

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



دانشکده علوم ریاضی

گروه ریاضی کاربردی

پایان نامه کارشناسی ارشد

یک کاربرد از سیستمهای دینامیکی برای حل رده ای از مسائل بهینه سازی ناهموار

عبدالحی نظری

استاد راهنما

دکتر علیرضا ناظمی

۱۳۹۴

تقدیم بہ
روح پاک مادرم و

پدر نزر کو ارم

سپاس‌گزاری

من لم یشکر المخلوق، لم یشکر الخالق حمد و سپاس یکتای بی همتا را که لطفش بر ما عیان است، ادای شکرش را هیچ زبان و دریای فضلش را هیچ کران نیست و اگر در این وادی هستیم، همه محبت اوست. الهی ای مهربانتر از ما به ما، از تو می‌خواهم همه کسانی را که حتی ذره‌ای در انجام این امر مرا یاری نموده‌اند، در سایه لطف و محبت بی‌کرانت، سلامت، شادکام و موفق بداری. بر خود لازم می‌دانم تا مراتب سپاس را از بزرگوارانی به جا آورم که اگر دست یاری‌گرشان نبود، هرگز این پایان نامه به انجام نمی‌رسید. ابتدا از استاد گرانقدرم جناب آقای دکتر علیرضا ناظمی که زحمت راهنمایی این پایان نامه را بر عهده داشتند، کمال سپاس را دارم. از جناب آقای دکتر جعفر فتحعلی و دکتر مهدی ایرانمنش که زحمت داوری این پایان نامه با ایشان بوده است، کمال سپاس را دارم. پدر، مادر، برادر، خواهر و دوستان عزیزم، از زحمات بی‌دریغ و بی‌منت شما متشکرم. همیشه نیازمند محبت، لطف و دعای خیر شما بوده و هستیم.

عبدالحی نظری
۱۳۹۴

تعمدنامه

اینجانب عبدالسخی نظری دانشجوی کارشناسی ارشد رشته ریاضی کاربردی دانشکده علوم ریاضی دانشگاه شاهرود، نویسنده پایان نامه با عنوان یک کاربرد از سیستمهای دینامیکی برای حل رده ای از مسائل بهینه سازی ناهموار، تحت راهنمایی دکتر 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 پژوهشگران، 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 شاهرود متعلق دارد، و مقالات مستخرج با نام “ دانشگاه شاهرود “ یا “ Shahrood University “ 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.

عبدالسخی نظری
۱۳۹۴

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 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

با پیشرفت فن‌آوری اطلاعات و ارتباطات و توسعه ارتباط درون سازمانی و بین سازمانی نیاز به استفاده از مدل‌های بهینه‌سازی را برای استفاده منطقی از داده‌ها و اطلاعات فراهم شده گسترش داده است. این مطلب متضمن بزرگ شدن اندازه مسائل بهینه‌سازی که در عمل وجود دارند خواهد بود. در این شرایط لزوم به کارگیری روش‌های کارآمدی که بتوانند با سرعت بالا مسائل بسیار بزرگ را با کیفیت قابل قبول حل کنند بیش از پیش احساس می‌شود.

در چند دهه اخیر روش‌های بهینه‌سازی که بر پایه رویکرد هوش مصنوعی توسعه یافته‌اند، موفقیت‌های چشم‌گیری در حل مؤثر و کارای مسائل بهینه‌سازی به‌دست آورده‌اند.

روش‌هایی چون الگوریتم ژنتیک، جست‌وجوی ممنوع، شبیه‌سازی تبریدی، شبکه عصبی و... قابلیت‌های خود را در حل مسائل بزرگ عملی به خوبی نشان داده‌اند.

امتیازات ویژه‌ی موجود در شبکه‌های عصبی امکان کاربرد آنها را در حوزه وسیعی از تحقیقات فراهم ساخته است. از جمله آن امتیازات می‌توان به امکان یادگیری و بهبود عملکرد براساس داده‌های ورودی اشاره کرد. همچنین امکان انجام محاسبات به صورت موازی در شبکه‌های عصبی امتیاز دیگری است که با توجه به گسترش سخت‌افزارهای موازی، امکان حل مسائل بسیار بزرگ را توسط این رویکرد ممکن می‌سازد.

در این پایان‌نامه یک مدل شبکه عصبی بازگشتی برای حل رده‌ای از مسائل بهینه‌سازی مشتق ناپذیر ارائه می‌شود. تحلیل وجود یکتایی، پایداری و همگرایی سراسری جواب‌ها مورد بررسی قرار می‌گیرند و عملکرد روش‌های ارائه شده با به کارگیری چند مثال از مسائل برنامه‌ریزی مشتق ناپذیر نشان داده می‌شود.

در انتها نتایج کار و پیشنهاداتی برای کارهای آتی ارائه می‌دهیم.

کلمات کلیدی: پایداری، همگرایی، برنامه ریزی مشتق ناپذیر، بهینه سازی، شبکه عصبی، مسائل مین ماکس، برنامه ریزی خطی و غیر خطی، آنروپی

فهرست مطالب

۱	مقدمه	۱
۱	تاریخچه تحقیق	۱.۱
۲	کاربرد مسائل برنامه ریزی مشتق ناپذیر	۲.۱
۲	تعاریف	۳.۱
۵	پایداری	۴.۱
۷	تابع انرژی	۵.۱
۹	تاریخچه مختصر از کاربرد شبکه های عصبی در بهینه سازی	۲
۹	مقدمه	۱.۲
۱۰	تاریخچه مختصر شبکه های عصبی	۱.۱.۲
۱۱	چرا از شبکه عصبی استفاده می کنیم؟	۲.۱.۲
۱۱	معایب شبکه ی عصبی	۳.۱.۲
۱۲	شبکه های عصبی در بهینه سازی	۲.۲
۱۳	برخی مدل های پیشین ارائه شده شبکه عصبی در بهینه سازی	۳.۲
۲۵	یک روش منظم سازی آنتروپی برای حل مسائل مین-ماکس با برخی کاربردها	۳
۲۵	مقدمه	۱.۳
۲۶	روش منظم سازی آنتروپی	۲.۳
۲۶	تابع آنتروپی شانون	۱.۲.۳
۲۷	خواص $F_p(x)$	۲.۲.۳
۲۹	کاربردهای مرتبط	۳.۳
۲۹	سیستم های نامساوی و مساوی	۱.۳.۳
۳۰	حل مسائل برنامه ریزی خطی	۲.۳.۳
۳۱	حل مسائل مین - ماکس مقید	۳.۳.۳
۳۳	یک مدل شبکه عصبی تصویری برای حل رده ای از مسائل بهینه سازی مشتق ناپذیر	۴
۳۳	مقدمه	۱.۴

۳۴		مدل شبکه عصبی	۲.۴
۳۶		تجزیه و تحلیل پایداری	۳.۴
۴۰		مثال شبیه سازی شده	۴.۴
۴۲		نتیجه گیری و پیشنهادات برای کارهای آتی	۵.۴
۴۳		پیشنهادهای برای کارهای آتی	۱۰.۵.۴

۴۵

مراجع

۴۹

واژه نامه فارسی به انگلیسی

۵۱

واژه نامه انگلیسی به فارسی

فصل ۱

مقدمه

در این پایان نامه یک مدل شبکه عصبی^۱ بازگشتی برای حل مسائل بهینه سازی مشتق ناپذیر ارائه می کنیم. شبکه های عصبی ارائه شده در این پایان نامه در مقایسه با شبکه های عصبی قبلی پیچیدگی کمتری برای پیاده سازی دارد. علاوه بر این، شبکه های عصبی پیشنهادی همگرایی سراسری به جواب بهینه مسأله هستند. نمونه های شبیه سازی بر اساس مشکلات پایه مورد بحث قرار گرفته اند تا عملکرد درست شبکه عصبی پیشنهادی را نشان دهند.

۱.۱ تاریخچه تحقیق

قدمت شبکه های عصبی مصنوعی به دهه ۵۰ میلادی بر می گردد. در سال ۱۹۵۶ بنیاد راکفلر برگزاری کنفرانسی را بر عهده داشت که چشم اندازش به صورت زیر بود:

امکان استفاده از کامپیوترها و شبیه سازی در هر زمینه از یادگیری و سایر حوزه های آموزش.

در همین کنفرانس بود که اصطلاح هوش مصنوعی مورد استفاده عمومی قرار گرفت. بطور کلی هوش مصنوعی را به صورت زیر می توان تعریف نمود:

فرایندهای کامپیوتری که سعی دارند تفکر انسان را تقلید نمایند، این فرایندها با فعالیت هایی که نیاز به استفاده از هوش دارند در ارتباطند.

عموماً این تعریف شامل زمینه های یادگیری خودکار، فهم زبان طبیعی، بینایی و تشخیص صدا، بازی کردن، حل مسائل ریاضی، رباتیک و سیستم های خبره است. در سال های اخیر برای محققان، شبکه های عصبی و فناوری های مرتبط با آن را به عنوان بخشی از هوش مصنوعی در نظر می گیرند. در حالی که بعضی دیگر با توجه به سر رشته خود در علوم پزشکی، شبکه های عصبی را زیرمجموعه هوش مصنوعی نمی دانند.

