب (۲۰۲۰)** پیشنهاد شده است که به دلیل پراکندگی داده‌ها در فضای با ابعاد بالا انجام شد. این کار با تقسیم کل مجموعه داده آموزشی به زیرمجموعه‌های مجزا بر اساس تقسیمات دودویی انجام شده است. در فرایند مدل‌سازی با حذف بعضی تقسیمات و تکرار این فرایند تنها دو تصمیم دودویی وجود خواهد داشت. آزمایشات نشان داد مدل آن‌ها عملکرد بهتری نسبت به بعضی از مدل‌های فعلی دارد. همچنین یک مدل چند-کلاسه با نام حداقل مربعات رگرسیون بردار پشتیبان با k کلاس توسط **موسایی و هِلادیک (۲۰۲۲)** پیشنهاد شده است. برای به حداقل رساندن متغیرهای کمکی در مسئله اولیه، محدودیت‌های نامساوی را با محدودیت‌های مساوی جایگزین کرده و ℓ_2 را به جای ℓ_1 اعمال کرده است. این تغییر باعث شد الگوریتم سرعت بیشتر و عملکرد بهتری داشته باشد.

المعصومه و همکاران (۲۰۲۲) داده‌های هیستوگرامی را با ماشین بردار پشتیبان احتمالی مبتنی بر امیدریاضی و مبتنی بر کمیت و با محدودیت‌های احتمالی ارائه کردند. روش آن‌ها برای هر دو حالت دودویی و چند-کلاسه انجام شد که عملکرد بهتری نسبت به الگوی معمولی داشت. یک مدل چند-کلاسه توسط **صالح و همکاران (۲۰۲۳)** با استفاده از طرح یک در برابر همه پیشنهاد شده است. این مدل، مسئله را به زیرمسئله‌های دودویی تقسیم می‌کند که در آن دو ابر صفحه غیرموازی را پیدا می‌کند. مدل آن‌ها به خوبی توانست در تقسیم‌بندی کلاس‌ها کارایی بیشتری نسبت به بعضی مدل‌های پیش از خود به دست آورد. در میان مدل‌های ارائه شده اخیر، **بنیتز و همکاران (۲۰۲۴)** خروجی‌های احتمالی را از طریق برآورد نقطه‌ای تولید کردند. مدل آن‌ها به گونه‌ای طراحی شده است که نسبت به هزینه حساس است. همچنین مقادیر احتمالات از طریق برآوردهای بوت استرپ انجام می‌شود. برای بهبود حساسیت، آن‌ها از دو رویکرد مختلف استفاده کردند، روش اول مبتنی بر بوت استرپ که اندازه‌گیری حساسیت توسط احتمالات کلاس پسین کنترل می‌شود. روش دوم مشابه اول انجام می‌شود، اما طبقه‌بندی متفاوتی را در نظر می‌گیرد. اخیراً مدلی توسط **موسوی و گلعلی‌زاده (۱۴۰۲)** ارائه شد که در آن به دلیل چالش ابعاد روش‌های تحلیل داده‌های ژنی برای ارزیابی سرطان از تکنیک‌های کاهش ابعاد و ماشین بردار پشتیبان برای حل این مشکل استفاده شد. در مدل آن‌ها به نام ماشین بردار پشتیبان تصادفی خوشه‌ای با نمایش تصادفی فضای ویژگی و ترکیب طبقه‌بندی‌کننده‌ها، دقت بالایی در شناسایی ژن‌های مهم و پیش‌بینی سرطان پروستات ارائه شد. بیشتر اوقات در زندگی واقعی با داده‌های معلوم و دقیقی مواجه نیستیم و در عمل به دلایلی مانند خطای اندازه‌گیری، خطای نمونه‌گیری با داده‌هایی که دقیق و قطعی نیستند، سروکار داریم. از سوی دیگر، در موارد مکرر، داده‌ها بیش از دو کلاس هستند، بنابراین تصمیم‌گرفتن هر دو مشکل را به طور همزمان برطرف کنیم. انگیزه اصلی این مقاله این است که برای حل این مشکلات، مدلی را برای ورودی‌های تصادفی و تقسیم‌بندی به سه کلاس و بیشتر ارائه کند.

در این مقاله یک مدل چند-کلاسه با ورودی‌های تصادفی و محدودیت‌های احتمالی در نظر گرفتیم. در بخش

۲ ساختار مدل را تعریف و با استفاده از قضایای آماری، آن را از حالت احتمالی خارج می‌کنیم. سپس با روش‌های آماری به برآورد پارامترهای مدل می‌پردازیم. در بخش ۳ به شبیه‌سازی و در بخش ۴ به ارائه مثال واقعی جهت بررسی عملکرد مدل با نمونه‌هایی از مجموعه داده‌های واقعی از پایگاه داده یادگیری ماشین **دوآ و گزف** (۲۰۱۹) می‌پردازیم. بطور کلی نتایج شبیه‌سازی و مثال واقعی نشان می‌دهد که کارایی مدل پیشنهادی به مراتب بهتر از مدل‌های با ورودی‌های قطعی است.

۲ ساختار مدل بر اساس محدودیت‌های احتمالی

این مقاله نسخه بهبودیافته مدل ارائه شده توسط **هیی و همکاران** (۲۰۱۲) را ارائه می‌کند که حالت بهبودیافته مدل **کرامر و سینگر** (۲۰۰۱) است، با این تفاوت که مدل **کرامر و سینگر** (۲۰۰۱) مجموعه فشرده‌ای از محدودیت‌ها است که تعداد متغیرها در مسئله دوگان $\ell \times k$ است، و این مقدار حتی برای مجموعه داده‌های کوچک پیچیدگی محاسباتی بالایی دارد. الگوی ارائه شده توسط **هیی و همکاران** (۲۰۱۲) روشی ساده برای کاهش محدودیت پیشنهاد می‌کند. در الگوی **هیی و همکاران** (۲۰۱۲) یک مسئله برنامه‌ریزی درجه دوم با ℓ متغیر خواهیم داشت. اکنون مسئله اولیه **کرامر و سینگر** (۲۰۰۱) را بیان می‌کنیم، مجموعه داده‌های آموزشی $S = \{(x_1, y_1), \dots, (x_\ell, y_\ell)\}$ را در نظر بگیرید، که در آن $x_i \in \mathbb{R}^n$ نمونه i ام با n ویژگی از کلاس $y_i = 1, \dots, k$ است. بنابراین مسئله اولیه به صورت

$$\begin{aligned} \min \quad & \frac{1}{2} \sum_{m=1}^k \mathbf{w}_m^T \mathbf{w}_m + C \sum_{i=1}^{\ell} \xi_i \\ \text{s.t.} \quad & \mathbf{w}_{y_i}^T \phi(x_i) - \mathbf{w}_t^T \phi(x_i) \geq 1 - \delta_{y_i, t} - \xi_i, \\ & \xi_i \geq 0, i = 1, \dots, \ell, t = 1, \dots, k \neq y_i. \end{aligned}$$

بیان می‌شود، که در آن $C > 0$ پارامتر جریمه و $\xi_i \geq 0$ متغیرهای کمکی هستند و

$$\delta_{ij} = \begin{cases} 1 & i = j, \\ 0 & i \neq j \end{cases}$$

دلتای کرونکر و $\mathbf{w}^T = [w_1, \dots, w_n]^T$ بردار وزن است. همچنین $\phi(\cdot) : \mathbb{R}^n \rightarrow \mathbb{R}^{n'}$ یک تبدیل غیرخطی است که از طریق ضرب نقطه‌ای، فضای ورودی را از n به فضای ویژگی با ابعاد بالاتر n' تبدیل می‌کند. تابع تصمیم^۱

