





دانشکده علوم ریاضی

رساله دکتری کنترل و بهینه سازی

حل رده ای از مسائل ماشین بردار پشتیبان با استفاده از روش های بهینه سازی دینامیکی

نگارنده: امیر فیضی

استاد راهنما
دکتر علیرضا ناظمی

استاد مشاور
دکتر محمدرضا ربیعی

بهمن ۱۳۹۹

شماره: ۴۵۵-۲۱۲-۲
تاریخ: ۱۴۰۰/۲/۷
ویرایش:

باسمه تعالی



پیوست شماره ۲

دانشکده: ریاضی

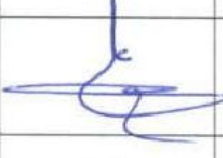
گروه: آمار

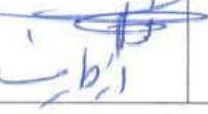
رساله دکتری: آقای امیر فیضی

تحت عنوان: حل رده ای از مسائل ماشین بردار پشتیبان با استفاده از روش های بهینه سازی دینامیکی

در تاریخ ۱۳۹۹/۱۱/۱۹ توسط کمیته تخصصی زیر جهت اخذ مدرک رساله دکتری ارزیابی گردید و با

درجه عالی مورد پذیرش قرار گرفت.

امضاء	اساتید مشاور	امضاء	اساتید راهنما
	نام و نام خانوادگی: دکتر محمدرضا ربیعی		نام و نام خانوادگی: دکتر علیرضا ناظمی
-----	نام و نام خانوادگی: -----	-----	نام و نام خانوادگی: -----

امضاء	نماینده تحصیلات تکمیلی	امضاء	اساتید داور
	نام و نام خانوادگی: دکتر سیدحیدر جعفری		نام و نام خانوادگی: دکتر جعفر فتحعلی
			نام و نام خانوادگی: دکتر مهرداد غزنوی
			نام و نام خانوادگی: دکتر مهدی روزبه



تقدیم بہ

پیشکامہ مقدس حضرت فاطمۃ الزہراء
سلام اللہ علیہا

بمسر مہربانم
و دختران عزیزم رقیہ و فاطمہ

پاس‌گزاری...

من لم يشكر المخلوق لم يشكر الخالق

حمد و ستایش خداوند سبحان را، که مرا توفیق عنایت فرمود تا در راه کسب علم و دانش گام بردارم. هم‌چنین درود و سلام خداوند به پیشگاه مقدس ائمه اطهار (علیهم السلام) خاصه ولی نعمت‌مان آقا علی بن موسی‌الرضا (علیه السلام)، که در تمامی عمر، به‌ویژه دوران تحصیل در مجاورت آن امام همام، از الطاف کریمانه ایشان برخوردار بوده‌ام.

و به پاس احترام به حرمت دانش بر خود لازم می‌دانم که از زحمات و راهنمایی‌های استاد گرانقدر و ارجمندم

جناب آقای دکتر علیرضا ناظمی

که نظارت این رساله را بر عهده داشته و در تمام مراحل از یاری و مساعدت ایشان برخوردار بوده‌ام، تشکر و قدرانی نمایم و موفقیت ایشان را از خدای متعال خواهانم.

از

جناب آقای دکتر محمدرضا ربیعی

که در امر مشاوره این رساله مساعدت نمودند و در این امر نهایت مراقبت، توجه و دقت خود را مبذول فرموده اند کمال تشکر و امتنان را دارم و برای ایشان از خداوند سلامت و سعادت ابدی را خواهانم.

هم‌چنین از اساتید فرهیخته جناب آقای دکتر جعفر فتحعلی، جناب آقای دکتر مهدی روزبه و جناب آقای دکتر مهرداد غزنوی که زحمت داوری این رساله را به عهده گرفتند با تمام وجود تشکر و قدردانی می‌نمایم.

امیر فیضی
بهمن ۱۳۹۹

تعهد نامه

اینجانب **امیر فیضی** دانشجوی دکتری رشته **ریاضی کاربردی علوم ریاضی** دانشگاه شاهرود، نویسنده پایان نامه با عنوان **حل رده ای از مسائل ماشین بردار پشتیبان با استفاده از روش های بهینه سازی دینامیکی**، تحت راهنمایی **علیرضا ناظمی** متعهد می شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش های دیگر پژوهش گران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان نامه، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا ارایه نشده است.
- حقوق معنوی این اثر، به دانشگاه صنعتی شاهرود تعلق دارد، و مقالات مستخرج با نام “دانشگاه صنعتی شاهرود” یا “Shahrood University of Technology” به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده اند، در مقالات مستخرج از پایان نامه رعایت می گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

امیر فیضی

بهمن ۱۳۹۹

مالکیت نتایج و حق نشر

- تمام حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه های رایانه ای، نرم افزارها و تجهیزات ساخته شده) متعلق به دانشگاه صنعتی شاهرود می باشد. این مطلب باید به نحو مقتضی، در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در این پایان نامه بدون ذکر منبع مجاز نمی باشد.

چکیده

ماشین های بردار پشتیبان ابزارهای قدرتمندی برای طبقه بندی داده ها و رگرسیون هستند. در سال های اخیر، الگوریتم های سریع بسیاری برای ماشین های بردار پشتیبان توسعه یافته اند. از ماشین بردار پشتیبان در بسیاری از کاربردهای نظامی و مهندسی، مسائل زمان واقعی، مانند طبقه بندی در محیط الکترومغناطیسی پیچیده، تشخیص در پزشکی و ... استفاده می شود. با توجه به این که روش های همگرایی عددی نیاز به محاسبات بیشتری داشته و نمی تواند نیاز ما در مسائل زمان واقعی را برآورده سازد، یک روند میسر و امیدوار کننده برای آموزش ماشین های بردار پشتیبان در زمان های واقعی بکارگیری شبکه های عصبی بازگشتی بر پایه اجرای دوره ای است.

در این رساله ما با معرفی شبکه عصبی به حل مسائلی از بردار پشتیبان مانند مساله رگرسیون بردار پشتیبان و نیز مساله رگرسیون بردار پشتیبان تصادفی با قیدهای احتمالی توسط شبکه عصبی می پردازیم. ثابت می کنیم نقطه تعادل شبکه عصبی با جواب بهینه مساله بردار پشتیبان معادل هستند. همچنین پایداری و همگرایی شبکه عصبی ارائه شده اثبات خواهد شد. در نهایت با ارائه مثال هایی نتایج نظری را تایید می کنیم.

کلمات کلیدی: بردار پشتیبان، شبکه عصبی، بهینه سازی، رگرسیون، متغیر تصادفی

لیست مقالات مستخرج از پایان نامه

۱. مقاله اول:

Feizi, Amir and Nazemi, Alireza. An application of a practical neural network model for solving support vector regression problems. *Intelligent Data Analysis*, 21:1443–1461, 2017.

۲. مقاله دوم:

Feizi, Amir, Nazemi, Alireza, and Mohammad Reza, Rabeie. Solving the stochastic support vector regression with probabilistic constraints by a high performance neural network model. *Engineering with Computers*, 2020.

فهرست مطالب

ف	فهرست تصاویر
ش	فهرست جداول
ث	پیشگفتار
۱	۱ ابرصفحه بهینه برای جداسازی نمونه‌ها و معرفی بردار پشتیبان
۱	۱.۱ شناسایی الگو و مفاهیم مقدماتی
۲	۱.۱.۱ نظریه یادگیری به وسیله سرپرست
۳	۲.۱.۱ مدل سازی مساله آموزش
۳	۳.۱.۱ رده بندی کننده خطی
۹	۲.۱ ماشین بردار پشتیبان
۱۱	۱.۲.۱ ابرصفحه بهینه برای نمونه های جدایی پذیر خطی
۱۵	۲.۲.۱ ابرصفحه بهینه برای نمونه های جدایی ناپذیر
۱۸	۳.۲.۱ رده بندی کننده های غیرخطی
۱۹	۴.۲.۱ هسته ضرب داخلی
۲۳	۲ شبکه عصبی
۲۳	۱.۲ مقدمه ای بر شبکه های عصبی
۲۳	۱.۱.۲ شبکه های عصبی بیولوژیکی
۲۴	۲.۱.۲ شبکه های عصبی مصنوعی
۲۸	۳.۱.۲ تاریخچه تکامل شبکه عصبی
۳۰	۴.۱.۲ تاریخچه حل مسائل بهینه سازی با استفاده از شبکه های عصبی
۳۳	۳ حل مسئله رگرسیون بردار پشتیبان با استفاده شبکه عصبی
۳۳	۱.۳ رگرسیون چیست؟
۳۴	۱.۱.۳ نمودار پراکندگی

۳۴	متغیرها و خطا	۲.۱.۳
۳۵	روش‌های رگرسیونی	۳.۱.۳
۳۵	استفاده از SVM در رگرسیون	۲.۳
۳۵	مساله رگرسیون خطی بردار پشتیبان	۱.۲.۳
۳۹	مساله رگرسیون غیرخطی بردار پشتیبان	۲.۲.۳
۴۱	معرفی شبکه عصبی	۳.۳
۴۳	مقایسه با برخی شبکه‌های عصبی موجود	۱.۳.۳
۴۷	تحلیل پایداری و همگرایی شبکه عصبی	۲.۳.۳
۵۳	مثال‌های عددی	۴.۳
۶۳	حل مسائل رگرسیون بردار پشتیبان تصادفی با قیده‌های احتمالی	۴
۶۳	مسائل رگرسیون بردار پشتیبان تصادفی با قیده‌های احتمالی	۱.۴
	مسائل رگرسیون بردار پشتیبان تصادفی با قیده‌های احتمالی در	۱.۱.۴
۶۳	حالت خطی	
	حل مسائل رگرسیون بردار پشتیبان تصادفی با قیده‌های احتمالی	۲.۱.۴
۶۶	در حالت غیرخطی	
۶۸	حل مساله رگرسیون به کمک شبکه عصبی با قیده‌های احتمالی	۳.۱.۴
۶۹	معرفی شبکه عصبی	۲.۴
۷۲	تجزیه و تحلیل پایداری شبکه عصبی	۳.۴
۷۶	مقایسه با برخی شبکه‌های عصبی موجود	۱.۳.۴
۷۸	مثال‌های عددی	۴.۴
۹۵	مراجع	
۱۰۱	واژه‌نامه فارسی به انگلیسی	
۱۰۵	واژه‌نامه انگلیسی به فارسی	
۱۰۹	پیوست	
۱۰۹	مقدمه‌ای بر بهینه‌سازی غیرخطی	۱.
۱۱۲	سیستم‌های دینامیکی	۲.
۱۱۷	نامساوی‌های وردشی	۳.

فهرست تصاویر

۴	۱.۱	کدام جداسازی بهتر است؟ حاشیه تقریباً صفر در A و یا حاشیه بزرگ در B ؟
	۲.۱	ابر صفحه جدا کننده بهینه با استفاده از بردارهای پشتیبان خاکستری بر روی ابرصفحه‌های حاشیه‌ای به دست می‌آید.
۹		
۱۲	۳.۱	ابرصفحه بهینه برای نمونه‌های جدایی پذیر خطی
۱۲	۴.۱	تصویر هندسی فاصله جبری نمونه‌ها تا ابرصفحه بهینه برای حالت دو بعدی
۱۶	۵.۱	داده x_i در داخل ناحیه جداکننده در سمت موافق صفحه تصمیم قرار دارد.
۱۶	۶.۱	داده x_i در داخل ناحیه جداکننده در سمت مخالف صفحه تصمیم قرار دارد.
۲۰	۷.۱	انتقال مساله از فضای دو بعدی به فضای سه بعدی
۲۵	۱.۲	ساختار یک نرون.
۲۵	۲.۲	مدل ریاضی یک نرون.
۲۶	۳.۲	ساختار یک نرون با لایه ورودی، لایه پنهان و لایه خروجی.
۲۸	۴.۲	شبکه‌های عصبی دو لایه بازگشتی.
۳۴	۱.۳	نمودار پراکندگی داده‌ها
۳۵	۲.۳	برآورد خط رگرسیونی
۳۶	۳.۳	تابع زیان ϵ - حساسیت
۴۳	۴.۳	یک نمودار ساده بلوکی شده برای شبکه عصبی (۲۰.۳) و (۲۱.۳)
۵۳	۵.۳	شبه کد اجرای روش $Leave - One - Out$ در اعتبار سنجی متقابل (CV)
۵۴	۶.۳	رفتار گذر $\alpha_1, \alpha_2, \dots, \alpha_9, \alpha'_1, \alpha'_2, \dots, \alpha'_9$ در شبکه عصبی (۲۰.۳) و (۲۱.۳) .
۵۵	۷.۳	تغییرات $MSPE$ نسبت به مقادیر مختلف C
	۸.۳	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)
۵۵		
۵۶	۹.۳	تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ
	۱۰.۳	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)
۵۶		

۵۸	۱۱.۳ تغییرات $MSPE$ نسبت به مقادیر مختلف C
	۱۲.۳ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از
۵۸	شبکه عصبی (۲۰.۳) و (۲۱.۳)
۵۹	۱۳.۳ تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ
	۱۴.۳ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه
۵۹	عصبی (۲۰.۳) و (۲۱.۳)
۶۰	۱۵.۳ تغییرات $MSPE$ نسبت به مقادیر مختلف C
	۱۶.۳ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از
۶۰	شبکه عصبی (۲۰.۳) و (۲۱.۳)
۶۱	۱۷.۳ تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ
	۱۸.۳ نتایج رگرسیون بردار پشتیبان به ازای مقادیر مختلف ϵ با استفاده از شبکه
۶۱	عصبی (۲۰.۳) و (۲۱.۳)
۷۱	۱.۴ یک نمودار ساده بلوکی شده برای شبکه عصبی (۲۴.۴)
۸۰	۲.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف C
	۳.۴ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از
۸۰	شبکه عصبی (۲۴.۴)
۸۱	۴.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف δ
	۵.۴ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف δ با استفاده از شبکه
۸۱	عصبی (۲۴.۴)
۸۲	۶.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ
	۷.۴ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه
۸۲	عصبی (۲۴.۴)
۸۳	۸.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف a
	۹.۴ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف a با استفاده از شبکه
۸۳	عصبی (۲۴.۴)
۸۵	۱۰.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف C
	۱۱.۴ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از
۸۵	شبکه عصبی (۲۴.۴)
۸۶	۱۲.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف δ
	۱۳.۴ نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف δ با استفاده از شبکه
۸۶	عصبی (۲۴.۴)
۸۷	۱۴.۴ تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ

۸۷	۱۵.۴	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۴.۴)
۸۸	۱۶.۴	تغییرات $MSPE$ نسبت به مقادیر مختلف a
۸۸	۱۷.۴	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف a با استفاده از شبکه عصبی (۲۴.۴)
۸۹	۱۸.۴	تغییرات $MSPE$ نسبت به مقادیر مختلف C
۸۹	۱۹.۴	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۴.۴)
۹۰	۲۰.۴	تغییرات $MSPE$ نسبت به مقادیر مختلف δ
۹۰	۲۱.۴	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف δ با استفاده از شبکه عصبی (۲۴.۴)
۹۱	۲۲.۴	تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ
۹۱	۲۳.۴	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۴.۴)
۹۲	۲۴.۴	تغییرات $MSPE$ نسبت به مقادیر مختلف a
۹۲	۲۵.۴	نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف a با استفاده از شبکه عصبی (۲۴.۴)
۱۱۲	۲۶	پایداری لیاپانوف.
۱۱۳	۲۷	پایداری مجانبی.
۱۱۸	۲۸	تعبیر هندسی مسأله نامساوی وردشی.

فهرست جداول

۲۷		۱.۲	برخی توابع فعال سازی نرون ها
۵۴		۱.۳	داده های رگرسیون
۵۷		۲.۳	داده های رگرسیون
۷۹		۱.۴	توزیع احتمال X_i و Y_i
۸۴		۲.۴	توزیع احتمال X_i و Y_i

پیشگفتار

ماشین بردار پشتیبان یکی از روش‌های یادگیری با نظارت است که از آن برای طبقه‌بندی و رگرسیون استفاده می‌کنند. این روش از جمله روش‌های نسبتاً جدیدی است که در سال‌های اخیر کارایی خوبی نسبت به روش‌های قدیمی‌تر برای طبقه‌بندی نشان داده است. مبنای کاری دسته‌بندی کننده‌ی SVM دسته‌بندی خطی داده‌ها است و در تقسیم خطی داده‌ها سعی می‌کنیم خطی را انتخاب کنیم که حاشیه اطمینان بیشتری داشته باشد. حل معادله پیدا کردن خط بهینه برای داده‌ها به وسیله روش‌های مسائل درجه دوم که روش‌های شناخته شده‌ای در حل مسائل محدودیت‌دار هستند صورت می‌گیرد. برای به دست آوردن بردار پشتیبان از قضیه دوگانگی لاگرانژ و شرایط کاروش کان تا کر استفاده می‌کنیم.

فصل اول این رساله به معرفی بردار پشتیبان با استفاده از مسائل کلاس بندی اختصاص دارد.

در فصل دوم این رساله شبکه عصبی بیولوژیکی و مصنوعی را توضیح می‌دهیم.

فصل سوم این رساله به معرفی یک شبکه عصبی بازگشتی برای حل مسائل رگرسیون می‌پردازد. در این فصل ثابت می‌شود که نقطه تعادل شبکه عصبی با جواب بهینه مساله رگرسیون بردار پشتیبان معادل است. همچنین همگرایی و پایداری شبکه عصبی معرفی شده، اثبات می‌گردد.

با توجه به این که در اغلب مسائل، کمیت‌های استفاده شده به دلیل انواع خطاهای ممکن دقیق نیستند، در فصل آخر با اصلاح شبکه عصبی استفاده شده در فصل سوم، مساله رگرسیون بردار پشتیبان با قیدهای احتمالی را به وسیله شبکه عصبی حل می‌کنیم.

فصل ۱

ابرصفحه بهینه برای جداسازی نمونه‌ها و معرفی بردار پشتیبان

ماشین بردار پشتیبان یکی از روش‌های یادگیری با نظارت است که از آن برای طبقه‌بندی و رگرسیون استفاده می‌شود. به عبارت دیگر، با دریافت داده‌های آموزشی برچسب خورده (آموزش نظارت شده)، الگوریتم یک ابرصفحه جدا کننده بهینه را خروجی می‌دهد که نمونه‌های جدید را طبقه‌بندی می‌کند. در فضای دو بعدی این ابرصفحه یک خط تقسیم کننده‌ی صفحه به دو بخش است که در آن هر کلاس در یک طرف خط قرار می‌گیرد. در این فصل به معرفی بردار پشتیبان با استفاده از مساله کلاس‌بندی می‌پردازیم. عمدتاً مطالب این فصل از مرجع [۱] گرفته شده است.

۱.۱ شناسایی الگو و مفاهیم مقدماتی

اگر تعدادی سیب و گلابی در اختیار شما قرار دهند و از شما بخواهند که تشخیص دهید هر یک به کدام دسته از دو نوع میوه تعلق دارند، چه می‌کنید؟ این مسأله ظاهراً زیاد پیچیده نیست، هر کسی می‌تواند بر اساس مشاهده به راحتی آن‌ها را از هم جدا کند. به علاوه، خیلی سخت نیست که تکه‌های فاسد، کثیف و بقیه اشیائی که نه سیب هستند و نه گلابی را تشخیص داد.

اگر چه این مساله تشخیص سیب‌ها و گلابی‌ها از یکدیگر به نظر مشکل نیست، اما ماشینی کردن کردن فرایند کم و بیش پیچیده است. مثلاً چه چیزی باید مبنای تصمیم‌گیری باشد که بگوییم یک شی «سیب» است یا «گلابی»؟ آیا وزن آن شی است، اندازه یا رنگ آن؟ شاید شکل آن، بو یا طعم آن باشد؟ احتمالاً ترکیبی از همه این ویژگی‌هاست. لحظه‌ای فکر کنید ما اندازه‌هایی از رنگ و بوی یک شی را داریم، چگونه یک سیب که روی آن کمی گلی است رده بندی می‌شود؟ آیا آن یک سیب است یا یک شی گلی؟

این همان مساله رده‌بندی است: تخصیص یک شی جدید به یک مجموعه از کلاس‌ها (در این جا سیب‌ها یا گلابی‌ها) که از پیش معلوم هستند. رده‌بندی کننده‌ای که باید عمل جداسازی را انجام دهد (یا به هر شی ورودی یک برچسب خروجی تخصیص دهد)، برپایه مجموعه‌ای از مثال‌ها عمل می‌کند. در این جا واژه‌ی شی کاملاً عمومی است، می‌تواند از «سیب» و «گلابی» تا «ارقام دست‌نویس»، از «سیگنال‌های گفتاری» تا «صدای موزیک» متغیر باشد [۴۸].

۱.۱.۱ نظریه یادگیری به‌وسیله سرپرست

آیا ماشین می‌تواند فقط دستوراتی را که ما به آن می‌دهیم انجام دهد؟ در موارد بسیاری ما از ماشین می‌خواهیم که وظایفی را انجام دهد که نمی‌توان آن‌ها را با الگوریتم شرح داد. این مساله هنگامی که مدل ریاضی مشخص و قابل دسترسی وجود ندارد یا موقعی که محاسبه راه حل نشدنی یا خیلی گران دارد، پیش می‌آید. حل کردن این مساله با استفاده از روش‌های برنامه‌نویسی کلاسیک غیرممکن است. به عنوان مثال شناسایی اعداد در دست‌نوشته‌های مختلف، را در نظر بگیرید (این موضوع یکی از مساله‌های کلاسیک در زمینه‌ی آموزش ماشین است). ما برای حل این مساله، نیاز داریم مساله را از یک دیدگاه دیگر بررسی کنیم. احتمالاً می‌توانیم ماشین را آموزش دهیم، به همان شیوه‌ای که کودکان را آموزش می‌دهیم، با این تفاوت که به آن‌ها خلاصه مفاهیم و تعاریف و مسائل تئوری را آموزش نمی‌دهیم، ولی توجه آن‌ها را به تابع‌هایی از ورودی و خروجی جلب می‌کنیم. چگونه بچه‌ها خواندن حروف الفبا را یاد می‌گیرند، معلم به آن‌ها تعریف دقیق و کاربرد حروف را به‌طور مجزا نمی‌گوید اما به آن‌ها مثال‌هایی را نشان می‌دهد که این حروف در آن‌ها به کار رفته‌اند. شاگرد با دقت و تمرین کردن زیاد این مثال‌ها یک استنباط کلی از مشخصات آن حرف پیدا می‌کند، در پایان اگر شاگردان به نوع حروف فکر کنند، می‌توانند لغات را در متن بخوانند و تشخیص دهند. از هدف‌های آشکار و اصلی آموزش، دسته‌بندی کردن صحیح مشاهدات آینده است. ما در این جا به یک سری روش‌های دسته‌بندی با تاثیرهای روشنی از اثرشان روی داده نیاز داریم. چگونه آموزش به عنوان مثال، در مساله ریاضی فرموله می‌شود؟

۲.۱.۱ مدل سازی مساله آموزش

مساله آموزش یک مدل کلی را شرح می دهد که شامل سه قسمت اصلی است. این سه قسمت عبارتند از: تولید کننده داده، یک سرپرست که یک برچسب y را به داده مربوط می کند و ماشین آموزش که خروجی y را تحت عنوان پاسخ سرپرست برمی گرداند. هدف تقلید از سرپرست می باشد، تلاش برای ساختن یک عملکرد، که بهترین پیشگویی را از خروجی های سرپرست به دست آورد، که برپایه ی داده های تولید شده است. این موضوع با این که تلاش کنیم تا عملکرد سرپرست یا پاسخ سرپرست را مشخص کنیم فرق دارد. تولید کننده محیطی که سرپرست ماشین آموزش باید در آن فعالیت کند را مشخص می کند. در واقع بردارهای $x \in \mathbb{R}^m$ را تولید می کند که مستقل از هم هستند و توزیع آن ها را با توجه به یک توزیع احتمال غیر مشخص $P(x)$ ، مشخص می کند. سرپرست مقدار واقعی را با توجه به توزیع احتمالی شرطی منصوب می کند. در این موضوع ما با مساله دسته بندی دوتایی سروکار داریم. ماشین آموزش با مجموعه ای ممکن از نگاشت های $x \rightarrow F(x, a)$ تعریف می شود، به طوری که a عضوی از فضای پارامتری A است که $A \equiv (w, b) \in \mathbb{R}^{m+1}$. به عنوان یک ماشین آموزش که با ابرصفحه جهت دار $\{x : \langle w, x \rangle + b = 0\}$ تعریف می شود، موقعیت یک نقطه در $\mathbb{R}^m$ را مشخص می کند، که در آن $\langle w, x \rangle$ ضرب داخلی دو بردار w و x است. این نتایج در ماشین آموزش به صورت زیر است:

$$f(x, (w, b)) = \text{sign}(\langle w, x \rangle + b) : (w, b) \in \mathbb{R}^{m+1}.$$

تابع $f : \mathbb{R}^m \rightarrow \{-1, 1\}$ که x را به $\{-1, 1\}$ نگاشت می کند، تابع تصمیم نام دارد.

تعریف ۱.۱.۱. انتخاب مشخص از پارامتر a که بر اساس مشاهده می باشد، مجموعه ی تمرینی نام دارد.

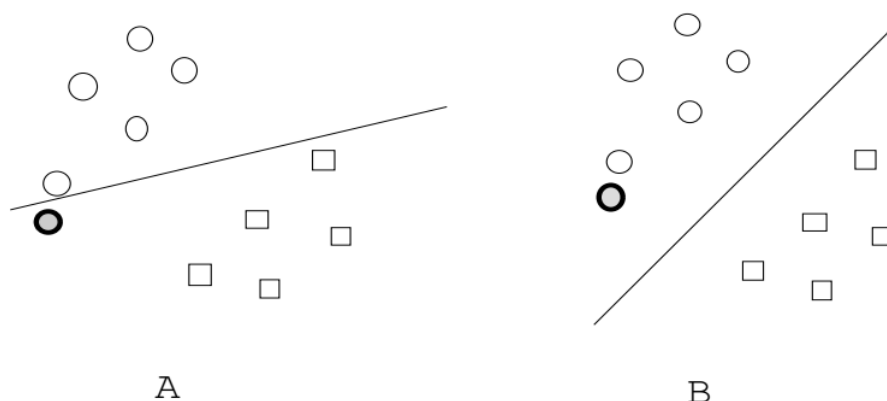
۳.۱.۱ رده بندی کننده خطی

به عنوان اولین قدم در ساخت ماشین بردار پشتیبان، دسته بندی کردن با ابرصفحه های خطی ساده را مورد مطالعه قرار می دهیم. رده ای از ابرصفحه های $\langle w, x \rangle + b = 0$ ، $x \in \mathbb{R}^n$ که نظیر توابع تصمیم زیر هستند را در نظر بگیرید:

$$f(x) = \text{sign}(\langle w, x \rangle + b).$$

برای شروع رده بندی مثال های جداشدنی (نقاط تمرینی جداشدنی) را در نظر می گیریم. راه های زیادی برای جدا کردن دو کلاس از هم وجود دارد. ایده یادگیری از مثال ها یعنی استفاده از اثر اعضای سازنده کلاس ها با امتحان کردن نقاط تمرینی که متعلق به آن کلاس هستند، تعریف می شود.

ابرصفحه‌ها باید به گونه‌ای انتخاب شوند که تغییرات خطی در اطلاعات، باعث تغییر پیش‌بینی‌ها نشوند. اگر فاصله بین جداکننده و نقاط تمرینی خیلی کوچک شوند، حتی مثال‌هایی که برای امتحان کردن نقاط بسیار نزدیک به نقاط تمرین زده‌ایم نیز به درستی دسته‌بندی نمی‌شوند. شکل (۱.۱) این موضوع را توضیح می‌دهد، دایره‌ها و مربع‌های سفید نقاط تمرینی هستند و دایره خاکستری یک نمونه برای تست نمودن ابرصفحه جداکننده است. رویکرد وپنیک^۱ و



شکل ۱.۱: کدام جداسازی بهتر است؟ حاشیه تقریباً صفر در A و یا حاشیه بزرگ در B ؟

لرنر در سال ۱۹۳۶ و وپنیک در سال ۱۹۶۴ ریشه در این مشاهده دارد که قدرت عمومی سازی بستگی به فاصله ابرصفحه تا نقاط تمرینی دارد. آن‌ها یک الگوریتم یادگیری برای مسائل جداسازی به نام تعمیم تصویر ارائه دادند، به طوری که ابرصفحه‌ای را می‌سازد که بیشترین فاصله را تا اعضای کلاس‌ها دارد،

$$\max_{w,b} \min\{\|x - x_i\| : x \in R^n, \langle w, x \rangle + b = 0, i = 1, \dots, n\}.$$

در این بخش ما یک راه مؤثر برای ساخت این ابرصفحه ارائه می‌دهیم.

تعریف ۲.۱.۱. (جداشوندگی) مجموعه تمرینی $\{x_i, y_i\}_{i=1}^n$ که $x_i \in \mathbb{R}^m$ و $y_i \in \{-1, 1\}$ ، جدا شونده توسط ابرصفحه $\langle w, x \rangle + b = 0$ نامیده می‌شود اگر بردار w و مقدار اسکالر b وجود داشته باشند به طوری که روابط زیر برقرار باشند:

$$\langle w, x_i \rangle + b > 0, \quad y_i = 1, \quad \text{برای برچسب } 1, \quad (1.1)$$

$$\langle w, x_i \rangle + b < 0, \quad y_i = -1. \quad \text{برای برچسب } -1. \quad (2.1)$$

ابرصفحه‌ای که توسط w و b تعریف می‌شود ابرصفحه جداکننده نامیده می‌شود.

می‌توان نامعادله‌های (۱.۱) و (۲.۱) را به صورت معادل زیر نوشت:

$$y_i(\langle w, x_i \rangle + b) > 0, \quad y_i \in \{-1, 1\}. \quad (3.1)$$

¹Vapnik

تعریف ۳.۱.۱. (حاشیه) ابرصفحه جداسازنده H را به صورت زیر در نظر بگیرید:

$$H : \langle w, x \rangle + b = 0.$$

حاشیه $\gamma_i(w, b)$ از یک نقطه تمرینی x_i به صورت فاصله بین H و x_i تعریف می شود:

$$\gamma_i(w, b) = y_i(\langle w, x_i \rangle + b).$$

حاشیه $\gamma_s(w, b)$ از یک مجموعه بردار $S = \{x_1, \dots, x_n\}$ به صورت حداقل فاصله از H تا بردارهای S تعریف می شود:

$$\gamma_s(w, b) = \min_{x_i \in S} \gamma_i(w, b).$$

اکنون بردار w^* و اسکالر b^* که حاشیه مجموعه ی تمرینی γ_i را تحت شرایط نامعادله های (۱.۱) و (۲.۱) ماکزیمم می کند را در نظر بگیرید.

زوج (w^*, b^*) ابرصفحه ای را تعیین می کند که مثال های مثبت را از مثال های منفی با حاشیه ماکسیمال جدا می کند. این ابرصفحه را **جداکننده بهینه** می نامیم.

تعریف ۴.۱.۱. (جدا شوندگی) ابرصفحه جداکننده ی بهینه برای مجموعه ی تمرینی χ به صورت زیر تعریف می شود:

$$(w^*, b^*) = \arg \max_{w, b} \gamma(w, b).$$

قضیه ۱.۱.۱. ابر صفحه جدا کننده بهینه یکتاست.

□

برهان. به مرجع [۵۰] مراجعه شود.

با استفاده از تعریف داریم:

$$\begin{aligned} (w^*, b^*) &= \arg \max_{w, b} \gamma_\chi(w, b) \\ &= \arg \max_{w, b} \min_{x_i \in \chi} \gamma_i(w, b) \\ &= \arg \max_{w, b} \min_{x_i \in \chi} y_i(\langle w, x_i \rangle + b). \end{aligned}$$

فرض می کنیم w^* نقطه ماکزیمم تابع پیوسته زیر باشد:

$$\gamma(w) = \min_{x_i \in \chi} y_i(\langle w, x_i \rangle),$$

که روی مجموعه $\|w\| \leq 1$ تعریف شده است. در این صورت داریم:

(۱) پیوستگی γ در ناحیه ی بسته $\|w\| \leq 1$ ، وجود ماکزیمم را نتیجه می دهد.

(۲) فرض کنید که ماکزیمم در نقطه‌ای داخلی مانند w' به دست می‌آید، که در آن $\|w'\| < 1$. اکنون بردار $w'' = \frac{w'}{\|w'\|}$ حاشیه‌ی بزرگ‌تری را تعریف می‌کند، چون

$$\gamma(w'') = \frac{\gamma(w')}{\|w'\|} > \gamma(w'),$$

و این با بهینگی w' تناقض دارد. پس $\|w\| = 1$ و ماکزیمم فقط می‌تواند روی مرز به دست آید. با داشتن w^* می‌توانیم b^* را این‌گونه تعریف کنیم.

$$b^* = \frac{1}{\gamma} (\min_{i \in I_+} \langle w^*, x_i \rangle - \max_{i \in I_-} \langle w^*, x_i \rangle),$$

که در آن $I_+ = \{i | y_i = 1\}$ و $I_- = \{i | y_i = -1\}$. بنابراین مسأله‌ای که برای پیدا کردن ابرصفحه بهینه به دست می‌آید به صورت زیر است:

$$\begin{aligned} & \text{maximize } \gamma_\chi(w, b), \\ & \text{subject to } \gamma_\chi(w, b) > 0, \\ & \|w\|^2 = 1. \end{aligned} \quad (4.1)$$

حل مساله (۴.۱) با توجه به این که تابع مورد نظر نه خطی است و نه درجه دوم و همچنین قیود مساله غیرخطی هستند، دشوار است.

برای پیدا کردن ابرصفحه جدا کننده بهینه می‌توان مساله زیر را در نظر گرفت:

$$\begin{aligned} & \text{minimize } \frac{1}{\gamma} \|w\|^2, \\ & \text{subject to } \langle w, x_i \rangle + b \geq +1, \quad y_i = -1, \quad \text{برای برچسب} \\ & \quad \langle w, x_i \rangle + b \leq -1, \quad y_i = +1, \quad \text{برای برچسب}. \end{aligned} \quad (5.1)$$

قضیه ۲.۱.۱. مسائل بهینه‌سازی (۴.۱) و (۵.۱) معادل هستند.

□

برهان. برای اثبات به مرجع [۳] مراجعه نمایید.

مساله بهینه‌سازی (۵.۱) را به روش لاگرانژ حل می‌کنیم. سعی می‌کنیم نقطه بهینه را با استفاده از تابع لاگرانژ زیر پیدا کنیم.

$$L_P(w, b, \alpha) = \frac{1}{\gamma} \|w\|^2 - \sum_{i=1}^n \alpha_i [y_i (\langle w, x_i \rangle + b) - 1],$$

که در آن $\alpha_i \geq 0$ ضرایب لاگرانژ هستند، تابع $L_P(w, b, \alpha)$ باید نسبت به متغیرهای اولیه w و b می‌نیمم و نسبت به متغیرهای دوگان (α_i ها) ماکزیمم گردد.

در نقطه زینی مشتق‌های L_P نسبت به متغیرهای اولیه باید صفر شود، یعنی:

$$\frac{\partial L_P}{\partial w}(w, b, \alpha) = 0,$$

$$\frac{\partial L_P}{\partial b}(w, b, \alpha) = 0.$$

بنابراین

$$w = \sum_{i=1}^n \alpha_i y_i x_i, \quad (6.1)$$

و

$$\sum_{i=1}^n \alpha_i y_i = 0. \quad (7.1)$$

با جایگزین کردن روابط فوق در L_P به دوگان ولف^۱ مساله بهینه‌سازی زیر می‌رسیم:

$$\begin{aligned} \text{maximize } L_D(\alpha) = & \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle - \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle \\ & + \sum_{i=1}^n \alpha_i = \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle. \end{aligned}$$

برای ساختن ابرصفحه بهینه باید ضرایب α_i^* که مساله‌ی فوق را حل می‌کند، بیابیم. P مخفف ابتدایی و D مخفف دوگان است. L_D و L_P از تابع هدف یکسان ولی با قیود متفاوت می‌باشند. جواب بهینه با می‌نیم کردن L_P و ماکزیمم کردن L_D به دست می‌آید.

لم ۱.۱.۱. فرض کنید بردار w^* ابرصفحه بهینه را تعریف می‌کند. مقدار ماکزیمم تابع $L_D(\alpha^*)$ برابر است با:

$$L_D(\alpha^*) = \frac{1}{2} \|w^*\|^2 = \frac{1}{2} \sum_{i=1}^n \alpha_i^*,$$

برهان. جواب‌های بهینه w^* و α^* باید در شرط مکمل کان – تاکر صدق کنند، یعنی:

$$\alpha_i^* [y_i (\langle w^*, x_i \rangle + b^*) - 1] = 0, \quad \forall i.$$

با جمع بر روی تمام (α_i) داریم:

$$\sum_{i=1}^n \alpha_i^* y_i \langle w^*, x_i \rangle + b^* \sum_{i=1}^n \alpha_i^* y_i - \sum_{i=1}^n \alpha_i^* = 0.$$

با استفاده از شرایط (۶.۱) و (۷.۱) داریم:

$$\sum_{i=1}^n \alpha_i^* = \sum_{i=1}^n \alpha_i^* y_i \langle w^*, x_i \rangle = \|w^*\|^2. \quad (8.1)$$

¹Wolf dual

از طرفی از رابطه‌ی (۶.۱) داریم:

$$\|w^*\|^2 = w^{*T} \cdot w^* = \sum_{i=1}^n \sum_{j=1}^n \alpha_i^* \alpha_j^* y_i y_j < x_i, x_j >.$$

با جای گذاری (۷.۱) و (۸.۱) در دوگان ولف مساله بهینه سازی، مقدار تابع هدف دوگان مساله به صورت زیر نوشته می شود:

$$L_D(\alpha^*) = \sum_{i=1}^n \alpha_i^* - \frac{1}{4} \sum_{i=1}^n \sum_{j=1}^n \alpha_i^* \alpha_j^* y_i y_j < x_i, x_j > = \|w^*\|^2 - \frac{1}{4} \|w^*\|^2 = \frac{1}{4} \|w^*\|^2,$$

و اثبات تمام است. $\square$

اکنون شرط مکمل کان – تاکر $(\alpha_i^* [y_i (< w^*, x_i > + b^*) - 1] = 0, \forall i)$ را دقیق تر بررسی می کنیم. معنی این شرط در مساله ماشین بردار پشتیبان این است که برای یک نقطه تمرینی x_i ، یا مؤلفه‌ی لاگرانژ α_i نظیر آن صفر است یا x_i بر روی یکی از ابرصفحه‌های H_1 یا H_2 است:

$$H_1 : < w, x > + b = +1,$$

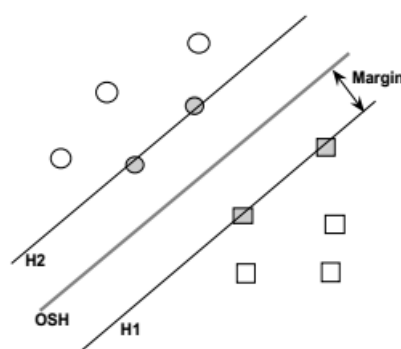
$$H_2 : < w, x > + b = -1.$$

ابرفصفحه‌های H_1 و H_2 ، ابرصفحه‌های حاشیه نامیده می شوند. آن‌ها دارای نقاط تمرینی با حداقل فاصله از ابرصفحه جداکننده بهینه هستند. ابرصفحه‌های حاشیه، مرز حاشیه را تشکیل می دهند و بردارهایی که α_i نظیر آن‌ها مثبت است، بر آن‌ها واقع هستند.

تعریف ۵.۱.۱. نقطه تمرینی x_i یک بردار پشتیبان نامیده می شود اگر و تنها اگر ضریب لاگرانژ نظیر آن مثبت باشد ($\alpha_i > 0$).

بردارهای پشتیبان نزدیک ترین نقاط تمرینی به ابرصفحه‌های جدا کننده هستند. بقیه نقاط تمرینی دارای $\alpha_i = 0$ هستند، یعنی یا بر یکی از ابرصفحه‌های حاشیه ای قرار دارند و یا به یکی از نیم فضاها تعریف (۲.۱.۱) تعلق دارند. توجه داشته باشید که شرط مکمل فقط می گوید که تمام بردارهای پشتیبان روی صفحه‌های حاشیه هستند، نه این که بردارهای پشتیبان تنها نقاط تمرینی روی آن ابرصفحه‌ها هستند. ممکن است موردی باشد که هر دوی α_i و $y_i (< w, x_i > + b)$ برابر صفر شوند. در این حالت x_i روی یکی از ابرصفحه‌های حاشیه ای است، بدون این که بردار پشتیبان باشد.

فرایند یافتن ابرصفحه بهینه جدا کننده‌ی دو کلاس همان مرحله یادگیری است، یعنی با استفاده از نقاط تمرینی ماشین را برای تشخیص الگو آموزش می دهیم. لذا از نقاط تمرینی به عنوان نقاط آموزشی نیز نام می بریم.



شکل ۲.۱: ابر صفحه جدا کننده بهینه با استفاده از بردارهای پشتیبان خاکستری بر روی ابر صفحه‌های حاشیه‌ای به دست می‌آید.

۲.۱ ماشین بردار پشتیبان

مساله و مفاهیم شناسایی الگو به عنوان یک عامل مهم در طراحی سیستم‌های اطلاعاتی بسیار مورد توجه قرار گرفته است. تئوری شناسایی مبتنی بر روش‌های تحلیلی قدرتمندی است که کاربردهای وسیعی در سیستم‌های کنترل و مخابرات دارد. میزان علاقه در این رشته (مخصوصاً در دهه‌های ۶۰ و ۷۰) در حال گسترش بوده و می‌باشد، به طوری که تبدیل به یک موضوع مشترک تحقیقاتی برای رشته‌های آکادمیک مختلف از قبیل مهندسی، علوم پایه مانند فیزیک، آمار و شیمی، زبان‌شناسی، روان‌شناسی، بیولوژی، تئوری اطلاعات، علم کامپیوتر و ... گشته است.

هر یک از حوزه‌های تحقیقاتی بر ویژگی خاصی از مساله شناسایی الگو از قبیل مدل‌سازی فرایندهای فیزیولوژی و یا مثلاً توسعه تکنیک‌های تحلیلی برای تصمیم‌گیری‌های اتوماتیک تاکید دارند.

مساله شناسایی الگو، عملاً چیزی جز جداسازی داده‌ها یا الگوهای ورودی بین دسته‌های مختلف نیست. در طراحی یک سیستم شناسایی الگو چند کار باید انجام گیرد: نخست اشیاء مورد شناسایی را باید به الگوها و یا بردارهای ورودی که به سیستم شناسایی اعمال می‌گردند، تبدیل نمود. مؤلفه‌های بردار ورودی از طریق اندازه‌گیری به دست می‌آیند. هر کمیت اندازه‌گیری شده، یک ویژگی از شیء مورد نظر را بیان می‌کند. از نظر هندسی هر شیء یا موضوع شناسایی را می‌توان به عنوان یک نقطه در فضای چندبعدی اقلیدسی در نظر گرفت.