^۱Neural network

۲.۱ کاربرد مسائل ریزی مشتق ناپذیر

مسئله برنامه ریزی مشتق ناپذیر در زمینه های مختلف علوم و مهندسی مانند رگرسیون، طبقه بندی داده ها، پردازش تصویر، مسائل ریاضی مالی و... کاربرد فراوانی دارد. بالاخص در مهندسی و بسیاری از کاربردهای نظامی که برنامه های کاربردی نیازمند پردازش در زمان واقعی داده ها است، اغلب نیازمند استفاده از شبکه های عصبی به عنوان یک ابزار کارا در یافتن جواب های زمان واقعی هستیم. چون زمان مورد نیاز برای حل مسئله بهینه سازی مشتق ناپذیر وابستگی زیادی به بعد و ساختار مسئله دارد لذا معمولاً روش های عددی مرسوم در کاربردهای واقعی چندان مؤثر نیستند. یک رهیافت امیدوار کننده عبارت است از استفاده از شبکه های عصبی بازگشتی مبنی بر پیاده سازی دوری با مدارهای الکترونیکی. در این فصل تعاریف مورد نیاز در فصل های بعدی به طور مختصر آورده شده است. ابتدا تعاریف مرتبط با تابع محدب^۲ را بیان نموده و در خصوص مسائل بهینه سازی مقید و نامقید به ذکر شرایط لازم و کافی بهینگی^۳ می پردازیم. در انتها مفاهیم پایداری^۴ و تابع انرژی در سیستم های دینامیکی را ارائه می دهیم.

۳.۱ تعاریف

تعریف ۱.۳.۱. تابع $f: \mathbb{R}^n \rightarrow \mathbb{R}^l$ را تابع آفین می گوئیم، هرگاه به صورت مجموع یک تابع خطی و یک تابع ثابت باشد. به عنوان مثال $f(x) = Ax + b$ که $A \in \mathbb{R}^{l \times n}$ و $b \in \mathbb{R}^l$ تابع آفین است.

تعریف ۲.۳.۱. مجموعه $E \subset \mathbb{R}^n$ را محدب می گوئیم اگر:

$$\forall x, y \in E, \lambda \in (0, 1) : \lambda x + (1 - \lambda)y \in E.$$

تعریف ۳.۳.۱. تابع $f: \mathbb{R}^n \rightarrow \mathbb{R}$ بر روی مجموعه نقاط واقع در مجموعه محدب $E \subset \mathbb{R}^n$ یک تابع محدب نامیده می شود، اگر:

$$\forall x, y \in E, \lambda \in (0, 1) : f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y).$$

تابع f را روی E مقعر گوئیم هرگاه $-f$ محدب باشد.

تعریف ۴.۳.۱. اگر $\bar{x} \in X$ و ε -همسایگی از $\bar{x}$ مثل $N_\varepsilon(\bar{x})$ موجود باشد که به ازای هر $x \in X \cap N_\varepsilon(\bar{x})$ داشته باشیم $f(\bar{x}) \leq f(x)$ آنگاه، $\bar{x}$ یک کمینه موضعی برای f است.

در توسعه برنامه ریزی غیرخطی، تابع محدب همواره به عنوان یک مفهوم اصلی ریاضی مورد نیاز است.

^۲Convex

^۳Optimality

^۴Stability

تعریف ۵.۳.۱. تابع $\|\cdot\| : \mathbb{R}^n \rightarrow \mathbb{R}$ را تابع نرم گوییم هرگاه به ازای هر $x, y \in \mathbb{R}$ و $\alpha \in \mathbb{R}$:

$$\begin{cases} \|x\| \geq 0, & \forall x, \\ \|x\| = 0 \Leftrightarrow x = 0, \\ \|\alpha x\| = |\alpha| \|x\|, \\ \|x + y\| \leq \|x\| + \|y\|. \end{cases}$$

تعریف ۶.۳.۱. یک زیر مجموعه E از فضای متریک X را یک مجموعه فشرده می‌گوییم اگر هر پوشش باز E یک زیر پوشش باز متناهی داشته باشد.

تعریف ۷.۳.۱. یک ماتریس $M(X)$ در اندازه $n \times n$ که عناصر آن m_{ij} ($i, j = 1, \dots, n$) توابعی روی $E \subset \mathbb{R}^n$ هستند،

الف. نیمه معین مثبت^۵ روی E نامیده می‌شود، اگر
 $\forall V \in \mathbb{R}^n, V \neq 0, X \in E : V^T M(X) V \geq 0.$

ب. معین مثبت^۶ روی E نامیده می‌شود، اگر
 $\forall V \in \mathbb{R}^n, V \neq 0, X \in E : V^T M(X) V > 0.$

شرایط لازم و کافی برای مسائل نامقید

قضیه ۸.۳.۱. (شرط لازم مرتبه اول [۳]): فرض کنید X یک مجموعه باز در $\mathbb{R}^n$ باشد. مسأله برنامه ریزی زیر را در نظر بگیرید:

$$\min \quad f(x) \quad (۱.۱)$$

$$\text{s.t} \quad x \in X \quad (۲.۱)$$

که در آن $f : \mathbb{R}^n \rightarrow \mathbb{R}$ در $\hat{x}$ مشتق پذیر است. اگر $\hat{x}$ یک کمینه موضعی باشد، آنگاه $\nabla f(\hat{x}) = 0$.

قضیه ۹.۳.۱. (شرط لازم مرتبه دوم [۳]): فرض کنید مشتق دوم $f : \mathbb{R}^n \rightarrow \mathbb{R}$ در $\hat{x}$ موجود باشد. اگر $\hat{x}$ یک کمینه موضعی باشد، آنگاه $\nabla f(\hat{x}) = 0$ و ماتریس هسین^۷ f در $\hat{x}$ نیمه معین مثبت است.

قضیه ۱۰.۳.۱. (شرط کافی [۳]): فرض کنید مشتق دوم $f : \mathbb{R}^n \rightarrow \mathbb{R}$ در $\hat{x}$ وجود باشد. اگر $\nabla f(\hat{x}) = 0$ و ماتریس هسین f در $\hat{x}$ معین مثبت باشد، آنگاه $\hat{x}$ یک کمینه موضعی اکید است.

قضیه ۱۱.۳.۱. ([۳]): فرض کنید مجموعه E بصورت زیر باشد:

$$E = \{x \in \mathbb{R}^n \mid g_k(x) \leq 0, k = 1, \dots, m, h_p(x) = 0, p = 1, \dots, l\}$$

اگر $g_k(x)$ ، $k = 1, \dots, m$ توابعی محدب و $h_p(x)$ ، $p = 1, \dots, l$ آفین باشند آنگاه، E یک مجموعه محدب است.

^۵Positive semi-definite

^۶Positive definite

^۷Hessian

قضیه ۱۲.۳.۱. [۳]: اگر $f(x)$ بر مجموعه محدب E تابعی محدب باشد آنگاه $f(x)$ دارای یک کمینه موضعی^۸ است. اگر چنین کمینه‌ای یافت شود، یک کمینه سراسری^۹ است و بر روی مجموعه محدب به دست آورده می‌شود.

شرایط لازم و کافی برای مسائل مقید

قضیه ۱۳.۳.۱. (شرایط لازم کروش-کان-تاکر^{۱۰}، (K, K, T)): فرض کنید X یک مجموعه باز در $\mathbb{R}^n$ باشد. مسأله کمینه‌سازی مقید زیر را در نظر بگیرید:

$$\min f(x) \quad (۳.۱)$$

$$\text{s.t.} \quad \begin{cases} g_k(x) \leq 0, & k = 1, \dots, m, \\ h_j(x) = 0, & j = 1, \dots, l, \\ x = (x_1, x_2, \dots, x_n)^T \in X, \end{cases} \quad (۴.۱)$$

که در آن $f: \mathbb{R}^n \rightarrow \mathbb{R}$ ، $g_k: \mathbb{R}^n \rightarrow \mathbb{R}$ ، $k = 1, \dots, m$ ، و $h_j: \mathbb{R}^n \rightarrow \mathbb{R}$ ، $j = 1, \dots, l$ فرض کنید $\hat{x}$ یک جواب شدنی این مسأله باشد و $K = \{k : g_k(\hat{x}) = 0\}$. همچنین فرض کنید f و g_k برای $k \in K$ در $\hat{x}$ مشتق پذیر، g_k برای $k \notin K$ در $\hat{x}$ پیوسته و h_j برای $j = 1, \dots, l$ دارای مشتق پیوسته باشند. به علاوه فرض کنید $\nabla g_k(\hat{x})$ برای $k \in K$ و $\nabla h_j(\hat{x})$ برای $j = 1, \dots, l$ مجموعاً مستقل خطی باشند. اگر $\hat{x}$ کمینه موضعی برای مسأله (۳.۱) و (۴.۱) باشد، آنگاه اسکالرهایی یکتای $\hat{u}_k \geq 0$ برای $k \in K$ و $\hat{v}_j \geq 0$ برای $j = 1, \dots, l$ موجود هستند به طوری که:

$$\begin{cases} \nabla f(\hat{x}) + \sum_{k \in K} \hat{u}_k \nabla g_k(\hat{x}) + \sum_{j=1}^l \hat{v}_j \nabla h_j(\hat{x}) = 0, \\ \hat{u}_k \geq 0, \quad k \in K. \end{cases}$$