^۱Decision function

بهینه در رابطه

$$\arg \max_m f_m(\mathbf{x}) = \mathbf{w}_m^T \phi(\mathbf{x}). \quad (۱)$$

تعریف می‌شود. اگرچه مدل چند-کلاسه **کرامر و سینگر** (۲۰۰۱) مجموعه فشرده‌ای از محدودیت‌ها را ارائه می‌دهد و تعداد متغیرها در مسئله دوگان آن‌ها $(\ell \times k)$ است که منجر به پیچیدگی محاسباتی زیادی می‌شود بنابراین مدل ارائه شده توسط **هیی و همکاران** (۲۰۱۲) شیوه **کرامر و سینگر** (۲۰۰۱) را دنبال می‌کند و یک روش ساده‌سازی شده برای کاهش محدودیت‌های **کرامر و سینگر** (۲۰۰۱) پیشنهاد می‌کند. با انجام این کار، به یک مسئله برنامه‌ریزی درجه دوم با ℓ متغیر تبدیل می‌شود. فرض کنید برای نمونه آموزشی x_i عبارت

$$\zeta_{i,m} = 1 - f_{y_i}(x_i) - f_m(x_i), \quad m \in \{1, \dots, k\} \setminus y_i.$$

نشان دهنده حاشیه زوجی نادرست بین کلاس واقعی y_i و هر کلاس دیگر m است (وستون و واتکینز، ۱۹۹۹) که در آن به جای "۱" از "۲" استفاده شده است. در تابع زیان وستون و واتکینز (۱۹۹۹) تمامی حاشیه‌های زوجی نادرست $(\zeta_{i,m} > 0)$ به شکل $\xi_i^{(1)} = \sum_{m \neq y_i} [\zeta_{i,m}]_+ \cdot$ جمع می‌شود، که در آن $[\cdot]_+ \equiv \max(\cdot, 0)$. اما در مدل کرامر و سینگر (۲۰۰۱) زیان نمونه تنها حداکثر حاشیه زوجی نادرست به حالت $\xi_i^{(2)} = [\max_{m \neq y_i} \zeta_{i,m}]_+$ بیان می‌شود. برای کاهش اندازه مسئله، تعداد محدودیت‌ها در **هیی و همکاران** (۲۰۱۲) به جای $(\ell \times k)$ باید متناسب با ℓ باشد. لذا آن‌ها کران جدید $\xi_i^{(3)} = [\frac{1}{k-1} \sum_{m=1, m \neq y_i}^k \zeta_{i,m}]_+$ را معرفی کردند. بنابراین رابطه بین این زیان‌ها $\xi_i^{(1)} \geq \xi_i^{(2)} \geq \xi_i^{(3)}$ است. در نهایت مدل ارائه شده توسط **هیی و همکاران** (۲۰۱۲) در رابطه

$$\begin{aligned} \min \quad & \frac{1}{\gamma} \sum_{m=1}^k \mathbf{w}_m^T \mathbf{w}_m + C \sum_{i=1}^{\ell} \xi_i \\ \text{s.t.} \quad & \mathbf{w}_{y_i}^T \phi(x_i) - \frac{1}{k-1} \sum_{m \neq y_i} \mathbf{w}_m^T \phi(x_i) \geq 1 - \xi_i, \\ & \xi_i \geq 0, i = 1, \dots, \ell. \end{aligned} \quad (۲)$$

بیان می‌شود. برای به دست آوردن مقادیر بهینه متغیرها، تابع دوگان لاگرانژ را به صورت

$$\begin{aligned} L(\mathbf{w}, \xi, \alpha, \lambda) = & \frac{1}{\gamma} \sum_{m=1}^k \mathbf{w}_m^T \mathbf{w}_m + C \sum_{i=1}^{\ell} \xi_i - \sum_{i=1}^{\ell} \alpha_i \xi_i - \sum_{i=1}^{\ell} \lambda_i \xi_i \\ & - \sum_{i: m=y_i} \alpha_i \mathbf{w}_{y_i}^T \phi(\mathbf{x}_i) + \frac{1}{k-1} \sum_{i: m \neq y_i} \alpha_i \mathbf{w}_m^T \phi(\mathbf{x}_i) + \sum_{i=1}^{\ell} \alpha_i \end{aligned}$$

۲۱۲ ماشین بردار پشتیبان چند-کلاسه با ورودی‌های تصادفی

تعریف می‌کنیم، که در آن $\alpha_i \geq 0$ و $\lambda_i \geq 0$ ضرایب لاگرانژ هستند. مقادیر بهینه متغیرهای مسئله اولیه به ترتیب

$$\mathbf{w}_m = \sum_{i:y_i=m} \alpha_i \mathbf{x}_i - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i \mathbf{x}_i, \quad (۳)$$

$$C\mathbf{e} = \boldsymbol{\alpha} + \boldsymbol{\lambda}, \quad (۴)$$

هستند، که در آن $\mathbf{e}$ بردار یکه است. با جایگذاری روابط (۳) و (۴) در (۱) تابع تصمیم

$$f_m(\mathbf{x}) = \sum_{i:y_i=m} \alpha_i K(\mathbf{x}_i, \mathbf{x}_j) - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i K(\mathbf{x}_i, \mathbf{x}_j) \quad (۵)$$

حاصل خواهد شد، که در آن $K(\mathbf{x}_i, \mathbf{x}_j) = \phi^T(\mathbf{x}_i)\phi(\mathbf{x}_j)$ تابع هسته است. با ساده کردن رابطه (۵) تابع تصمیم بهینه به صورت

$$\arg \max_m f_m(\mathbf{x}) = \arg \max_m \sum_{i:y_i=m} \alpha_i K(\mathbf{x}_i, \mathbf{x}_j), \quad (۶)$$

تبدیل می‌شود. اکنون بر اساس رابطه (۲) در حالتی که ورودی‌ها تصادفی هستند، مدل پیشنهادی ارائه می‌شود. پس مجموعه داده‌های آموزشی به شکل $\tilde{S} = \{(\mathbf{X}_1, y_1), \dots, (\mathbf{X}_\ell, y_\ell)\}$ تبدیل می‌شود، که در آن $\mathbf{X}_i = [X_1, \dots, X_n]^T$ بردار متغیر تصادفی ورودی است. همچنین در مسئله برنامه‌ریزی غیرقطعی یک حالت در نظر گرفتن محدودیت‌های مسئله به صورت احتمالی است. برای ضمانت تعمیم مدل (۲)، با اضافه کردن b_m^y به عبارت تنظیم‌کننده $\mathbf{w}_m^T \mathbf{w}$ تابع هدف مسئله اولیه در رابطه

$$\begin{aligned} \min \quad & \frac{1}{\gamma} \sum_{m=1}^k (\mathbf{w}_m^T \mathbf{w}_m + b_m^y) + C \sum_{i=1}^{\ell} \xi_i \\ \text{s.t.} \quad & Pr((\mathbf{w}_{y_i}^T \phi(\mathbf{X}_i) + b_{y_i}) - \frac{1}{k-1} \sum_{m \neq y_i} (\mathbf{w}_m^T \phi(\mathbf{X}_i) + b_m) \geq 1 - \xi_i) \geq \gamma, \\ & \xi_i \geq 0, i = 1, \dots, \ell, \end{aligned}$$

بیان می‌شود، که در آن γ احتمال به ازای اثر نمونه نام است.

قضیه ۱. (بارنت و همکاران، ۲۰۰۲) اگر U متغیری تصادفی در بازه $[c, a]$ با تابع توزیع تجمعی $F(u)$ باشد، نامساوی $P(U \leq u) - \frac{a-E(U)}{a-c} \leq \frac{1}{4} + \frac{|u - \frac{c+a}{2}|}{a-c}$ یا

$$\left| P(U \geq u) - \frac{E(U) - c}{a - c} \right| \leq \frac{1}{4} + \frac{|u - \frac{c+a}{2}|}{a - c}. \quad (7)$$

برقرار است.