کار دوم استخراج مشخصه‌های مهم از روی داده‌های ورودی اندازه‌گیری شده و کاهش ابعاد بردار الگوها می‌باشد، به این کار اصطلاحاً استخراج شاخص می‌گویند. مشخصه‌های یک طبقه آن‌هایی هستند که در تمامی الگوهای متعلق به آن طبقه خاص مشترک می‌باشند. عناصری از بردار مشخصه‌ها که در تمامی طبقه‌ها مشترک هستند حذف می‌گردند.

سومین کار در شناسایی الگو، این است که روندی جهت تصمیم‌گیری بهینه اتوماتیک، برای طبقه‌بندی الگوها به‌دست آوریم.

به بیان ریاضی مساله شناسایی الگو به شکل ساده، یعنی تخصیص الگوهای ورودی به یکی از طبقاتی که فضای چندبعدی اقلیدسی برای تصمیم‌گیری به تعداد متناهی از آن تقسیم شده است. یکی از ایده‌های جدید در شناسایی و دسته‌بندی الگوها، ماشین بردار پشتیبان یا *SVM* است. ماشین بردار پشتیبان دارای خواص ارزشمندی است که آن را برای شناسایی الگو مناسب می‌سازد، از جمله این که *SVM* در آموزش خود مشکل بهینه‌های محلی را ندارد، دسته‌بندی‌کننده‌ها را با حداکثر تعمیم بنا می‌کند، ساختار و توپولوژی خود را به‌صورت بهینه تعیین می‌کند و توابع متمایز غیرخطی را به راحتی و با محاسبات کم با استفاده از مفهوم حاصل‌ضرب داخلی در فضای هیلبرت تشکیل می‌دهد.

اولین الگوریتم برای طبقه‌بندی الگوها در سال ۱۹۳۶ توسط فیشر^۱ ارائه شد و معیار آن برای بهینه بودن، کم کردن خطای طبقه‌بندی الگوهای آموزشی بود. بسیاری از الگوریتم‌ها و روش‌هایی که تاکنون نیز برای طبقه‌بندی‌کننده‌های الگو ارائه شده است از همین استراتژی پیروی می‌کنند. در هیچ‌یک از این روش‌ها خاصیت تعمیم طبقه‌بندی‌کننده به‌طور مستقیم در تابع هزینه‌ی روش دخالت داده نشده است و طبقه‌بندی‌کننده طراحی شده نیز دارای خاصیت تعمیم‌دهنده‌گی کمی می‌باشد. اگر طراحی دسته‌بندی‌کننده‌ی الگو را به‌عنوان یک مساله بهینه‌سازی در نظر بگیریم بسیاری از این روش‌ها با مشکل بهینه‌های محلی در تابع هزینه مواجه هستند و در دام بهینه‌های محلی گرفتار می‌شوند.

شکل دیگری نیز وجود دارد و آن تعیین ساختار و توپولوژی دسته‌بندی‌کننده قبل از طراحی است که به‌عنوان مثال تعیین تعداد بهینه‌ی گره‌های لایه‌ی مخفی در شبکه‌های عصبی *MLP*^۲، تعداد توابع گاوسی در شبکه‌های عصبی *RBF*^۳ و یا تعداد بهینه‌ی حالت‌ها و توابع گاوسی در مدل پنهان مارکف می‌باشد. همه این عوامل باعث می‌شوند که در عمل با روش‌های پیشنهاد شده‌ی قبلی نتوانیم به یک دسته‌بندی‌کننده‌ی بهینه برسیم.

محقق روسی به‌نام ولادمیر وپنیک^۴ در سال ۱۹۶۵ گامی بسیار مهم در طراحی دسته‌بندی‌کننده‌ها برداشت و نظریه آماری یادگیری را به‌صورت مستحکم‌تری بنا نهاد و ماشین‌های بردار پشتیبان را بر این اساس ارائه داد. ماشین‌های بردار پشتیبان دارای خواص زیر هستند:

- (۱) طراحی دسته‌بندی‌کننده با حداکثر تعمیم
- (۲) رسیدن به بهینه سراسری تابع هزینه
- (۳) تعیین خوار ساختار و توپولوژی بهینه برای طبقه‌بندی‌کننده
- (۴) مدل کردن توابع متمایز غیرخطی با استفاده از هسته‌های غیرخطی و مفهوم حاصل‌ضرب داخلی در فضا‌های هیلبرت.

¹Fisher

²Multi Layer Perceptron

³Radial Basis Function

⁴Veladimir Vapnik

۱.۲.۱ ابرصفحه بهینه برای نمونه‌های جدایی‌پذیر خطی

نمونه آموزش $\{(x_i, y_i)\}_{i=1}^n$ را که در آن x_i نمونه‌ی ورودی و y_i پاسخ مطلوب نظیر (خروجی هدف) است را در نظر می‌گیریم. فرض می‌کنیم نمونه نمایش داده شده با $y_i = 1$ و نمونه نمایش داده شده با $y_i = -1$ به طور خطی جدایی‌پذیراند. معادله‌ی صفحه‌ی تصمیم به فرم ابرصفحه‌ای است که خاصیت جداسازی دارد، یعنی به صورت:

$$H : w^T x + b = 0, \quad (9.1)$$

که در آن x بردار ورودی و w بردار وزنی قابل تنظیم و b مقدار ثابت است. در این صورت داریم:

$$\begin{cases} w^T x + b \geq 0, & \text{برای برچسب } y_i = 1, \\ w^T x + b \leq 0, & \text{برای برچسب } y_i = -1. \end{cases} \quad (10.1)$$

فرض نمونه‌های جدایی‌پذیر خطی برای بیان ایده‌ی اصلی یک SVM به کار می‌رود. برای بردار وزنی w و ثابت b ، فضای جدایی‌پذیری بین ابرصفحه در معادله (۹.۱) و نزدیک‌ترین داده، **حاشیه‌ی جدایی‌پذیری** نامیده می‌شود که با ρ نمایش داده می‌شود. هدف از یافتن SVM یافتن ابرصفحه‌ای است که این حاشیه را ماکزیمم کند.

با هدف درک بهتر، تصویر حالتی از فضای ورودی دو بعدی را بررسی می‌کنیم. داده‌ها به طور خطی جدایی‌پذیر هستند و ابرصفحه‌های مختلفی می‌توانند دو کلاس را جداسازی کنند. در واقع برای $x \in \mathbb{R}^2$ جداسازی توسط صفحه‌های

$$w_1 x_1 + w_2 x_2 + b = 0,$$

انجام می‌شود. چگونه می‌توان بهترین ابرصفحه را یافت؟ از بین همه‌ی ابرصفحه‌های ممکن آن ابرصفحه‌ای را می‌یابیم که حاشیه‌ی جداکننده‌ی دو کلاس را ماکزیمم کند. این یک روش قابل قبول شهودی است.

اگر حاشیه جداسازی را با ρ نشان دهیم، هدف SVM این است که اندازه ρ ماکزیمم شود. شکل (۳.۱) ساختار هندسی یک ابرصفحه بهینه برای یک فضای ورودی دو بعدی را شرح می‌دهد.

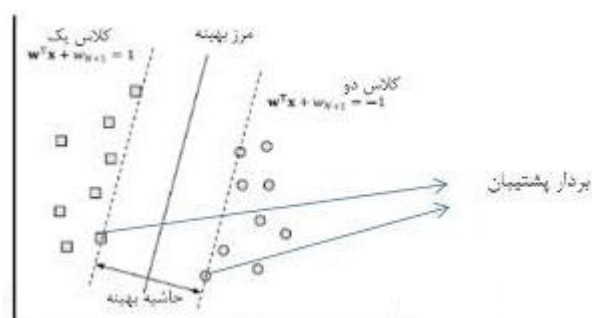
فرض کنی w_0 و b_0 به ترتیب مقادیر بهینه از بردار وزنی و مقدار ثابت باشند. ابرصفحه بهینه زیر

$$w_0^T x + b_0 = 0,$$

فضای تصمیم خطی چندبعدی را در فضای ورودی نمایش می‌دهد. تابع زیر

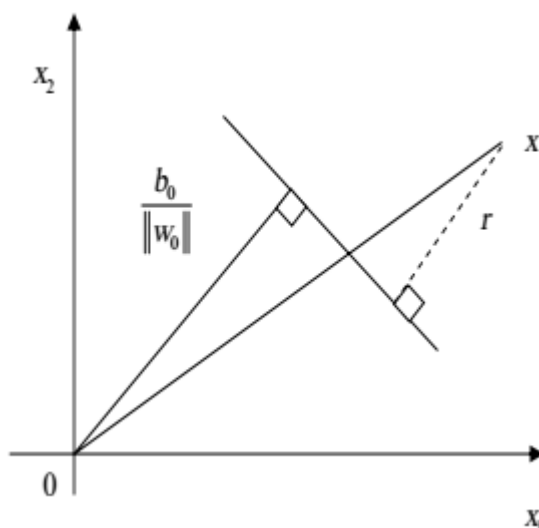
$$g(x) = w_0^T x + b_0,$$

یک اندازه‌ی جبری از فاصله‌ی x تا ابرصفحه بهینه را نشان می‌دهد. ساده‌ترین راه برای دیدن این مطلب ارائه x به صورت $x = x_\rho + r \frac{w_0}{\|w_0\|}$ است.



شکل ۳.۱: ابرصفحه بهینه برای نمونه‌های جدایی پذیر خطی

که در آن x_ρ تصویر نرمال x به توی ابرصفحه بهینه و r فاصله جبری مطلوب است. اگر x در طرف مثبت ابرصفحه بهینه باشد، r مثبت و اگر در طرف منفی ابرصفحه بهینه باشد، r منفی است. شکل (۴.۱) را ببینید. با تعریف $g(x_\rho) = 0$ نتیجه می‌شود:



شکل ۴.۱: تصویر هندسی فاصله جبری نمونه‌ها تا ابرصفحه بهینه برای حالت دو بعدی

$$g(x) = g(x_\rho + r \frac{w_\circ}{\|w_\circ\|}),$$

$$\Rightarrow g(x) = w_\circ^T x_\rho + b_\circ + r w_\circ^T \frac{w_\circ}{\|w_\circ\|},$$

$$\Rightarrow g(x) = w_\circ^T x + b_\circ = r \|w_\circ\|,$$

و یا

$$r = \frac{g(x)}{\|w_\circ\|},$$

به ویژه فاصله‌ی مبدا یعنی $(x = 0)$ تا ابرصفحه بهینه، $\frac{b_\circ}{\|w_\circ\|}$ است. اگر $b_\circ > 0$ ، مبدا در طرف ابرصفحه بهینه و اگر $b_\circ < 0$ ، مبدا در طرف منفی ابرصفحه بهینه و اگر $b_\circ = 0$ ، ابرصفحه بهینه از مبدا می‌گذرد.

مساله، یافتن پارامترهای w_0 و b_0 برای ابرصفحه بهینه با استفاده از مجموعه‌ی آموزشی $\{(x_i, y_i)\}_{i=1}^n$ است که زوج (w_0, b_0) در محدودیت‌های زیر صدق کند،

$$\begin{cases} w_0^T x + b_0 \geq 0, & \text{برای برچسب } y_i = 1, \\ w_0^T x + b_0 \leq 0, & \text{برای برچسب } y_i = -1. \end{cases} \quad (11.1)$$

اگر معادلات (۱۰.۱) برقرار باشند، نمونه‌ها به‌طور خطی جدایی پذیرند. بردارهایی که نامعادلات (۱۱.۱) را به تساوی‌هایی به ترتیب برابر ۱ و -۱ تبدیل می‌کنند SV نام دارند. این بردارها نقش اساسی در کاربرد ماشین‌های یادگیری دارند. SV ها داده‌هایی هستند که در نزدیک‌ترین فاصله تا صفحه تصمیم قرار می‌گیرند و طبقه بندی آن‌ها بسیار مشکل است و نسبت مستقیم با مکان صفحه تصمیم دارند. یک SV ، $x^{(s)}$ با $y^{(s)} = \pm 1$ را در نظر بگیرید، داریم:

$$g(x^{(s)}) = w_0^T x^{(s)} \pm b_0 = \pm 1, \quad y^{(s)} = \pm 1. \quad (12.1)$$

از معادله (۱۲.۱)، فاصله جبری SV یعنی فاصله $x^{(s)}$ تا ابرصفحه بهینه برابر

$$r = \frac{g(x^{(s)})}{\|w_0\|} = \begin{cases} \frac{1}{\|w_0\|}, & \text{برای برچسب } y_i = 1, \\ \frac{-1}{\|w_0\|}, & \text{برای برچسب } y_i = -1, \end{cases} \quad (13.1)$$

است. علامت + نشان می‌دهد که $x^{(s)}$ در طرف مثبت ابرصفحه بهینه و علامت منفی یعنی $x^{(s)}$ در طرف منفی ابرصفحه بهینه است. مقدار بهینه حاشیه از معادله (۱۳.۱) به‌صورت زیر به‌دست می‌آید:

$$\rho = \frac{1}{\|w_0\|} - \left(\frac{-1}{\|w_0\|} \right) = \frac{2}{\|w_0\|}.$$

بنابراین با توجه به مطالب ارائه شده‌ی قبل مساله پیدا کردن ابرصفحه بهینه جداساز را می‌توان به شکل زیر نوشت:

$$\begin{aligned} &\text{minimize } \phi(w) = \frac{1}{4} \|w\|^2, \\ &\text{subject to } y_i(\langle w, x_i \rangle + b) \geq 1, \quad i = 1, \dots, n. \end{aligned} \quad (14.1)$$

مساله (۱۴.۱) را با استفاده از ضرایب لاگرانژ حل می‌کنیم. ابتدا تابع لاگرانژ را به‌صورت زیر در نظر می‌گیریم:

$$L(w, b, \alpha) = \frac{1}{4} w^T w - \sum_{i=1}^n \alpha_i [y_i(w^T x_i + b) - 1].$$

که در آن متغیرهای نامنفی α_i ضرایب لاگرانژ نامیده می‌شوند. جواب مساله بهینه سازی مقید (۱۴.۱) توسط نقطه زینی تابع لاگرانژ $L(w, b, \alpha)$ که نسبت به w و b می‌نیم و نسبت به α

ماکزیمم می‌شود، به دست می‌آید. از $L(w, b, \alpha)$ نسبت به w و b مشتق گرفته و مساوی صفر قرار می‌دهیم و به شرایط بهینگی زیر می‌رسیم:

$$\frac{\partial L(w, b, \alpha)}{\partial w} = 0, \quad (15.1)$$

$$\frac{\partial L(w, b, \alpha)}{\partial b} = 0, \quad (16.1)$$

از شرایط بهینگی (۱۵.۱) و (۱۶.۱) داریم:

$$w - \sum_{i=1}^n \alpha_i y_i x_i = 0 \Rightarrow w = \sum_{i=1}^n \alpha_i y_i x_i, \quad (17.1)$$

$$\sum_{i=1}^n \alpha_i y_i = 0. \quad (18.1)$$

در هر نقطه زینی برای هر ضریب لاگرانژ α_i داریم:

$$\alpha_i [y_i (w^T x_i + b) - 1] = 0, \quad i = 1, \dots, n.$$

این خواص از شرایط کان-تاکر نتیجه می‌شود. حال از روی مساله بهینه سازی مقید (مساله اولیه) مساله دیگری می‌سازیم که مساله دوگان نامیده می‌شود. این مساله همان مقدار بهینه مساله اولیه را دارد با این تفاوت که با ضرایب لاگرانژ جواب بهینه به دست می‌آید. در ادامه قضیه دوگان را بیان می‌کنیم:

۱) شرط لازم و کافی برای آن که w_0 جواب بهینه مساله اولیه و α_0 جواب بهینه مساله دوگان باشد، این است که w_0 برای مساله اولیه شدنی و $\phi(w_0) = L(w_0, b_0, \alpha_0) = \min_w L(w, b_0, \alpha_0)$ باشد.

برای به دست آوردن مساله دوگان، تابع لاگرانژ را به صورت زیر جمله به جمله باز نویسی می‌کنیم:

$$L(w, b, \alpha) = \frac{1}{2} w^T w - \sum_{i=1}^n \alpha_i y_i w^T x_i - b \sum_{i=1}^n \alpha_i y_i + \sum_{i=1}^n \alpha_i. \quad (19.1)$$

با توجه به شرط بهینگی از معادله (۱۸.۱) جمله سوم در طرف راست معادله (۱۹.۱)، صفر است. بنابراین با توجه به معادله (۱۷.۱) داریم:

$$w^T w = \sum_{i=1}^n \alpha_i y_i w^T x_i = \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j x_i^T x_j.$$

به این ترتیب تابع هدف $L(w, b, \alpha) = Q(\alpha)$ به صورت زیر فرموله می‌شود:

$$Q(\alpha) = \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j x_i^T x_j, \quad (20.1)$$

که در آن α_i ها نامنفی هستند.

اکنون مساله دوگان را به صورت زیر بیان می کنیم:
 نمونه آموزشی $\{(x_i, y_i)\}_{i=1}^n$ داده شده، ضرایب لاگرانژ $\{\alpha_i\}_{i=1}^n$ را به گونه ای می یابیم که تابع هدف

$$Q(\alpha) = \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j x_i^T x_j,$$

با قیود زیر ماکزیمم شود:

$$\begin{cases} \sum_{i=1}^n \alpha_i y_i = 0, \\ \alpha_i \geq 0, \quad i = 1, \dots, n. \end{cases}$$

۲.۲.۱ ابرصفحه بهینه برای نمونه های جدایی ناپذیر

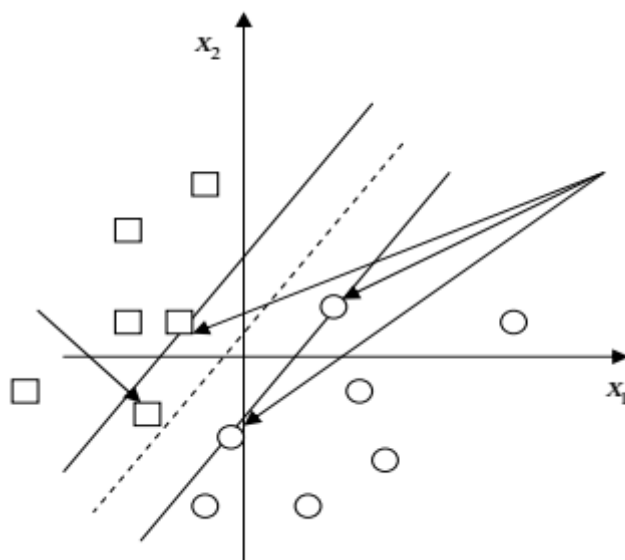
با یک مجموعه داده های آموزشی، ساخت ابرصفحه جداکننده بدون در نظر گرفتن خطاهای طبقه بندی ممکن نیست. با وجود این، ابرصفحه ای که خطای طبقه بندی را می نیمم کند، می یابیم.

در بخش قبل مساله SVM را برای حالت جدایی پذیر بررسی کردیم ولی در عمل تقریباً تمامی مسائل به صورت جدایی ناپذیر می باشند. برای حالت جدایی ناپذیر یک دسته از متغیرهای ζ_i را به عنوان متغیرهای کمبود تعریف می کنیم به طوری که شرایط زیر برقرار باشد:

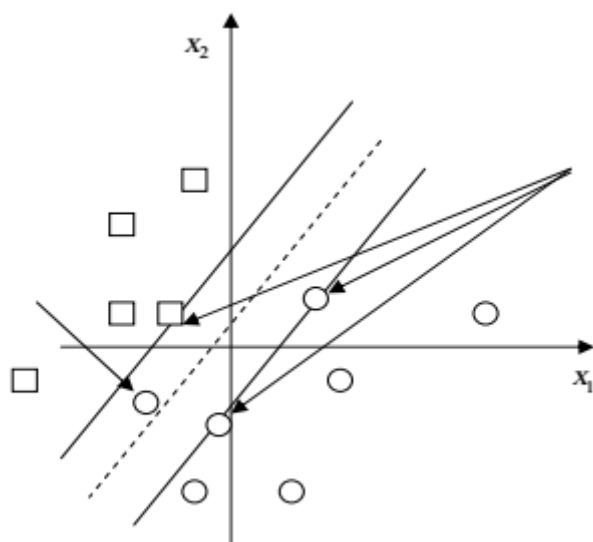
$$y_i(w^T x_i + b) \geq 1 - \zeta_i, \quad (21.1)$$

$$\zeta_i \geq 0, \quad i = 1, \dots, n. \quad (22.1)$$

در (۲۱.۱) و (۲۲.۱) اگر $0 \leq \zeta_i \leq 1$ ، داده x_i در داخل ناحیه جدا کننده در سمت موافق صفحه تصمیم قرار می گیرد (شکل (۵.۱)). و اگر $\zeta_i > 1$ ، رده بندی اشتباه است و داده در سمت مخالف ابر صفحه جداکننده قرار می گیرد (شکل (۶.۱)).



شکل ۵.۱: داده x_i در داخل ناحیه جداکننده در سمت موافق صفحه تصمیم قرار دارد.



شکل ۶.۱: داده x_i در داخل ناحیه جداکننده در سمت مخالف صفحه تصمیم قرار دارد.

همچنین اگر $\zeta_i = 0$ ، در این حالت برای نقطه تمرینی x_i خطای طبقه‌بندی وجود ندارد و داده خارج از ناحیه جدا کننده است و به درستی طبقه‌بندی می‌شود. بردارهای پشتیبان نمونه‌های خاصی هستند که در معادله‌های (۲۱.۱) و (۲۲.۱) صدق می‌کنند. واضح است که هرچه مقادیر متغیرهای ζ_i بیشتر شود از حالت بهینه دورتر خواهیم شد و خطا بیشتر می‌شود. پس مساله بهینه‌سازی مقید را به صورت زیر تعریف می‌کنیم:

$$\begin{aligned} & \text{minimize} \quad \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n \zeta_i, \\ & \text{subject to} \quad y_i(< w, x_i > + b) \geq 1 - \zeta_i, \quad i = 1, \dots, n, \\ & \quad \quad \quad \zeta_i \geq 0, \quad i = 1, \dots, n, \end{aligned}$$

که در آن C پارامتر جریمه است. تابع لاگرانژ این مساله به صورت زیر است:

$$L_P(w, b, \zeta, \alpha, \beta) = \frac{1}{2} w^T w + C \sum_{i=1}^n \zeta_i - \sum_{i=1}^n \alpha_i [y_i(< w, x_i > + b) - 1] - \sum_{i=1}^n \beta_i \zeta_i.$$

شرط لازم و کافی برای بهینگی به صورت زیر است:

$$\begin{aligned} \frac{\partial L}{\partial w} = 0 & \Rightarrow w = \sum_{i=1}^n \alpha_i y_i x_i, \\ \frac{\partial L}{\partial b} = 0 & \Rightarrow \sum_{i=1}^n \alpha_i y_i = 0, \\ \frac{\partial L}{\partial \zeta_i} = 0 & \Rightarrow C - \alpha_i - \beta_i = 0. \end{aligned}$$

برای این مساله نیز مشابه قبل، مساله دوگان به صورت زیر است:

$$\begin{aligned} & \text{maximize} \quad -\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j x_i^T x_j + \sum_{i=1}^n \alpha_i, \\ & \text{subject to} \quad \sum_{i=1}^n \alpha_i y_i = 0, \\ & \quad \quad \quad 0 \leq \alpha_i \leq C, \quad i = 1, \dots, n. \end{aligned}$$

همان‌طور که ملاحظه شد، حل مساله SVM در حالت جدایی ناپذیر مشابه حل آن در حالت جدایی پذیر است. با این تفاوت که محدوده تغییرات ضرایب لاگرانژ α_i ، فرق می‌کند. پس از به دست آوردن ضرایب لاگرانژ α_i ، نمونه‌هایی که ضرایب لاگرانژ آن‌ها در روابط $0 < \alpha_i \leq C$ صدق می‌کنند، **بردار پشتیبان** می‌باشند. بردارهای پشتیبان با $\alpha_i = C$ در جهت مخالف صفحه تصمیم قرار دارند. مقدار w و شکل تابع جداکننده مشابه حالت جدایی پذیر خواهد بود.

۳.۲.۱ رده‌بندی کننده‌های غیرخطی

در بخش‌های گذشته مشاهده کردیم که چگونه رده‌بندی کننده‌های خطی در داده‌های ورودی به آسانی توسط تکنیک‌های بهینه‌سازی استاندارد محاسبه می‌گردند. با این‌که کار کردن با رده‌بندی کننده‌های خطی آسان است اما آن‌ها محدودیت‌های شدیدی بر کار آموزش می‌گذارند.

در بیشتر مسائل زندگی داده‌ها به‌طور خطی جدا پذیر نیستند. در این موارد روش حل، استفاده از تبدیلات غیرخطی است. به‌عنوان مثال ایجاد یک بردار با بعد بالاتر به‌وسیله ضرب همه بردارها در بردار ویژگی و نگاشتی که داده‌ها را از فضای ورودی به فضایی با بعد بالاتر می‌برد، که یک تابع هسته غیرخطی است. ماشین‌های بردار پشتیبان ذکر شده در بخش‌های قبل، برای دسته‌بندی الگوهای یک مساله دوکلاسه، از مرزهای جداکننده خطی و از یک ابرصفحه استفاده می‌کند و در واقع حاصل ضرب داخلی بردار ورودی با هر کدام از بردارهای پشتیبان در فضای n بعدی ورودی محاسبه می‌گردد.

وینیک با استفاده از مفهوم ضرب داخلی در فضاها هیلبرت و قضیه هیلبرت-اشمیت نشان داد که ابتدا می‌توان بردار ورودی x را با یک تبدیل غیرخطی به یک فضای با بعد زیاد انتقال داد و در آن فضا حاصلضرب داخلی را انجام داد و ثابت کرد که اگر یک هسته متقارن، شرایط قضیه مرسر^۱ را داشته باشد، اعمال این هسته در فضای ورودی با بعد کم می‌تواند به‌عنوان حاصلضرب داخلی در یک فضای هیلبرت با بعد زیاد تلقی شود و محاسبات را به شدت کاهش دهد [۲].

تصمیم نهایی ما به‌عنوان یک ترکیب خطی از صفات داده شده ممکن است بسیار مشکل باشد. ایده کلی به این‌صورت است که اطلاعات ورودی را به فضایی با بعد بالاتر می‌نگاریم و سپس به کمک یک رده‌بندی کننده خطی نقاط تمرینی را جدا می‌کنیم، که این رده بندی کننده خطی یک جداکننده غیرخطی در فضای ورودی را نتیجه می‌دهد. به‌وجود آمدن این تکنیک به دهه ۶۰ برمی‌گردد و برای اولین بار در مقاله‌ای از آیزرمن^۲ مطرح شده است.

نحوه نمایش نمونه‌های آموزشی را با نگاشت به یک فضای بی‌نهایت بعدی مانند فضای هیلبرت انتقال می‌دهیم. معمولاً H بعد بسیار بالاتری نسبت به فضای اولیه H دارد (L و H به ترتیب به‌عنوان مخففی از کمترین و بیشترین بعد به کار می‌روند). این نگاشت قبل از یادگیری بر هر کدام از مثال‌ها اعمال می‌شود.

ابرصفحه جداکننده بهینه در فضای H ساخته می‌شود. i امین عضو $\Phi(x_i)$ مشخصه i ام تحت نگاشت $\Phi: \mathbb{R}^n \rightarrow H$ ویژگی فضا خوانده می‌شود. عمل انتخاب کردن بهترین نمایش اطلاعات نیز انتخاب مشخصه نامیده می‌شود.

¹Mercer

²Aizerman

۴.۲.۱ هسته ضرب داخلی

تعریف ۱.۲.۱. تعریف تابع کرنل (هسته ضرب داخلی) :
فرض کنید نگاشت $\Phi : L \rightarrow H$ داده‌های فضای ورودی L را به فضای مشخصه H می‌برد. تابع $K : L \times L \rightarrow H$ را یک تابع کرنل گویند اگر و تنها اگر داشته باشیم :

$$K(x, z) = \langle \Phi(x), \Phi(z) \rangle .$$

عملکرد تابع کرنل مانند یک ضرب داخلی در H است که می‌تواند مانند یک تابع در فضای اولیه L ارزیابی شود. در بخش قبل ما با معرفی نگاشت Φ قبل از اجرای الگوریتم یادگیری آغاز کردیم. حال برعکس شروع کرده و با انتخاب یک کرنل مانند K آغاز می‌کنیم که خود به‌طور غیرمستقیم یک نگاشت مانند Φ را معرفی خواهد کرد.

فرض کنید x برداری از فضای ورودی با بعد m_0 است و $\{\Phi_j(x)\}_{j=1}^{m_1}$ یک مجموعه از تبدیلات غیرخطی از فضای ورودی به فضای ویژگی است. با وجود این تبدیلات غیرخطی، ما می‌توانیم ابرصفحه‌ای به‌صورت زیر تعریف کنیم که مانند یک ابرصفحه تصمیم عمل می‌کند:

$$\sum_{j=1}^{m_1} w_j \Phi_j(x) + b = 0, \quad (23.1)$$

که در آن $\{w_j\}_{j=1}^{m_1}$ مجموعه‌ای از وزن‌های خطی است که فضای ویژگی را به فضای خروجی مرتبط می‌کند و b یک مقدار ثابت است. معادله (۲۳.۱) را به‌صورت خلاصه زیر می‌نویسیم:

$$\sum_{j=0}^{m_1} w_j \Phi_j(x) = 0,$$

که در آن $\Phi_0(x) = 1$ برای هر x و لذا w_0 مقدار ثابت b را نشان می‌دهد. $\Phi_j(x)$ نظیر وزن w_j در فضای تصویر است. $\Phi(x)$ تصویر تولید شده نظیر بردار ورودی x است.

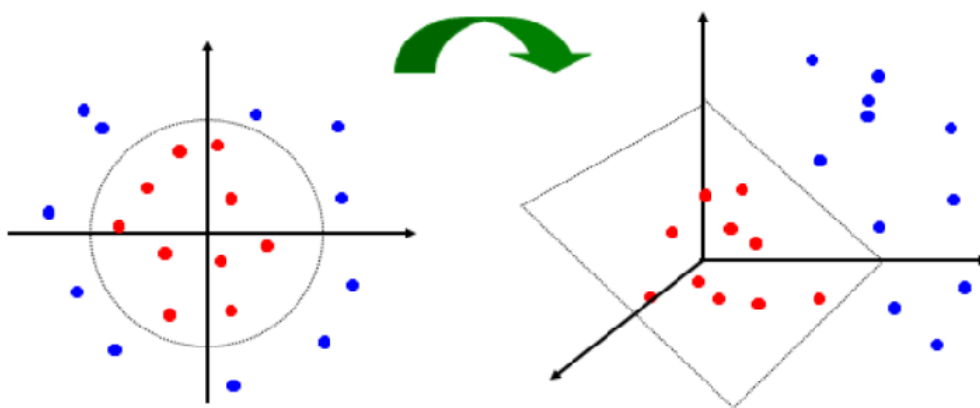
با تعریف $\Phi(x)$ به‌صورت $\Phi(x) = (\Phi_0(x)\Phi_1(x)\dots\Phi_{m_1}(x))$ که بنا به فرض $\Phi_0(x) = 1$ ، بردار $\Phi(x)$ تصویر القا شده بردار ورودی x در فضای ویژگی است، همان‌طور که در شکل (۷.۱) نشان داده شده است.

بنابراین می‌توانیم صفحه تصمیمی به شکل

$$w^T \Phi(x) = 0, \quad (24.1)$$

تعریف کنیم. با تطبیق فرمول (۱۷.۱) به موقعیت فعلی که اکنون انتظار داریم ترکیبی با قابلیت جدایی پذیری خطی داشته باشد، می‌توان نوشت:

$$w = \sum_{i=1}^n \alpha_i y_i \Phi(x_i), \quad (25.1)$$



شکل ۷.۱: انتقال مساله از فضای دو بعدی به فضای سه بعدی

که بردار ویژگی $\Phi(x_i)$ متناظر الگوی ورودی x_i در نمونه i ام است. بنابراین با جای‌گذاری معادله (۲۵.۱) در (۲۴.۱) می‌توانیم صفحه تصمیم محاسبه شده در فضای ویژگی را به صورت زیر تعریف کنیم:

$$\sum_{i=1}^n \alpha_i y_i \Phi^T(x_i) \Phi(x) = 0. \quad (26.1)$$

$\Phi^T(x_i) \Phi(x)$ ، ضرب داخلی دو بردار در فضای ویژگی توسط بردار ورودی x و الگوی ورودی x_i برای i امین نمونه است.

هسته ضرب داخلی $K(x, x_i)$ به صورت زیر تعریف می‌شود:

$$K(x, x_i) = \Phi^T(x_i) \Phi(x) = \sum_{j=1}^n \Phi_j(x) \Phi_j(x_i), \quad i = 1, \dots, n. \quad (27.1)$$

از این تعریف مشاهده می‌شود که هسته ضرب داخلی تابعی متقارن است:

$$K(x, x_i) = K(x_i, x), \quad \forall i. \quad (28.1)$$

از هسته ضرب داخلی برای ساخت ابرصفحه بهینه در فضای ویژگی بدون در نظر گرفتن شکل ضمنی فضای ویژگی استفاده می‌کنیم. از معادلات (۲۶.۱) و (۲۷.۱) داریم:

$$\sum_{i=1}^n \alpha_i y_i K(x, x_i) = 0. \quad (29.1)$$

تابع تصمیم در H به صورت زیر است:

$$f(x) = \text{sign}\left(\sum_{i=1}^n \alpha_i y_i < \Phi(x_i), \Phi(x) > + b\right).$$

تاکنون دیدیم که در دو مرحله یادگیری و آموزش تابع مشخصه فقط به مقدار ضرب داخلی در فضای ویژگی بستگی دارد. بنابراین آن‌ها می‌توانند با استفاده از توابع کرنل محاسبه شوند. لذا تابع تصمیم فوق را می‌توان به صورت زیر نوشت:

$$f(x) = \text{sign}\left(\sum_{i=1}^n \alpha_i y_i K(x, x_i) + b\right).$$

در نتیجه لازم نیست که نگاشت مشخصه را برای حل مساله یادگیری در فضای مشخصه بدانیم.

چه توابعی را می‌توان به عنوان تابع کرنل انتخاب کرد؟ به آسانی قابل مشاهده است که تابع هسته ضرب داخلی K متقارن است:

$$K(x, z) = \Phi(x)\Phi(z) = \Phi(z)\Phi(x) = K(z, x),$$

و در نامساوی کشی-شوارتز هم صدق می‌کند:

$$K(x, z)^2 = \langle \Phi(x), \Phi(z) \rangle^2 \leq \|\Phi(x)\|^2 \|\Phi(z)\|^2$$

$$= \langle \Phi(x), \Phi(x) \rangle \langle \Phi(z), \Phi(z) \rangle = K(x, x)K(z, z).$$

علاوه بر این باید در شرایط قضیه مرسر نیز صدق کند. به عنوان مثال تابع هسته کرنل می‌تواند به فرم‌های زیر باشد:

$$K(x, y) = (x \cdot y + 1)^p, \quad p = 2, 3, \dots \quad \text{هسته چندجمله‌ای}$$

$$K(x, y) = \exp\left(-\frac{\|x - y\|^2}{2\delta^2}\right), \quad \text{هسته گاوسی}$$

$$K(x, y) = \tanh(\beta \cdot x \cdot y + \beta_1), \quad \text{هسته تانژانت هیپربولیک}$$

فصل ۲

شبکه عصبی

شبکه‌های عصبی مصنوعی یا شبکه‌های عصبی صناعی یا به زبان ساده‌تر شبکه‌های عصبی، سیستم‌ها و روش‌های محاسباتی نوین برای یادگیری ماشینی، نمایش دانش و در انتها اعمال دانش به دست آمده در جهت بیش‌بینی پاسخ‌های خروجی از سامانه‌های پیچیده هستند. ایده‌ی اصلی این گونه شبکه‌ها تا حدودی الهام‌گرفته از شیوه‌ی کارکرد سیستم عصبی زیستی برای پردازش داده‌ها و اطلاعات به منظور یادگیری و ایجاد دانش می‌باشد. عنصر کلیدی این ایده، ایجاد ساختارهایی جدید برای سامانه‌ی پردازش اطلاعات است. در این فصل به‌طور خلاصه به معرفی شبکه عصبی و کاربرد آن در بهینه‌سازی می‌پردازیم.

۱.۲ مقدمه‌ای بر شبکه‌های عصبی

۱.۱.۲ شبکه‌های عصبی بیولوژیکی

واحد اصلی دستگاه عصبی، سلولی خاص به‌نام نرون می‌باشد و بدون شک نحوه کار مغز و راز شعور آدمی در آن نهفته است. نرون‌ها با وجود تفاوت‌های زیاد از نظر اندازه و شکل ظاهری، مشخصه‌های مشترکی دارند. همان‌طور که در شکل (۱.۲) نشان داده شده است از تنه سلولی نرون تعدادی شاخک کوتاه خارج می‌شود که دندریت نام دارند. دندریت‌ها و تنه سلولی، سیگنال‌ها را از نرون‌های مجاور دریافت می‌کنند و از طریق یک لوله باریک به نام آکسون به نرون‌های دیگر منتقل می‌گردد.

آکسون در انتهای خود به تعدادی رشته جانبی باریک تقسیم می‌گردد که پایانه‌های آکسونی نام دارد و با دندریته‌های سایر نرون‌ها مرتبط است. ارتباط بین آکسون یک نرون و دندریته نرون دیگر را سیناپس می‌نامند. یک سیناپس مرکب از پایانه قبل از سیناپسی، شکاف سیناپسی و پایانه بعد از سیناپسی می‌باشد. همه سیگنال‌های جمع‌بندی شده، در تنه نرون ترکیب می‌شوند و اگر وسعت سیگنال‌های ترکیب شده به آستانه نرون برسد، مرحله تحریک شدن فعال می‌شود و یک سیگنال خروجی تولید می‌شود. این سیگنال به صورت یک پالس منفرد یا بخشی از پالس‌ها در یک میزان خاص به موازات آکسون به پایانه‌های سیناپسی انتقال می‌یابد و موجب ترشح موادی به نام عصب-رسانه می‌شود. عصب-رسانه در داخل شکاف سیناپسی پخش می‌شود و نرون بعدی را تحریک می‌کند و به این ترتیب یک سیگنال از یک نرون به دیگری انتقال می‌یابد. تعداد بسیار زیادی آکسون از نرون‌های مختلف با دندریته‌های یک نرون به این صورت ارتباط برقرار می‌کنند.

نرون‌ها با توجه به عملکردشان به سه دسته تقسیم می‌شوند:

(۱) نرون‌های حسی: این نرون‌ها پیام‌ها را از گیرنده‌ها به دستگاه عصبی مرکزی منتقل می‌کنند. گیرنده سلول خاصی است که در اندام‌های حسی، عضلات، پوست و ... قرار دارد و تغییرات فیزیکی یا شیمیایی را در می‌یابد، این تغییرات را به سیگنال‌های عصبی تبدیل می‌کند.

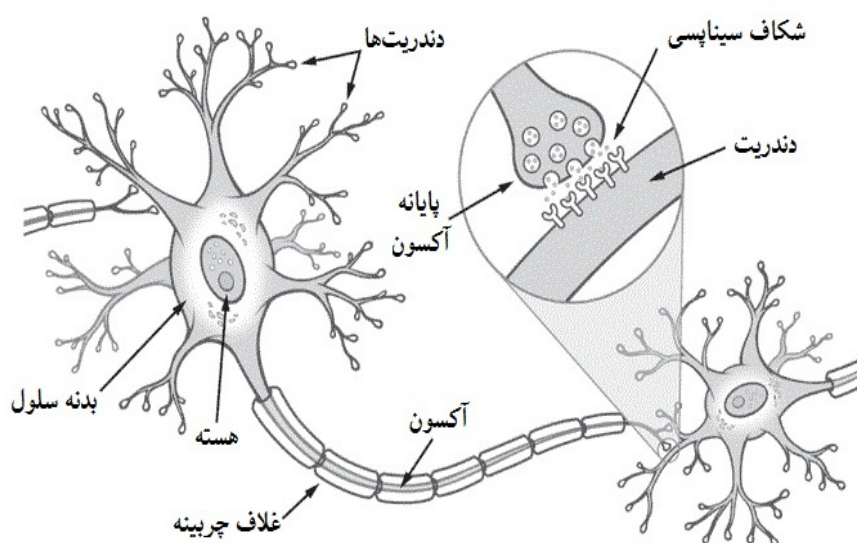
(۲) نرون‌های حرکتی: این نرون‌ها پیام‌ها را از مغز یا طناب نخاعی به اندام‌های عمل کننده می‌رسانند.

(۳) نرون‌های رابط: این نرون‌ها پیام‌ها را از نرون‌های حسی می‌گیرند و به یک نرون رابط دیگر انتقال می‌دهند و یا پیام‌ها را به یک نرون حرکتی می‌رسانند. نرون‌های رابط فقط در مغز، چشم و طناب نخاعی یافت می‌شوند.

حال منظور از یادگیری عبارت است از تغییری نسبتاً ساختاری در سیناپس (بعد از تحریک)، که این تغییر ساختاری باعث افزایش کارایی سیناپس می‌شود.

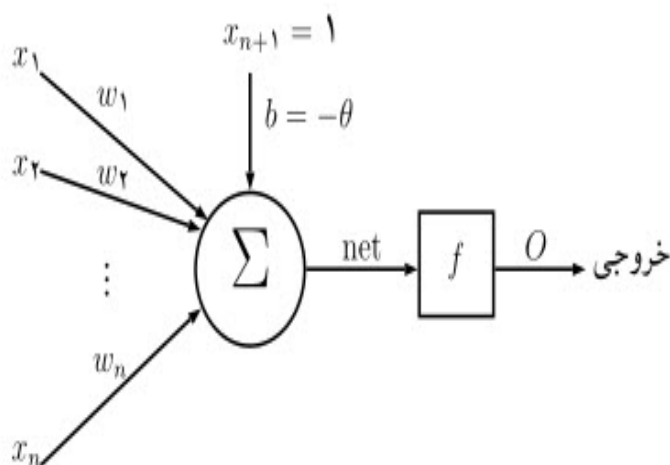
۲.۱.۲ شبکه‌های عصبی مصنوعی

بعد از اشاره‌ای گذرا بر شبکه عصبی بیولوژیکی و نشان دادن گوشه‌ای بسیار کوچک از این جهان مرموز و فوق‌العاده پیچیده به موضوع اصلی یعنی شبکه‌های عصبی مصنوعی می‌پردازیم. شبکه‌های عصبی مصنوعی مدل‌های سخت افزاری یا نرم افزاری الهام گرفته از ساختار و رفتار نرون‌های بیولوژیکی و سیستم‌های عصبی انسان هستند که شامل عناصر پردازشی (به نام نرون‌ها) و ارتباطات میان آنها با ضرایب (وزن‌ها) اختصاص یافته به ارتباطات است، که ساختار شبکه عصبی را تشکیل می‌دهند. در این رساله ما شبکه‌های عصبی مصنوعی را یک مدل



شکل ۱.۲: ساختار یک نرون.