علاوه بر فرض‌های بالا اگر g_k برای $k \notin K$ در $\hat{x}$ مشتق پذیر باشد، آنگاه شرایط کروش-کان-تاکر می‌تواند به صورت زیر نوشته شود:

$$\begin{cases} \nabla f(\hat{x}) + \sum_{k=1}^m \hat{u}_k \nabla g_k(\hat{x}) + \sum_{j=1}^l \hat{v}_j \nabla h_j(\hat{x}) = 0, \\ \hat{u}_k g_k(\hat{x}) = 0, \quad k = 1, \dots, m, \\ \hat{u}_k \geq 0, \quad k = 1, \dots, m. \end{cases}$$

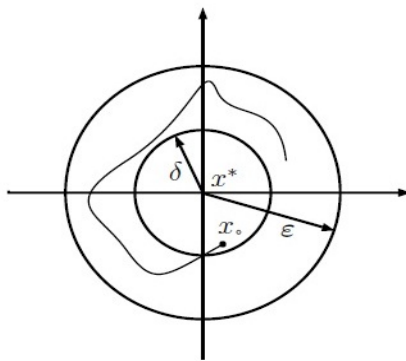
همچنین شرایط کروش-کان-تاکر می‌تواند به فرم ماتریسی زیر نوشته شود:

$$\begin{cases} \nabla f(\hat{x}) + \nabla g(\hat{x})^T u + \nabla h(\hat{x})^T v = 0, \\ u^T g(\hat{x}) = 0, \\ u \geq 0, \end{cases}$$

^۸Local minimum

^۹Global minimum

^{۱۰}Karush-Kuhn-Tucker



شکل ۱۰۱: پایداری لیاپانوف

که در آن $\nabla g(\hat{x})$ ماتریس ژاکوبی $m \times n$ و $\nabla h(\hat{x})$ ماتریس ژاکوبی $l \times n$ است و $\hat{u}$ یک بردار m تایی و $\hat{v}$ یک بردار l تایی است. $(\hat{u}^T, \hat{v}^T)^T$ بردار ضرایب لاگرانژ نامیده شده است.

قضیه ۱۴.۳.۱. (شرایط کافی کروش-کان-تاکر $(K.K.T)$ ، [۲]): فرض کنید X یک مجموعه باز در $\mathbb{R}^n$ باشد، مسأله (۳.۱) و (۴.۱) را در نظر بگیرید. فرض کنید $\hat{x}$ یک جواب شدنی باشد، همچنین فرض کنید f و g_k ها برای $k = 1, \dots, m$ توابعی محدب و h_j ها برای $j = 1, \dots, l$ توابعی آفین باشند و $K = \{k : g_k(\hat{x}) = 0\}$ ، اگر شرایط کروش-کان-تاکر در $\hat{x}$ برقرار باشد، یعنی اسکالرهایی $\hat{u}_k \geq 0$ برای $k \in K$ و $\hat{v}_j \geq 0$ برای $j = 1, \dots, l$ موجود باشند به طوری که:

$$\nabla f(\hat{x}) + \sum_{k \in K} \hat{u}_k \nabla g_k(\hat{x}) + \sum_{j=1}^l \hat{v}_j \nabla h_j(\hat{x}) = 0,$$

آنگاه $\hat{x}$ یک کمینه موضعی برای (۳.۱) و (۴.۱) خواهد بود.

۴.۱ پایداری

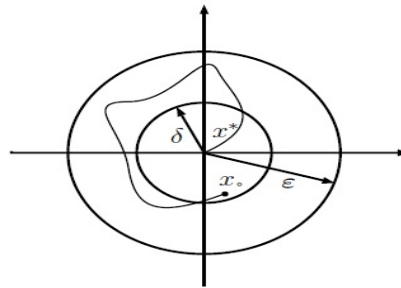
سیستم دینامیکی زیر را در نظر بگیرید:

$$\begin{cases} \dot{x} = f(x(t)), \\ x(t_0) = x_0 \in \mathbb{R}^n, \end{cases} \quad (5.1)$$

که در آن $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ است.

تعریف ۱۰.۴.۱. x^* یک نقطه تعادل^{۱۱} برای (۵.۱) نامیده می شود اگر $f(x^*) = 0$.

^{۱۱}Equilibrium point



شکل ۲.۱: پایداری مجانبی

تعریف ۲.۴.۱. فرض کنید که $x(t)$ یک جواب (۵.۱) باشد، نقطه تعادل تنها x^* ، پایدار به مفهوم لیاپانوف^{۱۲} است (شکل (۱.۱) را ببینید)، اگر برای هر $x(t_0) = x_0$ و هر $\varepsilon > 0$ ، یک $\delta > 0$ وجود داشته باشد به طوری که اگر $\|x^* - x_0\| \leq \delta$ آن گاه

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

تعریف ۳.۴.۱. سیستم دینامیکی (۵.۱) همگرای سراسری^{۱۳} به مجموعه

$$\Omega^* = \{x \mid \text{جواب سیستم (۵.۱) است.}\}$$

گفته می شود اگر هر جواب دلخواه $x(t)$ از سیستم در رابطه زیر صدق کند:

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

که در آن

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

تعریف ۴.۴.۱. سیستم دینامیکی در نقطه تعادل یکتا x^* پایدار مجانبی سراسری^{۱۴} نامیده می شود اگر x^* به مفهوم لیاپانوف پایدار باشد و در شرط زیر صدق کند:

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

شکل (۲.۱) را ببینید.

تعریف ۵.۴.۱. مجموعه $M \subseteq \mathbb{R}^n$ ، یک مجموعه پایدار^{۱۵} نسبت به سیستم (۵.۱) گفته می شود اگر $x(t_0) \in M$ و $t_0 \geq 0$ ، آن گاه به ازای هر $t \geq t_0$ ، $x(t) \in M$ باشد.

تعریف ۶.۴.۱. نگاشت F در شرط لپ شیتز^{۱۶} صدق می کند اگر عدد ثابت L وجود داشته باشد، به طوری که برای هر دو نقطه $x, y \in \mathbb{R}^n$ داشته باشیم:

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

^{۱۲}Lyapunov

^{۱۳}Global convergence

^{۱۴}Globally asymptotically stable

^{۱۵}Invariant set

^{۱۶}Lipschitz

تعریف ۷.۴.۱. F را لیپ شیتز محلی روی $\mathbb{R}^n$ نامیم اگر به ازای هر نقطه از $\mathbb{R}^n$ ، یک همسایگی مانند $D \subset \mathbb{R}^n$ وجود داشته باشد به طوری که نامساوی بالا برای هر دو نقطه $x, y \in D$ برقرار باشد.

تعریف ۸.۴.۱. نگاشت $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ یکنوا^{۱۷} گفته می‌شود، اگر برای هر زوج از نقاط $x, y \in \mathbb{R}^n$ داشته باشیم:

$$(x - y)^T (F(x) - F(y)) \geq 0.$$

همچنین روی $\mathbb{R}^n$ اکیداً یکنوا^{۱۸} است، اگر نامساوی فوق به صورت اکید برای هر $x \neq y$ برقرار باشد. و F روی $\mathbb{R}^n$ یکنوای قوی است اگر برای هر زوج از نقاط $x, y \in \mathbb{R}^n$ ثابت $\beta > 0$ وجود داشته باشد، به طوری که:

$$(x - y)^T (F(x) - F(y)) \geq \beta \|x - y\|^2.$$

قضیه ۹.۴.۱. ([۱]): فرض می‌کنیم که $F : K \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$ ، یک تابع به طور پیوسته مشتق پذیر روی K باشد و ماتریس ژاکوبین F که لزوماً متقارن نیست، نیمه معین مثبت (معین مثبت) باشد، آنگاه $F(x)$ یکنوا (یکنوای قوی) است.

قضیه ۱۰.۴.۱. ([۱]): در سیستم دینامیکی (۵.۱) فرض می‌کنیم که $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ یک تابع پیوسته باشد، آنگاه برای هر $t_0 > 0$ و $x_0 \in \mathbb{R}^n$ ، یک جواب محلی^{۱۹} $x(t)$ به ازای $t \in [t_0, \tau)$ که $\tau > t_0$ وجود دارد. علاوه بر این اگر f در x_0 در شرط لیپ شیتز محلی صدق کند آنگاه جواب یکتا خواهد بود و اگر f در $\mathbb{R}^n$ در شرط لیپ شیتز صدق کند آنگاه τ می‌تواند تا $+\infty$ توسعه داده شود.