قضیه ۲. فرض کنید $\mathbf{X}_i = [X_{i1}, \dots, X_{in}]^T$ برداری تصادفی با امیدریاضی معلوم $E(\mathbf{X}_i) = [E(X_{i1}), \dots, E(X_{in})]^T$ باشد. شرط کافی برای برقراری محدودیت

$$Pr((\mathbf{w}_{y_i}^T \phi(\mathbf{X}_i) + b_{y_i}) - \frac{1}{k-1} \sum_{m \neq y_i} (\mathbf{w}_m^T \phi(\mathbf{X}_i) + b_m) \geq 1 - \xi_i) \geq \gamma,$$

عبارت است از

$$(\mathbf{w}_{y_i}^T E(\phi(\mathbf{X}_i)) + b_{y_i}) - \frac{1}{k-1} \sum_{m \neq y_i} (\mathbf{w}_m^T E(\phi(\mathbf{X}_i)) + b_m) \geq 1 - \xi_i + 2a\gamma.$$

برهان: فرض کنید $U_i = (\mathbf{w}_{y_i}^T \phi(\mathbf{X}_i) + b_{y_i}) - \frac{1}{k-1} \sum_{m \neq y_i} (\mathbf{w}_m^T \phi(\mathbf{X}_i) + b_m)$ ، برای یک مقدار معین u بطوریکه $u \leq U_i + \xi_i \leq U_i - \xi_i \leq -u$ باشد و با استفاده از رابطه (۷) و $c = -a$ ، نامساوی

$$-\frac{1}{2} - \frac{|u|}{2a} + \frac{E(U_i + \xi_i) + a}{2a} \leq Pr(U_i + \xi_i \geq u) \leq \frac{1}{2} + \frac{|u|}{2a} + \frac{E(U_i + \xi_i) + a}{2a}.$$

به دست می‌آید. با در نظر گرفتن $u = 1$ ، داریم:

$$-\frac{1}{2a} + \frac{E(U_i) + \xi_i}{2a} \leq Pr(U_i + \xi_i \geq 1) \leq 1 + \frac{1}{2a} + \frac{E(U_i) + \xi_i}{2a},$$

و در نهایت چون $Pr(U_i + \xi_i \geq 1) \geq \gamma$ داریم $-\frac{1}{2a} + \frac{E(U_i) + \xi_i}{2a} \geq \gamma$ یا $E(U_i) \geq 1 - \xi_i + 2a\gamma$. بنابراین محدودیت مسئله اولیه از حالت احتمالی خارج و به

$$(\mathbf{w}_{y_i}^T E(\phi(\mathbf{X}_i)) + b_{y_i}) - \frac{1}{k-1} \sum_{m \neq y_i} (\mathbf{w}_m^T E(\phi(\mathbf{X}_i)) + b_m) \geq 1 - \xi_i + 2a\gamma.$$

تبدیل می‌شود.

اکنون برای به دست آوردن تابع دوگان با استفاده از تابع لاگرانژ^۱، مسئله اولیه نسبت به متغیرهای $\mathbf{w}$ ، b و ξ_i و ضرایب لاگرانژ مینیمم می‌شود. برای به دست آوردن مقادیر بهینه متغیرها، تابع

$$\begin{aligned} L(\mathbf{w}, b, \xi, \alpha, \lambda) = & \frac{1}{\gamma} \sum_{m=1}^k (\mathbf{w}_m^T \mathbf{w}_m + b_m^2) + C \sum_{i=1}^{\ell} \xi_i - \sum_{i=1}^{\ell} \alpha_i \xi_i - \sum_{i=1}^{\ell} \lambda_i \xi_i \\ & - \sum_{i:m=y_i} \alpha_i (\mathbf{w}_{y_i}^T E(\phi(\mathbf{X}_i)) + b_{y_i}) + \frac{1}{k-1} \sum_{i:m \neq y_i} \alpha_i (\mathbf{w}_m^T E(\phi(\mathbf{X}_i)) + b_m) \\ & + \sum_{i=1}^{\ell} \alpha_i + \gamma a \sum_{i=1}^{\ell} \alpha_i \gamma. \end{aligned}$$

تعریف می‌شود. از مشتق L نسبت به $\mathbf{w}_m$ ، b_m و ξ شرایط لازم بهینگی به صورت

$$\begin{aligned} \mathbf{w}_m &= \sum_{i:y_i=m} \alpha_i E(\phi(\mathbf{X}_i)) - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i E(\phi(\mathbf{X}_i)), \\ b_m &= \sum_{i:y_i=m} \alpha_i - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i = 0, \\ C\mathbf{e} &= \boldsymbol{\alpha} + \boldsymbol{\lambda} = 0. \end{aligned} \quad (8)$$

به دست می‌آید. با جایگذاری روابط (۸) در تابع لاگرانژ برای $i = 1, \dots, \ell$ و $i \neq j$ و بسط تابع لاگرانژ داریم

$$\begin{aligned} L(\alpha) = & \frac{1}{\gamma} \sum_{i:y_i=m} \alpha_i^2 E^T(\phi(\mathbf{X}_i)) E(\phi(\mathbf{X}_j)) \\ & - \frac{1}{k-1} \sum_{i:y_i=m} \alpha_i \sum_{i:y_i \neq m} \alpha_i E^T(\phi(\mathbf{X}_i)) E(\phi(\mathbf{X}_j)) \\ & + \frac{1}{\gamma(k-1)^2} \sum_{i:y_i \neq m} \alpha_i^2 E^T(\phi(\mathbf{X}_i)) E(\phi(\mathbf{X}_j)) \\ & + \frac{1}{\gamma} \sum_{i:y_i=m} \alpha_i^2 - \frac{1}{k-1} \sum_{i:y_i=m} \alpha_i \sum_{i:y_i \neq m} \alpha_i + \frac{1}{\gamma(k-1)^2} \sum_{i:y_i \neq m} \alpha_i^2, \end{aligned}$$

¹Lagrangian function

$$\begin{aligned}
 & - \sum_{i:y_i=m} \alpha_i^\gamma E^T(\phi(\mathbf{X}_i))E(\phi(\mathbf{X}_j)) + \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i^\gamma E^T(\phi(\mathbf{X}_i))E(\phi(\mathbf{X}_j)) \\
 & - \sum_{i:y_i=m} \alpha_i^\gamma + \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i^\gamma + \frac{1}{k-1} \sum_{i:y_i=m} \alpha_i^\gamma E^T(\phi(\mathbf{X}_i))E(\phi(\mathbf{X}_j)) \\
 & - \frac{1}{(k-1)^\gamma} \sum_{i:y_i \neq m} \alpha_i^\gamma E^T(\phi(\mathbf{X}_i))E(\phi(\mathbf{X}_j)) + \frac{1}{k-1} \sum_{i:y_i=m} \alpha_i^\gamma \\
 & - \frac{1}{(k-1)^\gamma} \sum_{i:y_i \neq m} \alpha_i^\gamma + \sum_{i=1}^{\ell} \alpha_i + \gamma \sum_{i=1}^{\ell} \alpha_i \gamma. \tag{۹}
 \end{aligned}$$