ریاضی فرض خواهیم کرد و به جای استفاده از عبارت شبکه‌های عصبی مصنوعی به طور خلاصه عبارت شبکه‌های عصبی را به کار خواهیم برد.

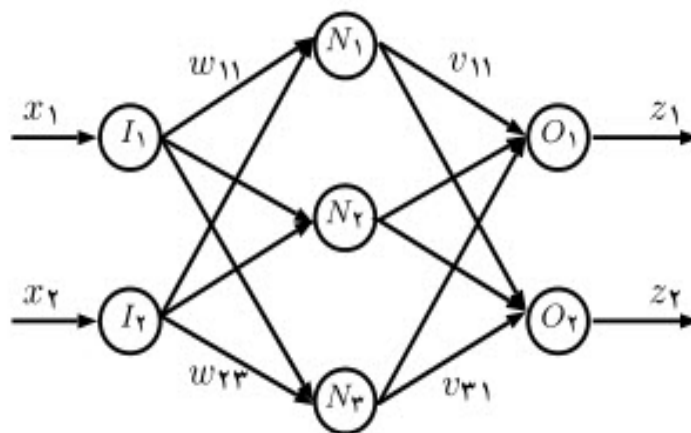


شکل ۲.۲: مدل ریاضی یک نرون.

مدل شبیه‌سازی شده یک نرون، در شکل (۲.۲) نمایش داده شده است. یک نرون در حالت کلی مجموعه‌ای از $n + 1$ ورودی x_j ($j = 1, \dots, n + 1$) دارد. این ورودی‌ها همان سیگنال‌های عصبی می‌باشند که از محیط یا از یک نرون دیگر فرستاده می‌شوند. w_j ($j = 1, \dots, n + 1$) که روی آکسون‌ها واقع هستند نقش ساختار سیناپسی را ایفا می‌کنند. این‌ها هستند که باید تغییر کنند تا فرآیند یادگیری اتفاق بیفتد حال این که این تغییر بر چه قاعده‌ای استوار باشد بستگی به الگوریتم آموزش یا تعلیم دارد. وزن آخرین ورودی یعنی $-\theta$ را با b نشان می‌دهند و آن را بایاس می‌نامند. سیگنال‌های عصبی x_1 و x_2 و ... و x_{n+1} به ترتیب در w_{n+1} ، ...، w_2 ، w_1 ضرب شده و مجموع این حاصل ضرب‌ها که در شکل آن را با net نشان می‌دهیم، وارد نرون می‌شوند. در این مرحله نرون همانند یک تابع روی net عمل می‌کند و

حاصل کار خود (برد تابع) را از طریق آکسون‌های طرف راست به نرون‌های دیگر می‌فرستد. این تابع را تابع فعال‌سازی می‌نامند.

البته خاطر نشان می‌کنیم که اگر نرون به‌عنوان تابع فرض شود این تابع از $\mathbb{R}^n$ به $\mathbb{R}$ تعریف می‌گردد. یعنی سیگنال خروجی از نرون در تمام آکسون‌های طرف راست یکسان می‌باشد که روی شکل با O نمایش داده‌ایم. همان‌طور که قبلاً اشاره شد نرون‌ها می‌توانند حسی، ارتباطی و یا حرکتی باشند. اطلاعات از طریق نرون‌های حسی دریافت می‌شود و از طریق نرون‌های ارتباطی انتقال می‌یابد و بالاخره نتیجه کار شبکه عصبی و یا واکنش شبکه عصبی به اطلاعات دریافت شده از طریق نرون‌های حرکتی به خارج منتقل می‌گردد (شکل ۳.۲).



شکل ۳.۲: ساختار یک نرون با لایه ورودی، لایه پنهان و لایه خروجی.
نرون‌های حرکتی نرون‌های ارتباطی نرون‌های حسی

نرون‌های حسی را لایه ورودی، نرون‌های حرکتی را لایه خروجی و نرون‌های ارتباطی را لایه پنهان نیز می‌گویند. لایه مخفی می‌تواند از چندین لایه تشکیل گردد البته در شکل (۳.۲) لایه مخفی تک لایه در نظر گرفته شده است.

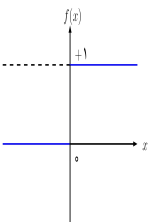
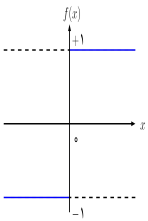
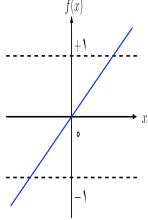
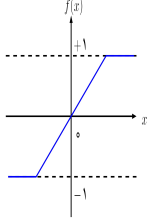
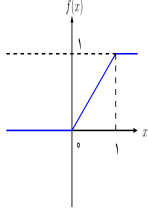
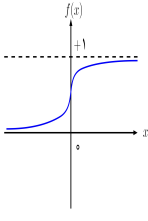
مشخصه‌های اصلی یک شبکه عصبی که در طراحی آن باید مدنظر قرار گیرند عبارت‌اند از

معماری شبکه عصبی: نحوه اتصالات بین نرون‌ها، تعداد آن‌ها و تعداد لایه‌های تشکیل دهنده بخش نرون‌های ارتباطی را معماری شبکه عصبی می‌گوییم.

تابع‌های فعال‌سازی: این‌که چه تابعی روی نرون قرار بگیرد تا بر ورودی‌های نرون اثر کرده و خروجی نرون را تولید کند تابع فعال‌سازی آن نرون گویند. در جدول (۱.۲) تعدادی از مهم‌ترین توابع فعال‌سازی نشان داده شده است.

الگوریتم آموزش: روشی که وزن‌های روی آکسون‌ها بر اساس آن تغییر می‌کنند (w_{ij} ها و v_{jk} ها در شکل (۳.۲)) تا خروجی شبکه عصبی به حالت مطلوب درآید را الگوریتم آموزش می‌گوییم.

جدول ۱.۲: برخی توابع فعال‌سازی نرون‌ها

ردیف	نام تابع	تعریف تابع	شکل تابع
۱	آستانه‌ای دو مقداره	$f(x) = \begin{cases} 0 & x < 0, \\ 1 & x \geq 0. \end{cases}$	
۲	آستانه‌ای دو مقداره متقارن	$f(x) = \begin{cases} -1 & x < 0, \\ 1 & x \geq 0. \end{cases}$	
۳	خطی	$f(x) = x.$	
۴	آستانه‌ای خطی متقارن	$f(x) = \begin{cases} -1, & x < -1, \\ x, & -1 \leq x \leq 1, \\ 1, & x > 1. \end{cases}$	
۵	آستانه‌ای خطی	$f(x) = \begin{cases} 0, & x < 0, \\ x, & 0 \leq x \leq 1, \\ 1, & x > 1. \end{cases}$	
۶	سیگموئید	$f(x) = \frac{1}{1+e^{-x}}.$	

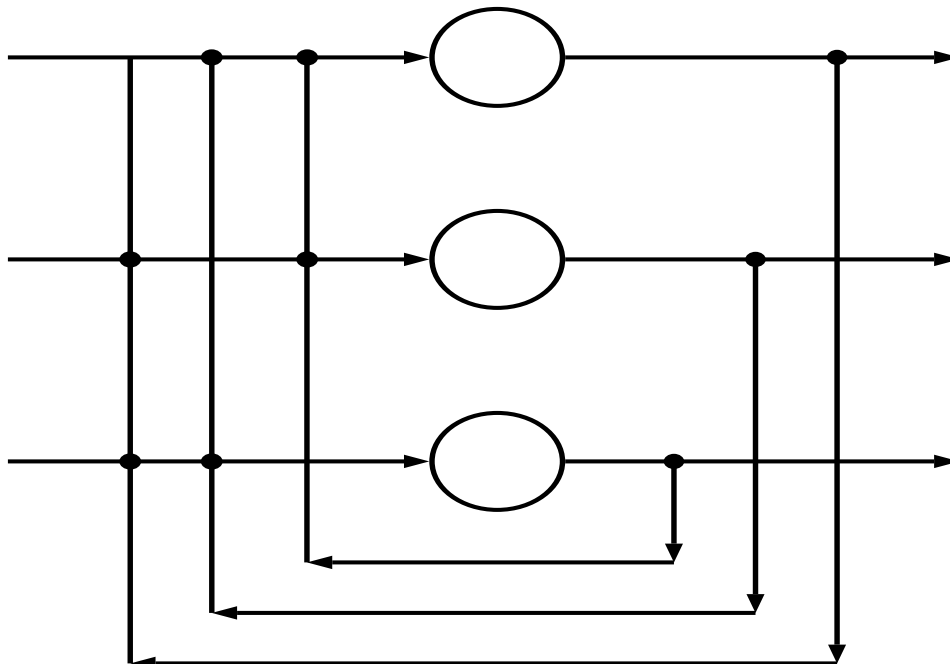
با توجه به مطالب قبل می‌توان برای شبکه‌های عصبی مصنوعی خصوصیات عمومی زیر را نسبت داد:

۱. پردازش اطلاعات در عناصر بسیار ساده به نام نرون (سلول عصبی) انجام می‌گیرد.

۲. سیگنال‌های عصبی بین سلول‌های عصبی از طریق اتصالات بین آن‌ها مبادله می‌شوند.
۳. به هر خط اتصال (آکسون) یک وزن نسبت داده می‌شود که در شبکه‌های عصبی معمولی این وزن‌ها در سیگنال‌های منتقل شده ضرب می‌گردند.
۴. هر نرون دارای یک تابع فعال‌سازی می‌باشد که این تابع در اکثر موارد غیرخطی است.

شبکه‌های عصبی بازگشتی

در شبکه‌های عصبی بازگشتی، حداقل یک سیگنال برگشتی از یک نرون به همان نرون یا نرون‌های همان لایه و یا لایه قبل وجود دارد. چنین شبکه‌ای حافظه‌ای را نگه می‌دارد و حالت بعد نه تنها به سیگنال‌های ورودی بلکه به حالات قبل شبکه نیز وابسته است. این شبکه‌های عصبی معمولاً در مدل ریاضی به صورت معادلات دیفرانسیل مرتبه اول ظاهر می‌شوند.



شکل ۴.۲: شبکه‌های عصبی دو لایه بازگشتی.

۳.۱.۲ تاریخچه تکامل شبکه عصبی

در سال ۱۹۴۳ یک نرولوژیست به نام وارن مک کلاچ^۱ به همراه یک متخصص آمار به نام والتر پیتز^۲ اولین مدل ریاضی نرون را ارائه دادند [۲۹]. این نرون ساده یک عنصر محاسبه‌گر بود که عناصری به نام ورودی را در مقادیر ثابتی به نام وزن ضرب می‌کرد و آنها را از یک عملگر خطی عبور می‌داد و حاصل خروجی نرون را تشکیل می‌داد. در این مدل وزن‌ها مقادیر ثابتی

¹Warren McCulloch

²Walter Pitts

بودند. در سال ۱۹۴۹ دونالد هب^۱ [۶] کتابی درباره نحوه یادگیری در مغز انسان نوشت. بر این مبنا قانون یادگیری برای نرون را ارائه داد که به قانون هب معروف است و بعدها در تعلیم شبکه‌های عصبی از آن استفاده شد. در سال ۱۹۵۴ فارلی و کلارک^۲ [۴] مدلی از نرون‌ها را به‌طور تصادفی به هم وصل کردند و قانون هب را پیاده‌سازی نمودند و نشان دادند که می‌توان دو الگوی ورودی را با این قانون تعلیم از هم تشخیص داد. با پیشرفت علم کامپیوتر اولین مدل از نرون عصبی مصنوعی توسط ناتانیال روچستر^۳ محقق شرکت IBM در سال ۱۹۵۸ با استفاده از کامپیوتر شبیه‌سازی شد. روزنبلات^۴ [۴۲، ۴۳] در سال ۱۹۵۸ نرون ساده مک‌کلاچ پیتر را اصلاح کرد و به آن قابلیت یادگیری و سازگاری را اضافه کرد و این نرون را پرسپترون نامید. قانون تعلیم پرسپترون اولین قانون رسمی برای شبکه‌های عصبی است.

برنارد ویدرو^۵ و مارسین هاف^۶ و همکارانشان یک مدل سه سطحی از نرون‌ها را ابداع کردند که آدلاین نام‌گذاری شد و مختصر شده حروف "عنصر خطی تطبیقی" بود [۵۲-۵۴] و برای آدلاین یک قانون تعلیم بر مبنای مشتق‌گیری بیان نمودند. این قانون تعلیم به $\alpha - LMS$ معروف است. از متصل کردن چند آدلاین به یکدیگر مادلاین ساخته شد. ویدرو و هاف از شبکه عصبی مادلاین برای حذف اکو از خطوط تلفن استفاده نمودند.

سال ۱۹۶۹ آغاز یک دوره رکود برای شبکه عصبی بود. پاپرت و مینسکی^۷ [۳۱] کتابی به نام پرسپترون نوشتند و نشان دادند که اگرچه گیت‌های AND و OR را می‌توان با استفاده از پرسپترون پیاده‌سازی کرد ولی گیت XOR قابل پیاده‌سازی نمی‌باشد و لذا نمی‌توان از شبکه عصبی برای پردازش اطلاعات استفاده کرد. الگوریتم پس انتشار پال ورباس^۸ [۵۱] در سال ۱۹۷۴ ارائه شد و بعدها توسط رامل هارت^۹ [۴۴] به‌طور مستقل کشف گردید. این الگوریتم از زمان پیدایش به‌طور گسترده به‌عنوان یک الگوریتم آموزش در شبکه‌های عصبی پیش‌خورد مورد استفاده قرار گرفته است. سال ۱۹۸۲ تا ۱۹۸۶ را می‌توان تولد دوباره شبکه عصبی دانست. در این سال‌ها دو اتفاق مهم در این زمینه رخ داد. اول ارائه شبکه عصبی بازگشتی و مفهوم تابع انرژی توسط جان هاپفیلد^{۱۰} [۲۰] بود که در سال ۱۹۸۲ منتشر شد. دومین اتفاق، حل مسأله فروشنده دوره‌گرد (TSP) توسط شبکه هاپفیلد بود [۲۱]. هاپفیلد نشان داد که از این شبکه بازگشتی می‌توان برای حل مسائل بهینه‌سازی استفاده کرد و افق جدیدی را در شبکه‌های عصبی مصنوعی گشود.

¹Donald Hebb

²Farley and Clark

³Nathanial Rochester

⁴Frank Rosenblatt

⁵Bernard Widrow

⁶Marcian Hoff

⁷Papert and Minsky

⁸Paul Werbos

⁹Rumelhart

¹⁰John Hopfield

۴.۱.۲ تاریخچه حل مسائل بهینه‌سازی با استفاده از شبکه‌های عصبی

طیف گسترده‌ای از مسائل مهم در رشته‌های علوم پایه و مهندسی از جمله کنترل بهینه، طراحی ساختار بهینه، پردازش تصویر، تقریب توابع، تحلیل رگرسیون و... قابل تبدیل به مسائل بهینه‌سازی خطی و غیرخطی هستند. در دهه‌های گذشته الگوریتم‌های بهینه‌سازی عددی فراوانی برای حل مسائل بهینه‌سازی خطی و غیرخطی ارائه شده‌اند. از آن‌جا که زمان مورد نیاز برای حل مسائل بهینه‌سازی و به‌خصوص مسائل بهینه‌سازی غیرخطی بسیار وابسته به بعد و ساختار مسأله است بنابراین الگوریتم‌های عددی کارایی کمی از خود نشان می‌دهند. یک رهیافت امیدوار کننده برای حل مسائل بهینه‌سازی با ابعاد بالا استفاده از شبکه‌های عصبی مصنوعی می‌باشد.

در سال ۱۹۸۶ تانک و هاپفیلد [۴۶] اولین شبکه عصبی را برای حل مسائل برنامه‌ریزی خطی معرفی کردند و نشان دادند که وضعیت شبکه در هر تکرار به‌گونه‌ای تغییر می‌کند که تابع انرژی متناظر با آن به‌طور یکنواخت کاهش می‌یابد تا جایی که به نقطه مینیمم خود می‌رسد و این نقطه مینیمم متناظر با نقطه تعادل شبکه عصبی می‌باشد. آن‌ها این شبکه عصبی را توسط یک مدار الکتریکی پیاده‌سازی کردند. هم‌چنین از این شبکه برای حل مسأله فروشنده دوره‌گرد با ۳۰ شهر استفاده نمودند. این شبکه دارای نقص‌هایی بود به‌خصوص این که نقطه تعادل شبکه در شرایط کاروش کان تاکر^۱ صدق نمی‌کرد و لذا جواب مطلوبی از مسأله حاصل نمی‌شد. با این وجود کارهای هاپفیلد انگیزه بسیار خوبی را برای محققین به‌وجود آورد تا در این زمینه فعالیت کنند. کندی و چا^۲ [۲۲] با افزودن یک پارامتر جریمه متناهی به شبکه هاپفیلد آن را توسعه دادند و از آن برای حل مسأله برنامه‌ریزی غیرخطی محدب استفاده نمودند. پارامترهای جریمه این شبکه در یافتن نقطه بهینه دقیق ناتوان است به‌خصوص اگر پارامتر جریمه بزرگ باشد این شبکه به‌سختی عمل می‌کند.

برای جلوگیری از بکار بردن پارامتر جریمه یک شبکه عصبی سوئیچ-خازن توسط رودریگز-وازکز^۳ و همکاران [۴۱] معرفی شد. هم‌چنین ما و شانبلات^۴ [۲۷] یک شبکه عصبی دو فازی ارائه کردند که فاز اول شبکه مشابه شبکه کندی و چا بود و در فاز دوم آن مسیر شبکه به جواب دقیق مسأله همگرا می‌شد بنابراین این روش جواب‌های دقیق‌تری نسبت به شبکه کندی و چا ارائه می‌داد. مشکل این شبکه در این بود که پایداری فاز دوم شبکه بستگی به انتخاب یک مقدار بزرگ از پارامتر جریمه داشت و لذا اگرچه تأثیر پارامتر جریمه در این روش کاهش یافته بود و جواب‌های حاصل دقیق‌تر بودند اما هنوز مستقل از پارامتر نبود.

ژانگ و کنستانتینیدس^۵ [۶۸] بر مبنای روش لاگرانژ یک شبکه عصبی ارائه کردند که کاملاً مستقل از پارامتر جریمه و قادر به حل مسائل غیرخطی بود. نقطه تعادل این شبکه عصبی در

^۱Karush-Kuhn-Tucker (KKT)

^۲Kennedy and Chua

^۳Rodriguez-Vazquez

^۴Maa and Shanblatt

^۵Zhang and Constantinides

شرایط بهینگی مرتبه اول و دوم صدق می‌کرد و همچنین شبکه حاصل همگرا بود. در سال ۱۹۹۳ بوزردوم و پتیسن^۱ [۱۲] شبکه‌ای را بر مبنای گرادیان و روش تصویر و مستقل از پارامتر جریمه ابداع کردند که تنها قادر به حل مسائل درجه دوم با متغیرهای کران‌دار بود. این روش در عمل روش کارایی بود اما نمی‌توانست مسائل کلی برنامه‌ریزی خطی و درجه دوم را حل کند.

زیا^۲ و همکاران [۵۷، ۶۲-۶۶] چندین مدل را برای حل مسائل بهینه‌سازی غیرخطی با قيود خطی و غیرخطی ارائه کردند. همچنین چندین مدل شبکه عصبی برای حل مسائل بهینه‌سازی غیرخطی توسط عفتی و همکاران [۱۴، ۱۵، ۱۷، ۳۲، ۳۳] ارائه شده است.

¹Bouzerdoum and Pattison

²Xia

فصل ۳

حل مسئله رگرسیون بردار پشتیبان با استفاده شبکه عصبی

۱.۳ رگرسیون چیست؟

تاریخچه: واژه رگرسیون در فرهنگ لغت به معنی «بازگشت» است و اغلب جهت رساندن مفهوم «بازگشت به یک مقدار متوسط یا میانگین» به کار می‌رود. به این معنی که برخی پدیده‌ها به مرور زمان از نظر کمی به طرف یک مقدار متوسط میل می‌کند.

بیش از ۱۰۰ سال پیش در سال ۱۸۷۷ فرانسیس گالتون^۱ در مقاله‌ای که در همین زمینه منتشر کرد اظهار داشت که متوسط قد پسران دارای پدران قدبلند، کمتر از قد پدرانشان می‌باشد. به نحو مشابه متوسط قد پسران دارای پدران کوتاه قد نیز بیشتر از قد پدرانشان گزارش شده است. به این ترتیب گالتون پدیده بازگشت به طرف میانگین را در داده‌هایش مورد تاکید قرار داد. برای گالتون رگرسیون مفهومی زیست شناختی داشت اما کارهای او توسط کارل پیرسون^۲ برای مفاهیم آماری توسعه داده شد. گرچه گالتون برای تاکید بر پدیده «بازگشت به سمت مقدار متوسط» از تحلیل رگرسیون استفاده کرد، اما به هر حال امروزه واژه تحلیل رگرسیون جهت اشاره به مطالعات مربوط به روابط بین متغیرها به کار برده می‌شود.

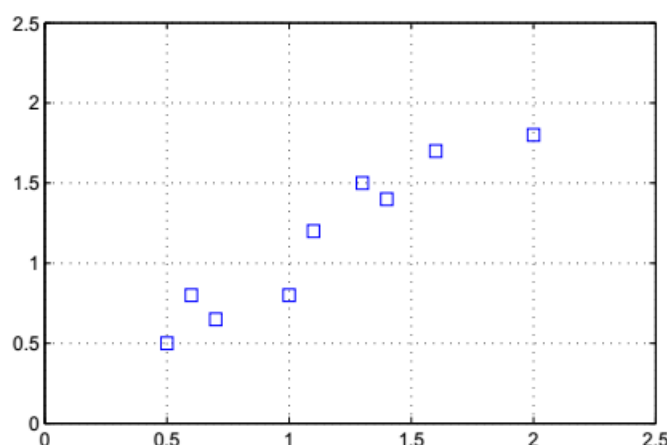
^۱Francis Galton

^۲Karl Pearson

۱.۱.۳ نمودار پراکندگی

در حقیقت تحلیل رگرسیونی فن و تکنیکی آماری برای بررسی و مدل سازی و ارتباط بین متغیرهاست. رگرسیون تقریباً در هر زمینه‌ای از جمله مهندسی، فیزیک، اقتصاد، مدیریت، علوم زیستی، بیولوژی و علوم اجتماعی برای برآورد و پیش بینی مورد نیاز است. می توان گفت تحلیل رگرسیونی، پرکاربردترین روش در بین تکنیک های آماری است. خلاصه ای از تحلیل رگرسیونی به صورت زیر است:

در ابتدا تحلیل گر حدس می زند که یک رابطه به شکل یک خط بین دو متغیر وجود دارد و سپس به جمع آوری اطلاعات کمی از دو متغیر می پردازد و این داده ها را به صورت نقاطی در یک نمودار دو بعدی رسم می کند (شکل (۱.۳)).



شکل ۱.۳: نمودار پراکندگی داده ها

این نمودار که به آن نمودار پراکندگی^۱ گفته می شود نقش بسیار مهمی در تحلیل های رگرسیونی و نمایش ارتباط بین متغیرها ایفا می کند. در صورتی که نمودار نشان دهنده این باشد که داده ها تقریباً (نه لزوماً دقیق) در امتداد یک خط مستقیم پراکنده شده اند، حدس تحلیل گر تایید شده (شکل (۲.۳)) و این ارتباط خطی به صورت زیر نمایش داده می شود:

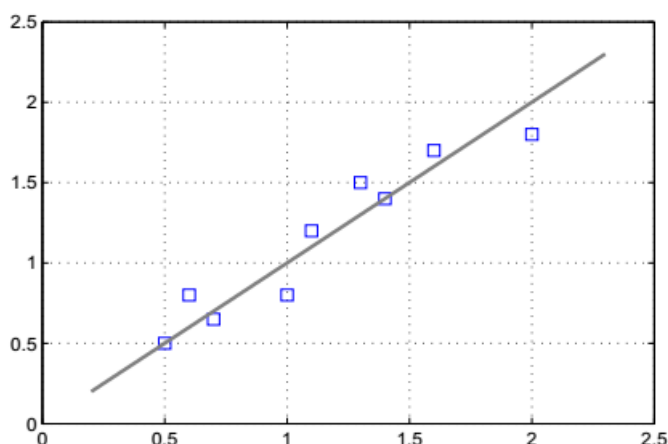
$$y = ax + b,$$

که در آن b عرض از مبدأ و a شیب این خط است.

۲.۱.۳ متغیرها و خطا

بین برخی از نقاط و تصویر آن ها بر روی خط رگرسیونی کمی تفاوت به چشم می خورد که از آن به عنوان خطای برآورد یاد می کنیم. بنابراین معادله اولیه را به صورت زیر اصلاح می کنیم:

^۱Scatter Plot



شکل ۲.۳: برآورد خط رگرسیونی

$$y = ax + b + \epsilon,$$

معادله فوق یک مدل رگرسیونی خطی نامیده می‌شود.

۳.۱.۳ روش‌های رگرسیونی

تا این مرحله مدل رگرسیونی معرفی شد و کافی است پارامترهای مجهول a و b برآورد شوند. برآورد پارامترها در مدل‌سازی با استفاده از روش‌های مختلف انجام می‌شود از جمله روش کم‌ترین مربعات خطا. روش کم‌ترین مربعات خطا که یکی از روش‌های مورد استفاده در تحلیل رگرسیونی است، اولین بار به صورت واضح و مختصر توسط لژاندر^۱ ریاضی‌دان فرانسوی در سال ۱۸۰۵ و گاوس^۲ ریاضی‌دان آلمانی در سال ۱۸۰۹ معرفی و در مطالعات نجومی به کار برده شد.

روش‌های ماشین بردار پشتیبان که در فصل اول آن را توضیح دادیم، در رگرسیون بسیار موفقیت آمیز عمل کرده است. این روش را رگرسیون بردار پشتیبان (SVR) می‌نامند.

۲.۳ استفاده از SVM در رگرسیون

۱.۲.۳ مساله رگرسیون خطی بردار پشتیبان

برای نمونه‌های آموزشی داده شده $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) \in X \times \mathbb{R}$ که $x_i \in X$ نقطه‌ای در فضای ورودی و y_i متناظر مقدار تابع $f(x_i)$ است. هدف پیدا کردن تابع تخمینی $f(x)$ است

^۱Legendre

^۲Gauss

که برای همه نمونه‌های آموزشی x_i ، حداکثر فاصله آن از مقدار هدف y_i ، ϵ باشد. در رگرسیون ϵ -بردار پشتیبان، هر انحراف که کمتر از ϵ است، مجاز می‌باشد.

تابع خطی $f(x) = \langle w, x \rangle + b$ را در نظر بگیرید که در آن $b \in \mathbb{R}$ و $w \in \mathbb{R}^n$. هدف SVR می‌نیمم سازی نرم اقلیدسی w است به قسمی که انحراف $f(x)$ از مقادیر خروجی y_i کمتر از ϵ باشد. این مساله بهینه سازی محدب به صورت زیر بیان می‌شود:

$$\text{minimize } \frac{1}{2} \|w\|^2, \quad (1.3)$$

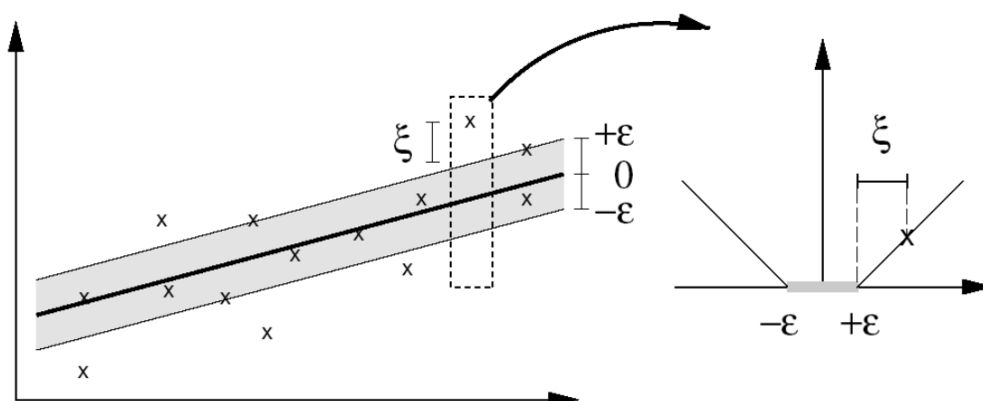
$$\text{subject to } \begin{cases} y_i - \langle w, x_i \rangle - b \leq \epsilon, & i = 1, \dots, n, \\ \langle w, x_i \rangle + b - y_i \leq \epsilon, & i = 1, \dots, n. \end{cases} \quad (2.3)$$

در مساله (۱.۳) و (۲.۳) فرض ضمنی بر این است که تابع $f(x)$ که همه اطلاعات ورودی $\{(x_i, y_i)\}_{i=1}^n$ را تخمین می‌زند، موجود است. برای ساختن یک ماشین بردار پشتیبان که هدف d را تخمین می‌زند، از تابع زیان زیر استفاده می‌کنیم:

$$L_\epsilon(d, y) = \begin{cases} |d - y| - \epsilon, & \text{اگر } |d - y| \geq \epsilon \\ 0, & \text{در غیر این صورت} \end{cases} \quad (3.3)$$

که ϵ پارامتری از پیش تعیین شده و مثبت است. تابع زیان (۳.۳) تابع زیان ϵ - حساسیت نامیده می‌شود.

این تعریف نتیجه می‌دهد که تابع زیان تنها زمانی دارای مقدار است که انحراف خروجی y از هدف مطلوب d بزرگ‌تر از پارامتر انحراف ϵ باشد، که به طور شهودی در شکل (۳.۳) نشان داده شده است.



شکل ۳.۳: تابع زیان ϵ - حساسیت

برای شذنی بودن محدودیت‌های (۲.۳)، متغیرهای کمکی ζ_i و ζ'_i معرفی می‌شوند تا مساله بهینه‌سازی را به شکل زیر تشکیل دهیم:

$$\text{minimize } \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n (\zeta_i + \zeta'_i), \quad (4.3)$$

$$\text{subject to } \begin{cases} y_i - \langle w, x_i \rangle - b \leq \epsilon + \zeta_i, & i = 1, \dots, n, \\ \langle w, x_i \rangle + b - y_i \leq \epsilon + \zeta'_i, & i = 1, \dots, n, \\ \zeta_i, \zeta'_i \geq 0, & i = 1, \dots, n. \end{cases} \quad (5.3)$$

که در آن ϵ پارامتر انحراف و C پارامتر جریمه است. برای ساختن دوگان مساله بهینه‌سازی (۴.۳) و (۵.۳) متغیرهای کمکی (α_i, α'_i) را معرفی می‌کنیم. α_i متناظر محدودیت‌های شامل ζ_i و α'_i متناظر محدودیت‌های کمکی ζ'_i است. این متغیرها در حقیقت ضرایب لاگرانژ برای مساله اولیه هستند. مشابه روش‌های به کار رفته در طراحی رده‌بندی کننده‌های SV ، مساله بهینه‌سازی مقید (۴.۳) و (۵.۳) را با تشکیل تابع لاگرانژین به صورت زیر حل می‌کنیم.

$$\begin{aligned} L_P(w, b, \zeta_i, \zeta'_i, \alpha_i, \alpha'_i, \beta_i, \beta'_i) = & \frac{1}{2} w^T w + C \sum_{i=1}^n (\zeta_i + \zeta'_i) - \sum_{i=1}^n (\beta_i \zeta_i + \beta'_i \zeta'_i) - \\ & \sum_{i=1}^n \alpha_i [w^T x_i + b - y_i + \epsilon + \zeta_i] - \\ & \sum_{i=1}^n \alpha'_i [y_i - w^T x_i - b + \epsilon + \zeta'_i]. \end{aligned}$$

تابع لاگرانژین $L_P(w, b, \zeta_i, \zeta'_i, \alpha_i, \alpha'_i, \beta_i, \beta'_i)$ باید نسبت به متغیرهای اولیه w, b, ζ_i و ζ'_i می‌نیم و نسبت به ضرایب لاگرانژ نامنفی $\alpha_i, \alpha'_i, \beta_i$ و β'_i ماکزیمم شود. بنابراین تابع لاگرانژین در نقطه جواب بهینه مساله (۴.۳) و (۵.۳) یعنی $(w^*, b^*, \zeta_i^*, \zeta'_i^*)$ دارای نقطه زینی است. در جواب بهینه مشتقات جزئی نسبت به متغیرهای اولیه صفر می‌شود، یعنی :

$$\begin{aligned} \frac{\partial L_P}{\partial w} = 0 & \Rightarrow w - \sum_{i=1}^n (\alpha_i - \alpha'_i) x_i = 0, \\ \frac{\partial L_P}{\partial b} = 0 & \Rightarrow \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0, \\ \frac{\partial L_P}{\partial \zeta_i} = 0 & \Rightarrow C - \alpha_i - \beta_i = 0, \quad i = 1, \dots, n, \\ \frac{\partial L_P}{\partial \zeta'_i} = 0 & \Rightarrow C - \alpha'_i - \beta'_i = 0. \quad i = 1, \dots, n. \end{aligned} \quad (6.3)$$

با استفاده از شرایط KKT دوگان مساله بهینه‌سازی برای رگرسیون خطی با استفاده از SVM به صورت زیر است:

$$\begin{aligned} \text{maximize } & -\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n (\alpha_i - \alpha'_i)(\alpha_j - \alpha'_j) x_i^T x_j + \sum_{i=1}^n y_i (\alpha_i - \alpha'_i) - \epsilon \sum_{i=1}^n (\alpha_i + \alpha'_i), \quad (7.3) \\ \text{subject to } & \begin{cases} \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0, \\ \alpha_i \in [0, C], \quad i = 1, \dots, n, \\ \alpha'_i \in [0, C], \quad i = 1, \dots, n. \end{cases} \quad (8.3) \end{aligned}$$

توجه شود که مساله (۷.۳) و (۸.۳) تنها به صورت عباراتی از ضرایب لاگرانژ α و α' بیان شده است.

مرحله‌ی آموزش n جفت ضریب لاگرانژ (α_i, α'_i) را نتیجه می‌دهد. بعد از آموزش تعداد پارامترهای ناصفر α_i یا α'_i مساوی تعداد بردار پشتیبان است. از آن جا که حداقل یک مولفه از هر جفت (α_i, α'_i) ($i = 1, \dots, n$) صفر است لذا حاصل ضرب α_i و α'_i همیشه صفر است، یعنی $\alpha_i \alpha'_i = 0$.

در بهینگی شرایط KKT زیر برقرار هستند:

$$\begin{aligned}\alpha_i(w^T x_i + b - y_i + \epsilon + \zeta_i) &= 0, \\ \alpha'_i(-w^T x_i - b + y_i + \epsilon + \zeta_i) &= 0, \\ \beta_i \zeta_i &= (C - \alpha_i) \zeta_i = 0, \\ \beta'_i \zeta'_i &= (C - \alpha'_i) \zeta'_i = 0.\end{aligned}\tag{۹.۳}$$

$$(۱۰.۳)$$

شرایط فوق نشان می‌دهد که برای $0 < \alpha_i < C$ ، شرط $\zeta_i = 0$ برقرار است. به طور مشابه اگر $0 < \alpha'_i < C$ آن گاه $\zeta'_i = 0$ و برای $0 < \alpha_i, \alpha'_i < C$ داریم:

$$\begin{aligned}\alpha_i(w^T x_i + b - y_i + \epsilon) &= 0, \\ \alpha'_i(-w^T x_i - b + y_i + \epsilon) &= 0.\end{aligned}$$

شرایط (۹.۳) و (۱۰.۳) در مورد داده‌هایی هستند که دارای انحراف بیشتر از ϵ نسبت به خط رگرسیون هستند. یعنی اگر $\zeta_i > 0$ آن گاه $\alpha_i = C$ و اگر $\zeta'_i > 0$ آن گاه $\alpha'_i = C$. به عبارتی دیگر برای داده‌هایی که بالای محدوده‌ی با خطای مجاز هستند، $\alpha_i = C$ و برای داده‌هایی که پایین محدوده با خطای مجاز هستند، $\alpha'_i = C$ ، این نقاط را بردارهای پشتیبان مقید می‌نامند. همچنین برای همه نقاط تمرینی داخل ناحیه مجاز، یا جایی که $|y - f(x)| < \epsilon$ هر دوی α_i و α'_i مساوی صفر هستند. لذا هیچ یک از آن‌ها بردار پشتیبان نیستند و تاثیری در صفحه تصمیم ندارند.

اما پس از حل مساله بهینه سازی (۷.۳) و (۸.۳) با داشتن ضرایب بهینه α_i و α'_i ، با استفاده از شرط بهینگی (۶.۳) بردار وزنی بهینه برای ابرصفحه رگرسیون به صورت زیر به دست می‌آید:

$$w_0 = \sum_{i=1}^n (\alpha_i - \alpha'_i) x_i = 0, \tag{۱۱.۳}$$

و ابرصفحه بهینه به صورت زیر به دست می‌آید:

$$f(x) = w_0 x + b = \sum_{i=1}^n (\alpha_i - \alpha'_i) x_i^T x + b.$$

۲.۲.۳ مساله رگرسیون غیرخطی بردار پشتیبان

حل مسائل رگرسیون غیرخطی بیشتر مورد توجه قرار می‌گیرد. یک تعمیم به حالت غیرخطی مشابه روش توسعه رده‌بندی کننده غیرخطی که در (۳.۲.۱) بیان شده است، استفاده از رده‌بندی کننده غیرخطی است، یعنی به وسیله یک نگاشت به فضای ویژگی (که معمولاً دارای بعد خیلی بالایی است) داده‌ها را به فضایی با قابلیت رگرسیون خطی می‌بریم و سپس تابع رگرسیون غیرخطی را متناظر معادله رگرسیون خطی در فضای ویژگی به دست می‌آوریم. بنابراین ابر صفحه بهینه رگرسیون بردار پشتیبان در حالت غیرخطی با حل مساله زیر به دست می‌آید:

$$\text{minimize } \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n (\zeta_i + \zeta'_i), \quad (12.3)$$

$$\text{subject to } \begin{cases} y_i - \langle w, \Phi(x_i) \rangle - b \leq \epsilon + \zeta_i, & i = 1, \dots, n, \\ \langle w, \Phi(x_i) \rangle + b - y_i \leq \epsilon + \zeta'_i, & i = 1, \dots, n, \\ \zeta_i, \zeta'_i \geq 0, & i = 1, \dots, n. \end{cases} \quad (13.3)$$

که در آن ϵ پارامتر انحراف و C پارامتر جریمه است. برای ساختن دوگان مساله بهینه‌سازی (۱۲.۳) و (۱۳.۳) متغیرهای کمکی (α_i, α'_i) را معرفی می‌کنیم. α_i متناظر محدودیت‌های شامل ζ_i و α'_i متناظر محدودیت‌های کمکی ζ'_i است. این متغیرها در حقیقت ضرایب لاگرانژ برای مساله اولیه هستند. مشابه روش‌های به کار رفته در رگرسیون خطی، مساله بهینه‌سازی مقید (۱۲.۳) و (۱۳.۳) را با تشکیل لاگرانژین متغیرهای اولیه حل می‌کنیم که تابع لاگرانژین به صورت زیر است:

$$\begin{aligned} L_P(w, b, \zeta_i, \zeta'_i, \alpha_i, \alpha'_i, \beta_i, \beta'_i) = & \frac{1}{2} w^T w + C \sum_{i=1}^n (\zeta_i + \zeta'_i) - \sum_{i=1}^n (\beta_i \zeta_i + \beta'_i \zeta'_i) - \\ & \sum_{i=1}^n \alpha_i [w^T \Phi(x_i) + b - y_i + \epsilon + \zeta_i] - \\ & \sum_{i=1}^n \alpha'_i [y_i - w^T \Phi(x_i) - b + \epsilon + \zeta'_i]. \end{aligned}$$

تابع لاگرانژین $L_P(w, b, \zeta_i, \zeta'_i, \alpha_i, \alpha'_i, \beta_i, \beta'_i)$ باید نسبت به متغیرهای اولیه w, b, ζ_i و ζ'_i می‌نیم شود و نسبت به ضرایب لاگرانژ نامنفی $\alpha_i, \alpha'_i, \beta_i$ و β'_i ماکزیمم گردد. بنابراین تابع لاگرانژین در نقطه جواب بهینه مساله (۱۲.۳) و (۱۳.۳) یعنی $(w^*, b^*, \zeta_i^*, \zeta_i'^*)$ دارای نقطه زینی است.