۵.۱ تابع انرژی

قضیه ۱۰.۵.۱. ([۱]): فرض کنید که $x = 0$ یک نقطه تعادل برای سیستم (۵.۱) و تابع $E : \Omega \subseteq \mathbb{R}^n \rightarrow \mathbb{R}$ به طور پیوسته مشتق پذیر باشد اگر

$$E(0) = 0. \quad ۱.$$

$$E(x) > 0, \quad x \in \Omega \subseteq \mathbb{R}^n \setminus \{0\} \quad ۲.$$

$$\dot{E}(x) \leq 0, \quad x \in \Omega \subseteq \mathbb{R}^n \setminus \{0\} \quad ۳.$$

آنگاه $x = 0$ نقطه پایداری لیاپانوف سیستم خواهد بود و $E(x)$ را «تابع لیاپانوف» یا «تابع انرژی» برای سیستم (۵.۱) نامیم.

قضیه ۲.۵.۱. [۱] اصل تغییرناپذیری (ناوردای) لازال^{۲۰}: فرض می‌کنیم تابع $V : \mathbb{R}^n \rightarrow \mathbb{R}$ به طور پیوسته مشتق پذیر باشد، همچنین فرض می‌کنیم

^{۱۷}Monotone

^{۱۸}Strictly monotonic

^{۱۹}Local solution

^{۲۰}Lasalle's invariance principle

۱. $M \subseteq \mathbb{R}^n$ یک مجموعه فشرده و پایدار نسبت به جواب سیستم (۵.۱) باشد،

۲. به ازای هر $x \in M$ ، $\dot{V}(x) \leq 0$ ،

۳. E مجموعه همه نقاط M باشد به طوری که $\dot{V}(x) = 0$ ،

۴. E بزرگترین مجموعه پایدار در $\mathbb{R}^n$ ، N باشد

آنگاه به ازای هر $x(t_0) \in M$ که $t_0 \geq 0$ باشد، داریم:

$$\lim_{t \rightarrow \infty} \text{dist}(x(t), N) = 0.$$

فصل ۲

تاریخچه مختصر از کاربرد شبکه های عصبی در بهینه سازی

۱.۲ مقدمه

دهه های آغازین سده بیستم میلادی و دوران پیشرفت شگرف صنعتی، همراه با تولید خودرو بود که انقلاب همه جانبه این در ترابری، افزایش شتاب جابجایی و صدها کار و پیشه جدید در رشته های بازرگانی بوجود آورده است. به نظر می رسد که سمبل دوران فراصنعتی و نماد فرآورده های بی همتای قرن آینده «هوش مصنوعی» است. امروزه موضوع هوش مصنوعی داغ ترین بحث میان کارشناسان دانش رایانه و اطلاعات و دیگر دانشمندان و تصمیم گیرندگان است. در سراسر تاریخ تا به امروز انسان از جنبه تن و روان، مرکز و محور بحث ها و پژوهش ها بوده است. ولی اکنون موجودی با رتبه ای پائین تر، بی جان و ساختگی می خواهد جانشین او شود، امری که بدون شك می توان ادعا نمود بیشتر انسان ها با آن مخالفتند. هوش مصنوعی چنانچه به هدف های والای خود برسد، جهش بزرگی در راه دستیابی بشر به رفاه بیشتر و حتی ثروت افزون تر خواهد بود. هم اکنون نمونه های خوبی از هوش مصنوعی در دنیای واقعی ما به کار افتاده است. چنین دستاوردهایی صرف منابع لازم در آینده را همچنان توجیه خواهد کرد. از سوی دیگر، منتقدین هوش مصنوعی چنین استدلال می کنند که صرف زمان و منابع ارزشمند دیگر در راه ساخت فرآورده ای که پر از نقص و کاستی و دست آوردهای مثبت اندکی است، مایه بدنام کردن و زیر پا گذاشتن توانمندی ها و هوشمندی های انسان می باشد. شبکه عصبی مصنوعی روشی عملی برای یادگیری توابع گوناگون نظیر توابع با مقادیر حقیقی، توابع با مقادیر گسسته و توابع با مقادیر برداری می باشد. یادگیری شبکه ی عصبی در برابر خطاهای داده های آموزشی مصون بوده و شبکه با موفقیت به مسائلی نظیر شناسائی گفتار، شناسائی و تعبیر تصاویر، و یادگیری روبات اعمال شده اند. لذا در این فصل تاریخچه ای از شبکه ی عصبی و کاربرد آن در زمینه های مختلف علوم اشاره کرده و در انتها مدلهایی از شبکه های عصبی که برای حل مسائل گوناگون بهینه سازی (مسائل خطی یا غیر خطی) به دست آمده است را معرفی می کنیم.

۱.۱.۲ تاریخچه مختصر شبکه های عصبی

دیدگاه جدید شبکه های عصبی در دهه ی چهل قرن بیستم شروع شد، زمانی که وارن مک کلوت^۱ و والتر پیتز^۲ نشان دادند که شبکه های عصبی در اصل می توانند هر تابع حسابی^۳ و منطقی^۴ را محاسبه نمایند. کار این افراد را می توان نقطه ی شروع حوزه ی علمی شبکه های عصبی مصنوعی نامید و این موضوع با دونالد هب^۵ ادامه یافت. نخستین کاربرد عملی شبکه های عصبی در اواخر دهه ی پنجاه قرن بیستم مطرح شد، زمانی که فرانک روزنبلات^۶ در سال ۱۹۵۸ شبکه ی پرسپترون^۷ را معرفی نمود. روزنبلات و همکارانش شبکه ای ساختند که قادر بود الگوها را از هم شناسایی نماید. در همین زمان بود که برنارد ویدر^۸ در سال ۱۹۶۰ شبکه ی عصبی تطبیقی خطی آدالین^۹ را با قانون یادگیری جدید مطرح نمود که از لحاظ ساختار شبیه شبکه ی پرسپترون می باشد. پیشرفت شبکه های عصبی تا دهه هفتاد قرن بیستم ادامه یافت. در ۱۹۷۲ تنوکوهونن^{۱۰} و جیمز اندرسون^{۱۱} به طور مستقل و بدون اطلاع از هم شبکه های عصبی جدید را معرفی نمودند که قادر بودند به عنوان عناصر ذخیره ساز عمل نمایند. استفان گروسبرگ^{۱۲} در این دهه روی شبکه های خود سازمانده^{۱۳} فعالیت می کرد. علاقه به کار کردن روی شبکه های عصبی در دهه شصت قرن بیستم در قیاس با دهه هشتاد به علت عدم بروز ایده های جدید و نبود کامپیوترهای سریع جهت پیاده سازی کمرنگ می نمود. لیکن در خلال دهه هشتاد، تحقیقات روی شبکه های عصبی دوباره افزایش یافت، در این زایش دوباره شبکه های عصبی دو نگرش جدید قابل تأمل می باشند. استفاده از مکانیسم تصادفی جهت توضیح عملکرد یک طبقه وسیع از شبکه های برگشتی^{۱۴} که می توان آن ها را جهت ذخیره سازی اطلاعات استفاده نمود. این ایده توسط جان هاپفیلد^{۱۵} در سال ۱۹۸۲ مطرح شد. دومین ایده ی مهم که کلید توسعه ی شبکه های عصبی در دهه هشتاد شد، الگوریتم پس انتشار خطا^{۱۶} می باشد که توسط دیوید راملهارت^{۱۷} و جیمز مک لند^{۱۸} مطرح گردید. با بروز این دو ایده شبکه های

^۱Warren McCulloch

^۲Walter Pitts

^۳Logical function

^۴Arithmetic

^۵Donald Hebb

^۶Frank Rosenblatt

^۷Perceptron

^۸Bernard Widrow

^۹Adaptive linear element

^{۱۰}Teo Kohonen

^{۱۱}James Anderson

^{۱۲}Stefan Grossberg

^{۱۳}Self-organizing

^{۱۴}Feedback (Recurrent)

^{۱۵}John Hopfield

^{۱۶}Error back-propagation

^{۱۷}David Rumelhart

^{۱۸}James Mclelland

عصبی متحول شدند و هزاران مقاله نوشته شد و شبکه‌های عصبی کاربردهای زیادی در رشته‌های مختلف علوم پیدا کردند. بیشتر پیشرفت‌ها در شبکه‌های عصبی به ساختارهای نوین و روش‌های یادگیری جدید مربوط می‌شوند.

۲.۱.۲ چرا از شبکه عصبی استفاده می‌کنیم؟

شبکه‌های عصبی با توانایی قابل توجه خود در استنتاج نتایج از داده‌های پیچیده می‌توانند در استخراج الگوها و شناسایی گرایش‌های مختلفی که برای انسان‌ها و کامپیوتر شناسایی آنها بسیار دشوار است استفاده شوند. از مزایای شبکه‌های عصبی می‌توان موارد زیر را نام برد:

۱. یادگیری تطبیقی: توانایی یادگیری اینکه چگونه وظایف خود را بر اساس اطلاعات داده شده به آن و یا تجارب اولیه انجام دهد در واقع اصلاح شبکه را گویند.