در نهایت با خلاصه کردن رابطه (۹)، مسئله دوگان به صورت

$$\begin{aligned}
 \max & - \frac{1}{\gamma} \frac{k}{k-1} \sum_{i:y_i=m} \alpha_i^\gamma (E^T(\phi(\mathbf{X}_i))E(\phi(\mathbf{X}_j)) + \mathbf{e}) \\
 & - \frac{1}{\gamma} \frac{-k}{(k-1)^\gamma} \sum_{i:y_i \neq m} \alpha_i^\gamma (E^T(\phi(\mathbf{X}_i))E(\phi(\mathbf{X}_j)) + \mathbf{e}) + \sum_{i=1}^{\ell} \alpha_i (1 + \gamma a \gamma) \\
 \text{s.t.} \quad & \mathbf{0} \leq \boldsymbol{\alpha} \leq C\mathbf{e}. \tag{۱۰}
 \end{aligned}$$

تبدیل می‌شود، که در آن $\boldsymbol{\alpha} = [\alpha_1, \dots, \alpha_n]^T$ بردار ضرایب لاگرانژ است. همچنین برای (۶) داریم:

$$\arg \max_m f_m(\mathbf{x}) = \arg \max_m \sum_{i:y_i=m} \alpha K(E(\mathbf{X}_i), E(\mathbf{X}_j)). \tag{۱۱}$$

که در آن

$$K(E(\mathbf{X}_i), E(\mathbf{X}_j)) = \phi^T(E(\mathbf{X}_i))\phi(E(\mathbf{X}_j)). \tag{۱۲}$$

تابع هسته است. با تعریف ماتریس هسین^۱ به صورت

$$\mathbf{G} = \begin{cases} \frac{k}{k-1} K(E(\mathbf{X}_i), E(\mathbf{X}_j)); & y_i = m \\ \frac{-k}{(k-1)^\gamma} K(E(\mathbf{X}_i), E(\mathbf{X}_j)); & y_i \neq m \end{cases}$$

¹Hessian matrix

مسئله دوگان (۱۰)، به صورت

$$\begin{aligned} \max \quad & -\frac{1}{\gamma} \alpha^T \mathbf{G} \alpha + \mathbf{e}(\gamma + \gamma a) \alpha \\ \text{s.t.} \quad & \mathbf{0} \leq \alpha \leq C \mathbf{e}. \end{aligned} \quad (13)$$

تبدیل می‌شود. برای تابع تصمیم (۱۱) داریم

$$\begin{aligned} f(\mathbf{x}) &= \mathbf{w}_m^T E(\phi(\mathbf{X}_i)) + b_m \\ &= \left(\sum_{i:y_i=m} \alpha_i E(\phi(\mathbf{X}_i)) - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i E(\phi(\mathbf{X}_i)) \right) E(\phi(\mathbf{X}_j)) \\ &\quad + \sum_{i:y_i=m} \alpha_i - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i, \end{aligned}$$

که در نهایت:

$$\begin{aligned} f(\mathbf{x}) &= \sum_{i:y_i=m} \alpha_i K(E(\mathbf{X}_i), E(\mathbf{X}_j)) - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i K(E(\mathbf{X}_i), E(\mathbf{X}_j)) \\ &\quad + \sum_{i:y_i=m} \alpha_i - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i. \end{aligned} \quad (14)$$

چون میانگین جامعه نامعلوم است، باروش گشتاوری برآورد می‌شود. اگر نمونه‌های

$$\mathbf{x}_1 = [x_{11}, \dots, x_{1n}]^T, \mathbf{x}_2 = [x_{21}, \dots, x_{2n}]^T, \dots, \mathbf{x}_t = [x_{t1}, \dots, x_{tn}]^T.$$

وجود داشته باشد، بردار برآورد امیدریاضی $\bar{\mathbf{x}}_j = \left[\frac{1}{N_j} \sum_{t=1}^{N_j} x_{t1}, \dots, \frac{1}{N_j} \sum_{t=1}^{N_j} x_{tn} \right]$ است، که در آن N_j ها حجم نمونه‌های ممکن است. بنابراین تابع هسته (۱۲)، به $K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) = \phi^T(\bar{\mathbf{x}}_i) \phi(\bar{\mathbf{x}}_j)$ و بردار برآورد وزن‌ها در رابطه (۸) به $\hat{\mathbf{w}}_m = \sum_{i:y_i=m} \alpha_i \phi(\bar{\mathbf{x}}_i) - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i \phi(\bar{\mathbf{x}}_i)$ تبدیل می‌شود.

$$\begin{aligned} \max \quad & -\frac{1}{\gamma} \frac{k}{k-1} \sum_{\{i:y_i=m\}} \alpha_i (K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) + e) \\ & - \frac{1}{\gamma} \frac{-k}{(k-1)^2} \sum_{\{i:y_i \neq m\}} \alpha_i (K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) + e) + \sum_{i=1}^{\ell} \alpha_i (\gamma + \gamma a_i) \\ \text{s.t.} \quad & \mathbf{0} \leq \alpha \leq C \mathbf{e}. \end{aligned}$$

بیان می‌شود و در نهایت برای برآورد مدل (۱۳) و همچنین برآورد تابع تصمیم بهینه (۱۴) به ترتیب داریم

$$\begin{aligned} \max \quad & -\frac{1}{\gamma} \alpha^T \hat{G} \alpha + e(1 + \gamma a) \alpha \\ \text{s.t.} \quad & 0 \leq \alpha \leq C e. \end{aligned}$$

و

$$\begin{aligned} f(\mathbf{x}) = & \sum_{i:y_i=m} \alpha_i K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) \\ & + \sum_{i:y_i=m} \alpha_i - \frac{1}{k-1} \sum_{i:y_i \neq m} \alpha_i. \end{aligned}$$

۳ مطالعه شبیه‌سازی

فرآیند شبیه‌سازی با روش مونت کارلو^۱ در (موونی، ۱۹۹۷) انجام شده است. روش بازنمونه‌گیری بوت استرپ به کار برده شد و از هر کلاس نمونه‌برداری شد، سپس میانگین داده‌ها محاسبه شده و دقت مدل به دست آمد. همچنین از سه تابع هسته خطی، چندجمله‌ای و گاوسی

$$K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) = \bar{\mathbf{x}}_i^T \bar{\mathbf{x}}_j, K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) = (\bar{\mathbf{x}}_i^T \bar{\mathbf{x}}_j + 1)^p, K(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) = \exp(-\mu \|\bar{\mathbf{x}}_i - \bar{\mathbf{x}}_j\|^2)$$

استفاده شد. فرایند شبیه‌سازی مطابق با گام‌های زیر انجام می‌شود.

گام ۱: با استفاده از روش بازنمونه‌گیری بوت استرپ از هر کلاس نمونه‌برداری می‌شود.

گام ۲: میانگین نمونه‌ها محاسبه و به عنوان ورودی در مدل پیشنهادی در نظر گرفته می‌شود.

گام ۳: محاسبه معیار دقت^۲ به عنوان یکی از معیارهای ماتریس آشفتگی^۳ با استفاده از رابطه

$$ACC = \frac{TP + TN}{TP + TN + FP + FN}, \quad (15)$$

انجام می‌شود، که در آن TP ، TN ، FP و FN به ترتیب مثبت درست، منفی درست، مثبت کاذب و منفی کاذب هستند. بر اساس رابطه (۱۵) دقت مدل‌ها تا چهار رقم اعشار به دست آمد.