در جواب بهینه مشتقات جزئی نسبت به متغیرهای اولیه صفر می شود، یعنی :

$$\frac{\partial L_P}{\partial w} = 0 \Rightarrow w - \sum_{i=1}^n (\alpha_i - \alpha'_i) \Phi(x_i) = 0, \quad (14.3)$$

$$\frac{\partial L_P}{\partial b} = 0 \Rightarrow \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0,$$

$$\frac{\partial L_P}{\partial \zeta_i} = 0 \Rightarrow C - \alpha_i - \beta_i = 0, \quad i = 1, \dots, n,$$

$$\frac{\partial L_P}{\partial \zeta_i} = 0 \Rightarrow C - \alpha_i - \beta_i = 0. \quad i = 1, \dots, n.$$

بنابراین مشابه حالت خطی دوگان مساله بهینه سازی برای رگرسیون غیرخطی با استفاده از SVR به صورت زیر است:

$$\begin{aligned} \text{maximize} \quad & -\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n (\alpha_i - \alpha'_i)(\alpha_j - \alpha'_j) \Phi(x_i)^T \Phi(x_j) + \sum_{i=1}^n y_i (\alpha_i - \alpha'_i) - \\ & \epsilon \sum_{i=1}^n (\alpha_i + \alpha'_i), \end{aligned} \quad (15.3)$$

$$\text{subject to} \quad \begin{cases} \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0, \\ \alpha_i \in [0, C], \quad i = 1, \dots, n, \\ \alpha'_i \in [0, C], \quad i = 1, \dots, n. \end{cases} \quad (16.3)$$

از شرط بهینگی (۱۴.۳)، w به صورت زیر به دست می آید:

$$w = \sum_{i=1}^n (\alpha_i - \alpha'_i) \Phi(x_i).$$

بنابراین تابع تخمین $f(x) = \langle w, x \rangle + b$ به صورت زیر بیان می شود:

$$f(x) = \sum_{i=1}^n (\alpha_i - \alpha'_i) \Phi(x_i)^T \Phi(x) + b.$$

می توان مساله (۱۵.۳) و (۱۶.۳) را به شکل خلاصه زیر نوشت،

$$\text{minimize} \quad \frac{1}{2} \theta^T Q \theta + d^T \theta,$$

$$\text{subject to} \quad \begin{cases} A\theta - b \leq 0, \\ e^T \theta = 0. \end{cases}$$

که در آن،

$$Q_{n \times n} = \begin{cases} \Phi(X_i)^T \Phi(X_j) & 1 \leq i, j \leq n \\ 0 & \text{در غیر این صورت} \end{cases},$$

$$d = (-y_1, -y_2, \dots, -y_n, \underbrace{\epsilon, \epsilon, \dots, \epsilon}_n, \underbrace{0, 0, \dots, 0}_{2n})^T,$$

$$e = (\underbrace{1, 1, \dots, 1}_n, \underbrace{0, 0, \dots, 0}_{2n})^T \in \mathbb{R}^{4n},$$

$$A = \begin{pmatrix} I_{n \times 2n} & O_{n \times 2n} \\ O_{2n \times 2n} & -I_{2n \times 2n} \end{pmatrix},$$

$$b = (\underbrace{C, C, \dots, C}_n, \underbrace{2C, 2C, \dots, 2C}_n, \underbrace{C, C, \dots, C}_n, \underbrace{0, 0, \dots, 0}_n)^T \in \mathbb{R}^{4n},$$

$$\theta = (\alpha_1 - \alpha'_1, \alpha_2 - \alpha'_2, \dots, \alpha_n - \alpha'_n, \alpha_1 + \alpha'_1, \alpha_2 + \alpha'_2, \dots, \alpha_n + \alpha'_n, \alpha_1 - \alpha'_1, \alpha_2 - \alpha'_2, \dots, \alpha_n - \alpha'_n, \alpha_1 + \alpha'_1, \alpha_2 + \alpha'_2, \dots, \alpha_n + \alpha'_n)^T.$$

در ادامه این فصل با استفاده از شبکه عصبی به حل مساله رگرسیون بردار پشتیبان می‌پردازیم. ثابت می‌شود که نقطه تعادل شبکه عصبی مورد نظر با نقطه بهینه مساله درجه دوم رگرسیون بردار پشتیبان برابر است. وجود داشتن و منحصر به فرد بودن نقطه تعادل ثابت شده و با ساخت تابع لیاپانوف مناسب پایداری مجانبی اثبات می‌شود. در ادامه نیز با ارائه مثال‌های عددی کارایی این شبکه بررسی می‌گردد.

۳.۳ معرفی شبکه عصبی

مساله برنامه‌ریزی غیرخطی زیر را در نظر بگیرید. در این رساله فرض بر این است که این مساله دارای جواب یکتا است.

$$\text{minimize } \frac{1}{2} \theta^T Q \theta + d^T \theta, \quad (17.3)$$

$$\text{subject to } \begin{cases} A\theta - b \leq 0, \\ e^T \theta = 0. \end{cases} \quad (18.3)$$

در این بخش می‌خواهیم یک شبکه عصبی با کارایی بالا را معرفی کرده و در مورد ویژگی‌های آن بحث کنیم.

در ادامه توضیح می‌دهیم که سیستم دینامیکی که در این بخش بررسی می‌شود، چگونه شکل می‌گیرد. همانند آنچه در [۳۸] نشان داده شد، مراحل ساخت شبکه عصبی شامل ساخت شبکه عصبی و ایجاد تابع لیاپانوف است. بدین منظور ابتدا شرایط KKT را به صورت زیر برای مساله (۱۷.۳) و (۱۸.۳) در نظر می‌گیریم.

$$\begin{cases} \nu^* \geq 0, & A\theta^* - b \leq 0, & \nu^{*T}(A\theta^* - b) = 0, \\ Q\theta^* + d + A^T \nu^* + e\vartheta^* = 0, \\ e^T \theta^* = 0. \end{cases} \quad (19.3)$$

قضیه ۱.۳.۳. [۱۱] $\theta^* \in \mathbb{R}^N$ یک نقطه بهینه (۱۷.۳) و (۱۸.۳) است اگر و فقط اگر نقاطی مانند $\nu^* \in \mathbb{R}^{2N}$ و $\vartheta^* \in \mathbb{R}$ وجود داشته باشد، به طوری که $(\theta^{*T}, \nu^{*T}, \vartheta^{*T})^T$ در شرایط KKT سیستم (۱۹.۳) صدق کند.

که در آن θ^* یک نقطه KKT معادله‌های (۱۷.۳) و (۱۸.۳) و زوج $(\nu^{*T}, \vartheta^{*T})^T$ بردار ضریب لاگرانژ متناظر با θ^* است.

لم ۱.۳.۳. $(\nu + A\theta - b)^+ = \nu$ اگر و فقط اگر $\nu^T(A\theta - b) = 0$ ، $A\theta - b \leq 0$ و $\nu \geq 0$ که در آن

$$(\nu + A\theta - b)^+ = [(\nu + A\theta - b)_1]^+, [(\nu + A\theta - b)_2]^+, \dots, [(\nu + A\theta - b)_m]^+]^T,$$

$$[(\nu + A\theta - b)_k]^+ = \max\{(\nu + A\theta - b)_k, 0\}, \quad k = 1, 2, \dots, m.$$

برهان. اگر $(\nu + A\theta - b)^+ = \nu$ ، آن گاه $\nu = 0$ یا $\nu > 0$. در این صورت یکی از دو حالت زیر را داریم:

۱. اگر $\nu = 0$ آن گاه $(\nu + A\theta - b)^+ = (A\theta - b)^+ = 0$ چون $A\theta - b \leq 0$ ، لذا داریم:

$$\nu^T(A\theta - b) = 0.$$

۲. اگر $\nu > 0$ چون $(\nu + A\theta - b)^+ = \nu + A\theta - b = \nu$ بنا براین $A\theta - b = 0$ لذا داریم:

$$\nu^T(A\theta - b) = 0.$$

عکس این رابطه نیز به راحتی قابل اثبات است.

□

اکنون فرض کنید $\theta(\cdot)$ ، $\nu(\cdot)$ و $\vartheta(\cdot)$ متغیرهایی وابسته به زمان باشند. هدف ما ساخت شبکه عصبی پیوسته‌ای است که به نقطه KKT مساله (۱۷.۳) و (۱۸.۳) همگرا باشد. شبکه‌ای که در این رساله پیشنهاد می‌شود به شکل زیر است.

$$\frac{dy}{dt} = \tau \Psi(y), \quad (20.3)$$

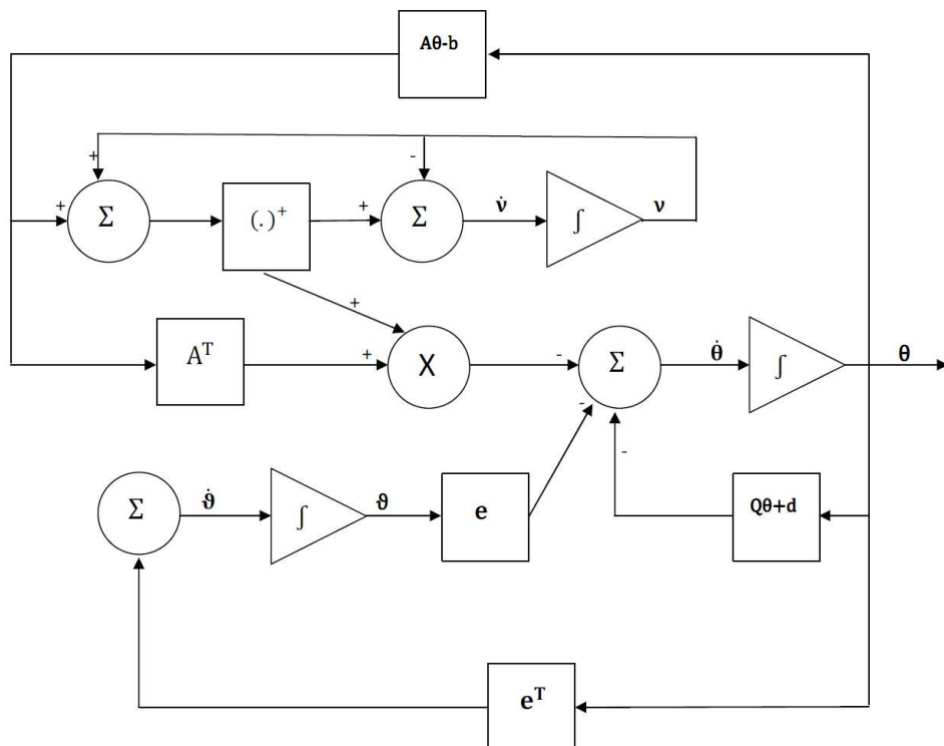
$$y(t_0) = y_0 = (\theta_0^T, \nu_0^T, \vartheta_0^T)^T, \quad \tau > 0, \quad (21.3)$$

که در آن $y = (\theta^T, \nu^T, \vartheta^T)^T \in \mathbb{R}^{3N+1}$ یک پارامتر مقیاس و

$$\Psi(y) = \begin{pmatrix} -(Q\theta + d + A^T(\nu + A\theta - b)^+ + e\vartheta) \\ (\nu + A\theta - b)^+ - \nu \\ e^T \theta \end{pmatrix}, \quad (22.3)$$

است.

لازم به ذکر است که τ نرخ همگرایی شبکه عصبی (۲۰.۳) و (۲۱.۳) است. تصویر (۴.۳) چگونگی اجرای این شبکه عصبی را توضیح می‌دهد.



شکل ۴.۳: یک نمودار ساده بلوکی شده برای شبکه عصبی (۲۰.۳) و (۲۱.۳)

۱.۳.۳ مقایسه با برخی شبکه‌های عصبی موجود

مساله (۱۷.۳) و (۱۸.۳) یک حالت خاص مساله درجه دوم محدب^۱ [۶۰] به شکل زیر است.

$$\text{minimize } \frac{1}{2} x^T Q x + D^T x, \quad (23.3)$$

$$\text{subject to } \begin{cases} Ax - b \leq 0, \\ Ex - f = 0. \end{cases} \quad (24.3)$$

که در آن $Q \in \mathbb{R}^{n \times n}$ یک ماتریس متقارن نیمه معین مثبت، $b \in \mathbb{R}^m$ ، $E \in \mathbb{R}^{l \times n}$ ، $f \in \mathbb{R}^l$ ، $x \in \mathbb{R}^n$ و $\text{rank}(A) = m$ ($0 < m < n$) است.

یک مدل کلی برای حل مساله (۲۳.۳) و (۲۴.۳) به شکل زیر است:

$$\frac{dy}{dt} = \tau \Psi(y), \quad (25.3)$$

$$y(t_0) = y_0 = (x_0^T, u_0^T, v_0^T)^T \in \mathbb{R}^{n+m+l}, \quad \tau > 0, \quad (26.3)$$

که در آن $\psi(y)$ به شکل زیر است:

$$\Psi(y) = \begin{pmatrix} -(Qx + D + A^T(u + Ax - b)^+ + E^T v) \\ (u + Ax - b)^+ - u \\ Ex - f \end{pmatrix}.$$

¹degenerate convex quadratic programming (DCQP)

واضح است که مساله (۲۰.۳) و (۲۱.۳) یک حالت خاص مساله (۲۵.۳) و (۲۶.۳) است. کارایی شبکه عصبی ارائه شده برای حل مسائل DCQP را با مقایسه‌ی این شبکه با برخی شبکه‌های موجود نشان می‌دهیم. مساله‌ی زیر را در نظر بگیرید:

$$\text{minimum } \frac{1}{2} x^T Q x + D^T x, \quad (27.3)$$

$$\text{subject to } Ax - b \leq 0, \quad x \in \Omega, \quad (28.3)$$

که در آن $\Omega = \{x \in \mathbb{R}^n, l_i \leq x_i \leq h_i, i = 1, 2, \dots, n\}$ مقادیر l_i و h_i می‌توانند بی کران بوده و $Ax - b \in \mathbb{R}^m$ است. در [۱۹] شبکه عصبی تصویر فریز^۱ برای حل مساله (۲۷.۳) و (۲۸.۳) به شکل زیر بیان شده است:

$$\frac{dx}{dt} = -(x - P_{\Omega}(x - (Qx + D) - A^T u)), \quad (29.3)$$

$$\frac{du}{dt} = -(u - (u + Ax - b)^+), \quad (30.3)$$

که در آن $\Omega \subset \mathbb{R}^n$ یک مجموعه‌ی محدب بسته و $P_{\Omega} : \mathbb{R}^n \rightarrow \Omega$ یک عملگر تصویر است. همان‌طور که در [۵۸] نشان داده شده است، همگرایی مجانبی مدل دینامیکی (۲۹.۳) و (۳۰.۳) با یک نگاشت یکنوای متقارن تضمین شده نیست. بنابراین سیستم توصیف شده توسط فریز نمی‌تواند مساله (۲۷.۳) و (۲۸.۳) را در حالتی که Q یک ماتریس نیمه معین مثبت است، مثلاً مساله برنامه ریزی خطی LP را حل کند. به عنوان مثال می‌توانید مثال LP را در [۳۷] ببینید.

در حالتی که ماتریس Q فقط معین مثبت باشد، از مدل شبکه لاگرانژی زیر که در [۳۵] برای حل مساله (۲۳.۳) و (۲۴.۳) ارائه شده، استفاده می‌کنیم.

$$\frac{dx}{dt} = -(Qx + D + \frac{1}{2} A^T u^2 + E^T v), \quad (31.3)$$

$$\frac{du}{dt} = \text{diag}(u_1, u_2, \dots, u_m)(Ax - b), \quad (32.3)$$

$$\frac{dv}{dt} = Ex - f, \quad (33.3)$$

با یک نقطه‌ی آغازین $(x_0^T, u_0^T, v_0^T)^T$ به‌طوری‌که $u_k(t_0) > 0$ ($k = 1, \dots, m$) هرچند همگرایی سراسری شبکه عصبی (۳۱.۳) – (۳۳.۳) اثبات شده است ولی این مدل قادر به حل مسائل LP نیست. برای درک بهتر این موضوع می‌توانید مثال LP در [۳۴] را ببینید. یک مدل شبکه عصبی گرادیانی به نام مدل کندی و چا^۲ برای حل مساله (۲۳.۳) و (۲۴.۳) در [۲۲] به شکل زیر ارائه شده است:

$$\frac{dx}{dt} = -(Qx + D + \varsigma[A^T(Ax - b)^+ + E^T(Ex - f)]),$$

¹Friesz

²Kennedy and Chau's

که در آن ϵ یک پارامتر جریمه است.

این مدل هنگامی که پارامتر جریمه آن متناهی است قادر به پیدا کردن جواب دقیق مساله نیست و زمانی که پارامتر جریمه آن بسیار زیاد شود، اجرای آن دشوار است [۱۶]. این مدل به ازای هر مقدار متناهی برای پارامتر جریمه به تقریبی از جواب مساله (۲۳.۳) و (۲۴.۳) همگرا است. همچنین می توان نشان داد که این مساله همگرای سراسری به یک جواب دقیق مسائل DCQP نیست. به عنوان مثال، مساله LP را در [۲۶] ببینید، در حالی که این مساله به صورت دقیق در [۳۷] حل شده است.

با استفاده از یک تابع مکمل غیرخطی و روش گرادیان در [۳۶] یک شبکه عصبی برای حل مسائل اکیداً محدب درجه دوم به شکل زیر ارائه شده است.

$$\frac{dy(t)}{dt} = -k \nabla E(y(t)), \quad k > 0,$$

$$y(0) = y_0,$$

که در آن

$$E(y) = \frac{1}{2} \|\eta(y)\|^2,$$

$$\eta(y) = \begin{bmatrix} Qx + D + E^T v + A^T u \\ f - Ex \\ \phi_{FB}^\epsilon(u, b - Ax) \end{bmatrix} = 0,$$

$$\phi_{FB}^\epsilon(a, b) = (a + b) - \sqrt{a^2 + b^2 - \epsilon}, \quad \epsilon \rightarrow 0_+.$$

این مدل دارای پایداری لیاپانوف و پایدار مجانبی است، ولی دارای نرون های بیشتری با پیچیدگی محاسباتی بالا است و به شرایط همگرایی قوی تر نیاز دارد. علاوه براین، این مدل نمی تواند مسائل برنامه ریزی خطی (LP) را حل کند، زیرا در مسائل برنامه ریزی درجه دوم (QP) فرض براین است که تابع هدف اکیداً محدب است در حالی که در مسائل (LP) تابع هدف فقط محدب است.

زیا^۱ در [۵۵] شبکه عصبی ای برای حل مساله زیر ارائه داد:

$$\text{minimize} \quad \frac{1}{2} x^T Q x + D^T x, \quad (34.3)$$

$$\text{subject to} \quad x \in \Omega_0, \quad (35.3)$$

$$Dx = b, \quad (36.3)$$

به طوری که :

$$\frac{dx}{dt} = (I + A_0)[P_0(x - A_0 x + D^T y - a) - x] - D^T(Dx - b),$$

$$\frac{dy}{dt} = -DP_0(x - A_0 x + D^T y - a) + b,$$

¹Xia

که در آن $A_{m \times m}$ ماتریس متقارن نیمه معین مثبت، D یک ماتریس $n \times m$ و $\Omega_0 \subset \mathbb{R}^m$ یک مجموعه‌ی محدب بسته است. این مساله ضرایب آنالوگ زیادی دارد و برای مسائل بهینه سازی با ابعاد بالا به صرفه نیست. برای حل این مشکل تاو و همکاران^۱ در [۴۷] مدل اصلاح شده‌ی زیر را بیان کردند:

$$\begin{aligned}\frac{dx}{dt} &= [P_0(x - A_0x + D^T y - a) - x] - D^T(Dx - b), \\ \frac{dy}{dt} &= -DP_0(x - A_0x + D^T y - a) + b.\end{aligned}$$

درباره‌ی این شبکه نیز می‌توان گفت هنوز همگرایی آن در زمان متناهی ثابت نشده است. اگر مدل ارائه شده در [۵۸] را برای حل مساله (۳۴.۳) – (۳۶.۳) استفاده کنیم، ابعاد شبکه بزرگ‌تر می‌شود، علاوه‌براین همگرایی این شبکه در زمان متناهی نیز ثابت نشده است. به‌منظور رسیدن به همگرایی در زمان متناهی و همچنین همگرایی نمایی، شیا و همکاران^۲ در [۵۶] شبکه عصبی‌ای را برای حل مسائل درجه دوم اکیداً محدب نشان دادند. گرچه شبکه مورد نظر دارای همگرایی نمایی در زمان متناهی به جواب یکتای مساله درجه دوم اکیداً محدب با قیدهای خطی است ولی این شبکه برای حل مساله (۲۳.۳) و (۲۴.۳) درحالی‌که ماتریس $Q \in \mathbb{R}^{n \times n}$ فقط متقارن و نیمه معین مثبت است، کارایی ندارد. لذا این شبکه نمی‌تواند مسائل LP را حل کند.

اخیراً زیو و بی‌ان^۳ در [۶۰]، بر اساس نتایج به‌دست آمده در [۵۶] برای مسائل درجه دوم اکیداً محدب، یک شبکه عصبی تصویر برای مساله (۳۷.۳) – (۳۹.۳) ارائه کردند،

$$\text{minimize} \quad \frac{1}{2}x^T Qx + c^T x, \quad (37.3)$$

$$\text{subject to} \quad b^1 \leq Bx \leq b^2, \quad (38.3)$$

$$d^\circ \leq h^\circ \leq h^\circ, \quad (39.3)$$

که در آن $Q \in \mathbb{R}^{n \times n}$ یک ماتریس متقارن نیمه معین مثبت، $B \in \mathbb{R}^{m \times n}$ ، $b^1, b^2 \in \mathbb{R}^m$ ، $c \in \mathbb{R}^n$ و $d^\circ, h^\circ \in \mathbb{R}^n$ است. در [۶۰] یک شبکه عصبی تصویر به‌گونه‌ای برای حل مساله (۳۷.۳) – (۳۹.۳) توسعه یافته که همگرایی کامل در زمان متناهی باشد. با این حال، مدل شبکه عصبی پیشنهادی دارای متغیرها و نرون‌های بیشتری است، که باعث می‌شود اجرای شبکه دشوارتر شود.

برای کاهش پیچیدگی شبکه ارائه شده در [۶۰]، زیو در [۵۹] شبکه عصبی جدیدی برای حل مساله (۳۷.۳) – (۳۹.۳) ارائه کرده و ادعا دارد که دامنه آن گسترده‌تر از مقاله قبلی است. ولی با دقت متوجه می‌شویم که ساختار عمگر تصویر برای حل مساله (۳۷.۳) – (۳۹.۳) پیچیده بوده و لذا شبکه عصبی پیشنهادی براساس آن نیز پیچیدگی محاسباتی بالایی دارد.

¹Tao et al

²Xia et al

³Xue and Bian

فورتی و همکاران در [۱۸] همگرایی مسیر دسته‌ای از شبکه‌های عصبی برای حل مسائل برنامه‌ریزی خطی و مسائل برنامه‌ریزی درجه دوم محدب را بررسی کردند. نتیجه اصلی این بررسی این بود که تمامی شبکه‌های مطالعه شده در این بررسی در زمان متناهی به یک مجموعه‌ی تک عضوی از زیرمجموعه‌های نقاط بحرانی قیود، همگرا هستند. با این حال، مسائل درجه دوم بررسی شده در [۱۸] فقط با محدودیت‌های آفین هستند و فضای شدنی آن‌ها یک چند وجهی محدب بسته، بدون نقاط داخلی است.

۲.۳.۳ تحلیل پایداری و همگرایی شبکه عصبی

در این بخش پایداری و همگرایی شبکه عصبی (۲۰.۳) و (۲۱.۳) را بررسی می‌کنیم.

قضیه ۲.۳.۳. فرض کنید $y^* = (\theta^{*T}, \nu^{*T}, \vartheta^{*T})^T$ نقطه تعادل شبکه عصبی (۲۰.۳) و (۲۱.۳) باشد، آن‌گاه θ^* یک نقطه KKT مساله (۱۷.۳) و (۱۸.۳) است. از طرف دیگر اگر $\theta^* \in \mathbb{R}^n$ یک نقطه بهینه مساله (۱۷.۳) و (۱۸.۳) باشد، آن‌گاه $\nu^* \in \mathbb{R}^N$ و $\vartheta^* \in \mathbb{R}$ وجود دارد به‌طوری که $y^* = (\theta^{*T}, \nu^{*T}, \vartheta^{*T})^T$ نقطه تعادل شبکه (۲۰.۳) و (۲۱.۳) است.

برهان. فرض کنید y^* نقطه تعادل شبکه (۲۰.۳) و (۲۱.۳) باشد. در این صورت $\frac{d\theta^*}{dt} = 0$ ، $\frac{d\nu^*}{dt} = 0$ و $\frac{d\vartheta^*}{dt} = 0$. با توجه به این روابط به راحتی نتیجه گرفته می‌شود که

$$Q\theta^* + d + A^T(\nu^* + A\theta^* - b)^+ + e\vartheta^* = 0, \quad (40.3)$$

$$(\nu^* + A\theta^* - b)^+ - \nu^* = 0, \quad (41.3)$$

$$e^T\theta^* = 0. \quad (42.3)$$

بنابر لم (۱.۳.۳)، اگر و فقط اگر $(\nu^* + A\theta^* - b)^+ = \nu^*$

$$\nu^* \geq 0, \quad A\theta^* - b \leq 0, \quad \nu^{*T}(A\theta^* - b) = 0.$$

با جای‌گذاری (۴۰.۳) در (۴۱.۳) داریم:

$$Q\theta^* + d + A^T\nu^* + e\vartheta^* = 0. \quad (43.3)$$

با توجه به روابط (۴۳.۳) – (۴۲.۳) مشخص است که $y^* = (\theta^{*T}, \nu^{*T}, \vartheta^{*T})^T$ در شرایط KKT ارائه شده در (۱۹.۳) صدق می‌کند. عکس این قضیه به‌وضوح مشخص است.

□

لم ۲.۳.۳. ماتریس ژاکوبین $(\nabla\Psi(y))$ برای نگاشت Ψ تعریف شده در (۲۲.۳) یک ماتریس نیمه معین منفی است.

برهان. بدون از دست دادن کلیت، فرض می‌کنیم عدد $0 < p < 2N$ وجود دارد به‌طوری که

$$(\nu + A\theta - b)^+ = ((\nu + A\theta - b)_1, (\nu + A\theta - b)_2, \dots, (\nu + A\theta - b)_p, \underbrace{0, 0, \dots, 0}_{2N-p})^T.$$

با یک محاسبه ساده، مشخص است که

$$\nabla \Psi(y) = \begin{pmatrix} -(Q + (A^p)^T A^p) & -(A^p)^T & -e \\ A^p & W_{2N \times 2N} & O_{2N \times 1} \\ e^T & O_{1 \times 2N} & 0 \end{pmatrix},$$

که در آن O ، ماتریس صفر،

$$A^p = \begin{pmatrix} U_{p \times N} \\ O_{(2N-p) \times N} \end{pmatrix} = \begin{pmatrix} A_1. \\ A_2. \\ \dots \\ \dots \\ A_p. \\ O_{1 \times N} \\ O_{1 \times N} \\ \dots \\ \dots \\ O_{1 \times N} \end{pmatrix},$$

9

$$W_{2N \times 2N} = \begin{pmatrix} O_{p \times p} & O_{p \times (2N-p)} \\ O_{(2N-p) \times p} & -I_{(2N-p) \times (2N-p)} \end{pmatrix}.$$

در [۴۰] می‌توان دید که $(A^p)^T A^p$ یک ماتریس نیمه معین مثبت است. ماتریس Q نیز فرض می‌شود نیمه معین مثبت باشد. علاوه براین به‌وضوح مشخص است ماتریس $W_{2N \times 2N}$ یک ماتریس نیمه معین منفی است.

با استفاده از مشاهدات بیان شده، می‌توان فهمید که ماتریس ژاکوبین $(\nabla \Psi(y))$ ، یک ماتریس نیمه معین منفی است.

اگر $p = 2N$ یعنی $(\nu + A\theta - b)^+ = ((\nu + A\theta - b)_1, (\nu + A\theta - b)_2, \dots, (\nu + A\theta - b)_{2N})^T$ آن‌گاه

$$\nabla \Psi(y) = \begin{pmatrix} -(Q + (A^T)A) & -A^T & -e \\ A & O_{2N \times 2N} & O_{2N \times 1} \\ e^T & O_{1 \times 2N} & 0 \end{pmatrix}.$$

شبیه قسمت قبلی به سادگی ثابت می‌شود که $\nabla \Psi(y)$ یک ماتریس نیمه معین منفی است.

در نهایت اگر $p = \circ$ ، یعنی $(\nu + A\theta - b)^+ = (\underbrace{\circ, \circ, \dots, \circ}_{2N})^T$ داریم:

$$\nabla \Psi(y) = \begin{pmatrix} -Q & O_{N \times 2N} & -e \\ O_{2N \times N} & -I_{2N \times 2N} & O_{2N \times 1} \\ e^T & O_{1 \times 2N} & \circ \end{pmatrix}.$$

در این حالت نیز به سادگی می‌توان ثابت کرد که ماتریس $\nabla \Psi(y)$ نیمه معین منفی است. $\square$

لم ۳.۳.۳. تابع $\|(\nu + A\theta - b)^+\|^2$ از $(\nu + A\theta - b)$ در (۲۲.۳)، روی $\mathbb{R}^N \times \mathbb{R}^{2N}$ محدب و به‌طور پیوسته مشتق پذیر است.

برهان. به‌وضوح، $\|(\nu + A\theta - b)^+\|^2 = \sum_{k=1}^{2N} ((\nu + A\theta - b)_k^+)^2$ و به‌ازای $k = 1, \dots, 2N$ داریم:

$$((\nu + A\theta - b)_k^+)^2 = \begin{cases} ((\nu + A\theta - b)_k)^2 & \text{اگر } (\nu + A\theta - b)_k \geq \circ \\ \circ & \text{در غیر این صورت} \end{cases},$$

با توجه به [۹]، از تحدب و مشتق‌پذیری $(A\theta - b)_k$ ($k = 1, \dots, 2N$) نتیجه حاصل می‌شود. هم‌چنین داریم:

$$\nabla(\|(\nu + A\theta - b)^+\|^2) = \begin{bmatrix} 2A^T(\nu + A\theta - b)^+ \\ 2(\nu + A\theta - b)^+ \end{bmatrix},$$

که این مطلب، اثبات را کامل می‌کند.

$\square$

اکنون نتایج اصلی خود را به شرح زیر بیان می‌کنیم.

قضیه ۳.۳.۳. مدل شبکه عصبی (۲۰.۳) و (۲۱.۳) به مفهوم لیاپانوف پایدار است.

برهان. تابع لیاپانوف $E : \mathbb{R}^{3N+1} \rightarrow \mathbb{R}$ را به‌صورت زیر در نظر بگیرید:

$$E(y) = E_1(y) + E_2(y), \quad (۴۴.۳)$$

که در آن $E_1(y) = \|\Psi(y)\|^2$ و $E_2(y) = \frac{1}{\gamma} \|y - y^*\|^2$ است. با توجه به لم (۳.۳.۳)، می‌دانیم $E_1(y)$ یک تابع مشتق‌پذیر است. بنابر (۲۲.۳) مشاهده می‌شود که

$$\frac{d\Psi}{dt} = \frac{\partial \Psi}{\partial y} \frac{dy}{dt} = \nabla \Psi(y) \Psi(y).$$

با محاسبه مشتق $E(t)$ نسبت به $y(t)$ در شبکه عصبی (۲۰.۳) و (۲۱.۳) داریم:

$$\begin{aligned} \frac{dE(y(t))}{dt} &= \left(\frac{d\Psi}{dt}\right)^T \Psi + \Psi^T \left(\frac{d\Psi}{dt}\right) + (y - y^*)^T \frac{dy(t)}{dt} \\ &= \Psi^T (\nabla \Psi(y)^T + \nabla \Psi(y)) \Psi + (y - y^*)^T \Psi(y). \end{aligned} \quad (۴۵.۳)$$

با استفاده از لم (۲.۳.۳) داریم:

$$\Psi^T(y)(\nabla\Psi(y)^T + \nabla\Psi(y))\Psi(y) \leq 0, \forall y \neq y^*.$$

علاوه بر این از تعریف ۲.۲ و لم ۳.۲ در [۳۸] داریم:

$$(y - y^*)^T(\Psi(y) - \Psi(y^*)) = (y - y^*)^T\Psi(y) \leq 0, \forall y \neq y^*,$$

بنابراین

$$\frac{dE(y(t))}{dt} \leq 0. \quad (۴۶.۳)$$

□ لذا شبکه عصبی (۲۰.۳) و (۲۱.۳) پایدار به مفهوم لیاپانوف است.

لم ۴.۳.۳

(i) به ازای هر نقطه آغازین $y(t_0) = (\theta(t_0)^T, \nu(t_0)^T, \vartheta(t_0)^T)^T$ ، یک جواب پیوسته $y(t) = (\theta(t)^T, \nu(t)^T, \vartheta(t)^T)^T$ برای سیستم (۲۰.۳) و (۲۱.۳) وجود دارد.

(ii) فرض کنید $y(t) = (\theta(t)^T, \nu(t)^T, \vartheta(t)^T)^T$ مسیر جواب سیستم (۲۰.۳) و (۲۱.۳) با نقطه آغازین $y(t_0) = (\theta(t_0)^T, \nu(t_0)^T, \vartheta(t_0)^T)^T$ باشد. اگر $\nu(t_0) \geq 0$ آن گاه $\nu(t) \geq 0$.

برهان.

(i) به راحتی قابل بررسی است که $(\nu + A\theta - b)^+ - \nu, Q\theta + d + A^T(\nu + A\theta - b)^+ + e\vartheta$ به صورت محلی، پیوسته لیپ شیتز هستند. با توجه به وجود جواب محلی معادلات دیفرانسیل معمولی در [۳۰]، شبکه عصبی (۲۰.۳) و (۲۱.۳) دارای یک جواب پیوسته یکتا $y(t)$ برای مقادیر $t \in [t_0, \eta]$ است.

با توجه به اثبات قضیه ۳.۳.۳، می دانیم که E یک تابع ناصعودی نسبت به t است، لذا

$$\frac{1}{2}\|y - y^*\|^2 \leq E(y(t)) \leq E(y(t_0)), \quad \forall t \geq t_0.$$

این نشان دهنده این مطلب است که مسیر جواب شبکه عصبی (۲۰.۳) و (۲۱.۳) کران دار است. بنابراین $\eta \rightarrow +\infty$.

(ii) برای هر نقطه آغازین $y(t_0)$ با $\nu(t_0) \geq 0$ داریم:

$$\begin{aligned} \frac{d\nu}{dt} + \nu &= (\nu + A\theta - b)^+, \\ \int_{t_0}^t \left(\frac{d\nu}{ds} + \nu\right) e^s ds &= \int_{t_0}^t e^s (\nu + A\theta - b)^+ ds. \end{aligned}$$

بنابراین

$$\nu(t) = e^{-(t-t_0)}\nu(t_0) + e^{-t} \int_{t_0}^t e^s(\nu + A\theta - b)^+ ds.$$

با توجه به این که $(\nu + A\theta - b)^+ \geq 0$ ، لذا $\nu(t) \geq 0$ به ازای هر $t \geq t_0$.

□

قضیه ۴.۳.۳. به ازای هر نقطه آغازین $y_0 = (\theta(t_0)^T, \nu(t_0)^T, \vartheta(t_0)^T)^T \in R^{3N+1}$ مسیر جواب شبکه عصبی (۲۰.۳) و (۲۱.۳) به یک نقطه تعادل همگرا است. به ویژه زمانی که D^* که مجموعه نقاط بهینه مساله (۱۷.۳) و (۱۸.۳) و دوگان آن است دارای نقطه تعادل یکتا باشد، شبکه عصبی (۲۰.۳) و (۲۱.۳) به ازای هر نقطه آغازین $y_0 = (\theta(t_0)^T, \nu(t_0)^T, \vartheta(t_0)^T)^T \in R^{3N+1}$ پایدار جانبی سراسری است.

برهان. با توجه به اثبات لم (۴.۳.۳) مسیر جواب $(\theta(t)^T, \nu(t)^T, \vartheta(t)^T)^T$ در شبکه عصبی (۲۰.۳) و (۲۱.۳) کران دار است. لذا دنباله‌ی صعودی مانند t_k به طوری که $t_k \rightarrow \infty$ وقتی $k \rightarrow \infty$ و یک نقطه‌ی حدی $(\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T$ وجود دارد به طوری که:

$$\lim_{k \rightarrow \infty} (\theta(t_k)^T, \nu(t_k)^T, \vartheta(t_k)^T)^T = (\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T.$$

با به کار بردن اصل تغییر ناپذیری لسل، زمانی که $t \rightarrow \infty$ ، $\{(\theta(t)^T, \nu(t)^T, \vartheta(t)^T)^T \rightarrow M\}$ که در آن M بزرگ‌ترین مجموعه‌ی تغییر ناپذیر در $K = \{(\theta(t)^T, \nu(t)^T, \vartheta(t)^T)^T \mid \frac{dE(y(t))}{dt} = 0\}$ است. از روابط (۲۰.۳) و (۲۱.۳) و (۴۶.۳) نتیجه می‌شود $\frac{d\theta}{dt} = 0$ ، $\frac{d\nu}{dt} = 0$ و $\frac{d\vartheta}{dt} = 0$ اگر و فقط اگر $\frac{dE(y(t))}{dt} = 0$. بنابراین $(\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T \in D^*$ که در آن $M \subseteq K \subseteq D^*$ است. ثانیاً ثابت می‌کنیم مسیر جواب $(\theta(t)^T, \nu(t)^T, \vartheta(t)^T)^T$ همگرای سراسری به نقطه تعادل $(\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T$ است. تابع لیاپانوف را به شکل زیر تعریف می‌کنیم:

$$\bar{E}(y) = \|\Psi(y)\|^2 + \frac{1}{\gamma} \|y - \bar{y}\|^2.$$

در رابطه (۴۴.۳)، $\theta^* = \bar{\theta}$ ، $\nu^* = \bar{\nu}$ و $\vartheta^* = \bar{\vartheta}$ را جایگزین می‌کنیم. با توجه به این که $\bar{E}(y)$ به طور پیوسته مشتق پذیر و $\bar{E}(\bar{y}) = 0$ است، لذا

$$\lim_{k \rightarrow \infty} (\theta(t_k)^T, \nu(t_k)^T, \vartheta(t_k)^T)^T = (\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T.$$

پس داریم $\lim_{k \rightarrow \infty} \bar{E}(\theta(t_k)^T, \nu(t_k)^T, \vartheta(t_k)^T)^T = \bar{E}(\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T$. بنابراین، برای هر $\epsilon > 0$ وجود دارد $q > 0$ به طوری که برای تمامی $t \geq t_q$ داریم $\bar{E}(y(t)) < \epsilon$. به طور مشابه، می‌توان به دست آورد $\frac{d\bar{E}(y(t))}{dt} \leq 0$. پس نتیجه می‌گیریم به ازای $t \geq t_q$ ،

$$\frac{1}{\gamma} \|y(t) - \bar{y}\|^2 \leq \bar{E}(y(t)) \leq \epsilon.$$

با توجه به مطالب بالا $\lim_{t \rightarrow \infty} \|y(t) - \bar{y}\| = 0$ و در نتیجه $\lim_{t \rightarrow \infty} y(t) = \bar{y}$. بنابراین سیستم (۲۰.۳) و (۲۱.۳) همگرای سراسری به یک نقطه تعادل مثل $\bar{y} = (\bar{\theta}^T, \bar{\nu}^T, \bar{\vartheta}^T)^T$ است، که در آن $\bar{\theta}$ جواب بهینه مساله (۱۷.۳) و (۱۸.۳) است.

در حالت خاص اگر $D^* = \{(\theta^{*T}, \nu^{*T}, \vartheta^{*T})^T\}$ ، آن گاه شبکه عصبی (۲۰.۳) و (۲۱.۳) برای حل مساله (۱۷.۳) و (۱۸.۳)، پایدار مجانبی سراسری به نقطه تعادل $y^* = \{(\theta^{*T}, \nu^{*T}, \vartheta^{*T})^T\}$ است. در این جا D^* مجموعه جواب بهینه مساله (۱۷.۳) و (۱۸.۳) و دوگان آن است. □

لم ۵.۳.۳. با افزایش τ سرعت همگرایی شبکه عصبی (۲۰.۳) و (۲۱.۳) افزایش می یابد. برهان. با توجه به روابط (۴۵.۳) – (۴۶.۳) به ازای هر $\tau > 0$ در (۲۰.۳) و (۲۱.۳) داریم:

$$\frac{dE(y)}{dt} \leq \tau \Psi(y)^T (\nabla \Psi(y))^T + \nabla \Psi(y) \Psi(y) \leq 0.$$

آن گاه

$$E(y(t)) \leq E(y(t_0)) + \tau \int_{t_0}^t \Psi(y(s))^T (\nabla \Psi(y(s)))^T + (\nabla \Psi(y(s))) \Psi(y(s)) ds.$$

چون $E(y) \geq \frac{1}{2} \|y - y^*\|^2$ که در آن y^* یک نقطه KKT مساله (۱۷.۳) و (۱۸.۳) است، داریم:

$$\|y(t) - y^*\|^2 \leq 2E(y(t_0)) + 2\tau \int_{t_0}^t \Psi(y(s))^T (\nabla \Psi(y(s)))^T + (\nabla \Psi(y(s))) \Psi(y(s)) ds.$$

بنابراین با افزایش τ سرعت همگرایی $y(t)$ افزایش پیدا می کند.

□

۴.۳ مثال‌های عددی

برای نشان دادن کارایی شبکه عصبی مورد نظر، مثال‌هایی را در این بخش ارائه می‌دهیم. شبیه سازی این مثال‌ها با استفاده از نرم افزار *Matlab 7* و حل کننده معادلات دیفرانسیل معمولی *ode45* انجام شده است.