۲. خود سازماندهی: یک شبکه عصبی مصنوعی به صورت خودکار سازماندهی و ارائه داده‌هایی که در طول آموزش دریافت کرده را انجام دهد. نرون‌ها با قاعده‌ی یادگیری سازگار شده و پاسخ به ورودی تغییر می‌یابد.

۳. عملگرهای بی‌درنگ: محاسبات در شبکه‌ی عصبی مصنوعی می‌تواند به صورت موازی و به وسیله سخت‌افزارهای مخصوصی که طراحی و ساخت آن برای دریافت نتایج بهینه قابلیت‌های شبکه عصبی مصنوعی است انجام شود.

۴. تحمل خطا: با ایجاد خرابی در شبکه مقداری از کارایی کاهش می‌یابد ولی برخی امکانات آن با وجود مشکلات بزرگ همچنان حفظ می‌شود.

۵. دسته بندی: شبکه‌های عصبی قادر به دسته بندی ورودی‌ها برای دریافت خروجی مناسب می‌باشند.

۶. تعمیم دهی: این خاصیت شبکه را قادر می‌سازد تا تنها با برخورد با تعداد محدودی نمونه، یک قانون کلی از آن را به دست آورده، نتایج آموخته‌ها را به موارد مشابه از قبل تعمیم دهد. توانایی که در صورت نبود آن سامانه باید بی نهایت واقعیت‌ها و روابط را به خاطر سپارد.

۷. پایداری-انعطاف پذیری: یک شبکه عصبی هم به حد کافی پایدار است تا اطلاعات فراگرفته خود را حفظ کند و هم قابلیت انعطاف و تطبیق را دارد و بدون از دست دادن اطلاعات قبلی می‌تواند موارد جدید را بپذیرد.

۳.۱.۲ معایب شبکه‌ی عصبی

با وجود برتری‌هایی که شبکه‌های عصبی نسبت به سامانه‌های مرسوم دارند، معایبی نیز دارند که پژوهشگران این رشته تلاش دارند که آنها را به حداقل برسانند، از جمله:

۱. قواعد یا دستورات مشخصی برای طراحی شبکه جهت یک کاربرد اختیاری وجود ندارد.
۲. در مورد مسائل مدل سازی، صرفاً نمی توان با استفاده از شبکه عصبی به فیزیک مساله پی برد. به عبارت دیگر مرتبط ساختن پارامترها یا ساختار شبکه به پارامترهای فرآیند معمولاً غیر ممکن است.
۳. دقت نتایج بستگی زیادی به اندازه مجموعه آموزش دارد.
۴. آموزش شبکه ممکن است مشکل و یا حتی غیر ممکن باشد.
۵. پیش بینی عملکرد آینده شبکه (عمومیت یافتن) آن به سادگی امکان پذیر نیست.

۲.۲ شبکه های عصبی در بهینه سازی

همانطور که گفته شد یکی از کاربردهای هوش مصنوعی در مسائل بهینه سازی می باشد، بهینه نمودن تابع هدف (ماکزیم یا مینیم کردن آن) با استفاده از روش های تحلیلی سر راست بسیار پیچیده می باشد. با وجود این امکان رسیدن به جواب با بکارگیری روش های عددی رایانه ای میسر می گردد. در این روش ها با تغییر متغیرهای طراحی در دوره های متوالی و با در نظر گرفتن قیود به جواب بهینه نزدیک می شویم. استفاده از روش های بهینه سازی متداول و تحلیل مستقیم سازه ها در دفعات مکرر موجب افزایش قابل ملاحظه ای زمان محاسبه می گردد. بحث و بررسی در مورد کاربرد شبکه های عصبی (شبکه های عصبی مصنوعی) در بهینه سازی از اوایل سال ۱۹۸۰ میلادی آغاز شده است. نتایج پژوهش های انجام شده از آن زمان تاکنون، مواردی چون برنامه ریزی خطی، برنامه ریزی درجه دوم، برنامه ریزی هندسی و برنامه ریزی غیرخطی را در بر می گیرد. ایده اصلی در استفاده از شبکه های عصبی مصنوعی برای مسائل بهینه سازی استفاده از یک تابع انرژی (نامنفی) و یک سیستم دینامیکی است که این دو بیان کننده ی مدل های شبکه های عصبی مصنوعی متناظر مسائل بهینه سازی هستند. سیستم دینامیکی بیان شده معمولاً یک دستگاه معادلات دیفرانسیل غیرخطی مرتبه اول است. انتظار می رود که برای یک نقطه ی آغازین نقطه ی پایداری دستگاه معادلات دیفرانسیل غیرخطی به دست آمده، جواب بهینه ی مسأله بهینه سازی اصلی باشد. یک اصل اساسی در استفاده از تابع انرژی این است که برای همگرایی دستگاه معادلات دیفرانسیل به نقطه ی پایداری آن، تابع انرژی متناظر با آن حتماً باید نامنفی باشد. البته می توان مدل شبکه های عصبی نظیر مسائل بهینه سازی را حتی بدون در نظر گرفتن تابع انرژی نیز بیان کرد. بنابراین برای یک مدل متناظر با مسائل بهینه سازی، اصل اساسی استفاده از شبکه های عصبی در اینگونه مسائل به صورت زیر بیان می شود:

برای یک نقطه ی آغازین دلخواه، نقطه ی پایداری دستگاه معادلات دیفرانسیل غیرخطی بدست آمده جواب بهینه مسأله اصلی است و برعکس.

مدل های مطرح شده متناظر مسائل مختلف بهینه سازی را می توان به دو قسمت مدل های دوگانی و مدل های جریمه ای تقسیم نمود. نظریه دوگانی و توابع جریمه ای دو نظریه بسیار مهم در مسائل بهینه سازی هستند که اکثر این مسائل بر مبنای این دو روش کلاسیک قابل حل هستند. در نظریه ی توابع جریمه ای

معمولاً از مدل های گرادینانی برای معرفی مدل مورد نظر استفاده می شود، ولی در نظریه دوگانی از مدل های اولیه - دوگان برای حل این مسائل استفاده می شود. اگر بتوان متناظر با هر مسئله بهینه سازی، روش مشخصی را ارائه نمود که آن روش برای حل آن مسئله شرایط لازم و کافی را برآورده سازد، آنگاه می توان متناظر با آن روش، یک مدل شبکه عصبی برای مسئله مورد نظر بسازیم. در زیر به بیان چند مدل که تاکنون برای مسائل برنامه ریزی خطی و درجه دوم و غیرخطی ارائه شده اند می پردازیم.

۳.۲ برخی مدل های پیشین ارائه شده شبکه عصبی در بهینه سازی

مدل اول

مسئله برنامه ریزی خطی زیر را در نظر بگیرید:

$$\min f(x) = a^T x \quad (۱.۲)$$

$$\text{s.t.} \begin{cases} g(x) = Dx - b = 0, \\ x \geq 0, \end{cases} \quad (۲.۲)$$

که در آن $D \in \mathbb{R}^{m \times n}$ ، $\text{rank}(D) = m$ ، $a \in \mathbb{R}^n$ و $b \in \mathbb{R}^m$ است. دوگان (۱.۲) و (۲.۲) به صورت زیر بیان می شود:

$$\max \bar{f}(y) = b^T y \quad (۳.۲)$$

$$\text{s.t.} \begin{cases} \bar{g}(y) = D^T y - a \leq 0, \\ y \in \mathbb{R}^m. \end{cases} \quad (۴.۲)$$

در [۲۵] مدل متناظر برای حل (۱.۲) و (۲.۲) به صورت زیر نمایش داده شده است:

$$\begin{cases} \frac{dx}{dt} = -[D^T(Dx - b) - \beta(D^T y - a)], \\ \frac{dy}{dt} = -\beta[(Dx - D^T y - a)^+ - b], \end{cases}$$

که در آن $\| (x + D^T y - a)^+ - x \|_2^2 = \beta$ ، $\beta = \| (x + D^T y - a)^+ - x \|_2^2$ و $(x, y)^T \in \{(x, y)^T \mid x \geq 0\}$ ، $\beta = \| (x + D^T y - a)^+ - x \|_2^2$ و $(Dx - D^T y - a)^+ = \max\{(Dx - D^T y - a), 0\}$.