گام ۴: مقادیر پارامترها با $\mu = 0.5$ ، $p = 2$ ، $\gamma = 0.7$ ، $C = 1$ تنظیم شده است. دقت به دست آمده از هر

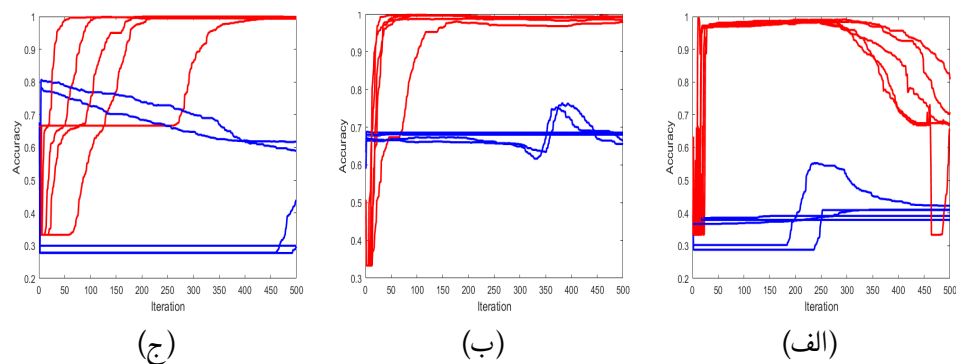
¹Monte Carlo

²Accuracy

³Confusion matrix

دو مدل در 50° تکرار و برای ۵ اجرا، در شکل ۱ برای سه تابع هسته خطی، چندجمله‌ای و گاوسی نشان داده شده است. رنگ قرمز، دقت مدل پیشنهادی و رنگ آبی، دقت مدل پایه را نشان می‌دهد. همان طور که دیده می‌شود تقریباً در تمام اجراها دقت مدل پیشنهادی برای هر سه تابع هسته در سطح بالاتری از مدل پایه قرار گرفته است.

گام ۵: برای ارزیابی سه مدل کرامر-سینگر، پایه و پیشنهادی، به ازای هر سه تابع هسته از آزمون فریدمن استفاده می‌شود. با توجه به نتایج این آزمون مشخص می‌شود که آیا تفاوتی بین مدل‌ها وجود دارد؟



شکل ۱. دقت مدل‌های پیشنهادی و پایه در داده‌های شبیه‌سازی شده برای تابع هسته الف-خطی، ب-چندجمله‌ای، ج-گاوسی.

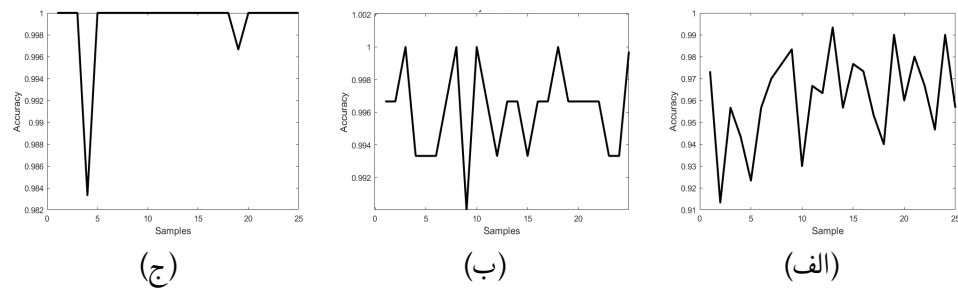
در جدول ۱ مقادیر مختلف γ در بازه $[0.1, 0.9]$ در نظر شده است و برای هر مقدار، ده تکرار برای هر دو مدل اجرا شده است و میانگین آن‌ها به دست آمده است. مقادیر سایر پارامترها با $C = 1, \mu = 0.5, p = 2$ تنظیم شده است. مشاهده می‌شود به ازای مقادیر کوچک γ ، مدل پیشنهادی دقت کم تری دارد و پیش‌بینی برچسب صحیح کلاس دارای خطای بیش تری است. بنابراین مدل ارائه شده به ازای مقادیر $\gamma \leq 0.5$ برتری قابل توجهی ندارد، اما به ازای مقادیر $\gamma > 0.5$ ، بهتر از مدل پایه عمل کرده است. مقدار آماره آزمون فریدمن برای مقایسه میانگین دقت‌های به دست آمده به ازای سه تابع هسته در جدول ۱ برابر 40.963 و p -مقدار کمتر از 0.1 بیانگر تفاوت بین دقت سه مدل است. در شکل ۲ نمودار دقت برای هر سه تابع هسته در مدل پیشنهادی و ۲۵ نمونه و به ازای مقادیر پارامترهای $\mu = [0.05, 0.09, 0.15, 0.5, 0.7, 0.9]$ و $C = [1, 5, 15, 25, 50]$ ، $\gamma = [0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9]$ ارائه شده است. بطور کلی با توجه به نمودارها در تابع هسته خطی (آ) کوچکترین مقدار دقت تقریباً 0.91 و بزرگ ترین آن بیش از 0.99 است. در تابع هسته چندجمله‌ای (ب) کوچک ترین مقدار دقت برابر 0.99 و بزرگ ترین آن 1 است. همچنین در تابع هسته گاوسی (ج) کم ترین مقدار دقت تقریباً 0.98 و بیشترین آن 1 است که نشان می‌دهد مدل با هسته گاوسی بهترین عملکرد را دارد.

جدول ۱. جدول مقادیر معیارهای آماری دقت داده‌های شبیه‌سازی در هر دو مدل به ازای سه تابع هسته.

پارامتر	مدل	تابع خطی	تابع چندجمله‌ای	تابع گاوسی		
γ		میانگین	انحراف معیار	میانگین	انحراف معیار	
۰/۱	کرامر-سینگر	۰/۷۸۲۴	۰/۱۳۲	۰/۹۱۱۸	۰/۲۳۱	۰/۲۱۶
۰/۱	پایه	۰/۷۹۲۱	۰/۱۲۴	۰/۸۹۸۷	۰/۲۱۲	۰/۱۲۷
۰/۱	پیشنهادی	۰/۸۱۳۲	۰/۲۴۱	۰/۹۵۱۳	۰/۳۰۱	۰/۳۱۸
۰/۲	کرامر-سینگر	۰/۸۲۲۱	۰/۲۲۹	۰/۹۰۱۳	۰/۳۱۷	۰/۳۰۹
۰/۲	پایه	۰/۸۱۴۵	۰/۲۲۱	۰/۹۱۱۱	۰/۲۰۱	۰/۲۰۱
۰/۲	پیشنهادی	۰/۸۳۰۴	۰/۲۳۸	۰/۹۴۵۴	۰/۲۶۲	۰/۳۱۵
۰/۳	کرامر-سینگر	۰/۸۱۱۹	۰/۲۶۴	۰/۹۰۲۹	۰/۲۸۱	۰/۲۷۵
۰/۳	پایه	۰/۸۰۲۲	۰/۱۸۳	۰/۹۲۲۵	۰/۲۱۱	۰/۱۹۷
۰/۳	پیشنهادی	۰/۸۴۲۵	۰/۲۳۷	۰/۹۵۰۱	۰/۲۲۵	۰/۲۱۴
۰/۴	کرامر-سینگر	۰/۷۹۲۵	۰/۱۹۸	۰/۸۶۴۳	۰/۲۲۵	۰/۱۵۱
۰/۴	پایه	۰/۸۱۳۳	۰/۲۶۸	۰/۸۸۸۶	۰/۲۳۱	۰/۳۱۷
۰/۴	پیشنهادی	۰/۸۵۱۱	۰/۳۲۱	۰/۹۶۳۶	۰/۲۰۵	۰/۲۱۱
۰/۵	کرامر-سینگر	۰/۸۵۱۱	۰/۲۸۶	۰/۹۴۱۶	۰/۱۹۷	۰/۳۰۲
۰/۵	پایه	۰/۸۳۲۵	۰/۲۴۱	۰/۹۳۲۱	۰/۳۱۹	۰/۲۵۷
۰/۵	پیشنهادی	۰/۸۹۳۱	۰/۳۹۷	۰/۹۸۲۵	۰/۲۱۲	۰/۱۲۹
۰/۶	کرامر-سینگر	۰/۸۲۹۰	۰/۲۶۹	۰/۹۳۲۱	۰/۳۱۱	۰/۱۴۸
۰/۶	پایه	۰/۷۹۶۱	۰/۲۵۴	۰/۹۰۹۱	۰/۳۷۱	۰/۲۳۵
۰/۶	پیشنهادی	۰/۹۱۴۸	۰/۳۸۵	۰/۹۹۱۶	۰/۱۹۸	۰/۱۱۹
۰/۷	کرامر-سینگر	۰/۷۷۳۶	۰/۲۲۱	۰/۹۳۶۷	۰/۲۸۹	۰/۲۷۱
۰/۷	پایه	۰/۷۹۴۴	۰/۱۳۴	۰/۹۱۱۱	۰/۲۰۸	۰/۳۲۹
۰/۷	پیشنهادی	۰/۹۲۳۳	۰/۱۸۷	۰/۹۹۷۶	۰/۱۰۳	۰/۱۵۳
۰/۸	کرامر-سینگر	۰/۸۲۱۹	۰/۱۹۲	۰/۹۱۲۰	۰/۲۴۶	۰/۱۲۵
۰/۸	پایه	۰/۸۴۲۱	۰/۲۲۷	۰/۹۱۷۵	۰/۲۴۵	۰/۲۷۸
۰/۸	پیشنهادی	۰/۹۲۷۴	۰/۱۹۶	۰/۹۹۸۹	۰/۱۰۸	۰/۱۱۱
۰/۹	کرامر-سینگر	۰/۸۰۳۸	۰/۱۲۴	۰/۹۱۶۶	۰/۲۷۸	۰/۱۲۶
۰/۹	پایه	۰/۸۲۹۷	۰/۲۱۳	۰/۹۲۲۵	۰/۲۲۹	۰/۲۶۹
۰/۹	پیشنهادی	۰/۹۱۲۶	۰/۱۷۴	۰/۹۹۹۲	۰/۱۱۷	۰/۱۰۹