برای مشخص کردن مقدار مناسب برای پارامترهای C و ϵ از روش *Leave – One – Out* در **اعتبارسنجی متقابل**^۱ استفاده می‌کنیم. در این روش، از مجموعه داده‌های آموزشی یک مشاهده خارج شده و براساس بقیه مشاهدات، پارامترها برآورد می‌شود. سپس میزان خطای مدل برای مشاهده خارج شده، محاسبه می‌شود. از آن جایی که در این روش در هر مرحله از فرآیند *CV* فقط یکی از مشاهدات خارج می‌شود، تعداد مراحل تکرار فرآیند *CV* با تعداد داده‌های آموزشی برابر است. همان‌طور که در شکل (۵.۳) نشان داده شده است، در مرحله i ام، مشاهده i ام از مجموعه داده‌های آموزشی خارج شده و مدل توسط تابع *Interpolate* برآورد شده و مقدار تخمینی برای متغیر پاسخ در مشاهده i ام توسط مدل به دست می‌آید. مربع فاصله مقدار واقعی متغیر پاسخ ($y[i]$) از مقدار برآورد شده (y_out)، نیز به عنوان برآورد خطای مدل در نظر گرفته شده و در انتها، میانگین خطای همه مدل‌ها به عنوان برآورد خطای کلی محاسبه می‌شود که به آن میانگین توان دوم خطای پیش‌بینی^۲ می‌گویند. واضح است که هرچه مقدار *MSPE* کمتر باشد، دقت برآورد بالاتر می‌رود.

```

1 Input:
2
3 x, {vector of length N with x-values of data points}
4
5 y, {vector of length N with y-values of data points}
6
7 Output:
8
9 err, {estimate for the prediction error}
10
11 Steps:
12
13 err = 0
14
15 for i = 1, . . . , N do
16
17 // define the cross-validation subsets
18
19 x_in = (x[1], . . . , x[i - 1], x[i + 1], . . . , x[N])
20
21 y_in = (y[1], . . . , y[i - 1], y[i + 1], . . . , y[N])
22
23 x_out = x[i]
24
25 interpolate(x_in, y_in, x_out, y_out)
26
27 err = err + (y[i] - y_out)^2
28
29 end for
30
31 err = err/N

```

شکل ۵.۳: شبه کد اجرای روش *Leave – One – Out* در اعتبارسنجی متقابل (*CV*)

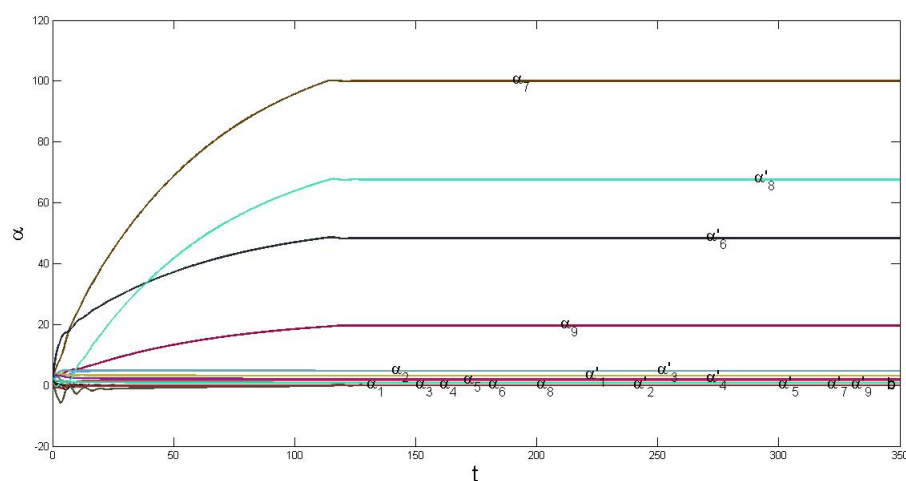
^۱CV

^۲MSPE

جدول ۱.۳: داده‌های رگرسیون

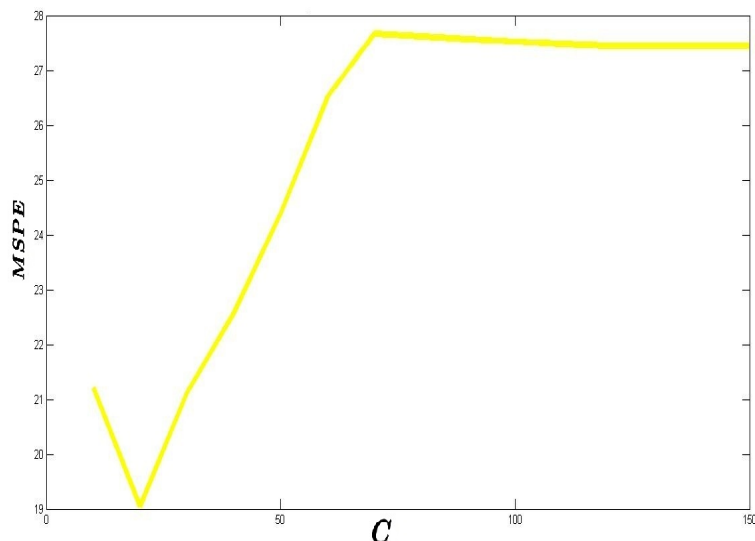
x	-۵	-۳	-۲	۱.۶	۳.۸	۱۰.۲	۱۱	۱۱.۵	۱۲.۷
y	-۱.۶	۱.۸	-۱	-۱.۲	۲.۲	-۶.۸	۹	۱۰	۱۰

مثال ۱.۴.۳. داده‌های رگرسیون جدول (۱.۳) را در نظر می‌گیریم. با توجه به غیرخطی بودن این مساله رگرسیون از کرنل تابع پایه شعاعی $(\Phi(x)^T \Phi(y) = K(x, y) = \exp(-\frac{\|x-y\|^2}{2\delta^2}))$ برای رگرسیون استفاده می‌کنیم. شکل (۶.۳) مسیرهای همگرایی $\alpha_1, \alpha_2, \dots, \alpha_9, \alpha'_1, \alpha'_2, \dots, \alpha'_9$ را نشان می‌دهد.

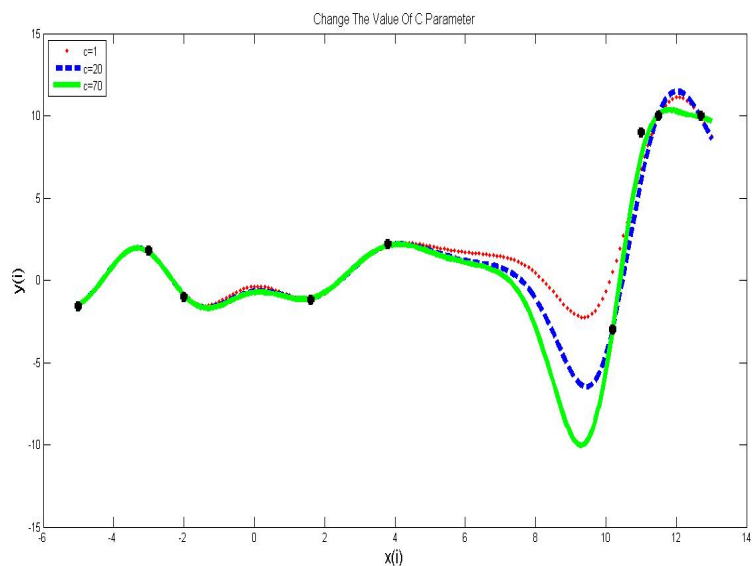


شکل ۶.۳: رفتار گذر $\alpha_1, \alpha_2, \dots, \alpha_9, \alpha'_1, \alpha'_2, \dots, \alpha'_9$ در شبکه عصبی (۲۰.۳) و (۲۱.۳)

در ادامه با تغییر یک پارامتر و ثابت نگه داشتن سایر پارامترها، رفتار $f(x)$ را بررسی می‌کنیم. در شکل (۷.۳) مقدار $MSPE$ به‌ازای مقادیر مختلف C نشان داده شده است. شکل (۸.۳) نیز $f(x)$ را به‌ازای سه مقدار مختلف C نشان می‌دهد. با توجه به نمودار شکل (۷.۳)، به‌ازای مقدار $C = 20$ دقت برآورد $f(x)$ که در شکل (۸.۳) نشان داده شده است، از دو مقدار دیگر C بیشتر است.

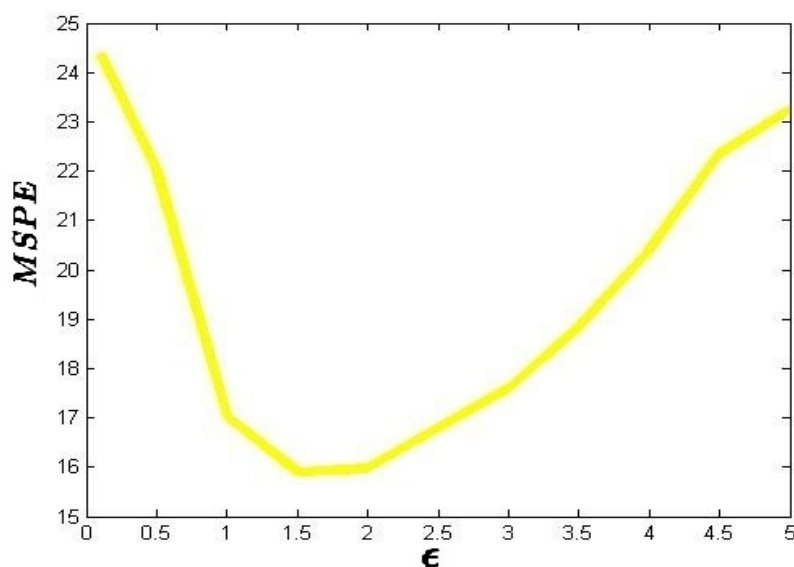


شکل ۷.۳: تغییرات $MSPE$ نسبت به مقادیر مختلف C

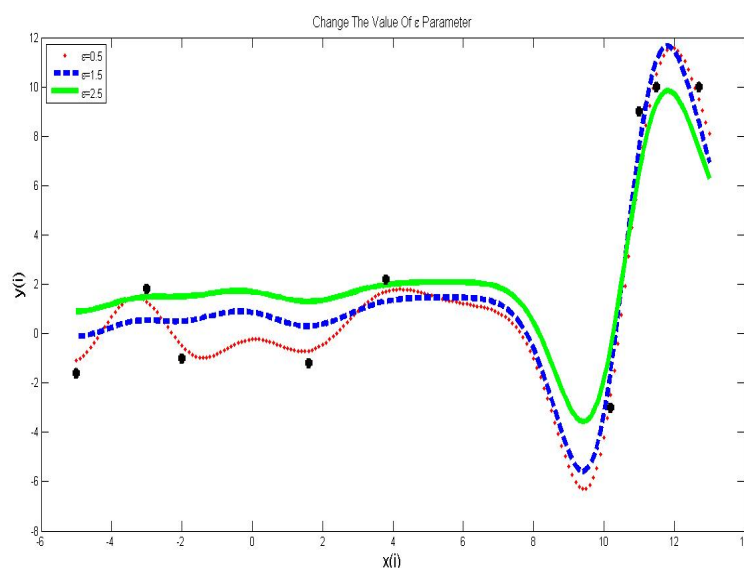


شکل ۸.۳: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)

در شکل (۹.۳) مقدار $MSPE$ به ازای مقادیر مختلف ϵ نشان داده شده است. شکل (۱۰.۳) نیز $f(x)$ را به ازای مقادیر مختلف ϵ نشان می دهد. با توجه به نمودار شکل (۹.۳)، به ازای مقدار $\epsilon = 1/5$ دقت برآورد $f(x)$ که در شکل (۱۰.۳) نشان داده شده است، از دو مقدار دیگر ϵ بیشتر است.



شکل ۹.۳: تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ



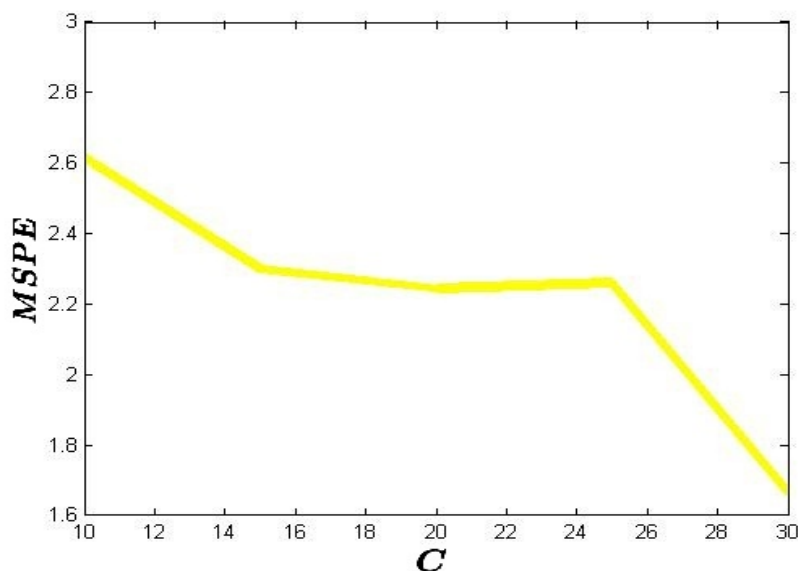
شکل ۱۰.۳: نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)

مثال ۲.۴.۳. در این مثال مساله رگرسیون خطی را با شبکه عصبی (۲۰.۳) و (۲۱.۳) حل می‌کنیم. داده‌های جدول (۲.۳) را در نظر بگیرید.

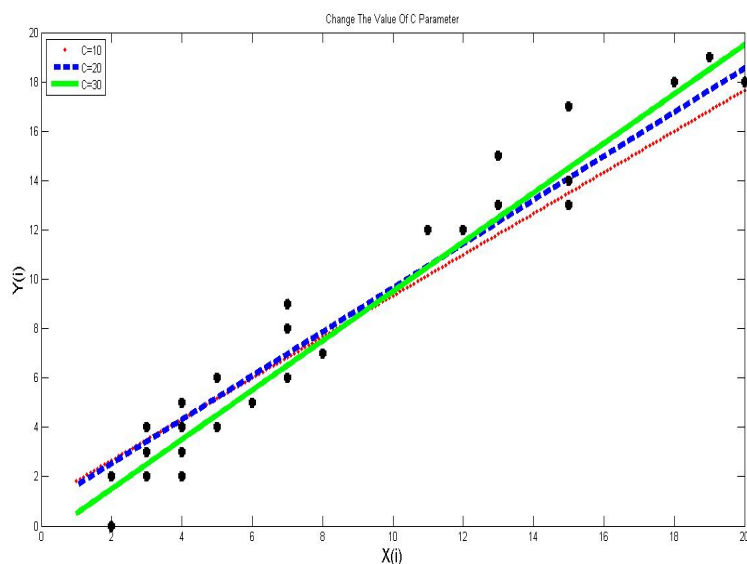
جدول ۲.۳: داده‌های رگرسیون

Y_i	X_i	i	Y_i	X_i	i
۱۷	۱۵	۱۶	۱۳	۱۵	۱
۱۲	۱۲	۱۷	۱۲	۱۱	۲
۳	۴	۱۸	۳	۳	۳
۴	۵	۱۹	۱۵	۱۳	۴
۲	۲	۲۰	۳	۳	۵
۱۹	۱۹	۲۱	۴	۳	۶
۱۴	۱۵	۲۲	۰	۲	۷
۱۲	۱۲	۲۳	۲	۳	۸
۸	۷	۲۴	۲	۴	۹
۴	۴	۲۵	۵	۴	۱۰
۱۳	۱۳	۲۶	۹	۷	۱۱
۱۸	۲۰	۲۷	۶	۷	۱۲
۲	۴	۲۸	۶	۵	۱۳
۵	۶	۲۹	۵	۶	۱۴
۷	۸	۳۰	۱۸	۱۸	۱۵

در شکل (۱۱.۳) مقدار $MSPE$ به ازای مقادیر مختلف C نشان داده شده است. شکل (۱۲.۳) نیز $f(x)$ را به ازای سه مقدار مختلف C نشان می دهد. با توجه به نمودار شکل (۱۱.۳)، به ازای مقدار $C = 30$ دقت برآورد $f(x)$ که در شکل (۱۲.۳) نشان داده شده است، از دو مقدار دیگر C بیشتر است.

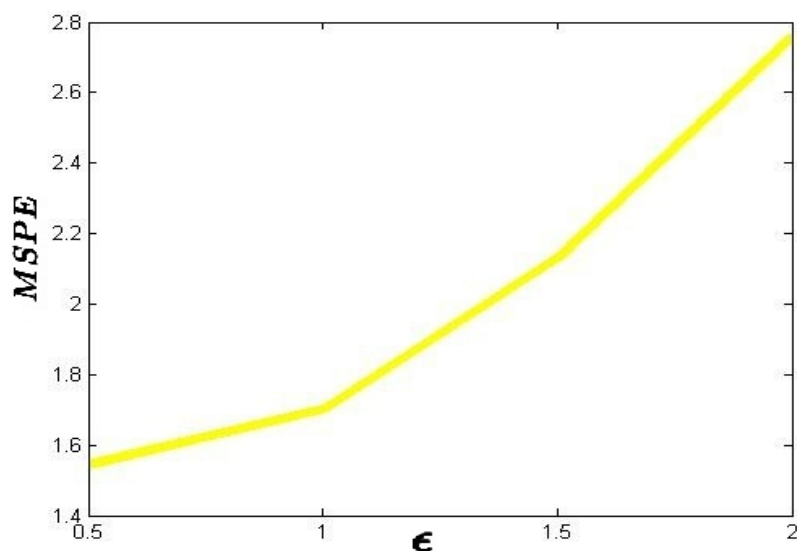


شکل ۱۱.۳: تغییرات $MSPE$ نسبت به مقادیر مختلف C

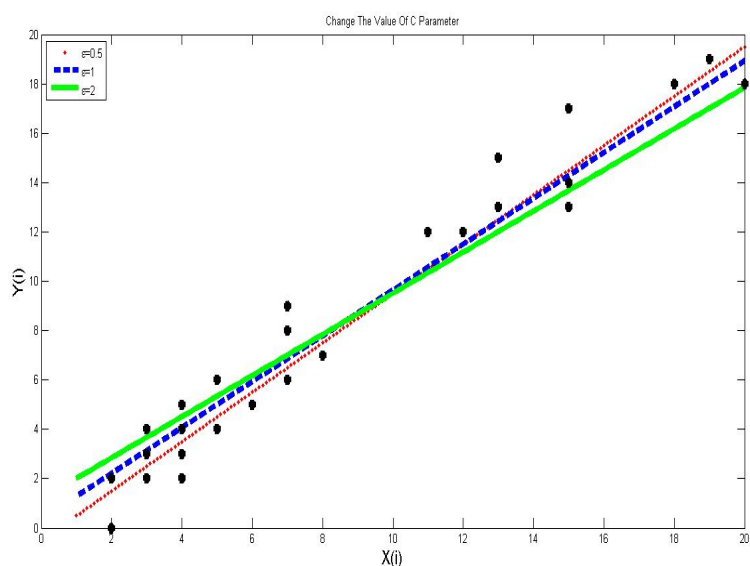


شکل ۱۲.۳: نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)

در شکل (۱۳.۳) مقدار $MSPE$ به‌ازای مقادیر مختلف ϵ نشان داده شده است. شکل (۱۴.۳) نیز $f(x)$ را به‌ازای مقادیر مختلف ϵ نشان می‌دهد. با توجه به نمودار شکل (۱۳.۳)، به‌ازای مقدار $\epsilon = 0.5$ دقت برآورد $f(x)$ که در شکل (۱۴.۳) نشان داده شده است، از دو مقدار دیگر ϵ بیشتر است.



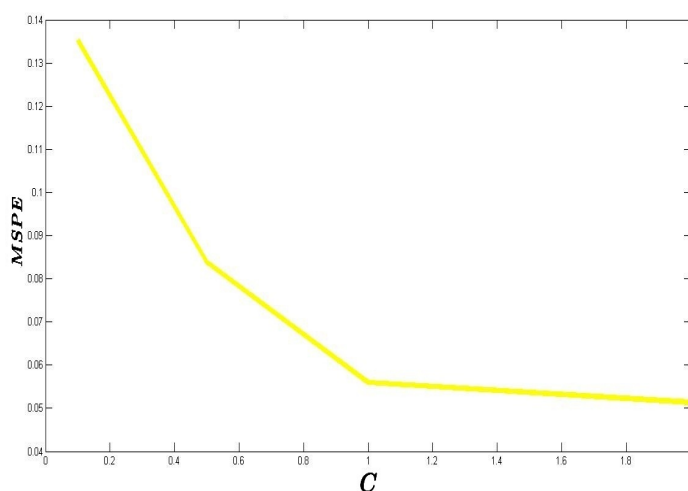
شکل ۱۳.۳: تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ



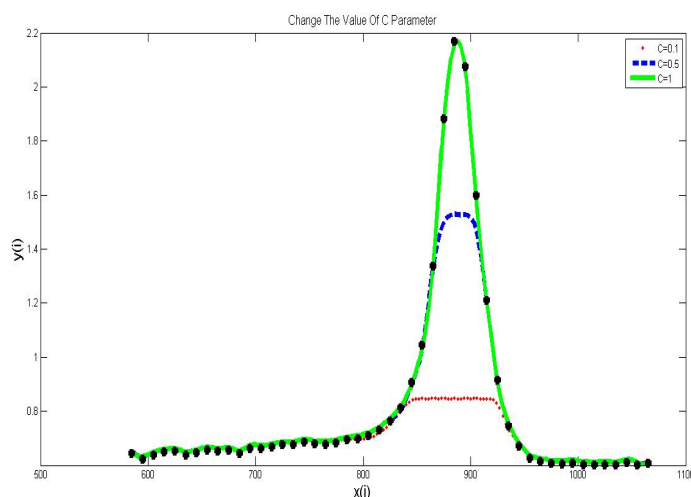
شکل ۱۴.۳: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)

مثال ۳.۴.۳. در این مثال، نتیجه رگرسیون را برای داده‌های رگرسیون تیتانیوم^۱ [۱۳] با استفاده از شبکه عصبی ارائه شده نشان می‌دهیم. کرنل استفاده شده مانند مثال (۱.۴.۳) است.

در شکل (۱۵.۳) مقدار $MSPE$ به‌ازای مقادیر مختلف C نشان داده شده است. شکل (۱۶.۳) نیز $f(x)$ را به‌ازای سه مقدار مختلف C نشان می‌دهد. با توجه به نمودار شکل (۱۵.۳)، به‌ازای مقدار $C = 1$ دقت برآورد $f(x)$ که در شکل (۱۶.۳) نشان داده شده است، از دو مقدار دیگر C بیشتر است.

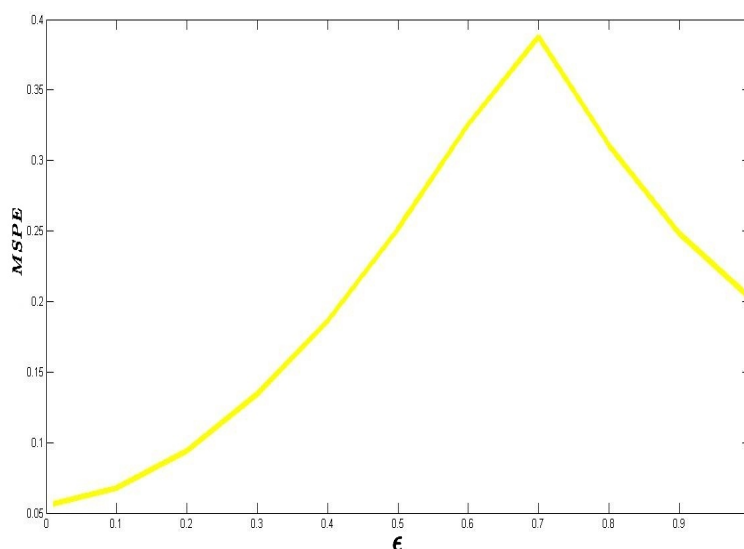


شکل ۱۵.۳: تغییرات $MSPE$ نسبت به مقادیر مختلف C

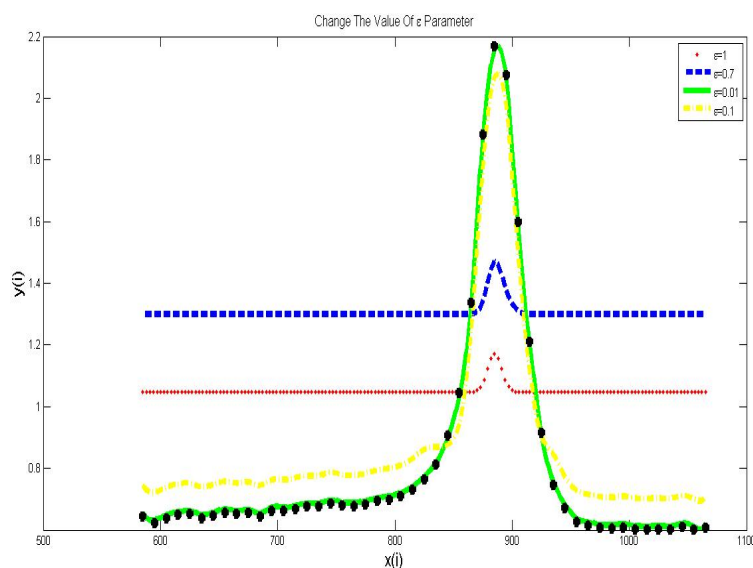


شکل ۱۶.۳: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)

در شکل (۱۷.۳) مقدار $MSPE$ به‌ازای مقادیر مختلف ϵ نشان داده شده است. شکل (۱۸.۳) نیز $f(x)$ را به‌ازای چهار مقدار مختلف ϵ نشان می‌دهد. با توجه به نمودار شکل (۱۷.۳)، به‌ازای مقدار $\epsilon = 0.01$ دقت برآورد $f(x)$ که در شکل (۱۸.۳) نشان داده شده است، از سه مقدار دیگر ϵ بیشتر است.



شکل ۱۷.۳: تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ



شکل ۱۸.۳: نتایج رگرسیون بردار پشتیبان به‌ازای مقادیر مختلف ϵ با استفاده از شبکه عصبی (۲۰.۳) و (۲۱.۳)

فصل ۴

حل مسائل رگرسیون بردار پشتیبان تصادفی با قیدهای احتمالی

در این فصل شبکه عصبی برای حل مساله رگرسیون بردار پشتیبان با قیدهای احتمالی را ارائه می‌دهیم. ثابت می‌شود که نقطه تعادل شبکه عصبی مورد نظر با نقطه بهینه مساله درجه دوم رگرسیون بردار پشتیبان برابر است. وجود داشتن و منحصر به فرد بودن نقطه تعادل ثابت شده و با ساخت تابع لیاپانوف مناسب پایداری مجانبی اثبات می‌گردد. در نهایت با ارائه مثال‌های عددی کارایی این شبکه بررسی می‌شود.

۱.۴ مسائل رگرسیون بردار پشتیبان تصادفی با قیدهای احتمالی

۱.۱.۴ مسائل رگرسیون بردار پشتیبان تصادفی با قیدهای احتمالی در حالت خطی

همان‌طور که در مساله (۴.۳) و (۵.۳) در قبل نشان داده شد، مساله رگرسیون به‌صورت زیر نوشته می‌شود:

$$\text{minimize } \frac{1}{\gamma} \|w\|^2 + C \sum_{i=1}^n (\zeta_i + \zeta'_i), \quad (1.4)$$

$$\text{subject to } \begin{cases} y_i - \langle w, x_i \rangle - b \leq \epsilon + \zeta_i, & i = 1, \dots, n, \\ \langle w, x_i \rangle + b - y_i \leq \epsilon + \zeta'_i, & i = 1, \dots, n, \\ \zeta_i, \zeta'_i \geq 0, & i = 1, \dots, n. \end{cases} \quad (2.4)$$

در رگرسیون عملی داده‌های آموزشی حاوی داده‌های ورودی و یا خروجی را به دلیل خطاهای نمونه‌گیری، خطاهای مدل‌سازی و یا خطاهای اندازه‌گیری، نمی‌توان به‌طور دقیق مشاهده کرد لذا آن‌ها را به‌صورت متغیر تصادفی بیان می‌کنند. فرض کنید نمونه‌های آموزشی $\{(X_1, Y_1), \dots, (X_n, Y_n)\}$ متغیرهای تصادفی هستند، همچنین $X_i = [X_i^1, \dots, X_i^m]^T$ بردار تصادفی در فضای ورودی با امید ریاضی $E(X_i) = [E(X_i^1), \dots, E(X_i^m)]^T$ و Y_i بردار تصادفی در فضای ورودی با امید ریاضی $E(Y_i)$ است. بنابراین قیود در رابطه (۲.۴) باید به‌صورت احتمالی نوشته شوند. لذا ابرصفحه بهینه رگرسیون با حل مساله‌ی زیر به‌دست می‌آید:

$$\text{minimize } \frac{1}{\gamma} \|w\|^2 + C \sum_{i=1}^n (\zeta_i + \zeta'_i), \quad (3.4)$$

$$\text{subject to } \begin{cases} P_r(Y_i - w^T X_i - b \leq \epsilon + \zeta_i) \geq \delta, \\ P_r(w^T X_i + b - Y_i \leq \epsilon + \zeta'_i) \geq \delta, \\ \zeta_i, \zeta'_i \geq 0, & i = 1, \dots, n. \end{cases} \quad (4.4)$$

که در آن $\delta \in [0, 1]$.

با توجه به این‌که حل مساله بهینه‌سازی (۳.۴) و (۴.۴) با قیود نامساوی احتمالی دشوار است، لذا در ادامه شرایط کافی برای برقراری این نامساوی احتمالی را به‌دست می‌آوریم و مساله بهینه‌سازی احتمالی را تبدیل به یک مساله بهینه‌سازی درجه دوم می‌کنیم.

قضیه ۱.۱.۴. فرض کنید T متغیر تصادفی در بازه $[c, a]$ باشد آن‌گاه قضیه زیر برقرار است:

$$\left| P_r(T \leq t) - \frac{a - E(T)}{a - c} \right| \leq \frac{1}{\gamma} + \frac{\left| t - \frac{c+a}{2} \right|}{a - c}, \quad (5.4)$$

به‌ازای هر t در $[c, a]$.

□

برهان. مرجع [۱۰] را ببینید.

ملاحظه ۱.۱.۴. با توجه به این واقعیت که $P_r(T \leq t) = 1 - P_r(T > t)$ ، از نامساوی (۵.۴) نامساوی معادل زیر را به‌ازای هر $t \in [c, a]$ به‌دست می‌آوریم:

$$\left| P_r(T \geq t) - \frac{E(T) - c}{a - c} \right| \leq \frac{1}{\gamma} + \frac{\left| t - \frac{c+a}{2} \right|}{a - c}.$$

قضیه ۲.۱.۴. فرض کنید $X_i = [X_i^1, \dots, X_i^m]^T$ بردار تصادفی با امید ریاضی $E(X_i) = [E(X_i^1), \dots, E(X_i^m)]^T$ و Y_i بردار تصادفی با امید ریاضی $E(Y_i)$ باشد، آن گاه یک شرط کافی برای برقراری

$$P_r(Y_i - w^T X_i - b \leq \varepsilon + \zeta_i) \geq \delta,$$

عبارت است از :

$$E(Y_i) - w^T E(X_i) - b \leq \zeta_i - 2a\delta - \varepsilon.$$

همچنین شرط کافی برای برقراری

$$P_r(w^T X_i + b - Y_i \leq \varepsilon + \zeta'_i) \geq \delta,$$

عبارت است از

$$w^T E(X_i) + b - E(Y_i) \leq \zeta'_i - 2a\delta - \varepsilon,$$

که در آن $a > \varepsilon$.

□

برهان. مرجع [۸] را ببینید.

بنابراین ابرصفحه بهینه رگرسیون را می توان با حل مساله زیر به دست آورد:

$$\text{minimize } \frac{1}{r} \|w\|^2 + C \sum_{i=1}^n (\zeta_i + \zeta'_i), \quad (6.4)$$

$$\text{subject to } \begin{cases} E(Y_i) - w^T E(X_i) - b \leq \zeta_i - 2a\delta - \varepsilon, \\ w^T E(X_i) + b - E(Y_i) \leq \zeta'_i - 2a\delta - \varepsilon, \\ \zeta_i, \zeta'_i \geq 0, \quad i = 1, \dots, n. \end{cases} \quad (7.4)$$

که در آن ε پارامتر انحراف و C پارامتر جریمه است. برای ساختن دوگان مساله بهینه سازی (۶.۴) و (۷.۴) متغیرهای کمکی (α_i, α'_i) را معرفی می کنیم. α_i متناظر محدودیت های شامل ζ_i و α'_i متناظر محدودیت های کمکی ζ'_i است. این متغیرها در حقیقت ضرایب لاگرانژ برای مساله اولیه هستند. مشابه روش های به کار رفته در طراحی رده بندی کننده های SV ، مساله بهینه سازی مقید (۶.۴) و (۷.۴) را با تشکیل لاگرانژین متغیرهای اولیه حل می کنیم که تابع لاگرانژین آن به صورت زیر است:

$$L(w, b, \zeta, \zeta', \alpha, \alpha', \beta, \beta') = \frac{1}{r} w^T w + C \sum_{i=1}^n (\zeta_i + \zeta'_i) - \sum_{i=1}^n (\beta_i \zeta_i + \beta'_i \zeta'_i) + \sum_{i=1}^n \alpha_i [E(Y_i) - w^T E(X_i) - b + (2a\delta + \varepsilon) - \zeta_i] + \sum_{i=1}^n \alpha'_i [w^T E(X_i) + b - E(Y_i) + (2a\delta + \varepsilon) - \zeta'_i].$$

تابع لاگرانژین $L_P(w, b, \zeta_i, \zeta'_i, \alpha_i, \alpha'_i, \beta_i, \beta'_i)$ باید نسبت به متغیرهای اولیه w, b, ζ_i و ζ'_i می‌نیمم شود و نسبت به ضرایب لاگرانژ نامنفی $\alpha_i, \alpha'_i, \beta_i$ و β'_i ماکزیمم گردد. بنابراین تابع لاگرانژین در نقطه جواب بهینه مساله (۶.۴) و (۷.۴) یعنی $(w^*, b^*, \zeta_i^*, \zeta'_i^*)$ دارای نقطه زینی است. در جواب بهینه مشتقات جزئی نسبت به متغیرهای اولیه صفر می‌شود، یعنی :

$$\frac{\partial L}{\partial w} = w - \sum_{i=1}^n (\alpha_i - \alpha'_i) E(x_i) = 0 \Rightarrow w = \sum_{i=1}^n (\alpha_i - \alpha'_i) E(x_i), \quad (۸.۴)$$

$$\frac{\partial L}{\partial b} = \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0,$$

$$\frac{\partial L}{\partial \zeta_i} = c - \alpha_i - \beta_i = 0, \quad i = 1, \dots, n,$$

$$\frac{\partial L}{\partial \zeta'_i} = c - \alpha'_i - \beta'_i = 0, \quad i = 1, \dots, n.$$

بنابراین همانند مساله رگرسیون که در فصل قبل بیان شد، دوگان مساله (۶.۴) و (۷.۴) به شکل زیر نوشته می‌شود:

$$\begin{aligned} \text{maximize} \quad & -\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n (\alpha_i - \alpha'_i)(\alpha_j - \alpha'_j) (E(X_i))^T E(X_j) + (\gamma a \delta + \varepsilon) \sum_{i=1}^n (\alpha_i + \alpha'_i) \\ & + \sum_{i=1}^n (\alpha_i - \alpha'_i) E(Y_i), \end{aligned} \quad (۹.۴)$$

$$\text{subject to} \quad \begin{cases} \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0, \\ \alpha_i, \alpha'_i \in [0, C], \quad i = 1, \dots, n. \end{cases} \quad (۱۰.۴)$$

با توجه به شرط بهینگی (۸.۴) تابع تخمین $f(x)$ به صورت زیر بیان می‌شود:

$$f(x) = \sum_{i=1}^n (\alpha_i - \alpha'_i) E(X_i)^T E(X) + b.$$

۲.۱.۴ حل مسائل رگرسیون بردار پشتیبان تصادفی با قیدهای احتمالی

در حالت غیرخطی

در این بخش نیز مانند حالت‌های غیرخطی که در (۳.۲.۱) بیان شد، با نگاشت $\Phi: \mathbb{R}^m \rightarrow \mathbb{R}^m$ داده‌ها را به فضایی با بعد بیشتر می‌بریم. بنابراین ابر صفحه بهینه رگرسیون بردار پشتیبان در حالت غیرخطی با قیدهای احتمالی با حل مساله زیر به دست می‌آید:

$$\text{minimize} \quad \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n (\zeta_i + \zeta'_i),$$

$$\text{subject to } \begin{cases} P_r(Y_i - w^T \Phi(X_i) - b \leq \varepsilon + \zeta_i) \geq \delta, \\ P_r(w^T \Phi(X_i) + b - Y_i \leq \varepsilon + \zeta'_i) \geq \delta, \\ \zeta_i, \zeta'_i \geq 0, \quad i = 1, \dots, n. \end{cases}$$

که در آن $\Phi(X_i) = [\Phi^1(X_i), \dots, \Phi^{m_1}(X_i)]^T$ یک بردار تصادفی است.

قضیه ۳.۱.۴. فرض کنید $\Phi(X_i) = [\Phi^1(X_i), \dots, \Phi^{m_1}(X_i)]^T$ بردار تصادفی با امید ریاضی $E(\Phi(X_i)) = [E(\Phi^1(X_i)), \dots, E(\Phi^{m_1}(X_i))]^T$ و Y_i یک بردار تصادفی با امید ریاضی $E(Y_i)$ باشد، آن گاه شرط کافی برای برقراری

$$P_r(Y_i - w^T \Phi(X_i) - b \leq \varepsilon + \zeta_i) \geq \delta,$$

عبارت است از:

$$E(Y_i) - w^T E(\Phi(X_i)) - b \leq \zeta_i - 2a\delta - \varepsilon.$$

همچنین یک شرط کافی برای برقراری

$$P_r(w^T \Phi(X_i) + b - Y_i \leq \varepsilon + \zeta'_i) \geq \delta,$$

به صورت

$$w^T E(\Phi(X_i)) + b - E(Y_i) \leq \zeta'_i - 2a\delta - \varepsilon,$$

می باشد که در آن $a > \varepsilon$.

برهان. مرجع [۸] را ببینید.

□

بنابراین ابرصفحه بهینه با حل مساله زیر به دست می آید:

$$\text{minimize } \frac{1}{2} w^T w + C \sum_{i=1}^n (\zeta_i + \zeta'_i), \quad (11.4)$$

$$\text{subject to } \begin{cases} E(Y_i) - w^T E(\Phi(X_i)) - b \leq \zeta_i - 2a\delta - \varepsilon, \\ w^T E(\Phi(X_i)) + b - E(Y_i) \leq \zeta'_i - 2a\delta - \varepsilon, \\ \zeta_i, \zeta'_i \geq 0, \quad i = 1, \dots, n. \end{cases} \quad (12.4)$$

که در آن ε پارامتر انحراف و C پارامتر جریمه است. برای ساختن دوگان مساله بهینه سازی (۱۱.۴) و (۱۲.۴) متغیرهای کمکی (α_i, α'_i) را معرفی می کنیم. α_i متناظر محدودیت های شامل ζ_i و α'_i متناظر محدودیت های کمکی ζ'_i است. این متغیرها در حقیقت ضرایب لاگرانژ برای مساله اولیه هستند. مشابه روش های به کار رفته در طراحی رده بندی کننده های SV ، مساله مقید

بهینه سازی (۱۱.۴) و (۱۲.۴) را با تشکیل لاگرانژین متغیرهای اولیه حل می کنیم که تابع لاگرانژین آن به صورت زیر است:

$$L(w, b, \zeta, \zeta', \alpha, \alpha', \beta, \beta') = \frac{1}{2} w^T w + C \sum_{i=1}^n (\zeta_i + \zeta'_i) - \sum_{i=1}^n (\beta_i \zeta_i + \beta'_i \zeta'_i) + \sum_{i=1}^n \alpha_i [E(Y_i) - w^T E(\phi(X_i)) - b + (\gamma a \delta + \varepsilon) - \zeta_i] + \sum_{i=1}^n \alpha'_i [w^T E(\phi(X_i)) + b - E(Y_i) + (\gamma a \delta + \varepsilon) - \zeta'_i].$$

تابع لاگرانژین $L_P(w, b, \zeta_i, \zeta'_i, \alpha_i, \alpha'_i, \beta_i, \beta'_i)$ باید نسبت به متغیرهای اولیه w, b, ζ_i و ζ'_i می نیمم شود و نسبت به ضرایب لاگرانژ نامنفی $\alpha_i, \alpha'_i, \beta_i$ و β'_i ماکزیمم گردد. بنابراین تابع لاگرانژین در نقطه جواب بهینه مساله (۱۱.۴) و (۱۲.۴) یعنی $(w^*, b^*, \zeta_i^*, \zeta'_i^*)$ دارای نقطه زینی است. در جواب بهینه مشتقات جزئی نسبت به متغیرهای اولیه صفر می شود، یعنی:

$$\frac{\partial L}{\partial w} = w - \sum_{i=1}^n (\alpha_i - \alpha'_i) E(\phi(X_i)) = 0 \Rightarrow w = \sum_{i=1}^n (\alpha_i - \alpha'_i) E(\phi(X_i)), \quad (13.4)$$

$$\frac{\partial L}{\partial b} = \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0,$$

$$\frac{\partial L}{\partial \zeta_i} = c - \alpha_i - \beta_i = 0, \quad i = 1, \dots, n,$$

$$\frac{\partial L}{\partial \zeta'_i} = c - \alpha'_i - \beta'_i = 0, \quad i = 1, \dots, n.$$

مشابه حالت خطی، دوگان مساله (۱۱.۴) و (۱۲.۴) به صورت زیر نوشته می شود:

$$\begin{aligned} \text{maximize} \quad & -\frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n (\alpha_i - \alpha'_i)(\alpha_j - \alpha'_j) E(\Phi(X_i))^T E(\Phi(X_j)) + \\ & (\gamma a \delta + \varepsilon) \sum_{i=1}^n (\alpha_i + \alpha'_i) + \sum_{i=1}^n (\alpha_i - \alpha'_i) E(Y_i), \end{aligned} \quad (14.4)$$

$$\text{subject to} \quad \begin{cases} \sum_{i=1}^n (\alpha_i - \alpha'_i) = 0, \\ \alpha_i, \alpha'_i \in [0, C], \quad i = 1, \dots, n. \end{cases} \quad (15.4)$$

با توجه به شرط بهینگی (۱۳.۴)، تابع تخمین $f(x)$ به صورت زیر بیان می شود:

$$f(x) = \sum_{i=1}^n (\alpha_i - \alpha'_i) E(\Phi(X_i))^T E(\Phi(X)) + b.$$

۳.۱.۴ حل مساله رگرسیون به کمک شبکه عصبی با قیدهای احتمالی

مساله درجه دوم (۱۴.۴) و (۱۵.۴) را می توان به شکل زیر نوشت:

$$\text{minimize} \quad \frac{1}{2} \theta^T Q \theta + d^T \theta, \quad (16.4)$$

$$\text{subject to } \begin{cases} A\theta - b \leq \circ, \\ e^T \theta = \circ. \end{cases} \quad (۱۷.۴)$$

که در آن :

$$Q_{f_n \times f_n} = \begin{cases} E(\Phi(X_i))^T E(\Phi(X_j)) & 1 \leq i, j \leq n \\ \circ & \text{در غیر این صورت} \end{cases},$$

$$d = (-E(Y_1), -E(Y_2), \dots, -E(Y_n), \underbrace{-2a\delta - \epsilon, -2a\delta - \epsilon, \dots, -2a\delta - \epsilon}_n, \underbrace{\circ, \circ, \dots, \circ}_{2n})^T,$$

$$e = (\underbrace{1, 1, \dots, 1}_n, \underbrace{\circ, \circ, \dots, \circ}_{2n})^T \in \mathbb{R}^{f_n},$$

$$A = \begin{pmatrix} I_{n \times 2n} & O_{n \times 2n} \\ O_{n \times 2n} & -I_{n \times 2n} \end{pmatrix},$$

$$b = (\underbrace{C, C, \dots, C}_n, \underbrace{2C, 2C, \dots, 2C}_n, \underbrace{C, C, \dots, C}_n, \underbrace{\circ, \circ, \dots, \circ}_n)^T \in \mathbb{R}^{f_n},$$

$$\theta = (\alpha_1 - \alpha'_1, \alpha_2 - \alpha'_2, \dots, \alpha_n - \alpha'_n, \alpha_1 + \alpha'_1, \alpha_2 + \alpha'_2, \dots, \alpha_n + \alpha'_n, \alpha_1 - \alpha'_1, \alpha_2 - \alpha'_2, \dots, \alpha_n - \alpha'_n, \alpha_1 + \alpha'_1, \alpha_2 + \alpha'_2, \dots, \alpha_n + \alpha'_n)^T.$$

۲.۴ معرفی شبکه عصبی

در این بخش ما یک شبکه عصبی برای حل مساله (۱۶.۴) و (۱۷.۴) ارائه می دهیم. ابتدا برای ساده سازی معماری شبکه لم زیر را بیان می کنیم.