مدل دوم

مسئله برنامه ریزی درجه دوم زیر را در نظر بگیرید:

$$\min f(x) = \frac{1}{2} x^T A x + a^T x \quad (۵.۲)$$

$$\text{s.t.} \begin{cases} Dx - b = 0, \\ x \geq 0, \end{cases} \quad (۶.۲)$$

که در آن A یک ماتریس متقارن معین مثبت است. دوگان (۵.۲) و (۶.۲) به صورت زیر است:

$$\max \quad \bar{f}(y) = b^T y - \frac{1}{\gamma} x^T A x \quad (7.2)$$

$$\text{s.t.} \quad \bar{D}y - Ax - a \leq 0. \quad (8.2)$$

در [۲۵] مدل متناظر برای حل (۵.۲) و (۶.۲) به صورت زیر نمایش داده شده است:

$$\begin{cases} \frac{dx}{dt} = -\{D^T(x - b) + \gamma(-D^T y + Ax + a) + \gamma A[x - (x + D^T y - Ax - a)^+]\}, \\ \frac{dy}{dt} = -\gamma\{Dx - b + D[(x + D^T y - Ax - a)^+ - x]\}, \end{cases}$$

که در آن $\gamma = \|(x + D^T y - Ax - a)^+ - x\|^2$ و $(x, y)^T \in \{(x, y)^T \mid x \geq 0\}$.

مدل سوم

مسأله برنامه ریزی غیرخطی زیر را در نظر بگیرید:

$$\min \quad f(x) \quad (9.2)$$

$$\text{s.t.} \quad g(x) = [g_1(x), g_2(x), \dots, g_m(x)]^T \leq 0, \quad (10.2)$$

که در آن $f(x), g_i(x) : \mathbb{R}^n \rightarrow \mathbb{R}^1$ و $f(x), g(x) \in C^2$. برای حل (۹.۲) و (۱۰.۲) هاپفیلد و تنک^{۱۹} در [۲۳]، یک مدل شبکه عصبی به صورت زیر ارائه کردند:

$$\dot{x} = C^{-1}\{-\nabla f(x) - \nabla g(x)g^+(x) - \frac{1}{s}Q^{-1}x\}s,$$

که در آن C یک ماتریس قطری $n \times n$ بیان کننده ظرفیت هر نرون و Q یک ماتریس قطری $n \times n$ است که اعضای آن بصورت زیر تعریف می شوند:

$$q_{ij} = \frac{1}{\frac{1}{\rho_i} + \sum_{j=1}^m -d_{ji}},$$

که در آن $\frac{1}{\rho_i}$ ضریب هدایت گرمایی هر نرون، d_{ji} وزن های تخصیص داده شده به نرون های داخلی و s پارامتر جریمه است. تابع انرژی متناظر به صورت زیر انتخاب شده است:

$$E_1(x) = f(x) + \sum_{j=1}^m [g_j^+]^2 + \sum_{i=1}^n \frac{x_i^2}{2sr_{ii}}.$$

مدل چهارم

برای حل (۹.۲) و (۱۰.۲)، کندی و چوآ^{۲۰} در [۲۴] مدلی دیگر به صورت زیر ارائه داده اند:

$$\dot{x} = C^{-1}[-\nabla f(x) - \nabla g(x)g^+(x)],$$

^{۱۹}Hopfield and Tank

^{۲۰}Kennedy and Chua

که در آن C و s شرایط مدل سوم را دارند. تابع انرژی متناظر این مدل به صورت زیر انتخاب شده است:

$$E_{\Psi}(x) = f(x) + \frac{s}{\Psi} \sum_{j=1}^m [g_j^+]^2.$$

مدل پنجم

مجدداً برای حل (۹.۲) و (۱۰.۲) توسط رودریگز-واسکوئز^{۲۱} در [۲۵] مدلی به صورت زیر ارائه شده است:

$$\dot{x} = -u_x \nabla f(x) - s \nabla g(x) g^+(x),$$

که در آن u_x متناظر اندیس های شدنی x است چنان که

$$u_x = \begin{cases} 1 & g(x) \leq 0, \\ 0 & \text{در غیر این صورت} \end{cases}$$

تابع انرژی متناظر این مدل نیز به صورت زیر انتخاب شده است:

$$E_{\Psi}(x) = -u_x f(x) + \frac{s}{\Psi} \sum_{j=1}^m [g_j^+(x)]^2.$$

همان طور که ملاحظه می کنید، مدل های بیان شده به دلیل مشکلاتی که هر کدام داشته اند به طور متوالی اصلاح شده اند.

مدل ششم

برنامه ریزی محدب درجه دوم^{۲۲} زیر را در نظر بگیرید:

$$\min \quad f(x) = \frac{1}{\Psi} x^T Q x + D^T x \quad (11.2)$$

$$\text{s.t.} \quad \begin{cases} g(x) = Ax - b \leq 0, \\ h(x) = Ex - f = 0, \end{cases} \quad (12.2)$$

که در آن $Q \in \mathbb{R}^{n \times n}$ یک ماتریس متقارن معین مثبت، $A \in \mathbb{R}^{m \times n}$ ، $b \in \mathbb{R}^m$ ، $E \in \mathbb{R}^{l \times n}$ ، $f \in \mathbb{R}^l$ و $x \in \mathbb{R}^n$ است. بر طبق شرایط کروش-کان-تاکر، $x^* \in \mathbb{R}^n$ جواب بهینه (۱۱.۲) و (۱۲.۲) است اگر و تنها اگر بردارهای $u^* \in \mathbb{R}^m$ و $v^* \in \mathbb{R}^l$ وجود داشته باشد که در رابطه زیر صدق کنند:

$$\begin{cases} u^* \geq 0, Ax^* - b \leq 0, u^{*T}(Ax^* - b) = 0, \\ Qx^* + D + A^T u^* + E^T v^* = 0, \\ Ex^* - f = 0. \end{cases} \quad (13.2)$$

^{۲۱}Rodríguez Vazquez

^{۲۲}Convex quadratic programming

برای حل مسأله (۱۳.۲) تابع زیر را که معروف به تابع فیشر-برمیستر^{۲۳} است، معرفی می کنیم:

$$\begin{cases} \phi: \mathbb{R}^2 \rightarrow \mathbb{R} \\ \phi(a, b) = \sqrt{a^2 + b^2} - a - b. \end{cases}$$

قضیه ۱.۳.۲. [۲۵]:

۱. $\phi(a, b) = 0$ اگر و فقط اگر $a \geq 0, b \geq 0, ab = 0$.

۲. مربع $\phi(a, b)$ به طور پیوسته مشتق پذیر باشد.

۳. $\phi(a, b)$ همه جا به جز در مبدأ دو بار به طور پیوسته مشتق پذیر است، ولی در مبدأ قویاً^{۲۴} هموار است.

تابع ϕ که در شرط ۱ قضیه ی (۱.۳.۲) صدق کند، مساله مکمل غیر خطی^{۲۵} نامیده می شود.

متناظر با دستگاه (۱۳.۲) دستگاه زیر را در نظر می گیریم:

$$\phi(x, u, v) = \begin{bmatrix} \phi(\alpha, Qx + D + A^T u + E^T v) \\ \phi(u, Ax - b) \\ \phi(\beta, Ex - f) \end{bmatrix} = 0, \quad (14.2)$$

که در آن $\alpha = (\alpha_1, \dots, \alpha_m)^T > 0$ و $u = (u_1, \dots, u_m)$ ، $\beta = (\beta_1, \dots, \beta_l)^T > 0$ است. مدل شبکه عصبی متناظر با مساله ی (۱۴.۲) به صورت زیر بیان شده است:

$$\begin{cases} \frac{dy}{dt} = -\tau \nabla E(y(t)), \quad \tau > 0 \\ y(t_0) = y_0 \in \mathbb{R}^{n+m+l}. \end{cases} \quad (15.2)$$

که در آن τ ضریبی برای افزایش سرعت همگرایی و کاهش تعداد تکرارها در روش عددی انتخاب شده برای حل دستگاه دینامیکی (۱۵.۲) است. تابع انرژی متناظر این مدل به صورت زیر انتخاب شده است:

$$E(y) = \frac{1}{\tau} \|\phi(y)\|^2.$$

مدل هفتم

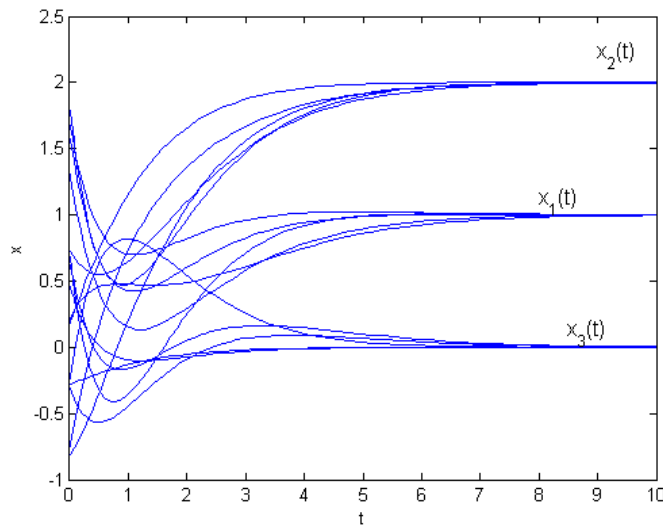
برای حل (۱۱.۲) و (۱۲.۲) تعریف کنید:

$$U(x, u, v) = \begin{bmatrix} -(Qx + D + A^T u + E^T v) \\ (u + Ax - b)^+ - u \\ (Ex - f) \end{bmatrix}. \quad (16.2)$$

^{۲۳}Fischer-Burmeister

^{۲۴}Strongly smooth

^{۲۵}Nonlinear complementarity problem



شکل ۱.۲: مسیر جواب در مثال ۲.۳.۲

مدل شبکه عصبی برای حل (۱۱.۲) و (۱۲.۲) در [۲۵] به صورت زیر بیان شده است:

$$\begin{cases} \frac{dy}{dt} = \kappa U(y), \\ y(t_0) = y_0, \end{cases} \quad (۱۷.۲)$$

که در آن κ نرخ همگرایی (نرخ همگرایی پارامتری است که سرعت همگرایی مدل شبکه عصبی با آن قابل تنظیم است) می باشد.