۴ مثال کاربردی

در این بخش با استفاده از مجموعه داده های دانه‌ها و اکولی در **دوآ و گزف (۲۰۱۹)** و همچنین مجموعه داده سرطان خون در <http://csse.szu.edu.cn/staff/zhuzx/Datasets.html>، دقت هر سه مدل مقایسه



شکل ۲. نمودار دقت به ازای ۲۵ نمونه برای سه تابع هسته الف- خطی، ب- چندجمله‌ای، ج- گاوسی

شده است.

(آ) جدول ۲ نتایج دقت برای مجموعه داده دانه‌ها به ازای طیف مختلف پارامترهای C ، μ ، γ و $p = 2$ است. در توابع هسته خطی و چندجمله‌ای مدل پایه به علت مستقل بودن از پارامتر γ تغییر زیادی دیده نمی‌شود، اما این تغییرات بر روی مدل پیشنهادی به علت وابستگی به γ مشهود است. به ازای برخی مقادیر کوچک γ ، دقت مدل پیشنهادی کم تر از خود در مقایسه با مقادیر بزرگ γ است که می‌تواند به علت فرایند نمونه‌گیری باشد. همان طور که مشاهده می‌شود در تابع هسته گاوسی به ازای مقادیر مختلف پارامترهای μ و γ ، دقت مدل پیشنهادی بهتر از مدل پایه است. برای مقادیر بزرگ γ ، مدل ارائه شده دقت بهتری نسبت به مدل پایه دارد. پس به طور کلی الگوی پیشنهادی، عملکرد بهتری نسبت به الگوی پایه و نیز نسبت به خودش در مقایسه با مقادیر کوچک γ به ازای هر سه تابع هسته دارد.

(ب) جدول ۳ مقادیر دقت را برای مجموعه داده اکولی نشان می‌دهد که بر مبنای پارامترهای C ، μ ، γ و $p = 2$ است. در تابع هسته خطی، مدل پیشنهادی نسبت به مدل پایه تقریباً عملکرد بهتری دارد. به ازای تابع هسته چندجمله‌ای، مدل پیشنهادی به طور کلی دقت بیش تری نسبت به مدل پایه دارد. برای تابع هسته گاوسی، همان طور که از جدول مشهود است، مقادیر دقت مدل پیشنهادی هم از مدل پایه و هم از دو تابع هسته خطی و چندجمله‌ای بهتر است. بنابراین می‌توان گفت تابع هسته گاوسی برای مدل پیشنهادی بهتر از مدل پایه است.

(ج) جدول ۴ مقادیر دقت را برای مجموعه داده سرطان خون به ازای پارامترهای C ، μ ، γ و $p = 2$ نشان می‌دهد. به ازای مقادیر کوچک γ و تابع هسته خطی، مدل پیشنهادی عملکرد ضعیف تری نسبت به مدل پایه دارد، اما برای مقادیر بزرگ تر پارامتر γ ، مدل پیشنهادی دقت بهتری دارد. برای تابع هسته چندجمله‌ای، مدل پیشنهادی بطور کلی از مدل پایه عملکرد بهتری دارد. به ازای تابع هسته گاوسی نیز، مدل پیشنهادی عملکرد بسیار بهتری نسبت به مدل پایه دارد.

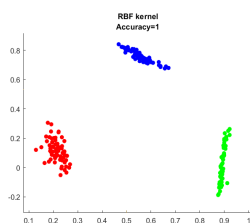
(ح) شکل ۳ برای نحوه طبقه‌بندی مجموعه داده دانه‌ها به ترتیب برای توابع هسته خطی، چندجمله‌ای و گاوسی با پارامترهای $C = 1$ ، $\mu = 0.1$ ، $\gamma = 0.7$ و $p = 2$ رسم شده است. برای توابع هسته خطی و گاوسی دقت کامل به دست آمده است و به ازای هسته چندجمله‌ای، دقت برابر ۰.۹۹۳۴ است.

(د) شکل ۴ متعلق به دسته‌بندی مجموعه داده اکولی است. مقادیر پارامترها $C = 1$, $\mu = 0.5$, $\gamma = 0.5$ و $p = 3$ هستند. برای تابع هسته خطی، دقت مدل برابر 0.9983 به دست آمده است. مقادیر دقت برای هسته‌های چندجمله‌ای و گاوسی به ترتیب 0.9846 و 0.9775 به دست آمده است.

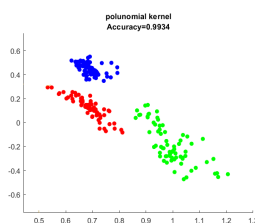
(ر) شکل ۵ مربوط به دسته‌بندی مجموعه داده سرطان خون برای توابع هسته خطی، چندجمله‌ای و گاوسی با پارامترهای $\mu = 0.7$, $\gamma = 0.65$ و $p = 3$ است. مقادیر دقت مدل برای توابع هسته خطی، چندجمله‌ای و گاوسی به ترتیب برابر 0.9972 , 0.9989 و 1 به دست آمده است.

جدول ۲. جدول مقادیر دقت مجموعه داده دانه‌ها برای هر دو مدل به ازای سه تابع هسته.