لم ۱.۲.۴. θ^* یک جواب بهینه مساله (۱۶.۴) و (۱۷.۴) است اگر و فقط اگر وجود داشته باشد $\circ \leq u^*$ به طوری که $(\theta^{*T}, u^{*T})^T$ در شرایط زیر صدق کند:

$$(I - \tilde{P})[Q\theta^* + d + A^T u^*] + \tilde{Q}(e^T \theta^*) = \circ, \quad (۱۸.۴)$$

$$(u^* + A\theta^* - b)^+ - u^* = \circ, \quad (۱۹.۴)$$

که در آن $\tilde{P} = e(e^T e)^{-1} e^T$ ، $\tilde{Q} = e(e^T e)^{-1}$ و $(u)^+ = ([u_1]^+, \dots, [u_p]^+)^T$ که $[u_i]^+ = \max\{\circ, u_i\}$ است.

۷۰ حل مسائل رگرسیون بردار پشتیبان تصادفی با قیدهای احتمالی

برهان. اگر θ^* یک جواب بهینه (۱۶.۴) و (۱۷.۴) باشد، بنابر شرایط KKT برای بهینه سازی محدب [۱۱]، $(\theta^{*T}, u^{*T}, v^{*T})^T$ و $u^* \geq 0$ وجود دارد به طوری که معادله‌های زیر برقرار باشند.

$$Q\theta^* + d + ev^* + A^T u^* = 0, \quad (20.4)$$

$$e^T \theta^* = 0, \quad (21.4)$$

$$A\theta^* \leq b, \quad u^* \geq 0, \quad (u^*)^T (A\theta^* - b) = 0. \quad (22.4)$$

واضح است که (۱۹.۴) و (۲۲.۴) معادل اند، یعنی:

$$(A\theta^* - b) \leq 0, \quad u^* \geq 0, \quad (u^*)^T (A\theta^* - b) = 0 \iff (u^* + A\theta^* - b)^+ - u^* = 0.$$

در ادامه ما بایستی ثابت کنیم جواب‌های (۱۸.۴) و جواب‌های (۲۰.۴) و (۲۱.۴) معادل اند. از (۲۰.۴) داریم:

$$-e^T (Q\theta^* + d + ev^* + A^T u^*) = 0.$$

از جمع رابطه بالا با (۲۱.۴) داریم:

$$e^T \theta^* - e^T (Q\theta^* + d) - e^T ev^* - e^T A^T u^* = 0.$$

بنابراین

$$ev^* = e(e^T e)^{-1} (e^T \theta^*) - e(e^T e)^{-1} e^T (Q\theta^* + d + A^T u^*). \quad (23.4)$$

با جایگذاری (۲۳.۴) در (۲۰.۴) داریم:

$$[I - e(e^T e)^{-1} e^T] (Q\theta^* + d + A^T u^*) + e(e^T e)^{-1} (e^T \theta^*) = 0.$$

اکنون قرار دهید $\tilde{P} = e(e^T e)^{-1} e^T$ و $\tilde{Q} = e(e^T e)^{-1} e^T$ آن گاه معادله‌ی بالا را می‌توان به شکل زیر نوشت:

$$(I - \tilde{P})[Q\theta^* + d + A^T u^*] + \tilde{Q}(e^T \theta^*) = 0.$$

برعکس اگر u^* وجود داشته باشد به طوری که $(\theta^{*T}, u^{*T})^T$ در معادله (۱۸.۴) صدق کند، با ضرب e^T در دو طرف (۱۸.۴)، داریم:

$$e^T (I - \tilde{P})[Q\theta^* + d + A^T u^*] + e^T \tilde{Q}(e^T \theta^*) = 0.$$

توجه شود که $e^T (I - \tilde{P}) = 0$ و $e^T \tilde{Q} = I$ ، بنابراین داریم:

$$e^T \theta^* = 0,$$

$$(I - \tilde{P})(Q\theta^* + d + A^T u^*) = 0.$$

قرار دهید $v^* = -(e^T e)^{-1} e^T (Q\theta^* + d + A^T u^*)$. بنابراین

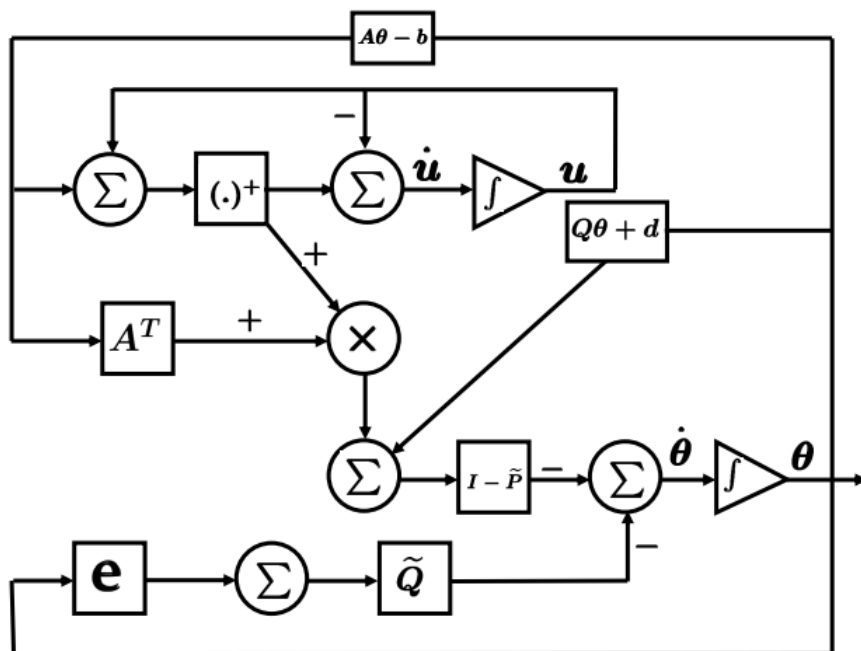
$$\begin{aligned} Q\theta^* + d + e v^* + A^T u^* &= Q\theta^* + d + A^T u^* - e(e^T e)^{-1} e^T (Q\theta^* + d + A^T u^*) \\ &= (I - \tilde{P})(Q\theta^* + d + A^T u^*) = 0. \end{aligned}$$

در نتیجه $(\theta^{*T}, u^{*T}, v^{*T})^T$ وجود دارد به‌طوری که (۲۰.۴) و (۲۱.۴) در شرایط KKT صدق می‌کنند و این اثبات را کامل می‌کند. $\square$

با توجه به لم ۱.۲.۴ ما یک شبکه عصبی کارا به‌صورت زیر برای حل مساله (۱۶.۴) و (۱۷.۴) پیشنهاد می‌کنیم:

$$\begin{cases} \frac{d\theta}{dt} = \eta \{ -(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+] - \tilde{Q}(e^T \theta) \}, \\ \frac{du}{dt} = \eta \{ -u + (u + A\theta - b)^+ \}, \end{cases} \quad (24.4)$$

که در آن η یک ضریب مقیاس است و نرخ همگرایی شبکه عصبی (۲۴.۴) را مشخص می‌کند. در این رساله $\eta = 1$ در نظر گرفته شده است. چگونگی پیاده‌سازی سخت افزاری شبکه عصبی (۲۴.۴) در شکل (۱.۴) نشان داده شده است.



شکل ۱.۴: یک نمودار ساده بلوکی شده برای شبکه عصبی (۲۴.۴)

ملاحظه ۱.۲.۴. با توجه به لم ۱.۲.۴ به راحتی فهمیده می‌شود که θ^* یک جواب بهینه (۱۶.۴) و (۱۷.۴) است اگر و فقط اگر $u^* \geq 0$ وجود داشته باشد به‌طوری که $(\theta^{*T}, u^{*T})^T$ یک نقطه تعادل شبکه (۲۴.۴) باشد. بنابراین زمانی که شبکه عصبی به نقطه تعادل همگراست، مسیر $\theta(t)$ نیز به جواب بهینه مساله (۱۶.۴) و (۱۷.۴) همگراست.

۳.۴ تجزیه و تحلیل پایداری شبکه عصبی

در این بخش ما همگرایی سراسری شبکه عصبی (۲۴.۴) را ثابت می‌کنیم. فرض کنید Ω^e مجموعه نقاط تعادل شبکه عصبی (۲۴.۴) باشد.

لم ۱.۳.۴. فرض کنید $(\theta^{*T}, u^{*T})^T \in \Omega^e$ یک نقطه تعادل (۲۴.۴) است و

$$\begin{aligned} V(\theta, u) = & \frac{1}{\gamma} \theta^T Q \theta + d^T \theta - \left(\frac{1}{\gamma} \theta^{*T} Q \theta^* + d^T \theta^* \right) \\ & + \frac{1}{\gamma} \|(u + A\theta - b)^+\|^2 - \frac{1}{\gamma} \|(u^* + A\theta^* - b)^+\|^2 \\ & - (\theta - \theta^*)^T (Q\theta^* + d + A^T u^*) - (u - u^*)^T u^* + \frac{1}{\gamma} \|\theta - \theta^*\|^2 + \frac{1}{\gamma} \|u - u^*\|^2. \end{aligned}$$

آن‌گاه:

$$\begin{aligned} (I) \quad V(\theta, u) & \geq \frac{1}{\gamma} \|\theta - \theta^*\|^2 + \frac{1}{\gamma} \|u - u^*\|^2, \\ (II) \quad \frac{dV}{dt} & \leq -\|(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+]\| + \tilde{Q}(e^T \theta)\|^2 - \frac{1}{\gamma} \|u - (u + A\theta - b)^+\|^2. \end{aligned}$$

برهان. (I) به سادگی می‌توان بررسی کرد که $\frac{1}{\gamma} \theta^T Q \theta + d^T \theta + \frac{1}{\gamma} \|(u + A\theta - b)^+\|^2$ یک تابع محدب مشتق پذیر است [۱۱] و

$$\nabla \left(\frac{1}{\gamma} \theta^T Q \theta + d^T \theta + \frac{1}{\gamma} \|(u + A\theta - b)^+\|^2 \right) = \begin{bmatrix} Q\theta + d + A^T(u + A\theta - b)^+ \\ (u + A\theta - b)^+ \end{bmatrix}.$$

بنابراین

$$\begin{aligned} & \frac{1}{\gamma} \theta^T Q \theta + d^T \theta - \left(\frac{1}{\gamma} \theta^{*T} Q \theta^* + d^T \theta^* \right) + \frac{1}{\gamma} \|(u + A\theta - b)^+\|^2 - \frac{1}{\gamma} \|(u^* + A\theta^* - b)^+\|^2 \\ & \geq (\theta - \theta^*)^T (Q\theta + d + A^T u^*) + (u - u^*)^T u^*. \end{aligned}$$

یعنی

$$V(\theta, u) \geq \frac{1}{\gamma} \|\theta - \theta^*\|^2 + \frac{1}{\gamma} \|u - u^*\|^2.$$

(II) اکنون مشتق V را نسبت به متغیر t محاسبه می‌کنیم،

$$\begin{aligned} \frac{dV}{dt} & = \frac{\partial V}{\partial \theta} \frac{d\theta}{dt} + \frac{\partial V}{\partial u} \frac{du}{dt} \\ & = -[Q\theta + d + A^T(u + A\theta - b)^+ - (Q\theta^* + d) - A^T u^* + \theta - \theta^*]^T \end{aligned}$$

$$\begin{aligned}
 & \{(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+] + \tilde{Q}(e^T\theta)\} \\
 & - \frac{1}{\Upsilon}[u - (u + A\theta - b)^+]^T[(u + A\theta - b)^+ - \Upsilon u^* + u] \\
 = & -[Q(\theta - \theta^*) + \theta - \theta^* + A^T(u + A\theta - b)^+ - A^T u^*]^T \\
 & \{(I - \tilde{P})[Q(\theta - \theta^*) + A^T(u + A\theta - b)^+ - A^T u^*] + \tilde{P}(\theta - \theta^*)\} \\
 & - \frac{1}{\Upsilon}[u - (u + A\theta - b)^+]^T[u - (u + A\theta - b)^+ + \Upsilon(u + A\theta - b)^+ - \Upsilon u^*] \\
 = & -\{[Q(\theta - \theta^*)]^T(I - \tilde{P})[Q(\theta - \theta^*)] \\
 & + [Q(\theta - \theta^*)]^T(I - \tilde{P})[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & + [Q(\theta - \theta^*)]^T \tilde{P}(\theta - \theta^*) \\
 & + (\theta - \theta^*)^T(I - \tilde{P})[Q(\theta - \theta^*)] \\
 & + (\theta - \theta^*)^T(I - \tilde{P})[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & + (\theta - \theta^*)^T \tilde{P}(\theta - \theta^*) \\
 & + [A^T(u + A\theta - b)^+ - A^T u^*]^T(I - \tilde{P})[Q(\theta - \theta^*)] \\
 & + [A^T(u + A\theta - b)^+ - A^T u^*]^T(I - \tilde{P})[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & + [A^T(u + A\theta - b)^+ - A^T u^*]^T \tilde{P}(\theta - \theta^*)\} \\
 & - \frac{1}{\Upsilon}\|u - (u + A\theta - b)^+\|^{\Upsilon} + [(u + A\theta - b)^+ - u]^T[(u + A\theta - b)^+ - u^*].
 \end{aligned}$$

با توجه به این که

$$(I - \tilde{P})^{\Upsilon} = I - \tilde{P}, \quad \tilde{P}^{\Upsilon} = \tilde{P}, \quad \tilde{P}(I - \tilde{P}) = \circ, \quad (u + A\theta - b)^+ - u = A\theta - b + (-u + b - A\theta)^+$$

داریم:

$$\begin{aligned}
 \frac{dV}{dt} = \frac{dV}{dt} = & -[Q(\theta - \theta^*)]^T(I - \tilde{P})^{\Upsilon}[Q(\theta - \theta^*)] \\
 & - \Upsilon[Q(\theta - \theta^*)]^T(I - \tilde{P})^{\Upsilon}[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & - [A^T(u + A\theta - b)^+ - A^T u^*]^T(I - \tilde{P})^{\Upsilon}[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & - (\theta - \theta^*)^T \tilde{P}^{\Upsilon}(\theta - \theta^*) - (\theta - \theta^*)^T[Q(\theta - \theta^*)] \\
 & - (\theta - \theta^*)^T[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & - \frac{1}{\Upsilon}\|u - (u + A\theta - b)^+\|^{\Upsilon} + [A\theta - b + (-u + b - A\theta)^+]^T[(u + A\theta - b)^+ - u^*] \\
 = & -\|(I - \tilde{P})[Q(\theta - \theta^*)] + (I - \tilde{P})[A^T(u + A\theta - b)^+ - A^T u^*] + \tilde{P}(\theta - \theta^*)\|^{\Upsilon} \\
 & - (\theta - \theta^*)^T[Q(\theta - \theta^*)] - (\theta - \theta^*)^T[A^T(u + A\theta - b)^+ - A^T u^*] \\
 & - \frac{1}{\Upsilon}\|u - (u + A\theta - b)^+\|^{\Upsilon} + (A\theta - b)^T(u + A\theta - b)^+ - (A\theta - b)^T u^*
 \end{aligned}$$

$$\begin{aligned}
 & + [(-u + b - A\theta)^+]^T (u + A\theta - b)^+ - [(-u + b - A\theta)^+]^T u^* \\
 = & - \|(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+] + \tilde{Q}(e^T \theta)\|^2 \\
 & - (\theta - \theta^*)^T [Q(\theta - \theta^*)] - (\theta - \theta^*)^T [A^T(u + A\theta - b)^+ - A^T u^*] \\
 & - \frac{1}{\gamma} \|u - (u + A\theta - b)^+\|^2 + (A\theta - b)^T [(u + A\theta - b)^+] - (A\theta - b)^T u^* \\
 & + [(-u + b - A\theta)^+]^T (u + A\theta - b)^+ - [(-u + b - A\theta)^+]^T u^* \\
 = & - \|(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+] + \tilde{Q}(e^T \theta)\|^2 \\
 & - \frac{1}{\gamma} \|u - (u + A\theta - b)^+\|^2 - (\theta - \theta^*)^T [Q(\theta - \theta^*)] - [(-u + b - A\theta)^+]^T u^* \\
 & + [A(\theta^* - \theta) + A(\theta - \theta^*)]^T u^* - (A\theta^* - b)^T u^* \\
 & + [A(\theta - \theta^*) + A(\theta^* - \theta)]^T (u + A\theta - b)^+ + (A\theta^* - b)^T (u + A\theta - b)^+.
 \end{aligned}$$

به راحتی می‌توان بررسی کرد که

$$\begin{aligned}
 & [(-u + b - A\theta)^+]^T (u + A\theta - b)^+ = 0, \\
 & (A\theta^* - b)^T u^* = 0, \\
 & [(-u + b - A\theta)^+]^T u^* \geq 0, \\
 & (A\theta^* - b)^T [(u + A\theta - b)^+] \leq 0.
 \end{aligned}$$

چون ماتریس Q نیمه معین مثبت است، داریم:

$$(\theta - \theta^*)^T [Q(\theta - \theta^*)] \geq 0.$$

بنابراین

$$\frac{dV}{dt} \leq -\|(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+] + \tilde{Q}(e^T \theta)\|^2 - \frac{1}{\gamma} \|u - (u + A\theta - b)^+\|^2 \leq 0.$$

□

که این مطلب اثبات را کامل می‌کند.

ملاحظه ۱.۳.۴. به آسانی می‌توان دید که $V(\theta, u)$ یک تابع لیپانوف است. با توجه به اثبات **لم ۱.۳.۴**، داریم $\frac{dV}{dt} \leq 0$. بنابراین شبکه عصبی (۲۴.۴) پایدار به مفهوم لیپانوف است.

قضیه ۱.۳.۴. به ازای هر نقطه اولیه $(\theta(0)^T, u(0)^T)^T \in \mathbb{R}^{\Lambda n}$ ، یک جواب پیوسته $(\theta(t)^T, u(t)^T)^T \in \mathbb{R}^{\Lambda n}$ از (۲۴.۴) برای $t \geq 0$ وجود دارد.

برهان. عملگر تصویر $(\cdot)^+$ انقباضی است. واضح است که

لیپ شیتز پیوسته هستند. بنابراین با توجه به قضیه وجود جواب محلی معادله دیفرانسیل معمولی [۳۰]، جواب یکتای پیوسته‌ای برای (۲۴.۴) در $[0, T)$ وجود دارد.

فرض کنید $(\circ, T)$ بازه ماکسیمال وجود جواب باشد. با توجه به لم ۱.۳.۴، می‌دانیم V نسبت به t تابع ناصعودی است، بنابراین

$$\frac{1}{4}\|\theta(t) - \theta^*\|^2 + \frac{1}{4}\|u(t) - u^*\|^2 \leq V(\theta(t), u(t)) \leq V(\theta(\circ), u(\circ)), \quad \forall t \geq \circ. \quad (25.4)$$

این نشان می‌دهد که مسیر (۲۴.۴) کران‌دار است، بنابراین $T = +\infty$ و این اثبات را کامل می‌کند. $\square$

قضیه ۲.۳.۴. به ازای هر نقطه اولیه $(\theta(\circ)^T, u(\circ)^T)^T \in \mathbb{R}^{2n}$ ، مسیر شبکه عصبی (۲۴.۴) به یک نقطه تعادل همگرا است. به‌ویژه زمانی که Ω^e نقطه تعادل یکتا داشته باشد، مسیر شبکه به‌ازای هر نقطه اولیه پایدار سراسری مجانبی است.

برهان. تعریف می‌کنیم

$$D(\theta, u) = \|(I - \tilde{P})[Q\theta + d + A^T(u + A\theta - b)^+] + \tilde{Q}(e^T\theta)\|^2 + \frac{1}{4}\|u - (u + A\theta - b)^+\|^2.$$

آن‌گاه $D(\theta, u) = \circ$ اگر و فقط اگر $(\theta^T, u^T)^T$ یک نقطه تعادل شبکه عصبی (۲۴.۴) باشد. با توجه به اثبات قضیه (۱.۳.۴) مسیر $(\theta(t)^T, u(t)^T)^T$ از شبکه عصبی (۲۴.۴) کران‌دار است. بنابراین یک دنباله صعودی $\{t_n\}$ با $\lim_{n \rightarrow \infty} t_n \rightarrow \infty$ و یک نقطه حدی به‌صورت $(\hat{\theta}^T, \hat{u}^T)^T$ که $\lim_{n \rightarrow \infty} \theta(t_n) \rightarrow \hat{\theta}$ و $\lim_{n \rightarrow \infty} u(t_n) \rightarrow \hat{u}$ وجود دارد. در ادامه ثابت می‌کنیم که $(\hat{\theta}^T, \hat{u}^T)^T \in \Omega^e$ و مسیر $(\theta(t)^T, u(t)^T)^T$ همگرای سراسری به نقطه تعادل $(\hat{\theta}^T, \hat{u}^T)^T$ است. ابتدا ثابت می‌کنیم که $D(\hat{\theta}, \hat{u}) = \circ$ ، یعنی یک نقطه تعادل شبکه عصبی (۲۴.۴) است. اگر این‌چنین نباشد، یعنی $D(\hat{\theta}, \hat{u}) > \circ$ ، چون $D(\theta, u)$ نسبت به θ و u پیوسته است، به‌ترتیب $\varepsilon > \circ$ ، $q > \circ$ و یک همسایگی ε از $(\hat{\theta}^T, \hat{u}^T)^T$ وجود دارد به‌طوری که $D(\theta, u) > q$ به‌ازای تمام $(\theta^T, u^T)^T \in B((\hat{\theta}, \hat{u}), \varepsilon)$. با توجه به این‌که $\lim_{n \rightarrow \infty} \theta(t_n) \rightarrow \hat{\theta}$ و $\lim_{n \rightarrow \infty} u(t_n) \rightarrow \hat{u}$ ، عدد صحیح مثبت N وجود دارد، به‌طوری که به‌ازای تمامی $n \geq N$ ، $\|\theta(t_n) - \hat{\theta}\| \leq \frac{1}{4}\varepsilon$ و $\|u(t_n) - \hat{u}\| \leq \frac{1}{4}\varepsilon$.

از شبکه (۲۴.۴) و کران‌داری $(\theta(t)^T, u(t)^T)^T$ ، مشتقات $\dot{\theta}$ و $\dot{u}$ را داریم که کران‌دار اند و کران بالای آن‌ها را با M با نشان می‌دهیم. در نظر بگیرید $n \geq N$ ، $t \in [t_n - \frac{\varepsilon}{\lambda M}, t_n + \frac{\varepsilon}{\lambda M}]$ لذا داریم

$$\begin{aligned} \|\theta(t) - \hat{\theta}\| + \|u(t) - \hat{u}\| &\leq \|\theta(t) - \theta(t_n)\| + \|u(t) - u(t_n)\| + \|\theta(t_n) - \hat{\theta}\| + \|u(t_n) - \hat{u}\| \\ &= \|\dot{\theta}(\xi_1)\| \times |t - t_n| + \|\dot{u}(\xi_2)\| \times |t - t_n| \\ &\quad + \|\theta(t_n) - \hat{\theta}\| + \|u(t_n) - \hat{u}\| \\ &\leq 2M|t - t_n| + \frac{\varepsilon}{4} \leq \varepsilon. \end{aligned}$$

یعنی برای هر $n \geq N$ ، $t \in [t_n - \frac{\varepsilon}{\lambda M}, t_n + \frac{\varepsilon}{\lambda M}]$ داریم $(\theta^T, u^T)^T \in B((\hat{\theta}, \hat{u}), \varepsilon)$. بنابراین به‌ازای هر $n \geq N$ ، $t \in [t_n - \frac{\varepsilon}{\lambda M}, t_n + \frac{\varepsilon}{\lambda M}]$ داریم $D(\theta(t), u(t)) > q$.

چون $\lim_{n \rightarrow \infty} t_n = \infty$ ، و اندازه لبگ مجموعه $t \in \cup_{n \geq N} [t_n - \frac{\varepsilon}{\lambda M}, t_n + \frac{\varepsilon}{\lambda M}]$ نامتناهی است بنابراین

$$\int_0^\infty D(\theta(t), u(t)) dt = \infty. \quad (۲۶.۴)$$

در حالی که

$$\begin{aligned} \int_0^\infty D(\theta(t), u(t)) dt &= \lim_{s \rightarrow \infty} \int_0^s D(\theta(t), u(t)) dt \\ &\leq - \lim_{s \rightarrow \infty} \int_0^s \dot{V}(\theta(t), u(t)) dt \\ &= \lim_{s \rightarrow \infty} [V(\theta(\circ), u(\circ)) - V(\theta(s), u(s))] \\ &\leq V(\theta(\circ), u(\circ)), \end{aligned}$$

که با (۲۶.۴) تناقض دارد. لذا $D(\hat{\theta}, \hat{u}) = \circ$ ، یعنی $(\hat{\theta}^T, \hat{u}^T)^T \in \Omega^e$.
دوماً ثابت می‌کنیم مسیر $(\theta(t)^T, u(t)^T)^T$ همگرای سراسری به نقطه تعادل $(\hat{\theta}, \hat{u})$ است.
برای این منظور تابع لیاپانوف را به شکل زیر تعریف می‌کنیم

$$\begin{aligned} \hat{V}(\theta, u) &= \frac{1}{\gamma} \theta^T Q \theta + d^T \theta - \left(\frac{1}{\gamma} \hat{\theta}^T Q \hat{\theta} + d^T \hat{\theta} \right) + \frac{1}{\gamma} \|(u + A\theta - b)^+\|^2 - \frac{1}{\gamma} \|(\hat{u} + A\hat{\theta} - b)^+\|^2 \\ &\quad - (\theta - \hat{\theta})^T (Q\hat{\theta} + d + A^T \hat{u}) - (u - \hat{u})^T \hat{u} + \frac{1}{\gamma} \|\theta - \hat{\theta}\|^2 + \frac{1}{\gamma} \|u - \hat{u}\|^2. \end{aligned}$$

داریم: $\hat{V}(\hat{\theta}, \hat{u}) = \circ$. چون $\lim_{n \rightarrow \infty} \theta(t_n) \rightarrow \hat{\theta}$ و $\lim_{n \rightarrow \infty} u(t_n) \rightarrow \hat{u}$ ، $\forall \varepsilon > \circ$ ، وجود دارد یک t_k ، به طوری که $\hat{V}(\theta(t_k), u(t_k)) < \varepsilon$. با توجه به لم ۱.۳.۴، داریم،

$$\frac{1}{\gamma} \|\theta(t) - \hat{\theta}\|^2 + \frac{1}{\gamma} \|u(t) - \hat{u}\|^2 \leq \hat{V}(\theta(t), u(t)),$$

و $\hat{V}$ ناصعودی است. بنابراین به ازای تمامی $t \geq t_k$ ، داریم

$$\frac{1}{\gamma} \|\theta(t) - \hat{\theta}\|^2 + \frac{1}{\gamma} \|u(t) - \hat{u}\|^2 \leq \hat{V}(\theta(t), u(t)) \leq \hat{V}(\theta(t_k), u(t_k)) \leq \varepsilon.$$

یعنی $\lim_{n \rightarrow \infty} \theta(t) = \hat{\theta}$ و $\lim_{n \rightarrow \infty} u(t) = \hat{u}$. بنابراین مسیر حالت شبکه عصبی (۲۴.۴) همگرای سراسری به نقطه تعادل $(\hat{\theta}^T, \hat{u}^T)^T$ است. به ویژه اگر $\Omega^e = \{(\theta^{*T}, u^{*T})^T\}$ ، آن گاه مسیر حالت $(\theta(t)^T, u(t)^T)^T$ به ازای هر نقطه اولیه $(\theta(\circ)^T, u(\circ)^T)^T$ با توجه به تجزیه و تحلیل بالا به نقطه تعادل $(\theta^{*T}, u^{*T})^T$ میل می‌کند. لذا شبکه عصبی پایدار مجانبی سراسری است. $\square$

۱.۳.۴ مقایسه با برخی شبکه‌های عصبی موجود

در این بخش برای این که ببینیم شبکه عصبی (۲۴.۴) چگونه مساله (۱۶.۴) و (۱۷.۴) را حل می‌کند، این شبکه را با برخی از شبکه‌های عصبی موجود مقایسه می‌کنیم.

نوعی از شبکه عصبی به نام شبکه عصبی مدل گرادیان وجود دارد که برای استفاده از این شبکه یک مساله بهینه سازی مقید توسط یک مساله بهینه سازی نامقید تخمین زده می شود سپس تابع انرژی با استفاده از روش تابع جریمه ساخته می شود. مزیت مدل شبکه عصبی گرادیان این است که می توان مدل شبکه عصبی را به طور مستقیم توسط گرادیان تابع انرژی تعریف کرد ولی نقص آن در تضمین نشدن همگرایی به ویژه در مورد مجموعه های جواب بی کران است [۶۱]. علاوه بر این در شبکه عصبی گرادیان بر اساس تابع جریمه، برای هر پارامتر قابل تنظیم نیاز به پارامتر جریمه دارد. به عنوان نمونه با استفاده از روش جریمه، مساله بهینه سازی مقید (۱۶.۴) و (۱۷.۴) را می توان با مساله بهینه سازی نامقید زیر تخمین زد.

$$\text{minimize } E_1(x) = \frac{1}{2} \theta^T Q \theta + d^T \theta + \frac{\gamma}{2} \left\{ \sum_{k=1}^{f_n} [(A\theta - b)_k^+]^2 + (e^T \theta)^2 \right\},$$

که در آن γ پارامتر جریمه است. مدل شبکه عصبی ارائه شده برای مساله فوق به نام مدل شبکه عصبی کندی و چا^۱ به شکل زیر است [۲۲]

$$\frac{dx}{dt} = -\nabla E_1(x) = -(Q\theta + d + \gamma[A^T(A\theta - b)^+ + e^T(e\theta)]). \quad (27.4)$$

شبکه (۲۷.۴) به ازای مقادیر متناهی پارامتر جریمه قادر به پیدا کردن جواب دقیق مساله بهینه سازی نیست و به ازای مقادیر بسیار زیاد برای پارامتر جریمه، پیاده سازی این شبکه دشوار است [۲۴]. بنابراین این شبکه به ازای هر مقدار متناهی از پارامتر جریمه به تقریبی از جواب بهینه مساله (۱۶.۴) و (۱۷.۴) همگرا است. هم چنین می توان نشان داد که شبکه عصبی (۲۷.۴) در مساله بهینه سازی محدب، به یک جواب بهینه دقیق همگرای سراسری نیست. به عنوان نمونه مثال ۱.۴ را در [۳۴] ببینید.

ناظمی و همکار در [۳۶] شبکه عصبی گرادیانی به شکل زیر برای حل مساله (۱۶.۴) و (۱۷.۴) ارائه دادند.

$$\frac{d(y(t))}{dt} = -\nabla E_Y(y(t)), \quad (28.4)$$

$$y(\circ) = y_\circ, \quad y(t) = (x(t), u(t), v(t))^T, \quad (29.4)$$

که در آن

$$E_Y(y) = \frac{1}{2} \|\rho(y)\|^2,$$

$$\rho(y) = \begin{bmatrix} Q\theta + d + e^T v + A^T u \\ -e^T \theta \\ \phi_{FB}^\varepsilon(u, b - A\theta) \end{bmatrix} = \circ,$$

¹Kennedy and Chua's

$$\phi_{\text{FB}}^{\varepsilon}(a, b) = (a + b) - \sqrt{a^2 + b^2 + \varepsilon}, \quad \varepsilon \rightarrow 0_+.$$

همچنین ناظمی در [۳۵] مدل شبکه عصبی گرادینانی به شکل زیر برای حل مساله (۱۶.۴) و (۱۷.۴) بیان کرده است،

$$\frac{dy}{dt} = \Phi(y), \quad (۳۰.۴)$$

$$y(t_0) = y_0, \quad y(t) = (x(t), u(t), v(t))^T, \quad u(t_0) > 0, \quad (۳۱.۴)$$

که در آن

$$\Phi(y) = \begin{bmatrix} -\left(Q\theta + d + \frac{1}{\gamma}A^T u^* + e^T v\right) \\ \text{diag}(u_1, \dots, u_p)(A\theta - b) \\ e^T \theta \end{bmatrix}.$$

در مدل‌های (۲۹.۴)–(۲۸.۴) و (۳۰.۴)–(۳۱.۴) متغیرهای حالت $x \in \mathbb{R}^{fn}$ ، $u \in \mathbb{R}^{fn}$ و $v \in \mathbb{R}$ است. بنابراین این شبکه‌های عصبی دارای $fn+1$ نرون است. در حالی تعداد نرون‌های شبکه عصبی (۲۴.۴) fn است، که زمانی کارایی خود را نشان می‌دهد که تعداد قیود تساوی مساله بهینه‌سازی زیاد باشد. مهم‌تر این که مدل‌های (۲۹.۴)–(۲۸.۴) و (۳۰.۴)–(۳۱.۴) در حالتی که به‌ازای تمامی $x \in \mathbb{R}^{fn}$ و $u_k \geq 0$ ماتریس Q معین مثبت است، همگراست. لذا این مدل قادر به حل مسائل برنامه‌ریزی خطی که کاربرد آن را محدود می‌کند، نیست. در حالی که شبکه عصبی (۲۴.۴) وقتی که ماتریس Q به‌ازای تمامی $x \in \mathbb{R}^{fn}$ و $u_k \geq 0$ نیمه معین مثبت است، به‌جواب بهینه همگراست.

شبکه عصبی ارائه شده در این فصل دارای fn نرون است در حالی که شبکه عصبی که در بخش قبل با آن مساله رگرسیون بردار پشتیبان را حل کردیم، دارای $fn+1$ نرون بود. البته قابل ذکر است که در شبکه عصبی (۲۰.۳) و (۲۱.۳) ما نیازی به محاسبه ماتریس معکوس نداشتیم ولی در شبکه عصبی (۲۴.۴) بایستی ماتریس معکوس را محاسبه کرد.

۴.۴ مثال‌های عددی

برای نشان دادن کارایی شبکه عصبی مورد نظر، مثال‌هایی را در این بخش ارائه می‌دهیم. شبیه سازی این مثال‌ها با استفاده از نرم افزار *Matlab 7* و حل کننده معادلات دیفرانسیل معمولی *ode45* انجام شده است.

برای مشخص کردن مقدار مناسب برای پارامترهای C, ϵ, a و δ از روش *Leave-One-Out* در *اعتبارسنجی متقابل (CV)* استفاده می‌کنیم.

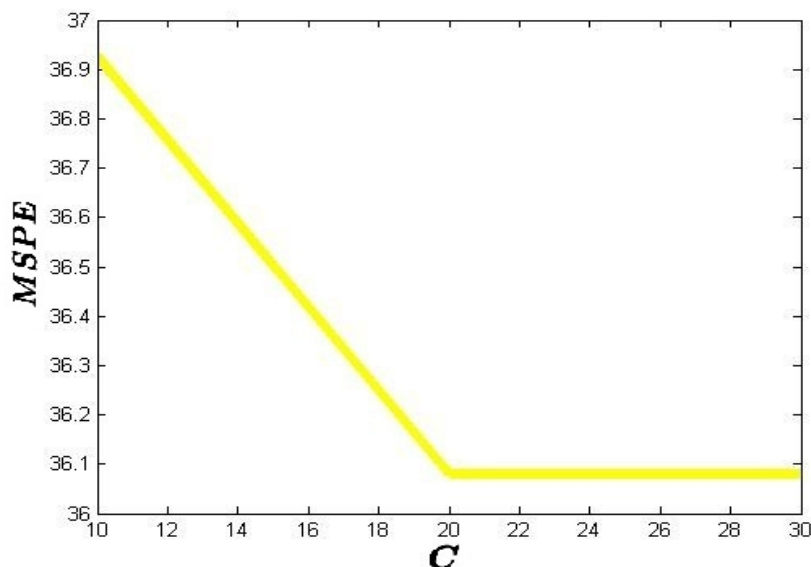
مثال ۱.۴.۴. داده‌های رگرسیون جدول (۱.۴) را در نظر می‌گیریم. می‌دانیم اگر $X_i \sim u(a, b)$ آن‌گاه $E(X_i) = \frac{a+b}{2}$.

در ادامه با تغییر یک پارامتر و ثابت نگه‌داشتن سایر پارامترها، رفتار $f(x)$ را بررسی می‌کنیم.

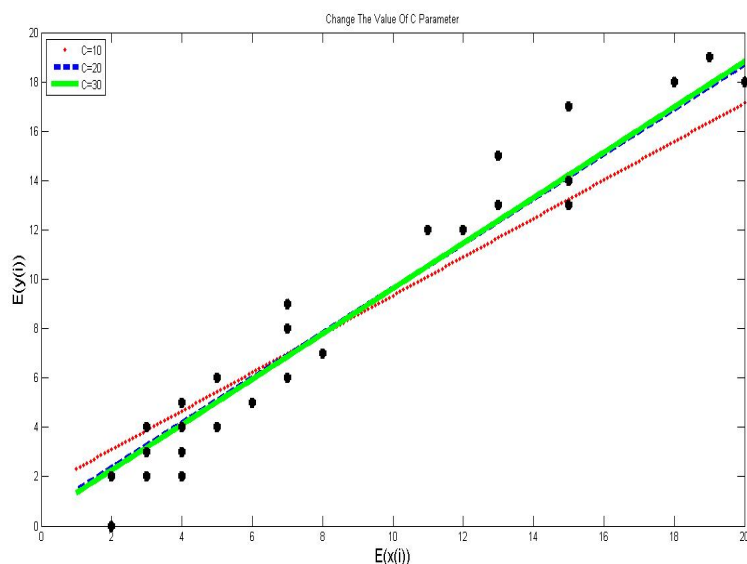
جدول ۱.۴: توزیع احتمال X_i و Y_i

Y_i	X_i	i	Y_i	X_i	i
$u(۱۶, ۱۸)$	$u(۱۳, ۱۷)$	۱۶	$u(۱۲, ۱۴)$	$u(۱۴, ۱۶)$	۱
$u(۱۰, ۱۴)$	$u(۱۱, ۱۳)$	۱۷	$u(۱۱, ۱۳)$	$u(۱۰, ۱۲)$	۲
$u(۲/۵, ۳/۵)$	$u(۲, ۶)$	۱۸	$u(۲, ۴)$	$u(۲, ۴)$	۳
$u(۲, ۶)$	$u(۳, ۷)$	۱۹	$u(۱۴, ۱۶)$	$u(۱۲, ۱۴)$	۴
$u(۱۷/۷۵, ۲/۲۵)$	$u(۱/۵, ۲/۵)$	۲۰	$u(۱, ۵)$	$u(۱, ۵)$	۵
$u(۱۸, ۲۰)$	$u(۱۸, ۲۰)$	۲۱	$u(۳, ۵)$	$u(۲/۵, ۳/۵)$	۶
$u(۱۳, ۱۵)$	$u(۱۴/۵, ۱۵/۵)$	۲۲	$u(۰, ۰)$	$u(۱, ۳)$	۷
$u(۱۲/۵, ۱۳/۵)$	$u(۱۰, ۱۴)$	۲۳	$u(۱, ۳)$	$u(۱/۵, ۴/۵)$	۸
$u(۷, ۹)$	$u(۶/۵, ۷/۵)$	۲۴	$u(۱/۵, ۲/۵)$	$u(۳/۵, ۴/۵)$	۹
$u(۳/۵, ۴/۵)$	$u(۲/۵, ۵/۵)$	۲۵	$u(۴, ۶)$	$u(۳, ۵)$	۱۰
$u(۱۲/۵, ۱۳/۵)$	$u(۱۱, ۱۵)$	۲۶	$u(۸, ۱۰)$	$u(۶, ۸)$	۱۱
$u(۱۶, ۲۰)$	$u(۱۹, ۲۱)$	۲۷	$u(۵, ۷)$	$u(۵, ۹)$	۱۲
$u(۱/۲۵, ۲/۷۵)$	$u(۳/۵, ۴/۵)$	۲۸	$u(۴, ۸)$	$u(۴, ۶)$	۱۳
$u(۳, ۷)$	$u(۴, ۸)$	۲۹	$u(۳, ۷)$	$u(۵, ۷)$	۱۴
$u(۶, ۸)$	$u(۷, ۹)$	۳۰	$u(۱۷, ۱۹)$	$u(۱۷, ۱۹)$	۱۵

در شکل (۲.۴) مقدار $MSPE$ به ازای مقادیر مختلف C نشان داده شده است. شکل (۳.۴) نیز $f(x)$ را به ازای سه مقدار مختلف C نشان می‌دهد. با توجه به نمودار شکل (۲.۴)، به ازای مقدار $C = 30$ دقت برآورد $f(x)$ که در شکل (۳.۴) نشان داده شده است، از دو مقدار دیگر C بیشتر است.

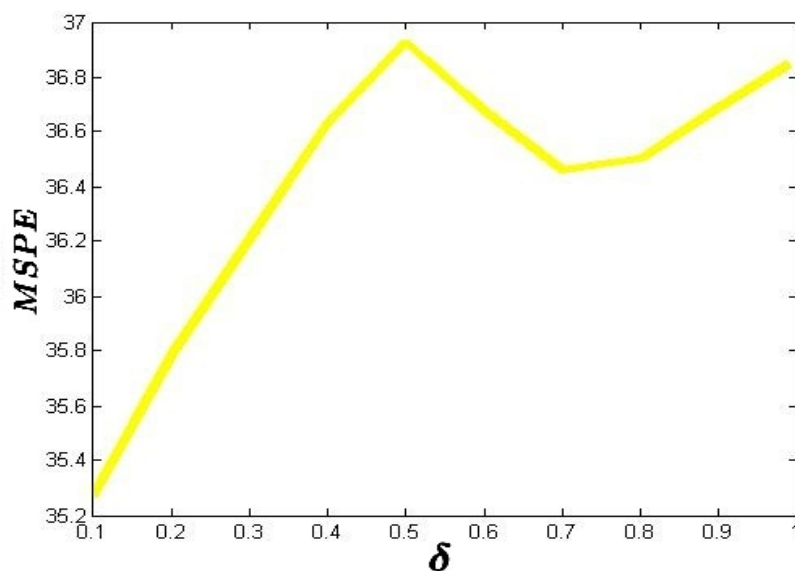


شکل ۲.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف C

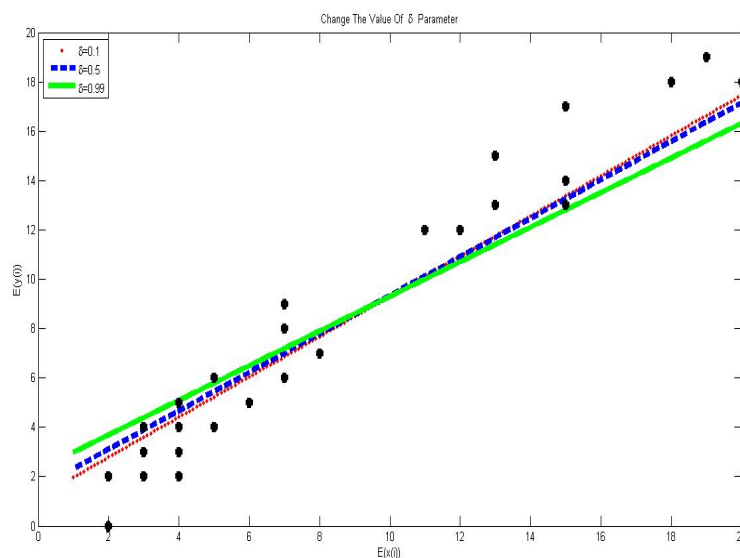


شکل ۳.۴: نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۴.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف δ نشان داده شده است. شکل (۵.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف δ نشان می‌دهد. با توجه به نمودار شکل (۴.۴)، به‌ازای مقدار $\delta = 0.1$ دقت برآورد $f(x)$ که در شکل (۵.۴) نشان داده شده است، از دو مقدار دیگر δ بیشتر است.