مثال ۲.۳.۲.

$$\begin{aligned} \min \quad & x_1^2 + x_2^2 + x_3^2 \\ \text{s.t.} \quad & \begin{cases} 2x_1 + x_2 - 5 \leq 0, \\ x_1 + x_3 - 2 \leq 0, \\ -x_1 + 1 \leq 0, \\ -x_2 + 2 \leq 0, \\ -x_3 \leq 0, \end{cases} \end{aligned}$$

جواب بهینه به ازای $\kappa = 1$ ، $x = (1, 2, 0)^T$ است. با استفاده از شبکه عصبی (۱۷.۲) نقطه تعادل بدست آمده جواب بهینه مساله می باشد که مسیرهای جواب در شکل (۱.۳) رسم شده است.

مدل هشتم

مسأله‌ی برنامه ریزی درجه دوم (۱۱.۲) و (۱۲.۲) را در نظر بگیرید که در آن $Q \in \mathbb{R}^{n \times n}$ یک ماتریس نیمه معین مثبت است. برای حل مساله (۱۱.۲) و (۱۲.۲) در [۲۹] یک مدل شبکه عصبی به صورت

زیر ارائه می شود:

$$\frac{dy}{dt} = \kappa(I + M^T) \begin{bmatrix} -(Qx + D + A^T u + E^T v) \\ (u + Ax - b)^+ - u \\ (Ex - f) \end{bmatrix}, \quad (18.2)$$

که در آن κ نرخ همگرایی است و

$$M = \begin{bmatrix} -Q & A^T & E^T \\ A & \circ_{m \times m} & \circ_{m \times l} \\ E & \circ_{l \times m} & \circ_{l \times l} \end{bmatrix}.$$

مدل نهم

مسأله برنامه ریزی درجه دوم زیر را در نظر بگیرید:

$$\min \quad f(x) = \frac{1}{2} x^T A x + a^T x \quad (19.2)$$

$$\text{s.t.} \quad \begin{cases} Dx \leq b, \\ x \geq \circ. \end{cases} \quad (20.2)$$

دوگان آن عبارت است از:

$$\max \quad \bar{f}(w) = b^T w - \frac{1}{2} x^T A x \quad (21.2)$$

$$\text{s.t.} \quad \begin{cases} \bar{D}^T w - Ax \leq a, \\ w \geq \circ. \end{cases} \quad (22.2)$$

که در آن $A \in \mathbb{R}^{n \times n}$ متقارن و نیمه معین مثبت است، $a \in \mathbb{R}^n$ و $b \in \mathbb{R}^m$ و $D \in \mathbb{R}^{m \times n}$ است. با توجه به شرایط شرایط فروش-کان-تاکر برای مسائل برنامه ریزی محدب، $(x^{*T}, w^{*T})^T$ به ترتیب یک جواب بهینه برای مسائل (۱۹.۲) - (۲۲.۲) است اگر و فقط اگر $(x^{*T}, w^{*T})^T$ در شرایط فروش-کان-تاکر زیر صدق کند:

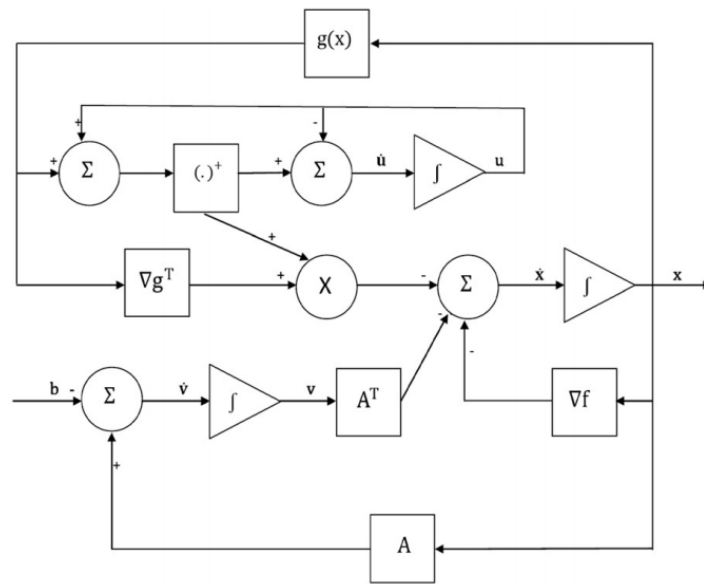
$$\begin{cases} w^* \geq \circ, Dx^* - b \leq \circ, w^{*T}(Dx^* - b) = \circ, \\ Ax^* + a + D^T w^* \geq \circ, x \geq \circ, \\ x^{*T}(Ax^* + a + D^T w^*) = \circ. \end{cases}$$

مدل متناظر با (۱۹.۲) و (۲۰.۲) در [۲۹] به صورت زیر بیان شده است:

$$\dot{u} = B(I + M^T)\{(u - (Mu + q)^+ - u)\},$$

که در آن

$$u = \begin{bmatrix} x \\ w \end{bmatrix}, q = \begin{bmatrix} a \\ b \end{bmatrix}, M = \begin{bmatrix} A & D^T \\ -D & \circ \end{bmatrix}.$$



شکل ۲۰۲: دیاگرام مدل شبکه عصبی ۲۵۰۲

M یک ماتریس نیمه معین مثبت است زیرا $u^T M u = x^T A x \geq 0$ همچنین $B = \lambda I$ که $\lambda > 0$ و با افزایش λ سرعت همگرایی افزایش می یابد.

مدل دهم [۲]

فرم کلی یک مساله ی برنامه ریزی غیر خطی محدب^{۲۶} را به صورت زیر در نظر بگیرید:

$$\min f(x) \quad (23.2)$$

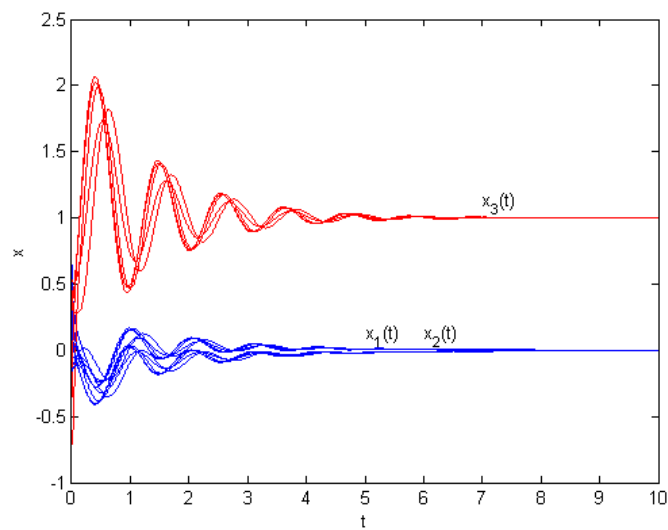
$$\text{s.t.} \begin{cases} g(x) \leq 0, \\ h(x) = 0, \end{cases} \quad (24.2)$$

که در آن $g(x) = (g_1(x), g_2(x), \dots, g_m(x))^T$, $f: \mathbb{R}^n \rightarrow \mathbb{R}$, $x \in \mathbb{R}^n$ ، تابع برداری $-m$ بعدی است. توابع $f, g_1, g_2, \dots, g_m$ محدب و دو بار مشتق پذیر فرض شده اند. $h(x) = Ax - b$. در اینجا فرض می کنیم که مساله (۲۳.۲) و (۲۴.۲) یک جواب بهینه دارد. متناظر با مساله (۲۳.۲) و (۲۴.۲) شبکه عصبی زیر را معرفی می کنیم:

$$\begin{cases} \frac{dx}{dt} = -(\nabla f(x) + \nabla g(x)^T (u + g(x))^+ + \nabla h(x)^T v), \\ \frac{du}{dt} = (u + g(x))^+ - u, \\ \frac{dv}{dt} = h(x), \end{cases} \quad (25.2)$$

با نقطه ی آغازی $y_0 = (x_0^T, u_0^T, v_0^T)^T$ شکل (۲۰۲) معماری این مدل را نشان می دهد.