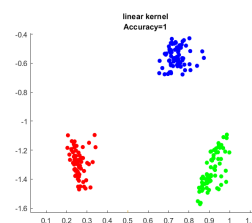
پارامترها			تابع خطی		تابع چندجمله‌ای		تابع گاوسی	
C	μ	γ	پایه	پیشنهادی	پایه	پیشنهادی	پایه	پیشنهادی
۵	۰/۰۵	۰/۳	۰/۹۵۴۲	۰/۹۷۱۶	۰/۹۸۷۳	۰/۹۷۵۳	۰/۹۷۱۰	۰/۹۶۴۴
۱۰	۰/۰۹	۰/۵	۰/۹۵۹۱	۰/۹۷۷۸	۰/۹۸۵۲	۰/۹۷۳۶	۰/۹۸۰۸	۰/۹۹۲۵
۲۵	۰/۱	۰/۶	۰/۹۶۱۷	۰/۹۸۶۸	۰/۹۷۸۳	۰/۹۹۱۹	۰/۹۸۶۷	۰/۹۹۴۷
۱	۰/۱۵	۰/۷	۰/۹۵۲۲	۰/۹۷۴۳	۰/۹۸۰۹	۰/۹۸۳۸	۰/۹۸۰۱	۰/۹۹۵۸
۱	۰/۲	۰/۸	۰/۹۷۲۹	۰/۹۸۹۱	۰/۹۸۹۴	۱	۰/۹۵۳۶	۰/۹۷۴۸
۱۵	۰/۷	۰/۹	۰/۹۸۱۴	۱	۰/۹۸۷۸	۰/۹۹۷۳	۰/۹۷۱۵	۰/۹۸۹۱
۵	۰/۰۱	۰/۸	۱	۰/۹۹۳۸	۰/۹۷۶۳	۰/۹۹۲۵	۰/۹۹۰۱	۱
۵۰	۰/۹	۰/۷	۰/۹۸۸۱	۰/۹۹۵۴	۰/۹۶۱۲	۱	۰/۹۸۲۳	۱
۳۰	۱/۷	۰/۴	۰/۹۹۵۱	۰/۹۸۶۱	۰/۹۷۰۴	۰/۹۸۰۴	۰/۹۶۵۷	۰/۹۷۱۳
۱۰	۲/۵	۰/۹	۰/۹۶۴۸	۰/۹۷۹۱	۰/۹۸۵۶	۱	۰/۹۳۷۱	۰/۹۵۱۵



(ج)



(ب)



(الف)

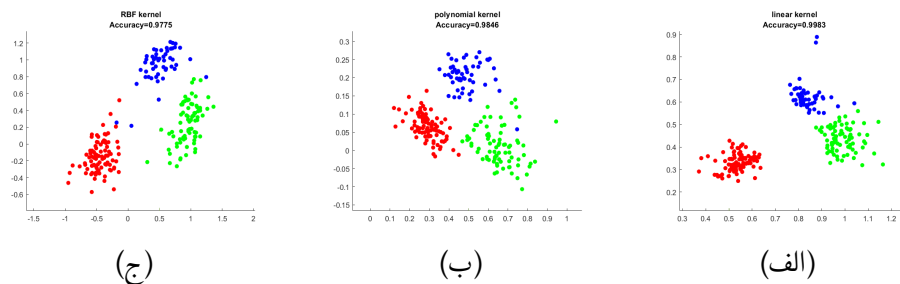
شکل ۳. دسته‌بندی کلاس‌ها برای مجموعه داده دانه‌ها. الف- خطی، ب- چندجمله‌ای، ج- گاوسی

جدول ۳. جدول مقادیر دقت مجموعه داده اکولی برای هر دو مدل به ازای سه تابع هسته.

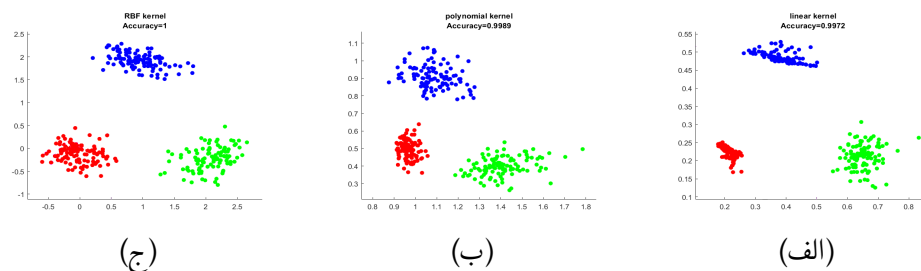
پارامترها		تابع خطی		تابع چندجمله‌ای		تابع گاوسی	
C	μ	γ	پایه	پیشنهادی	پایه	پیشنهادی	پیشنهادی
۵	۰/۰۵	۰/۳	۰/۹۲۱۱	۰/۹۳۶۴	۰/۹۷۷۵	۰/۹۶۳۲	۰/۹۷۳۳
۱۰	۰/۰۹	۰/۵	۰/۹۴۵۱	۰/۹۵۳۶	۰/۹۸۲۵	۰/۹۷۱۲	۰/۹۷۱۰
۲۵	۰/۱	۰/۶	۰/۹۲۴۱	۰/۹۶۱۷	۰/۹۶۱۸	۰/۹۸۳۱	۰/۹۸۲۱
۱	۰/۱۵	۰/۷	۰/۹۴۳۶	۰/۹۸۱۹	۰/۹۷۳۹	۰/۹۹۵۸	۱
۱	۰/۲	۰/۸	۰/۹۸۳۳	۰/۹۹۴۵	۰/۹۷۳۴	۱	۰/۹۹۷۱
۱۵	۰/۷	۰/۹	۰/۹۹۳۱	۱	۰/۹۸۲۴	۰/۹۸۷۸	۱
۵	۰/۰۱	۰/۸	۱	۰/۹۹۱۲	۰/۹۸۲۶	۱	۰/۹۹۵۴
۵۰	۰/۹	۰/۷	۰/۹۹۰۱	۰/۹۹۵۶	۱	۰/۹۸۹۹	۰/۹۹۶۸
۳۰	۱/۷	۰/۴	۰/۹۸۱۱	۰/۹۹۲۳	۰/۹۸۷۷	۰/۹۹۶۴	۰/۹۸۸۵
۱۰	۲/۵	۰/۹	۰/۹۸۱۹	۰/۹۹۵۷	۰/۹۸۱۶	۰/۹۹۶۵	۰/۹۶۵۲

جدول ۴. جدول مقادیر دقت مجموعه داده سرطان خون برای هر دو مدل به ازای سه تابع هسته.

پارامترها		تابع خطی		تابع چندجمله‌ای		تابع گاوسی	
C	μ	γ	پایه	پیشنهادی	پایه	پیشنهادی	پیشنهادی
۵	۰/۰۵	۰/۳	۰/۹۷۱۶	۰/۹۵۹۹	۰/۹۸۱۲	۰/۹۷۲۳	۰/۹۶۹۵
۱۰	۰/۰۹	۰/۵	۰/۹۷۴۲	۰/۹۵۲۰	۰/۹۷۱۷	۰/۹۶۱۹	۰/۹۸۱۱
۲۵	۰/۱	۰/۶	۰/۹۹۰۱	۰/۹۷۹۸	۰/۹۸۱۳	۰/۹۸۱۳	۰/۹۹۱۸
۱	۰/۱۵	۰/۷	۰/۹۶۱۸	۱	۰/۹۸۵۹	۰/۹۹۸۱	۰/۹۹۰۱
۱	۰/۲	۰/۸	۰/۹۷۷۶	۰/۹۸۳۲	۰/۹۷۷۹	۱	۰/۹۹۱۳
۱۵	۰/۷	۰/۹	۰/۹۹۶۶	۰/۹۹۲۵	۰/۹۸۹۲	۱	۱
۵	۰/۰۱	۰/۸	۰/۹۹۰۳	۰/۹۹۵۱	۰/۹۸۹۱	۰/۹۹۹۰	۰/۹۹۲۸
۵۰	۰/۹	۰/۷	۰/۹۷۶۳	۱	۰/۹۸۹۹	۰/۹۹۹۳	۱
۳۰	۱/۷	۰/۴	۰/۹۹۴۶	۰/۹۸۱۱	۰/۹۸۰۱	۰/۹۹۳۵	۰/۹۹۶۱
۱۰	۲/۵	۰/۹	۰/۹۸۸۸	۰/۹۹۸۶	۰/۹۹۱۴	۰/۹۹۸۹	۰/۹۸۱۵