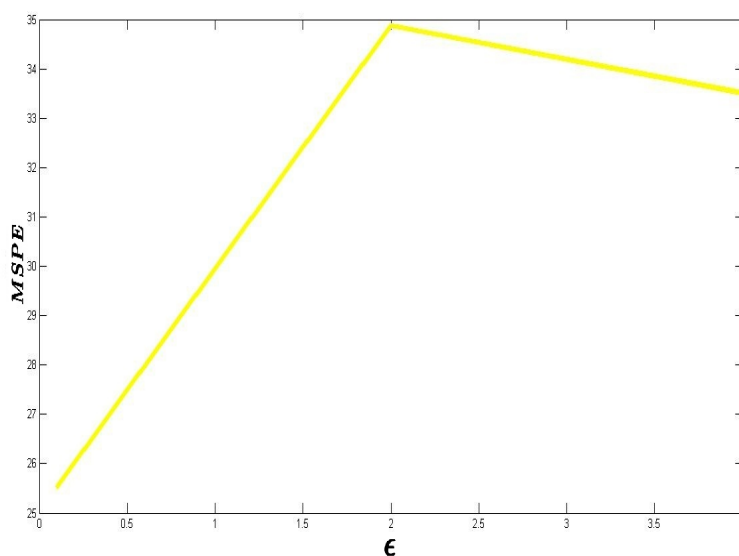


شکل ۴.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف δ

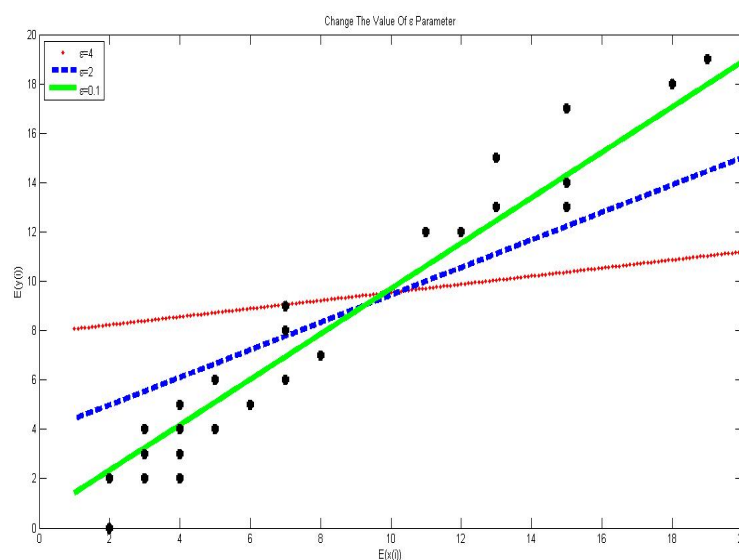


شکل ۵.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف δ با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۶.۴) مقدار $MSPE$ به ازای مقادیر مختلف ϵ نشان داده شده است. شکل (۷.۴) نیز $f(x)$ را به ازای سه مقدار مختلف ϵ نشان می‌دهد. با توجه به نمودار شکل (۶.۴)، به ازای مقدار $\epsilon = 1$ دقت برآورد $f(x)$ که در شکل (۷.۴) نشان داده شده است، از دو مقدار دیگر ϵ بیشتر است.

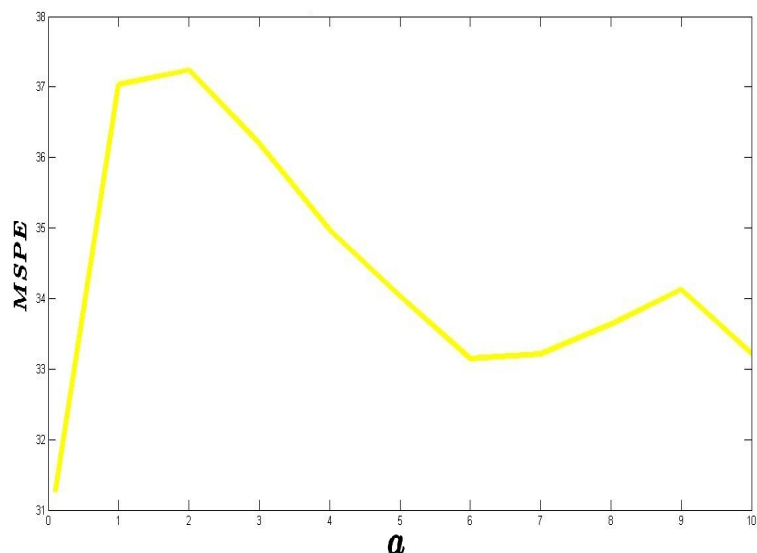


شکل ۶.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ

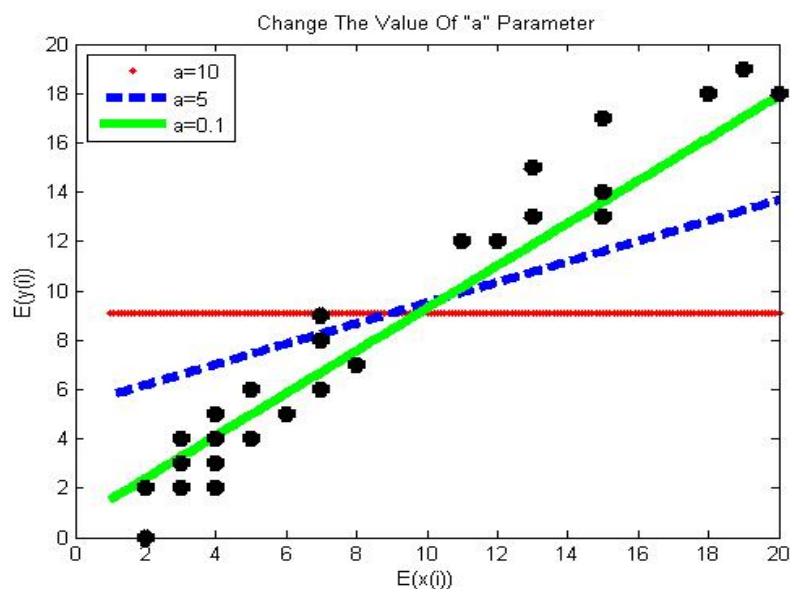


شکل ۷.۴: نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۸.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف a نشان داده شده است. شکل (۹.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف a نشان می‌دهد. با توجه به نمودار شکل (۸.۴)، به‌ازای مقدار $a = 0.1$ دقت برآورد $f(x)$ که در شکل (۹.۴) نشان داده شده است، از دو مقدار دیگر a بیشتر است.



شکل ۸.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف a



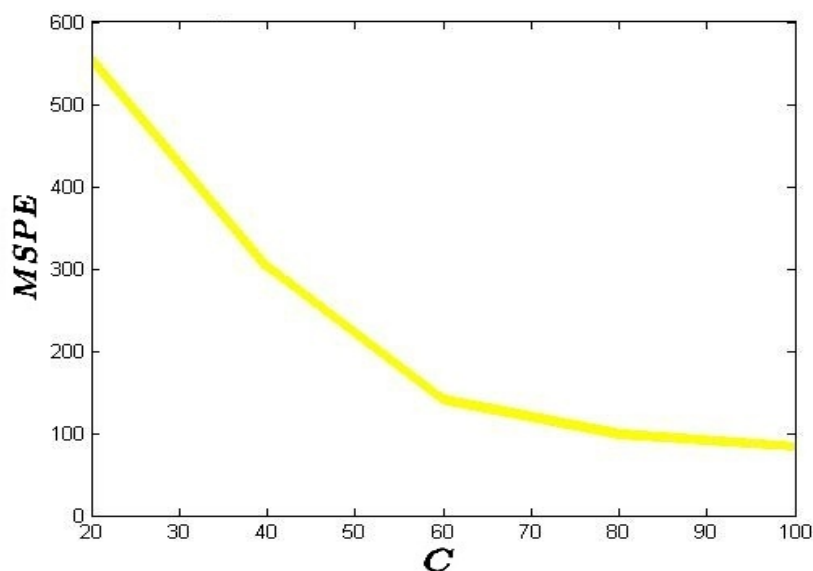
شکل ۹.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف a با استفاده از شبکه عصبی (۲۴.۴)

مثال ۲.۴.۴. نقاط رگرسیون جدول (۲.۴) را که دارای توزیع یکنواخت هستند را در نظر می‌گیریم. با توجه به غیرخطی بودن این مساله رگرسیون از کرنل تابع پایه شعاعی $(\Phi(x)^T \Phi(y) = K(x, y) = \exp(-\frac{\|x-y\|^2}{2\delta^2}))$ برای رگرسیون استفاده می‌کنیم. در ادامه با تغییر یک پارامتر و ثابت نگه‌داشتن سایر پارامترها، رفتار $f(x)$ را بررسی می‌کنیم.

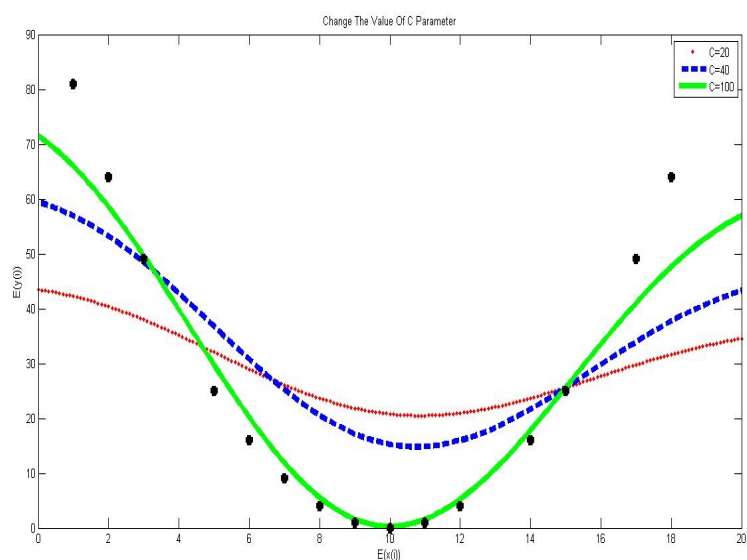
جدول ۲.۴: توزیع احتمال X_i و Y_i

Y_i	X_i	i	Y_i	X_i	i
$u(0, 2)$	$u(8, 10)$	۱۶	$u(24, 26)$	$u(14, 16)$	۱
$u(47, 51)$	$u(1, 5)$	۱۷	$u(23, 27)$	$u(13, 17)$	۲
$u(0, 2)$	$u(7, 11)$	۱۸	$u(63, 65)$	$u(17, 19)$	۳
$u(58, 70)$	$u(16, 20)$	۱۹	$u(62, 66)$	$u(1, 3)$	۴
$u(1, 7)$	$u(6, 10)$	۲۰	$u(8, 10)$	$u(6, 8)$	۵
$u(46, 52)$	$u(2, 4)$	۲۱	$u(80, 82)$	$u(0, 2)$	۶
$u(2, 6)$	$u(11, 13)$	۲۲	$u(48, 50)$	$u(16, 18)$	۷
$u(57, 71)$	$u(15, 21)$	۲۳	$u(0, 2)$	$u(10, 12)$	۸
$u(7, 11)$	$u(5, 9)$	۲۴	$u(7, 1)$	$u(7, 9)$	۹
$u(79, 83)$	$u(0, 2)$	۲۵	$u(22, 28)$	$u(4, 6)$	۱۰
$u(45, 53)$	$u(0, 6)$	۲۶	$u(0, 2)$	$u(9, 13)$	۱۱
$u(56, 62)$	$u(1, 3)$	۲۷	$u(15, 17)$	$u(4, 8)$	۱۲
$u(21, 29)$	$u(12, 18)$	۲۸	$u(61, 67)$	$u(0, 4)$	۱۳
$u(0, 0)$	$u(9, 11)$	۲۹	$u(60, 68)$	$u(1, 3)$	۱۴
$u(14, 18)$	$u(13, 15)$	۳۰	$u(59, 69)$	$u(0, 4)$	۱۵

در شکل (۱۰.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف C نشان داده شده است. شکل (۱۱.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف C نشان می‌دهد. با توجه به نمودار شکل (۱۰.۴)، به‌ازای مقدار $C = ۱۰۰$ دقت برآورد $f(x)$ که در شکل (۱۱.۴) نشان داده شده است، از دو مقدار دیگر C بیشتر است.

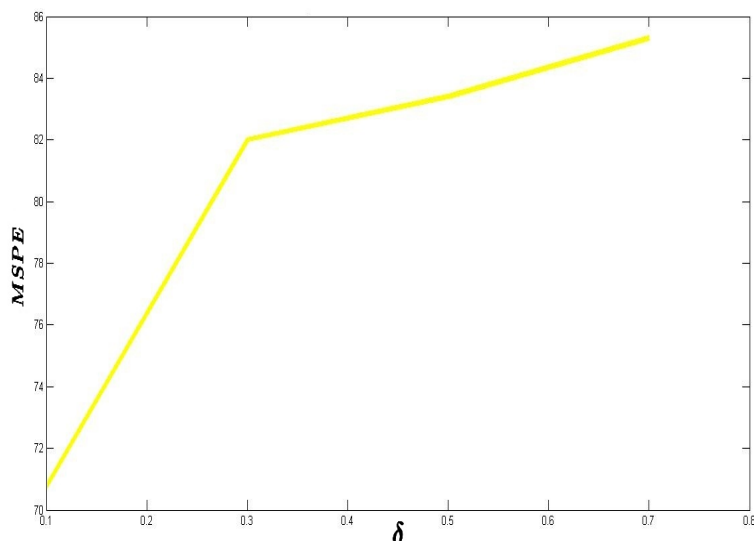


شکل ۱۰.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف C

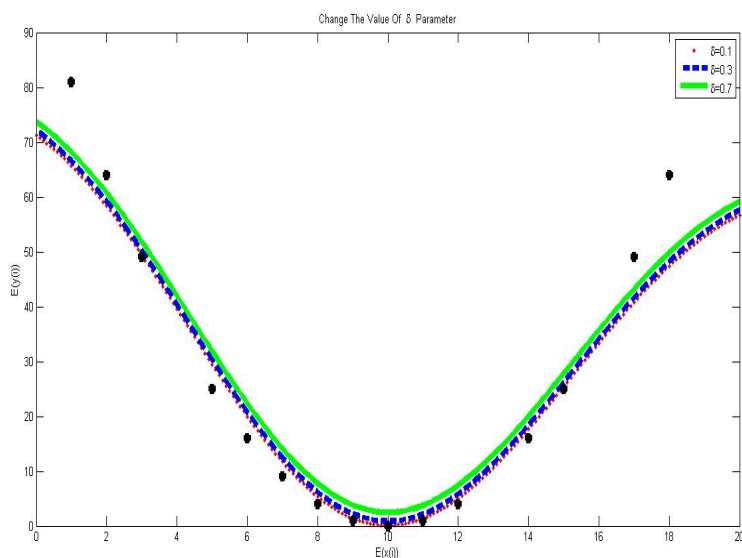


شکل ۱۱.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۱۲.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف δ نشان داده شده است. شکل (۱۳.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف δ نشان می‌دهد. با توجه به نمودار شکل (۱۲.۴)، به‌ازای مقدار $\delta = 0.1$ دقت برآورد $f(x)$ که در شکل (۱۳.۴) نشان داده شده است، از دو مقدار دیگر δ بیشتر است.

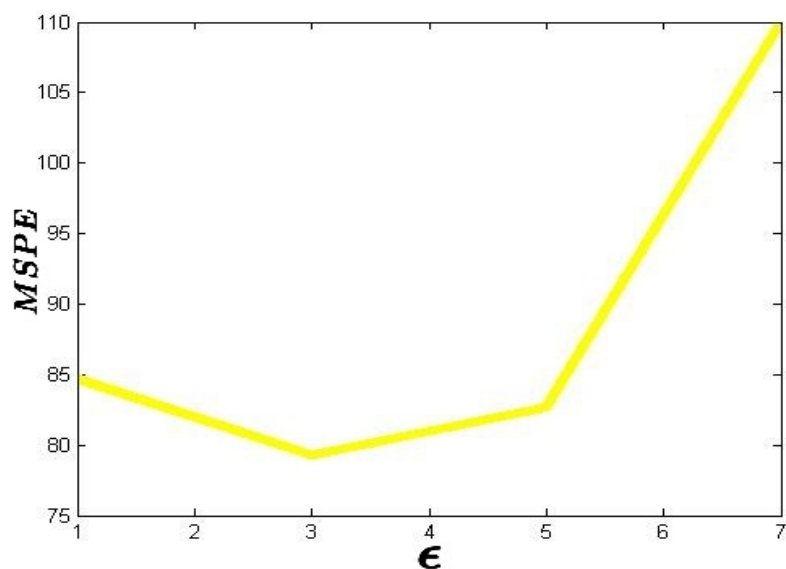


شکل ۱۲.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف δ

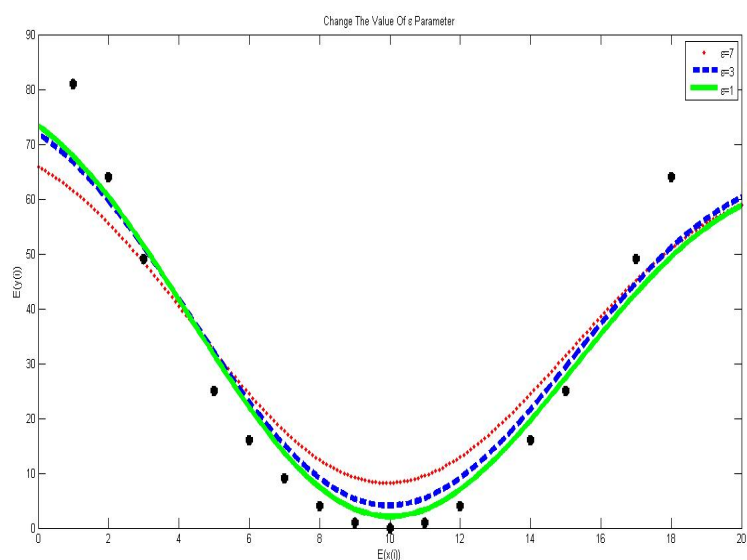


شکل ۱۳.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف δ با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۱۴.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف ϵ نشان داده شده است. شکل (۱۵.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف ϵ نشان می‌دهد. با توجه به نمودار شکل (۱۴.۴)، به‌ازای مقدار $\epsilon = 3$ دقت $f(x)$ که در شکل (۱۵.۴) نشان داده شده است، از دو مقدار دیگر ϵ بیشتر است.

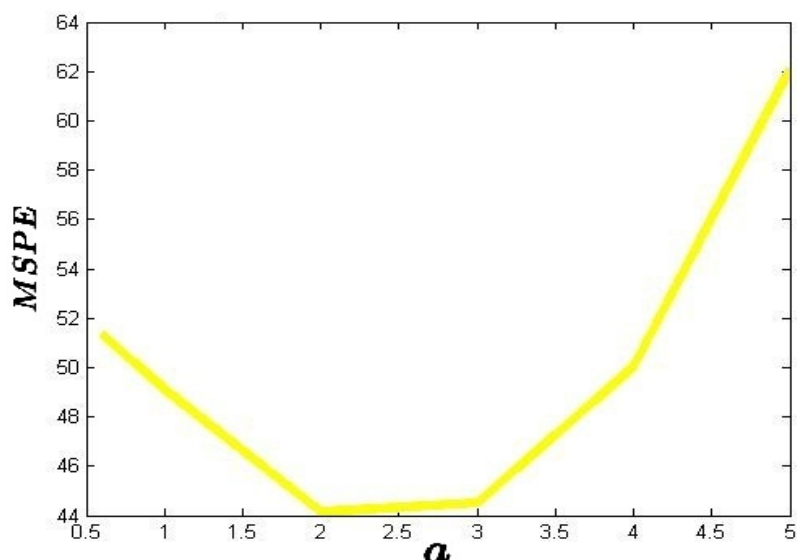


شکل ۱۴.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ

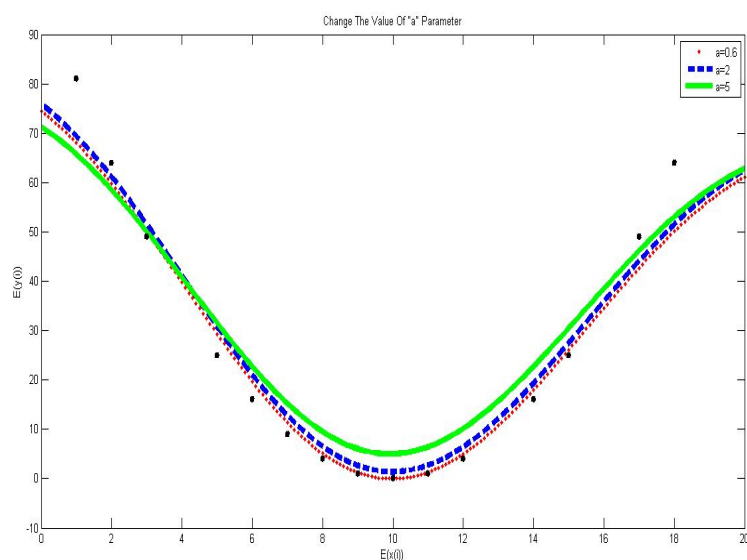


شکل ۱۵.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۱۶.۴) مقدار $MSPE$ به ازای مقادیر مختلف a نشان داده شده است. شکل (۱۷.۴) نیز $f(x)$ را به ازای سه مقدار مختلف a نشان می‌دهد. با توجه به نمودار شکل (۱۶.۴)، به ازای مقدار $a = 2$ دقت برآورد $f(x)$ که در شکل (۱۷.۴) نشان داده شده است، از دو مقدار دیگر a بیشتر است.



شکل ۱۶.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف a

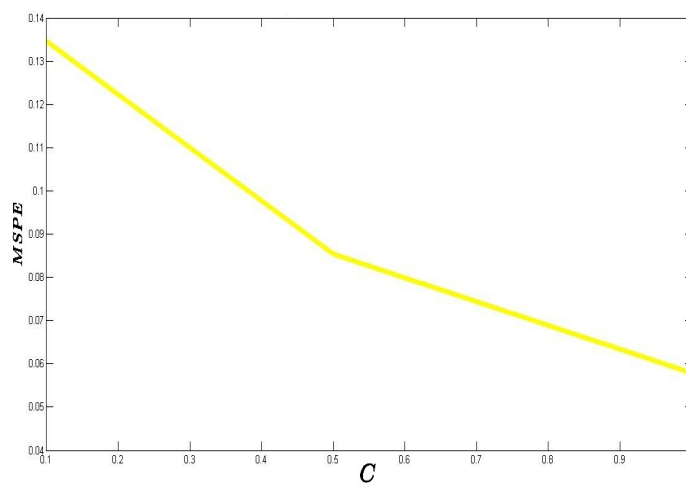


شکل ۱۷.۴: نتایج رگرسیون بردار پشتیبان به ازای سه مقدار مختلف a با استفاده از شبکه عصبی (۲۴.۴)

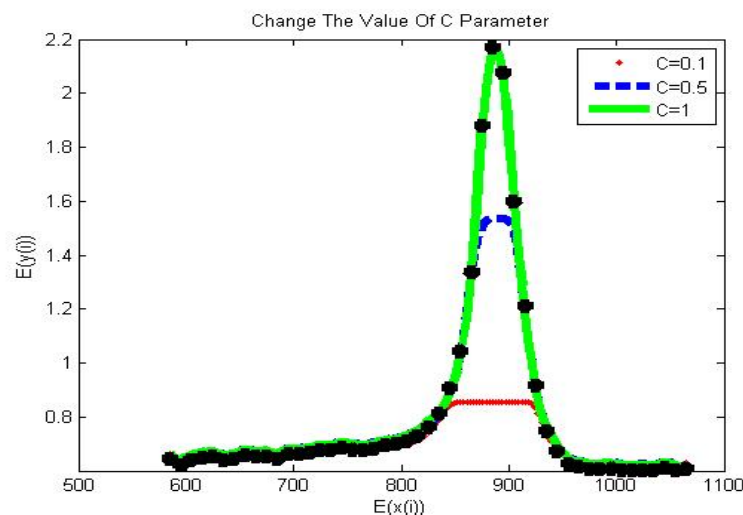
مثال ۳.۴.۴. در این مثال نقاط تیتانیوم مثال (۳.۴.۳) را طوری در نظر می‌گیریم که X_i ها و Y_i ها امید ریاضی نقاط تصادفی باشند.

در ادامه همانند مثال‌های قبل با تغییر یک پارامتر و ثابت نگه‌داشتن سایر پارامترها، رفتار $f(x)$ را بررسی می‌کنیم.

در شکل (۱۸.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف C نشان داده شده است. شکل (۱۹.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف C نشان می‌دهد. با توجه به نمودار شکل (۱۸.۴)، به‌ازای مقدار $C = 1$ دقت برآورد $f(x)$ که در شکل (۱۹.۴) نشان داده شده است، از دو مقدار دیگر C بیشتر است.

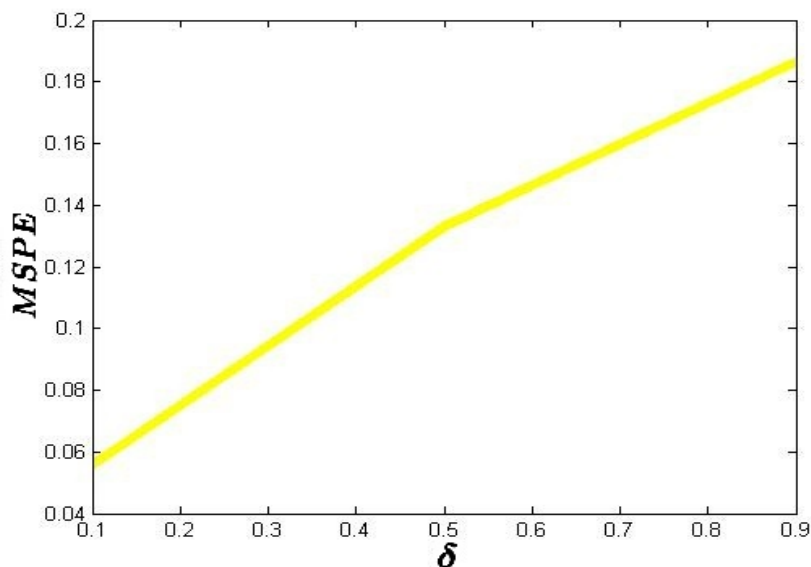


شکل ۱۸.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف C

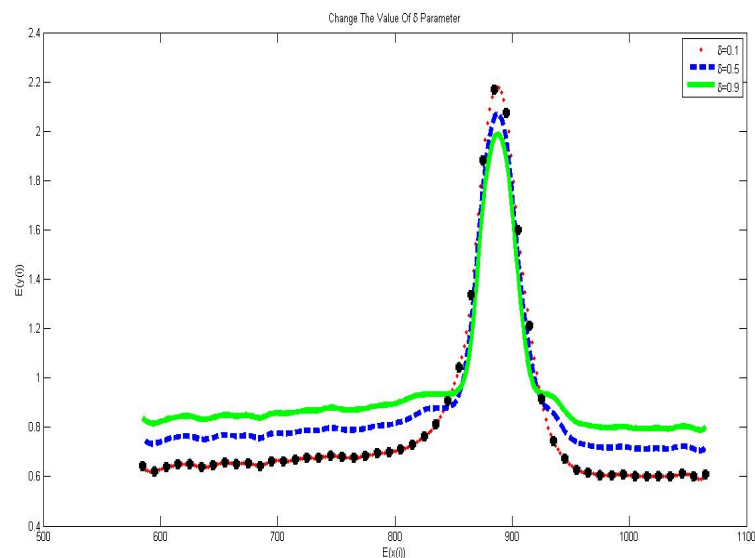


شکل ۱۹.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف C با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۲۰.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف δ نشان داده شده است. شکل (۲۱.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف δ نشان می‌دهد. با توجه به نمودار شکل (۲۰.۴)، به‌ازای مقدار $\delta = 0.1$ دقت برآورد $f(x)$ که در شکل (۲۱.۴) نشان داده شده است، از دو مقدار دیگر δ بیشتر است.

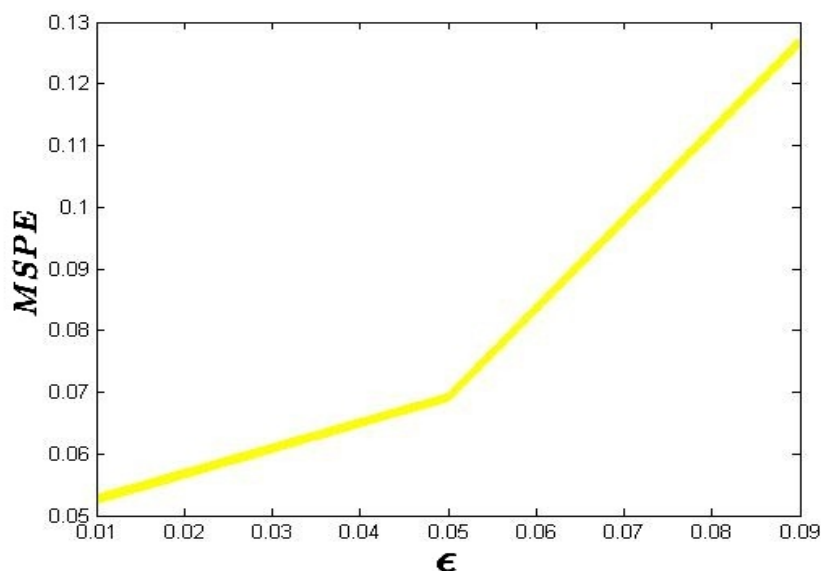


شکل ۲۰.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف δ

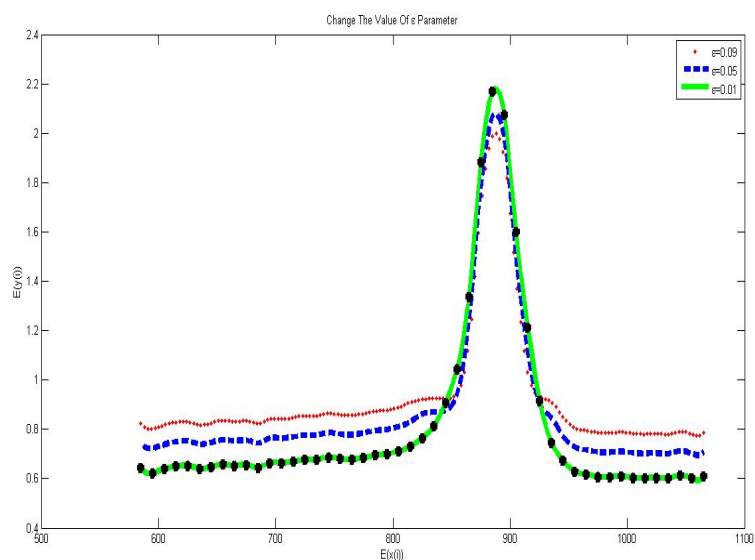


شکل ۲۱.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف δ با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۲۲.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف ϵ نشان داده شده است. شکل (۲۳.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف ϵ نشان می‌دهد. با توجه به نمودار شکل (۲۲.۴)، به‌ازای مقدار $\epsilon = 0.01$ دقت برآورد $f(x)$ که در شکل (۲۳.۴) نشان داده شده است، از دو مقدار دیگر ϵ بیشتر است.

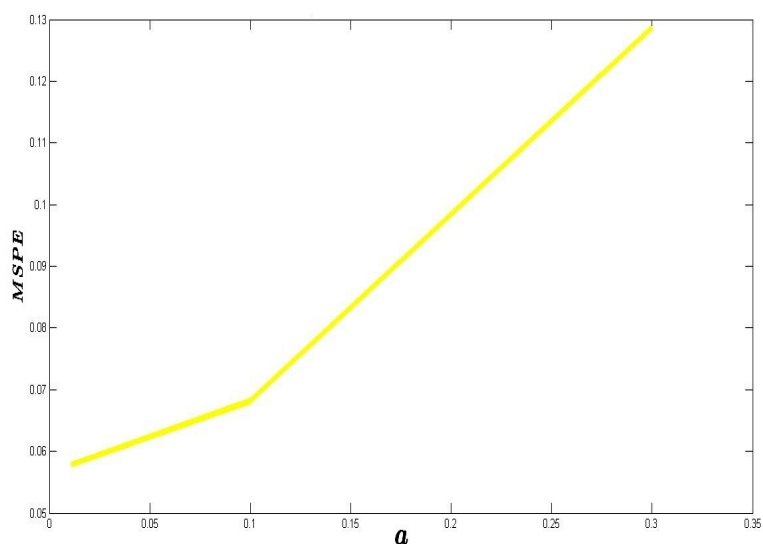


شکل ۲۲.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف ϵ

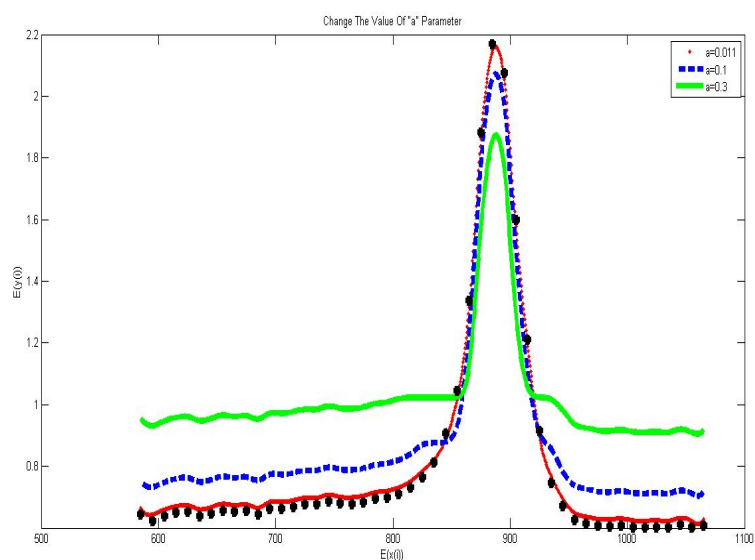


شکل ۲۳.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف ϵ با استفاده از شبکه عصبی (۲۴.۴)

در شکل (۲۴.۴) مقدار $MSPE$ به‌ازای مقادیر مختلف a نشان داده شده است. شکل (۲۵.۴) نیز $f(x)$ را به‌ازای سه مقدار مختلف a نشان می‌دهد. با توجه به نمودار شکل (۲۴.۴)، به‌ازای مقدار $a = 0.11$ دقت برآورد $f(x)$ که در شکل (۲۵.۴) نشان داده شده است، از دو مقدار دیگر a بیشتر است.



شکل ۲۴.۴: تغییرات $MSPE$ نسبت به مقادیر مختلف a



شکل ۲۵.۴: نتایج رگرسیون بردار پشتیبان به‌ازای سه مقدار مختلف a با استفاده از شبکه عصبی (۲۴.۴)

نتایج و پیشنهادات

با استفاده از دانش برنامه‌نویسی رایانه می‌توان ساختار داده‌ای طراحی کرد که همانند یک نرون عمل نماید. سپس با ایجاد شبکه‌ای از این نرون‌های مصنوعی به هم پیوسته، ایجاد یک الگوریتم آموزشی برای شبکه و اعمال این الگوریتم به شبکه آن را آموزش داد. این شبکه‌ها کارایی بسیار بالایی از خود نشان داده‌اند. گستره کاربرد این مدل‌های ریاضی بر گرفته از عملکرد مغز انسان، بسیار وسیع می‌باشد که به عنوان چند نمونه کوچک می‌توان استفاده از این ابزار ریاضی در پردازش سیگنال‌های بیولوژیکی، مخابراتی و الکترونیکی تا کمک در نجوم و فضاوردی را نام برد. ما در این رساله یک مدل شبکه عصبی برای حل مساله رگرسیون بردار پشتیبان بر اساس شرایط بهینگی KKT ارائه دادیم. شبکه عصبی پیشنهاد شده ساختاری ساده دارد. علاوه بر آن نشان دادیم که شبکه عصبی ارائه شده پایدار لیاپانوف و همگرایی سراسری به جواب بهینه مساله رگرسیون بردار پشتیبان است. همچنین کارایی مدل شبکه عصبی ارائه شده با حل مثال‌های متنوع نشان داده شد. همچنین یک مدل شبکه عصبی اصلاح شده برای حل مساله رگرسیون بردار پشتیبان با قیدهای احتمالی ارائه دادیم. نشان دادیم که مدل شبکه عصبی ارائه شده پایدار مجانبی و همگرایی سراسری به جواب بهینه مساله رگرسیون بردار پشتیبان با قیدهای احتمالی است. با ارائه مثال‌هایی کارایی مدل شبکه عصبی پیشنهاد شده تایید شد.

با توجه به این که تاکنون مدل شبکه عصبی برای حل مساله رگرسیون بردار پشتیبان با قیدهای احتمالی ارائه نشده است، در این رساله برای اولین بار یک مدل شبکه عصبی مبتنی بر شرایط KKT برای حل این مسئله بیان شد. با توجه به کاربردهای بسیار گسترده رگرسیون بردار پشتیبان با قیدهای احتمالی نتایج بدست آمده در این رساله می‌تواند مورد توجه قرار گیرد.

به عنوان پیشنهادی برای کارهای آتی می‌توان به ارائه مدل‌هایی از شبکه عصبی برای مسائلی چون مسائل کلاس‌بندی که در آن نقاط، متغیرهایی تصادفی هستند و یا مسائل رگرسیون بردار پشتیبان با متغیرهای بازه‌ای که در آن به‌جای یک نقطه با بازه‌ای از نقاط سرو کار داریم و یا مسائل رگرسیون بردار پشتیبان با نقاط خروجی تصادفی اشاره کرد.

مراجع

- [۱] ارجمندزاده، آمنه. ماشین بردار پشتیبان و رده بندی داده های بازه ای. دانشگاه تربیت معلم سبزوار، دانشکده علوم پایه، پایان نامه کارشناسی ارشد، ۱۳۸۸.
- [۲] کبودیان، جهانشاه، رحمتی، محمد، و همایون پور، محمد مهدی. استفاده از ماشین بردار پشتیبان (svm) در سه مساله شناسایی الگو. اولین کنفرانس بین المللی فناوری اطلاعات و دانش، تهران، ۱۳۸۲.
- [۳] ساده، محمدجواد. رده بندی در داده های ریزآرایه ها. پایان نامه کارشناسی ارشد، دانشگاه فردوسی مشهد، ۱۳۸۴.
- [۴] ضمیری، آذر. شبکه های عصبی مصنوعی. پایان نامه کارشناسی ارشد، دانشگاه فردوسی مشهد، ۱۳۸۷.
- [۵] عباس، قماش لنگرودی. حل مسائل بهینه سازی با استفاده از شبکه های عصبی تصویر. دانشگاه تربیت معلم سبزوار، دانشکده علوم پایه، پایان نامه کارشناسی ارشد، ۱۳۸۳.
- [6] Hebb, D.O. the organization of behavior: A neuropsychological theory. new york: John wiley and sons, inc., 1949. *Science Education*, 34:336–337, 1950.
- [7] Widrow, B. generalization and information storage in networks of adaline neurons. in self organizing systems. *Spartan Books*, 1959.
- [8] Abaszade, M. and Effati, S. Stochastic support vector regression with probabilistic constraints. *Applied Intelligence*, 48:243–256, 2018.
- [9] Mordecai, A. *Nonlinear programming : analysis and methods / Mordecai Avriel*. Englewood Cliffs (N.J.) : Prentice-Hall, 1976.
- [10] Barnett, N.S., Dragomir, S.S. and Agarwal, R.P. Some inequalities for probability, expectation, and variance of random variables defined over a finite interval. *Computers and Mathematics with Applications*, 43:1319–1357, 2002.

- [11] Bazaraa, M.S., Sherali, H.D., and Shetty, C.M. Nonlinear programming, theory and algorithms (2nd edn), 2013.
- [12] Bouzerdoun, A. and Pattison, T.R. Neural network for quadratic optimization with bound constraints. *IEEE Transactions on Neural Networks*, 4:293–304, 1993.
- [13] Dierckx, P. Curve and surface fitting with splines, chapter univariate splines, 1993.
- [14] Effati, S. and Baymani, M. A new nonlinear neural network for solving convex nonlinear programming problems. *Applied Mathematics and Computation*, 168:1370 – 1379, 2005.
- [15] Effati, S., Ghomashi, A. and Nazemi, A.R. Application of projection neural network in solving convex programming problems. *Applied Mathematics and Computation*, 188:1103 – 1114, 2007.
- [16] Effati, S. and Nazemi, A.R. Neural network models and its application for solving linear and quadratic programming problems. *Applied Mathematics and Computation*, 172:305 – 331, 2006.
- [17] Effati, S., Ghomashi, A. and Abbasi, M. A novel recurrent neural network for solving mlcps and its application to linear and quadratic programming. *Asia-Pacific Journal of Operational Research*, 28:523–541, 2011.
- [18] Forti, M., Nistri, P., and Quincampoix, M. Convergence of neural networks for programming problems via a nonsmooth lojasiewicz inequality. *IEEE Transactions on Neural Networks*, 17:1471–1486, 2006.
- [19] Friesz, T.L., Bernstein, D., Mehta, N.J., Tobin, R.L. and Ganjalizadeh, S. Day-to-day dynamic network disequilibria and idealized traveler information systems. *Operations Research*, 42:1120–1136, 1994.
- [20] Hopfield, J.J. Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences*, 79:2554–2558, 1982.
- [21] Hopfield, J.J. and Tank, D.W. "neural" computation of decisions in optimization problems. *Biological cybernetics*, 52:141–152, 1985.
- [22] Kennedy, M.P. and Chua, L.O. Neural networks for nonlinear programming. *IEEE Transactions on Circuits and Systems*, 35:554–562, 1988.