شکل ۴. دسته‌بندی کلاس‌ها برای مجموعه داده اکولی. الف- خطی، ب- چندجمله‌ای، ج- گاوسی



شکل ۵. دسته‌بندی کلاس‌ها برای مجموعه داده سرطان خون. الف- خطی، ب- چندجمله‌ای، ج- گاوسی

۵ بحث و نتیجه‌گیری

در دنیای واقعی اغلب با داده‌هایی سروکار داریم که به علت خطای ناشی از اندازه‌گیری یا نمونه‌گیری یا مدل‌سازی دقیق نیستند. در نتیجه نیاز به ارائه مدل‌هایی بر پایه الگوی غیرقطعی است. در این مقاله مدلی برای ورودی‌های غیرقطعی پیشنهاد شده است. در الگوی پیشنهادی با اضافه کردن مقداری اربیبی در عبارت تنظیم کننده، تابع هدف اولیه به یک مسئله برنامه‌ریزی درجه دوم تبدیل گردید و اضافه کردن مقدار اربیبی در محدودیت آن، سبب راحتی فرایند دسته‌بندی کلاس‌ها شده است. در این رویکرد از میانگین نمونه‌ای به عنوان برآوردی مناسب برای امیدریاضی استفاده شد. با توجه به نتایج بخش‌های ۳ و ۴، کارایی بهتر مدل با ورودی‌های غیرقطعی نسبت به مدل با ورودی‌های دقیق نشان داده شد. ضعف مدل پیشنهاد شده، سرعت آن در هنگام پردازش داده‌های با حجم بالا است. همچنین تأثیرپذیری میانگین از داده‌های پرت از دیگر ضعف‌های مدل پیشنهادی است که می‌تواند در مطالعات آینده مدنظر قرار گیرد.

۶ تقدیر و تشکر

نویسندگان این مقاله صمیمانه از داوران، سردبیر و ویراستار محترم به خاطر بازخورد ارزشمند و نظرات سازنده تشکر می‌کنند. پیشنهادات آن‌ها در ارتقای کیفیت نسخه نهایی این مقاله بسیار مؤثر بوده است. تحقیق نویسنده مسئول

از حمایت مالی دانشگاه فردوسی مشهد (۳/۶۰۷۵۴) برخوردار بوده است.

مراجع

موسوی، ن. و گلعلی‌زاده م. (۱۴۰۲)، رویکردی نوین در بکارگیری روش دسته ماشین بردار پشتیبان تصادفی در تحلیل داده‌های بیان ژن سرطان پروستات، *مجله علوم آماری*، ۱۷(۲)، ۴۵۹-۴۷۶.

Abaszade, M., and Effati, S. (2018), Stochastic Support Vector Regression with Probabilistic Constraints, *Applied Intelligence*, **48**, 243-256.

Abaszade, M., and Effati, S. (2019), A New Method for Classifying Random Variables Based on Support Vector Machine, *Journal of Classification*, **36**, 152-174.

Al-Ma'shumah, F., Razmkhah, M., and Effati, S. (2022), Expectation-Based and Quantile-Based Probabilistic Support Vector Machine Classification for Histogram-Valued Data, *International Journal on Electrical Engineering and Informatics*, **14**(1).

Barnett, N. S., Dragomir, S. S., and Agarwal, R. P. (2002), Some Inequalities for Probability, Expectation, and Variance of Random Variables Defined over a Finite Interval, *Computers & Mathematics with Applications*, **43**(10-11), 1319-1357. Pergamon.

Benítez-Peña, S., Blanquero, R., Carrizosa, E., and Ramírez-Cobo, P. (2024), Cost-sensitive Probabilistic Predictions for Support Vector Machines, *European Journal of Operational Research*, **314**(1), 268-279.

Bosch, P., López, J., Ramírez, H., and Robotham, H. (2013), Support Vector Machine under Uncertainty: An Application for Hydroacoustic Classification of Fish-schools in Chile, *Expert Systems with Applications*, **40**(10), 4029-4034.

Chen, W. J., Shao, Y. H., Li, C. N., and Deng, N. Y. (2016), MLTSVM: A Novel Twin Support Vector Machine to Multi-label Learning, *Pattern Recognition*, **52**, 61-74.

- Cortes, C., and Vapnik, V. (1995), Support Vector Networks, *Machine Learning*, **20**, 273-297.
- Crammer, K., and Singer, Y. (2001), On the Algorithmic Implementation of Multi-class Kernel-based Vector Machines, *Journal of Machine Learning Research*, **2**, 265-292.
- Don, D. R., and Iacob, I. E. (2020), DCSVM: Fast Multi-class Classification using Support Vector Machines, *International Journal of Machine Learning and Cybernetics*, **11**(2), 433-447.
- Dua, D., and Graff, C. (2019), UCI Machine Learning Repository. *University of California, Irvine, School of Information and Computer Sciences*.
- Gonen, M., Tanugur, A. G., and Alpaydin, E. (2008), Multi-class Posterior Probability Support Vector Machines, *IEEE Transactions on Neural Networks*, **19**(1), 130-139.
- Hastie, T., and Tibshirani, R. (1997), Classification by Pairwise Coupling, *The Annals of Statistics*, **26**(2), 451-471.
- He, X., Wang, Z., Jin, C., Zheng, Y., and Xue, X. (2012), A Simplified Multi-class Support Vector Machine with Reduced Dual Optimization, *Pattern Recognition Letters*, **33**(1), 71-82.
- Knerr, S., Personnaz, L., and Dreyfus, G. (1990), Single-layer Learning Revisited: A Stepwise Procedure for Building and Training a Neural Network, *Neurocomputing: Algorithms, Architectures and Applications*, 41-50.
- Li, H. X., Yang, J. L., Zhang, G., and Fan, B. (2013), Probabilistic Support Vector Machines for Classification of Noise Affected Data, *Information Sciences*, **221**, 60-71.
- Madjarov, G., Gjorgjevikj, D., and Chorbev, I. (2009), A Multi-class SVM Classifier Utilizing Binary Decision Tree, *Informatica*, **33**, 225-233.

- Mooney, C. Z. (1997), Monte Carlo Simulation (No. 116), *Sage*.
- Moosaei, H., and Hladik, M. (2022), Least Squares Approach to K-SVCR Multi-class Classification with Its Applications, *Annals of Mathematics and Artificial Intelligence*, **90**(7), 873-892.
- Platt, J., Cristianini, N., and Shawe-Taylor, J. (1999), Large Margin DAGs for Multi-class Classification, *NIPS'99: Proceedings of the 13th International Conference on Neural Information Processing Systems*, **12**, 547-553..
- Sahleh, A., Salahi, M., and Eskandari, S. (2023), Multi-class Nonparallel Support Vector Machine, *Progress in Artificial Intelligence*, **12**(4), 349-361.
- Schiilkop, P. B., Burgest, C., and Vapnik, V. (1995), Extracting Support Data for a Given Task, *Proceedings of the First International Conference on Knowledge Discovery and Data Mining*. AAAI Press, Menlo Park, CA, 252-257.
- Vapnik, V. (1995), The Nature of Statistical Learning Theory, *NY: Springer-Verlag*.
- Vapnik, V. (1999), An Overview of Statistical Learning Theory, *IEEE Transactions on Neural Networks*, (**10**)5, 988-999.
- Weston, J., and Watkins, C. (1999), Support Vector Machines for Multi-class Pattern Recognition. *In Esann*, **99**, 219-224.
- Yazdi, H. S., Effati, S., and Saberi, Z. (2007), The Probabilistic Constraints in the Support Vector Machine *Applied Mathematics and Computation*, **194**(2), 467-479.