- [23] Khalil, H.K. and Grizzle, J.W. *Nonlinear systems*, volume 3. Prentice hall Upper Saddle River, NJ, 2002.
- [24] Lillo, W.E., Loh, M.H., Hui, S. and Zak, S.H. On solving constrained optimization problems with neural networks: a penalty method approach. *IEEE Transactions on Neural Networks*, 4(6):931–940, 1993.
- [25] Luenberger, D. and Ye, Y. *Linear and Nonlinear Programming*. Springer US, 2008.
- [26] Maa, C.Y. and Shanblatt, M.A. Linear and quadratic programming neural network analysis. *IEEE transactions on neural networks*, 3:580–594, 1992.
- [27] Maa, C.Y. and Schanblatt, M.A. A two-phase optimization neural network. *IEEE transactions on neural networks*, 3:1003—1009, 1992.
- [28] Maugeri, A. and Zichichi, A. *Variational analysis and applications*. New York, Springer, 2005.
- [29] Warren, S.M. and Walter, P. A logical calculus of the ideas immanent in nervous activity. *The bulletin of mathematical biophysics*, 5:115–133, 1943.
- [30] Miller, R.K. and Michel, A.N. Ordinary differential equations. academic press, new york et. al. 1982.
- [31] Minsky, M.L. and Perceptrons, S.P. An introduction to computational geometry. *Mit Press*, 1972.
- [32] Moghaddas, M. and Eshaghnejad, M. A recurrent neural network for solving non-convex nonlinear optimization problem. *The 44th Annual Iranian Mathematics Conference, Mashhad, Iran,, 2013*.
- [33] Moghaddas, M. and Effati, S. A novel recurrent neural network based on ncp function for solving convex nonlinear optimization problems. *The 5th Iranian Conference on Applied Mathematics, Hamadan, Iran,, 2013*.
- [34] Nazemi, A.R. Solving general convex nonlinear optimization problems by an efficient neurodynamic model. *Engineering Applications of Artificial Intelligence*, 26:685–696, 2013.
- [35] Nazemi, A.R. A neural network model for solving convex quadratic programming problems with some applications. *Engineering Applications of Artificial Intelligence*, 32:54–62, 2014.

- [36] Nazemi, A.R. and Nazemi, M. A gradient-based neural network method for solving strictly convex quadratic programming problems. *Cognitive Computation*, 6:484–495, 2014.
- [37] Nazemi, A.R. and Omid, F. A capable neural network model for solving the maximum flow problem. *Journal of Computational and Applied Mathematics*, 236:3498 – 3513, 2012.
- [38] Nazemi, A.R. A dynamic system model for solving convex nonlinear optimization problems. *Communications in Nonlinear Science and Numerical Simulation*, 17:1696 – 1705, 2012.
- [39] Perko, L. *Differential equations and dynamical systems*. Springer, 3rd ed, 2001.
- [40] Quarteroni, A., Sacco, R. and Saleri, F. *Numerical mathematics*, volume 37. Springer-Verlag, Berlin, second edition, 2007.
- [41] Rodriguez-Vazquez, A., Castro, R., Rueda, A., Huertas, J.L. and Sanchez-Sinencio, E. Nonlinear switched-capacitor ‘neural network’ for optimisation problems. *Circuits and Systems, IEEE Transactions on*, 37:384 – 398, 1990.
- [42] Rosenblatt, F. The perception: a probabilistic model for information storage and organization in the brain, neurocomputing: foundations of research, 1988.
- [43] Rosenblatt, F. Principles of neurodynamics. perceptrons and the theory of brain mechanisms. Technical report, Cornell Aeronautical Lab Inc Buffalo NY, 1961.
- [44] Rumelhart, D.E., Hinton, G.E. and Williams, R.J. Learning representations by back-propagating errors. *nature*, 323:533–536, 1986.
- [45] Schölkopf, B., Smola, A.J. and et al. Learning with kernels: support vector machines, regularization. *Optimization, and Beyond. MIT press*, 1(2), 2002.
- [46] Tank, D.W. and Hopfield, J.J. Simple ‘neural’ optimization networks: An a/d converter, signal decision circuit, and a linear programming circuit. *IEEE Transactions on Circuits and Systems*, CAS-33:533–541, 1986.
- [47] Tao, Q., Cao, J. and Sun, D. A simple and high performance neural network for quadratic programming problems. *Applied Mathematics and Computation*, 124:251–260, 2001.

- [48] Tax, D. One-class classification: Concept learning in the absence of counter-examples. 2002.
- [49] Ortega, T.M. and Rheinboldt, W.C. *Iterative solution of nonlinear equations in several variables*. Academic Press, New York, 1970.
- [50] Vapnik, V. Statistical learning theory wiley. *New York*, 1:624, 1998.
- [51] Werbos, P.J. Beyond regression: New tools for prediction and analysis in the behavioral sciences. ph. d. thesis, harvard university, cambridge, ma, 1974. 1974.
- [52] Widrow, B. and Hoff, M.E. Adaptive switching circuits (pp. 96–104). In *Institute of Radio Engineers (IRE) WESCON Convention record*, 1960.
- [53] Widrow, B. and Lehr, M.A. 30 years of adaptive neural networks: perceptron, madaline, and backpropagation. *Proceedings of the IEEE*, 78:1415–1442, 1990.
- [54] Widrow, B. and Winter, R. Neural nets for adaptive filtering and adaptive pattern recognition. *Computer*, 21:25–39, 1988.
- [55] Xia, Y. A new neural network for solving linear and quadratic programming problems. *IEEE transactions on neural networks*, 7:1544–1548, 1996.
- [56] Xia, Y., Feng, G., and Wang, J. A recurrent neural network with exponential convergence for solving convex quadratic program and related linear piecewise equations. *Neural Networks*, 17:1003–1015, 2004.
- [57] Xia, Y., Feng, G., and Wang, J. A novel recurrent neural network for solving nonlinear optimization problems with inequality constraints. *IEEE transactions on neural networks / a publication of the IEEE Neural Networks Council*, 19:1340–53, 2008.
- [58] Xia, Y. and Wang, J. A recurrent neural network for solving linear projection equations. *Neural Networks*, 13:337–350, 2000.
- [59] Xu, H. Projection neural networks for solving constrained convex and degenerate quadratic problems. *2010 IEEE International Conference on Intelligent Computing and Intelligent Systems*, 3:91–96, 2010.
- [60] Xue, X. and Bian, W. A project neural network for solving degenerate convex quadratic program. *Neurocomputing*, 70:2449–2459, 2007.

-
- [61] Yang, Y. and Cao, J. A feedback neural network for solving convex constraint optimization problems. *Applied Mathematics and Computation*, 201(1):340 – 350, 2008.
- [62] Xia, Y. and Wang, J. A general methodology for designing globally convergent optimization neural networks. *IEEE Transactions on Neural Networks*, 9:1331–1343, 1998.
- [63] Xia, Y. and Wang, J. A general projection neural network for solving monotone variational inequalities and related optimization problems. *IEEE Transactions on Neural Networks*, 15:318–328, 2004.
- [64] Xia, Y. and Wang, J. A recurrent neural network for nonlinear convex optimization subject to nonlinear inequality constraints. *IEEE Transactions on Circuits and Systems I: Regular Papers*, 51:1385–1394, 2004.
- [65] Xia, Y. and Wang, J. A recurrent neural network for solving nonlinear convex programs subject to linear constraints. *IEEE Transactions on Neural Networks*, 16:379–386, 2005.
- [66] Xia, Y., Leung, H. and Wang, J. A projection neural network and its application to constrained optimization problems. *IEEE Transactions on Circuits and Systems I: Fundamental Theory and Applications*, 49:447–458, 2002.
- [67] Zabczyk, J. Mathematical control theory. *An Introduction*, 1992.
- [68] Zhang, G., Shen, Y., Yin, Q. and Sun, J. Global exponential periodicity and stability of a class of memristor-based recurrent neural networks with multiple delays. *Inf. Sci.*, 232:386–396, 2013.

واژه‌نامه فارسی به انگلیسی

Back propagation algorithm	الگوریتم پس انتشار
Cross validation	اعتبار سنجی متقابل
ϵ - insensitive	ϵ - حساسیت
Optimum separating hyperplane	ابر صفحه بهینه جدا کننده
Adeline	آدلاین
LaSalle principle of invariance	اصل تغییرناپذیری لاسال
Axon	آکسون
Strictly monotone	اکیداً یکنوا
Learning	آموزش
Nonexpansive	انقباضی
Radially unbounded	به طور شعاعی بی کران
Convex optimization	بهینه سازی محدب
Nonconvex optimization	بهینه سازی نامحدب
Supervisor response	پاسخ سرپرست
Postsynaptic terminal	پایانه بعد از سیناپسی
Presynaptic terminal	پایانه قبل از سیناپسی
Globally asymptotically stable	پایدار مجانبی سراسری
Stability in the sense of Lyapunov	پایداری به مفهوم لیاپانوف
Globally exponential stability	پایداری نمایی سراسری
Perceptron	پرسپترون
Lipschitz continuous	پیوسته لیپشیتز
Radial basis function	تابع پایه شعاعی
Activation function	تابع فعال سازی
Risk function	تابع مخاطره
Firing	تحریک شدن
Generalization portrait	تعمیم تصویر

Generator	تولید کننده
Feasible solution	جواب شدنی
Margin	حاشیه
Binary classification	دسته‌بندی دوتایی
Dendrite	دندریت
Local duality	دوگانی موضعی
ϵ - SV Regression	رگرسیون ϵ - بردار پشتیبان
Supervisor	سرپرست
Dynamic system	سیستم دینامیکی
Synapse	سیناپس
Feedback (Recurrent) neural network	شبکه عصبی بازگشتی
Feed forward neural network	شبکه عصبی پیش‌خورد
Switched-Capacitor neural network	شبکه عصبی سوئیچ-خازن
Artificial neural network	شبکه عصبی مصنوعی
Second-order sufficient conditions	شرایط کافی مرتبه دوم
Synaptic cleft	شکاف سیناپسی
Pattern recognition	شناسایی الگو
Lagrangian multiplier	ضرب‌گر لاگرانژ
Projection operator	عملگر تصویر
Adaptive linear element	عنصر خطی تطبیقی
Feature space	فضای ویژگی
Brouwer's fixed point theorem	قضیه نقطه ثابت بروئر
Strongly monotone	قویاً یکنوا
Hessian matrix	ماتریس هسین
Madaline	مادلاین
Support vector machine	ماشین بردار پشتیبان
Dual variable	متغیر دوگان
Invariant set	مجموعه تغییر ناپذیر
Level set	مجموعه سطح
Lagrangian dual problem	مساله دوگان لاگرانژ
Linear complementarity problem (LCP)	مساله مکمل خطی
Nonlinear complementarity problem (NCP)	مساله مکمل غیرخطی
Mean squared prediction error	میانگین توان دوم خطای پیش‌بینی
Strongly positive definite	معین مثبت قوی

Positive definite	معین مثبت
Variational inequality	نامساوی وردشی
Neuron	نرون
Isolated equilibrium point	نقطه تعادل تنها
Equilibrium point	نقطه تعادل
Regular point	نقطه منظم
Strict local minimum point	نقطه مینیمم محلی اکید
Positive semidefinite	نیمه معین مثبت
Kernel	هسته
Globally convergent	همگرای سراسری
Monotone	یکنوا

واژه‌نامه انگلیسی به فارسی

Artificial neural network	شبکه عصبی مصنوعی
Binary classification	دسته‌بندی دوتایی
ϵ - insensitive	ϵ - حساسیت
ϵ - SV Regression	رگرسیون ϵ - بردار پشتیبان
Activation function	تابع فعال‌سازی
Adaptive linear element	عنصر خطی تطبیقی
Adeline	آدلین
Axon	آکسون
Back propagation algorithm	الگوریتم پس انتشار
Brouwer's fixed point theorem	قضیه نقطه ثابت بروئر
Cross validation	اعتبار سنجی متقابل
Convex optimization	بهینه‌سازی محدب
Dendrite	دندریت
Dual variable	متغیر دوگان
Dynamic system	سیستم دینامیکی
Equilibrium point	نقطه تعادل
Feasible solution	جواب شدنی
Feature space	فضای ویژگی
Feed forward neural network	شبکه عصبی پیش‌خورد
Feedback (Recurrent) neural network	شبکه عصبی بازگشتی
Firing	تحریک شدن
Generalization portrait	تعمیم تصویر
Generator	تولید کننده
Globally asymptotically stable	پایدار مجانبی سراسری
Globally convergent	همگرای سراسری
Globally exponential stability	پایداری نمایی سراسری

Hessian matrix	ماتریس هسین
Invariant set	مجموعه تغییر ناپذیر
Isolated equilibrium point	نقطه تعادل تنها
Kernel	هسته
Lagrangian dual problem	مساله دوگان لاگرانژ
Lagrangian multiplier	ضرب‌گر لاگرانژ
LaSalle principle of invariance	اصل تغییرناپذیری لاسال
Learning	آموزش
Level set	مجموعه سطح
Linear complementarity problem (LCP)	مساله مکمل خطی
Lipschitz continuous	پیوسته لیپ‌شیتز
Local duality	دوگانی موضعی
Madaline	مادلاین
Margin	حاشیه
Mean squared prediction error	میانگین توان دوم خطای پیش‌بینی
Monotone	یکنوا
Neuron	نرون
Nonconvex optimization	بهینه‌سازی نامحدب
Nonexpansive	انقباضی
Nonlinear complementarity problem (NCP)	مساله مکمل غیرخطی
Optimum separating hyperplane	ابرصفحه بهینه جدا کننده
Pattern recognition	شناسایی الگو
Perceptron	پرسپترون
Positive definite	معین مثبت
Positive semidefinite	نیمه معین مثبت
Postsynaptic terminal	پایانه بعد از سیناپسی
Presynaptic terminal	پایانه قبل از سیناپسی
Projection operator	عملگر تصویر
Radial basis function	تابع پایه شعاعی
Radially unbounded	به‌طور شعاعی بی‌کران
Regular point	نقطه منظم
Risk function	تابع مخاطره
Second-order sufficient conditions	شرایط کافی مرتبه دوم
Stability in the sense of Lyapunov	پایداری به مفهوم لیاپانوف

Strict local minimum point	نقطه مینیمم محلی اکید
Strictly monotone	اکیداً یکنوا
Strongly monotone	قویاً یکنوا
Strongly positive definite	معین مثبت قوی
Supervisor response	پاسخ سرپرست
Supervisor	سرپرست
Support vector machine	ماشین بردار پشتیبان
Switched-Capacitor neural network	شبکه عصبی سوئیچ-خازن
Synapse	سیناپس
Synaptic cleft	شکاف سیناپسی
Variational inequality	نامساوی وردشی

پیوست

۱. مقدمه‌ای بر بهینه‌سازی غیرخطی

در پیوست رساله مطالبی را که خواننده احیاناً برای مطالعه بهتر رساله نیاز به آن پیدا می‌کند، آورده شده است.

تعریف ۱.۱. [۲۵] تابع $f : \mathbb{R}^n \rightarrow \mathbb{R}$ روی مجموعه محدب و ناتهی $S \subseteq \mathbb{R}^n$ محدب نامیده می‌شود اگر برای هر $x, y \in S$ و $\alpha \in (0, 1)$ داشته باشیم

$$f(\alpha x + (1 - \alpha)y) \leq \alpha f(x) + (1 - \alpha)f(y).$$

تابع f روی S به طور اکید محدب نامیده می‌شود اگر نامساوی بالا زمانی که $x \neq y$ است، به‌طور اکید برقرار باشد.

یک مسأله برنامه‌ریزی غیرخطی در حالت کلی به‌صورت زیر در نظر گرفته می‌شود

$$\text{minimize} \quad f(x), \quad (32)$$

$$\text{subject to} \quad g_i(x) \leq 0, \quad (33)$$

$$h_j(x) = 0, \quad (34)$$

$$x \in X \subset \mathbb{R}^n, \quad (35)$$

که $f : \mathbb{R}^n \rightarrow \mathbb{R}$ ، $g_i : \mathbb{R}^n \rightarrow \mathbb{R}$ ($i = 1, \dots, m$) و $h_j : \mathbb{R}^n \rightarrow \mathbb{R}$ ($j = 1, \dots, p$) توابع دو بار مشتق پذیر پیوسته و X زیر مجموعه‌ای باز از $\mathbb{R}^n$ است. برای تسهیل در نمادگذاری، توابع برداری $g = (g_1, g_2, \dots, g_m)^T$ و $h = (h_1, h_2, \dots, h_p)^T$ را معرفی می‌کنیم و (۳۲)–(۳۵) را به‌صورت زیر می‌نویسیم

$$\text{minimize} \quad f(x), \quad (36)$$

$$\text{subject to} \quad g(x) \leq 0, \quad (37)$$

$$h(x) = 0, \quad (38)$$

$$x \in X \subset \mathbb{R}^n, \quad (39)$$

که محدودیت‌های $g(x) \leq 0$ و $h(x) = 0$ ، محدودیت‌های تابعی و محدودیت $x \in X \subset \mathbb{R}^n$ ، محدودیت مجموعه‌ای نامیده می‌شود. هرگاه توابع برداری f و g روی X محدب باشند و h آفین باشد یعنی $h = Ax - b$ ، $A \in \mathbb{R}^{p \times n}$ ، $rank(A) = p$ ($0 \leq p < n$) و $b \in \mathbb{R}^p$ باشد، آنگاه (۳۶)–(۳۹) یک مسأله بهینه‌سازی محدب و در غیر این صورت یک مسأله بهینه‌سازی نامحدب نامیده می‌شود.

تعریف ۲.۱. [۲۵] نقطه x جواب شدنی مسأله (۳۶)–(۳۹) نامیده می‌شود هرگاه در محدودیت‌های مسأله صدق کند. نقطه شدنی x یک نقطه منظم نامیده می‌شود اگر بردارهای گرادیان $\nabla g_i(x)$ ، $i \in \{i | g_i(x) = 0\}$ و $j = 1, \dots, p$ ، $\nabla h_j(x)$ مستقل خطی باشند.

قضیه ۱.۱. (شرایط کافی مرتبه دوم) [۲۵] فرض کنید x^* یک نقطه شدنی برای مسأله (۳۶)–(۳۹) باشد. اگر $\mu \in \mathbb{R}^m$ و $\lambda \in \mathbb{R}^p$ وجود داشته باشند به‌طوری که

$$\begin{aligned}\mu &\geq 0, \\ \mu^T g(x^*) &= 0, \\ \nabla f(x^*) + \lambda^T \nabla h(x^*) + \mu^T \nabla g(x^*) &= 0,\end{aligned}$$

و ماتریس هسین

$$\nabla^2 L(x^*, \lambda, \mu) = \nabla^2 f(x^*) + \sum_{i=1}^m \lambda_i \nabla^2 g_i(x^*) + \sum_{j=1}^p \mu_j \nabla^2 h_j(x^*),$$

بر زیرفضای

$$M(x^*) = \{d | \nabla h(x^*)d = 0, \nabla g_j(x^*)d = 0, \forall j \in J(x^*)\},$$

که در آن

$$J(x^*) = \{j | g_j(x^*) = 0, \mu_j > 0\},$$

معین مثبت باشد، آنگاه x^* یک نقطه مینیمم محلی اکید برای مسأله (۳۶)–(۳۹) است.

مسأله دوگان لاگرانژ

مسأله برنامه‌ریزی غیرخطی (۳۶)–(۳۹) را در نظر بگیرید، این مسأله را مسأله اولیه (P) نیز می‌نامیم. مسأله دوگان لاگرانژ (D) به صورت زیر تعریف می‌شود

$$\max \quad \theta(u, v), \quad (40)$$

$$\text{subject to } u \geq 0, \quad (41)$$

که در آن

$$\theta(u, v) = \inf_x \left\{ f(x) + \sum_{i=1}^m u_i g_i(x) + \sum_{j=1}^p v_j h_j(x), x \in X \right\},$$

$$u = (u_1, u_2, \dots, u_m)^T, \quad v = (v_1, v_2, \dots, v_p)^T,$$

u_i -امین مؤلفه از u ، متغیر دوگان یا ضرب‌گر لاگرانژ مربوط به قید $g_i(x) \leq 0$ می‌باشد و v_j -امین مؤلفه v ، متغیر دوگان یا ضرب‌گر لاگرانژ مربوط به قید $h_j(x) = 0$ می‌باشد. مسأله دوگان لاگرانژ (۴۰) و (۴۱) در شکل برداری به‌صورت زیر نوشته می‌شود

$$\text{maximize } \theta(u, v), \quad (42)$$

$$\text{subject to } u \geq 0, \quad (43)$$

که در آن

$$\theta(u, v) = \inf_x \{ f(x) + u^T g(x) + v^T h(x), x \in X \},$$

هم چنین $f: \mathbb{R}^n \rightarrow \mathbb{R}$ ، $g = (g_1, \dots, g_m)^T: \mathbb{R}^n \rightarrow \mathbb{R}^m$ و $h = (h_1, \dots, h_p): \mathbb{R}^n \rightarrow \mathbb{R}^p$ هستند.

قضیه ۲.۱.۰. (دوگانی موضعی) [۲۵] فرض کنید مسأله

$$\text{minimize } f(x),$$

$$\text{subject to } h(x) = 0,$$

$$x \in X \subset \mathbb{R}^n,$$

که $h(x) \in \mathbb{R}^m$ و $f, h \in C^2$ ، دارای یک جواب موضعی در x^* با مقدار متناظر r^* و ضریب لاگرانژ v^* است. هم‌چنین فرض کنید که x^* یک نقطه منظم برای قیود است و هسین لاگرانژ متناظر $\nabla^2 L(x^*, v^*)$ معین مثبت است. در این صورت مسأله دوگان

$$\max \theta(v),$$

که در آن

$$\theta(v) = \inf_x \{ f(x) + v^T h(x), x \in X \},$$

دارای یک جواب موضعی در v^* با مقدار متناظر r^* است و x^* جواب متناظر با v^* در تعریف θ است.

به‌سادگی می‌توان نتایج فوق را به مسائلی که علاوه بر قیود تساوی قیود نامساوی هم دارند تعمیم داد [۲۵].

۲. سیستم‌های دینامیکی

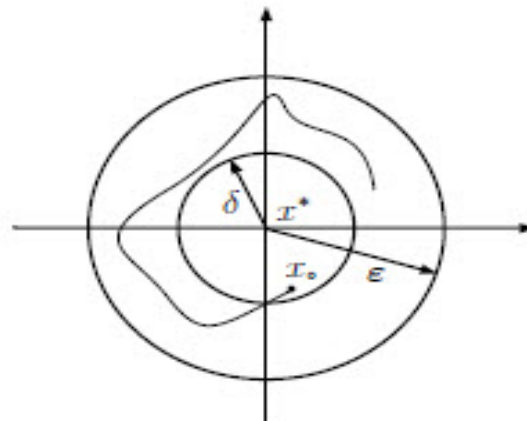
تعریف ۱.۲. [۳۹] سیستم دینامیکی زیر را در نظر بگیرید:

$$\dot{x} = f(x(t)), \quad x(t_0) = x_0 \in \mathbb{R}^n, \quad (44)$$

که در آن f یک تابع برداری از $\mathbb{R}^n$ به $\mathbb{R}^n$ می‌باشد. x^* یک نقطه تعادل (۴۴) نامیده می‌شود اگر $f(x^*) = 0$. اگر همسایگی $\Omega^* \subseteq \mathbb{R}^n$ از x^* وجود داشته باشد که $f(x^*) = 0$ و $\forall x \in \Omega^* \setminus \{x^*\}$ ، $f(x) \neq 0$ ، آن‌گاه x^* یک نقطه تعادل تنها نامیده می‌شود.

تعریف ۲.۲. (پایداری به مفهوم لیاپانوف) [۳۹] فرض می‌کنیم که $x(t)$ یک جواب (۴۴) باشد، نقطه تعادل تنها x^* ، پایدار به مفهوم لیاپانوف است اگر برای هر $x_0 = x(t_0)$ و هر $\varepsilon > 0$ ، یک $\delta > 0$ وجود داشته باشد به‌طوری که اگر $\|x(t_0) - x^*\| < \delta$ آن‌گاه

$$\|x(t) - x^*\| < \varepsilon, \quad \forall t \geq t_0.$$



شکل ۲۶: پایداری لیاپانوف.

تعریف ۳.۲. [۳۹] سیستم دینامیکی (۴۴) همگرای سراسری به مجموعه جواب‌های (۴۴) که با Ω^* نمایش داده می‌شود، گفته می‌شود اگر برای هر نقطه آغازین دلخواه، مسیر $x(t)$ در رابطه زیر صدق کند:

$$\lim_{t \rightarrow \infty} \text{dist}(x(t), \Omega^*) = 0,$$

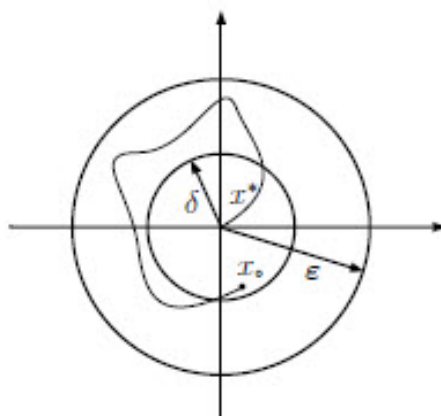
که در آن

$$\text{dist}(x, \Omega^*) = \inf_{y \in \Omega^*} \|x - y\|,$$

و $\|\cdot\|$ نرم l_2 را نشان می‌دهد.

تعریف ۴.۲. [۳۹] سیستم دینامیکی (۴۴) در نقطه تعادل x^* که یکتا می‌باشد، پایدار مجانبی سراسری نامیده می‌شود اگر x^* به مفهوم لیاپانوف پایدار باشد و در شرط زیر صدق کند:

$$\lim_{t \rightarrow \infty} x(t) = x^*.$$



شکل ۲۷: پایداری مجانبی.

تعریف ۵.۲. [۳۹] سیستم دینامیکی (۴۴) در نقطه x^* دارای پایداری نمایی سراسری با نرخ β است اگر برای هر نقطه آغازین دلخواه، مسیر $x(t)$ در رابطه زیر صدق کند

$$\|x(t) - x^*\| \leq \alpha \|x(t_0) - x^*\| e^{-\beta(t-t_0)}, \quad \forall t \geq t_0,$$

که در آن α و β ثابت‌های مثبت و مستقل از نقاط آغازین هستند. واضح است که پایداری نمایی سراسری پایداری مجانبی سراسری را نتیجه می‌دهد.

تعریف ۶.۲. [۳۹] مجموعه $M \subseteq \mathbb{R}^n$ یک مجموعه تغییر ناپذیر نسبت به سیستم (۴۴) گفته می‌شود اگر $x(t_0) \in M$ و $t_0 \geq 0$ ، به ازای هر $t \geq t_0$ ، $x(t) \in M$ باشد.

تعریف ۷.۲. [۳۹] فرض می‌کنیم که $\Omega \subset \mathbb{R}^l$ یک مجموعه بسته و محدب باشد که ممکن است بیکران باشد. $P_\Omega(x): \mathbb{R}^l \rightarrow \Omega$ یک عملگر تصویر نامیده می‌شود و به صورت زیر تعریف می‌شود

$$P_\Omega(x) = \arg \min_{v \in \Omega} \|x - v\|,$$

که $\|\cdot\|$ نرم l_2 را نشان می‌دهد.

تعریف ۸.۲. [۳۹] تابع $F: E \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$ روی E در شرط لیپ شیتز صدق می‌کند (پیوسته لیپ شیتز است) اگر ثابت L وجود داشته باشد به طوری که برای هر دو نقطه $x, y \in E$ داشته باشیم:

$$\|F(x) - F(y)\| \leq L \|x - y\|.$$

تابع F پیوسته لیپ شیتز محلی روی E نامیده می‌شود اگر برای نقطه $x_0 \in E$ ، یک $\varepsilon -$ همسایگی مانند $N_\varepsilon(x_0) \subset E$ و یک ثابت L_0 وجود داشته باشد به طوری که برای هر دو نقطه $x, y \in N_\varepsilon(x_0)$ داشته باشیم:

$$\|F(x) - F(y)\| \leq L_0 \|x - y\|.$$

قضیه ۱.۲. [۳۹] تابع $F : E \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$ را در نظر بگیرید. اگر $F \in C^1(E)$ ، آن گاه F روی E لیپ شیتز محلی است.

تعریف ۹.۲. نگاشت $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$

(الف) یکنوا نامیده می شود اگر برای هر دو نقطه $x, y \in \mathbb{R}^n$ داشته باشیم

$$\langle x - y, F(x) - F(y) \rangle \geq 0.$$

(ب) اکیداً یکنوا نامیده می شود اگر برای هر دو نقطه متمایز $x, y \in \mathbb{R}^n$ داشته باشیم

$$\langle x - y, F(x) - F(y) \rangle > 0.$$

(پ) قویاً یکنوا نامیده می شود اگر یک ثابت $\beta > 0$ وجود داشته باشد که برای هر دو نقطه $x, y \in \mathbb{R}^n$ داشته باشیم

$$\langle x - y, F(x) - F(y) \rangle \geq \beta \|x - y\|^2.$$

که $\langle \cdot, \cdot \rangle$ ضرب داخلی را نشان می دهد.

تعریف ۱۰.۲. [۴۹] اگر $g : \Omega \subset \mathbb{R}^l \rightarrow \mathbb{R}$ ، آن گاه هر مجموعه غیر تهی به شکل زیر

$$L(r) = \{u \in \Omega \mid g(u) \leq r, r \in \mathbb{R}\},$$

یک مجموعه سطح از g نامیده می شود.

تعریف ۱۱.۲. [۴۹] یک ماتریس $n \times n$ ، $M(x)$ که عناصر آن m_{ij} ، $i = 1, \dots, n$ و $j = 1, \dots, n$ توابعی هستند که روی مجموعه $S \subset \mathbb{R}^n$ تعریف شده اند، روی S نیمه معین مثبت نامیده می شود اگر

$$v^T M(x) v \geq 0, \forall v \in \mathbb{R}^n, \forall x \in S.$$

ماتریس $M(x)$ روی S معین مثبت قوی نامیده می شود اگر عددی مثبت مانند α وجود داشته باشد به قسمی که:

$$v^T M(x) v \geq \alpha \|v\|^2, \forall v \in \mathbb{R}^n, \forall x \in S.$$

اگر $\gamma(x)$ کوچکترین مقدار ویژه ماتریس $M(x)$ باشد سه حالت زیر را می توان نتیجه گرفت:

۱. $M(x)$ روی S نیمه معین مثبت است اگر و فقط اگر به ازای هر $x \in S$ ، $\gamma(x) \geq 0$.

۲. $M(x)$ روی S معین مثبت است اگر و فقط اگر به ازای هر $x \in S$ ، $\gamma(x) > 0$.

۳. $M(x)$ روی S قویاً معین مثبت است اگر و فقط اگر به ازای هر $x \in S$ ، $\gamma(x) \geq \alpha > 0$.

تعریف ۱۲.۲. [۴۹] فرض می‌کنیم $V: \mathbb{R}^n \rightarrow \mathbb{R}$ یک تابع به‌طور پیوسته مشتق پذیر باشد. $V(x)$ به‌طور شعاعی بی‌کران گفته می‌شود اگر

$$\|x\| \rightarrow \infty \Rightarrow V(x) \rightarrow \infty.$$

لم ۱.۲. [۴۹] فرض کنیم که $g: \Omega_1 \subset \mathbb{R}^l \rightarrow \mathbb{R}$ باشد که در آن Ω_1 بی‌کران است. آن‌گاه همه مجموعه‌های سطح g کراندار می‌باشند اگر و فقط اگر

$$\lim_{k \rightarrow \infty} g(u^k) = +\infty,$$

که در آن $\{u^k\} \subset \Omega_1$ و $\lim_{k \rightarrow \infty} \|u^k\| = +\infty$.

لم ۲.۲. [۴۹] فرض کنیم که $g: \Omega_1 \subset \mathbb{R}^l \rightarrow \mathbb{R}$ روی مجموعه بسته Ω_1 پیوسته باشد آن‌گاه همه مجموعه‌های سطح g کراندار می‌باشند اگر و فقط اگر مجموعه مینیمم‌کننده‌های g غیرتهی و کراندار باشند.

قضیه ۲.۲. (قضیه نقطه ثابت بروئر) [۴۹] تابع پیوسته $f: \Omega \rightarrow \Omega$ ، که Ω یک مجموعه فشرده و محدب می‌باشد حداقل یک نقطه ثابت دارد.

قضیه ۳.۲. [۴۹] فرض کنیم که تابع $F: K \subset \mathbb{R}^l \rightarrow \mathbb{R}^l$ روی K به‌طور پیوسته مشتق‌پذیر باشد، آن‌گاه F روی K یکنوا (اکیداً یکنوا) است اگر و فقط اگر ماتریس ژاکوبین آن $\nabla F(x)$ برای هر $x \in K$ نیمه‌معین مثبت (معین مثبت) باشد.

قضیه ۴.۲. [۴۹] فرض کنیم که تابع $F(x)$ روی K به‌طور پیوسته مشتق‌پذیر باشد و ژاکوبین $F(x)$ قویاً معین مثبت باشد آن‌گاه $F(x)$ قویاً یکنوا می‌باشد.

قضیه ۵.۲. [۳۹] در سیستم دینامیکی (۴۴) فرض کنیم که f یک تابع پیوسته از $\mathbb{R}^n$ به $\mathbb{R}^n$ باشد، آن‌گاه برای هر $t_0 \geq 0$ و $x_0 \in \mathbb{R}^n$ ، یک جواب محلی $x(t)$ به ازای $t \in [t_0, \tau)$ که $\tau > t_0$ وجود دارد. علاوه براین اگر f در x_0 پیوسته لیپ‌شیتز محلی باشد آن‌گاه جواب یکتاست و اگر f در $\mathbb{R}^n$ پیوسته لیپ‌شیتز باشد آن‌گاه τ می‌تواند تا $+\infty$ توسعه داده شود.

تعریف ۱۳.۲. [۳۹] اگر یک جواب محلی تعریف شده در بازه $[t_0, \tau)$ نتواند به یک جواب محلی روی یک بازه بزرگتر $[t_0, \tau_1)$ که $\tau_1 > \tau$ می‌باشد گسترش یابد، آن‌گاه یک جواب ماکسیمال نامیده می‌شود و بازه $[t_0, \tau)$ ، بازه ماکسیمال وجود جواب نامیده می‌شود. بازه ماکسیمال وجود جواب مربوط به x_0 اغلب به‌صورت $[t_0, \tau(x_0))$ تعریف می‌شود.

قضیه ۶.۲. [۶۷] فرض کنیم $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ یک تابع پیوسته باشد. اگر $x(t)$ ، $t \in [t_0, \tau(x_0))$ یک جواب ماکسیمال باشد و $\tau(x_0) < +\infty$ آن‌گاه

$$\lim_{t \rightarrow \tau(x_0)} \|x(t)\| = +\infty.$$

قضیه ۷.۲. (قضیه پایداری لیاپانوف) [۲۳] فرض کنید که $x = 0$ یک نقطه تعادل سیستم (۴۴) باشد و $V: \mathbb{R}^n \rightarrow \mathbb{R}$ یک تابع به طور پیوسته مشتق پذیر باشد به طوری که:

$$1. \quad V(0) = 0$$

$$2. \quad \text{به ازای هر } x \in \mathbb{R}^n \setminus \{0\}, \quad V(x) > 0$$

$$3. \quad \text{به ازای هر } x \in \mathbb{R}^n \setminus \{0\}, \quad \dot{V}(x) \leq 0$$

آن گاه $x = 0$ نقطه پایداری سیستم خواهد بود.

قضیه ۸.۲. (قضیه پایداری مجانبی) [۲۳] تحت شرایط قضیه ۷.۲ اگر $V(\cdot)$ در شرایط زیر صدق کند

$$1. \quad V(0) = 0$$

$$2. \quad \text{به ازای هر } x \in \mathbb{R}^n \setminus \{0\}, \quad V(x) > 0$$

$$3. \quad \text{به ازای هر } x \in \mathbb{R}^n \setminus \{0\}, \quad \dot{V}(x) < 0$$

آن گاه $x = 0$ پایدار مجانبی خواهد بود.

قضیه ۹.۲. (قضیه پایداری مجانبی سراسری) [۲۳] تحت شرایط قضیه ۷.۲ اگر $V(\cdot)$ به طور شعاعی بی کران باشد آن گاه نقطه $x = 0$ پایدار مجانبی سراسری است.

قضیه ۱۰.۲. (اصل تغییرناپذیری لیسال) [۲۳] فرض کنیم یک تابع به طور پیوسته مشتق پذیر باشد، همچنین فرض کنیم

$$1. \quad M \subset \mathbb{R}^n \text{ یک مجموعه فشرده و پایدار نسبت به جواب سیستم (۴۴) باشد.}$$

$$2. \quad \text{به ازای } x \in M, \quad \dot{V}(x) \leq 0$$

$$3. \quad E \text{ مجموعه همه نقاط } M \text{ باشد به طوری که } \dot{V}(x) = 0$$

$$4. \quad N \text{ بزرگترین مجموعه پایدار در } E \text{ باشد.}$$

آن گاه به ازای هر $x(t_0) \in M$ که $t_0 \geq 0$ می باشد داریم:

$$\lim_{t \rightarrow \infty} \text{dist}(x(t), N) = 0$$

۳. نامساوی‌های وردشی

مقدمه‌ای بر نامساوی وردشی

نظریه نامساوی وردشی اولین بار در سال ۱۹۶۶ به وسیله هارتمن^۱ و استمپکیا^۲ به عنوان ابزاری برای مطالعه معادلات دیفرانسیل جزئی با کاربردهایی در مکانیک معرفی شد [۲۸]. این نامساوی‌ها با بعد نامتناهی بودند. نظریه نامساوی وردشی با بعد متناهی در سال ۱۹۸۰ که به وسیله دافرموس^۳ معرفی شد، دافرموس متوجه شد که شرایط تعادل شبکه ترافیک که به وسیله اسمیت^۴ در سال ۱۹۷۹ بیان شد ساختار نامساوی وردشی دارد. نامساوی‌های وردشی در اقتصاد، علم مدیریت، تحقیق در عملیات و همچنین در مهندسی کاربردهای فراوانی دارد.

تعریف ۱.۳. فرض کنید $F : K \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$ تابعی پیوسته روی مجموعه بسته و محدب K باشد، مسأله نامساوی وردشی که آن را با نماد $VI(F, K)$ نشان می‌دهند یافتن یک بردار $x^* \in K$ است به‌طوری که

$$F(x^*)^T (x - x^*) \geq 0, \quad \forall x \in K,$$

یا به‌طور معادل

$$\langle F(x^*), x - x^* \rangle \geq 0, \quad \forall x \in K, \quad (۴۵)$$

که در آن $\langle \cdot, \cdot \rangle$ ضرب داخلی در فضای اقلیدسی n بعدی است.

از دیدگاه هندسی، نامساوی وردشی (۴۵) بیان می‌کند که $F(x^*)$ بر ناحیه شدنی K در نقطه x^* عمود است (شکل ۲۸).

قضیه ۱.۳. [۱۱] مسأله بهینه‌سازی زیر را در نظر بگیرید:

$$\begin{aligned} &\text{minimize } f(x), \\ &\text{subject to } x \in K, \end{aligned}$$

که در آن f یک تابع محدب و به‌طور پیوسته مشتق پذیر و K یک مجموعه محدب و غیرتهی در $\mathbb{R}^n$ است آن‌گاه x^* جواب بهینه مسأله فوق است اگر و تنها اگر x^* جواب مسأله نامساوی وردشی

$$\nabla f(x^*)^T (x - x^*) \geq 0, \quad \forall x \in K$$

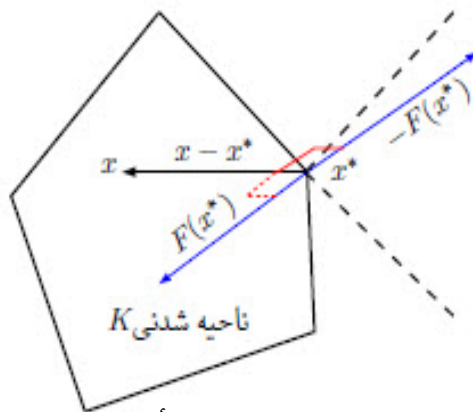
باشد.

¹Hartman

²Stampacchia

³Dafermos

⁴Smith



شکل ۲۸: تعبیر هندسی مسأله نامساوی وردشی.

قضیه ۲.۳. [۵] فرض کنید تابع $F : K \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$ روی K به طور پیوسته مشتق پذیر باشد و ماتریس ژاکوبین F متقارن و نیمه معین مثبت باشد، آن گاه تابع محدب حقیقی مقدار $f : K \rightarrow \mathbb{R}$ وجود دارد که در رابطه

$$\nabla f(x) = F(x),$$

صدق می کند و x^* جواب نامساوی وردشی $VI(F, K)$ است اگر و فقط اگر x^* جواب بهینه مسأله زیر باشد:

$$\begin{aligned} &\text{minimize } f(x), \\ &\text{subject to } x \in K, \end{aligned}$$

تعریف ۲.۳. [۵] فرض می کنیم $\mathbb{R}_+^n$ ناحیه نامنفی روی $\mathbb{R}^n$ را نشان دهد و فرض کنیم $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ باشد. مسأله مکمل غیرخطی روی $\mathbb{R}_+^n$ یک سیستم از معادلات و نامعادلات است به صورت زیر است :

یافتن $x^* \geq 0$ به طوری که

$$F(x^*) \geq 0, \quad F(x^*)^T x^* = 0. \quad (۴۶)$$

هرگاه نگاشت F آفین باشد یعنی $F(x) = Mx + b$ که در آن M یک ماتریس $n \times n$ و b یک بردار $n \times 1$ می باشد مسأله (۴۶) یک مسأله مکمل خطی نامیده می شود.

قضیه زیر رابطه بین مسأله مکمل و مسأله نامساوی وردشی را نشان می دهد.

قضیه ۳.۳. [۵] مسأله نامساوی وردشی $VI(F, \mathbb{R}_+^n)$ و مسأله مکمل (۴۶) جواب های یکسان دارند.

Aabstract

Support vector machines are powerful tools for data classification and regression. In recent years, many fast algorithms have been developed for support vector machines. Support vector machines are used in many military and engineering applications, real-time issues, such as classification in complex electromagnetic environments, diagnostics in medicine, and so on. Given that numerical convergence methods require more calculations and can not meet our needs in real-time problems. One promising process for training real-time support vector machines is the use of recurrent neural networks based on periodic execution.

In this thesis, we introduce two neural networks to solve support vector problems such as support vector regression problem and stochastic support vector regression with probabilistic constraints . We prove that the equilibrium points of the neural networks are equivalent to the optimal solution of the support vector problems. The stability and convergence of the proposed neural networks will also be demonstrated. We confirm the theoretical results by providing examples.

Keywords: Support vector, Neural network, Optimization, Regression, Random variable .



Faculty Of Mathematical Sciences

PhD Thesis in Control and Optimization

**Solving a class of support vector machine
problems using dynamic optimization
methods**

By: Amir Feizi

Supervisor:

Dr Alireza Nazemi

Advisor:

Dr Mohammad Reza Rabiei

February 2021