^{۲۶}General convex nonlinear programming

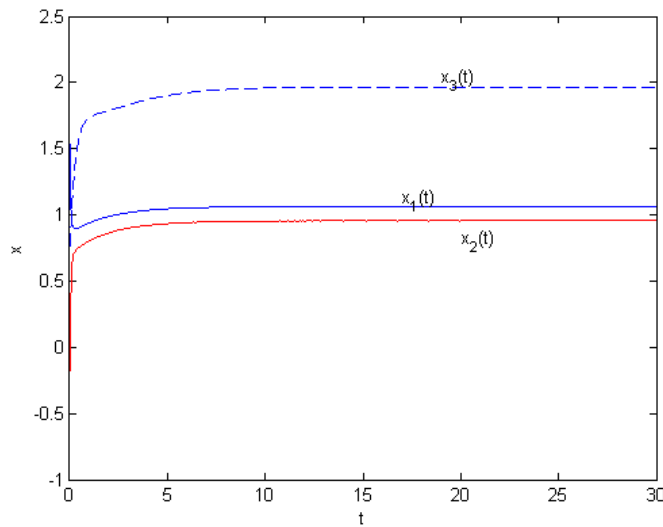


شکل ۳.۲: مسیر جواب در مثال ۳.۳.۲

مثال ۳.۳.۲.

$$\begin{aligned} \min \quad & (x_1 + 3x_2 + x_3)^2 + 4(x_1 - x_2)^2 \\ \text{s.t.} \quad & \begin{cases} x_1^3 - 6x_2 - 4x_3 \leq 0, \\ x_1 + x_2 + x_3 - 1 = 0, \\ x_1, x_2, x_3 \geq 0. \end{cases} \end{aligned}$$

جواب بهینه $x^* = (0, 0, 1)^T$ است. مسیرهای جواب در شکل (۳.۲) نشان داده شده اند.



شکل ۴.۲: مسیر جواب در مثال ۴.۳.۲

مثال ۴.۳.۲

$$\min \frac{1}{4}((x_1 - x_2)^4 + (x_2 + x_3)^2 + (x_1 + x_3)^2)$$

$$\text{s.t.} \begin{cases} x_1^2 + x_2^2 - x_3 \leq 0, \\ (2 - x_1)^2 + (2 - x_2)^2 - x_3 \leq 0, \\ 2e^{-x_1+x_2} - x_3 \leq 0, \\ x_1^2 + x_2^2 - 2x_1 + x_2 - 4 \leq 0, \\ |x_1| \leq 2, |x_2| \leq 2, x_3 \geq 0, \end{cases}$$

جواب بهینه $x^* = (1/109, 0/922, 1/954)^T$ است. مسیرهای جواب در شکل (۴.۲) نشان داده شده‌اند.

مدل یازدهم [۱]

مساله برنامه‌ریزی غیر خطی محدب^{۲۷} زیر را در نظر بگیرید:

$$\min f(x) \quad (26.2)$$

$$\text{s.t.} \begin{cases} g(x) \leq 0, \\ h(x) = 0. \end{cases} \quad (27.2)$$

^{۲۷}Convex nonlinear programming

که در آن $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$, $x \in \mathbb{R}^n$, $g(x) = (g_1(x), g_2(x), \dots, g_m(x))^T$, یک تابع برداری m -بعدی است. توابع $f, g_1, g_2, \dots, g_m$ محدب و دو بار مشتق پذیر فرض شده اند. $h(x) = Ax - b$, $A \in \mathbb{R}^{l \times n}$, $\text{rank}(A) = l$, $0 \leq l < n$, $b \in \mathbb{R}^l$. ابتدا تابع لاگرانژ زیر را در نظر می گیریم:

$$L(x, u, v) = f(x) + \frac{1}{r} \sum_{k=1}^m u_k^r g_k(x) + \sum_{p=1}^l v_p h_p(x)$$

که در آن $u \in \mathbb{R}^m$ و $v \in \mathbb{R}^l$ را ضرایب لاگرانژ می نامند. $x^* \in \mathbb{R}^n$ یک جواب بهینه (۲۶.۲) و (۲۷.۲) اگر و فقط اگر $u^* \in \mathbb{R}^m$ و $v^* \in \mathbb{R}^l$ وجود داشته باشد به طوری که (x^{*T}, u^{*T}, v^{*T}) در شرایط کروش-کان-تاکر زیر صدق کند:

$$\begin{cases} u^* \geq 0, g(x^*) \leq 0, u^{*T} g(x^*) = 0, \\ \nabla f(x^*) + \nabla g(x^*)^T u^* + \nabla h(x^*)^T v^* = 0, \\ h(x^*) = 0, \end{cases}$$

برای حل مساله ی (۲۶.۲) و (۲۷.۲) از سیستم دینامیکی زیر استفاده می کنیم:

$$\begin{cases} \frac{dx_i}{dt} = -\nabla_{x_i} L(x, u, v) = -\left(\frac{\partial f}{\partial x_i} + \frac{1}{r} \sum_{k=1}^m u_k^r \frac{\partial g_k}{\partial x_i} + \sum_{p=1}^l v_p \frac{\partial h_p}{\partial x_i}\right), \quad i = 1, \dots, n, \\ \frac{du}{dt} = \nabla_u L(x, u, v) = \text{diag}(u_1, \dots, u_m) g(x), \\ \frac{dv_p}{dt} = \nabla_{v_p} L(x, u, v) = h_p(x), \quad i = 1, \dots, l, \end{cases} \quad (28.2)$$

با یک نقطه آغازین $(x_0^T, u_0^T, v_0^T)^T$ و $u_k(t_0) > 0$ ($k = 1, 2, \dots, m$)

مثال ۵.۳.۲.

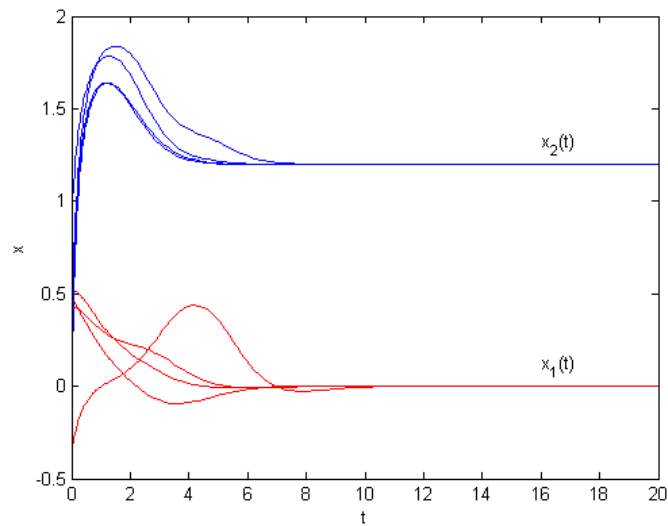
$$\begin{aligned} \min \quad & x_1^2 + x_2^2 + 2x_1x_2 + (x_1 - 1)^4 + (x_2 - 3)^4 \\ \text{s.t.} \quad & \begin{cases} x_1^2 + x_2^2 - 64 \leq 0, \\ (x_1 + 3)^2 + (x_2 + 4)^2 - 36 \leq 0, \\ (x_1 - 3)^2 + (x_2 + 4)^2 - 36 \leq 0, \\ x_1 \geq 0, \quad x_2 \geq 0, \end{cases} \end{aligned}$$

جواب بهینه $x^* = (0, 1/96)^T$ است و مسیرهای خروجی $x(t)$ در شکل (۵.۲) نشان داده شده است.

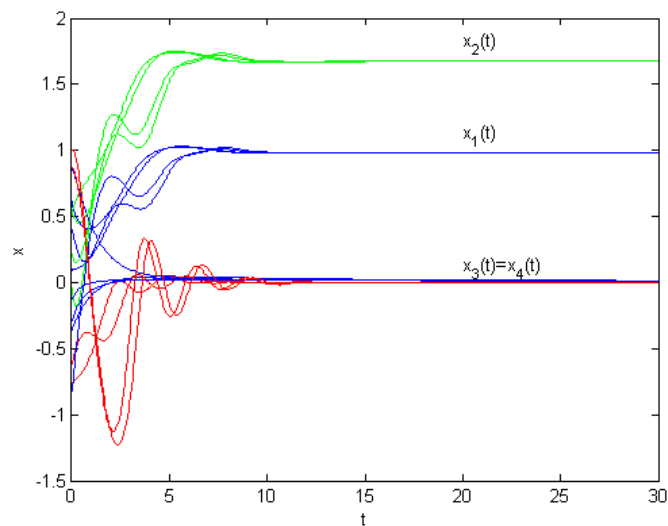
مثال ۶.۳.۲.

$$\begin{aligned} \min \quad & 0.4x_1 + x_1^2 + x_2^2 - x_1x_2 + 0.5x_3^2 + 0.5x_4^2 + \frac{1}{30}x_3^3 \\ \text{s.t.} \quad & \begin{cases} -x_1 + x_2 - x_3 \leq 2, \\ 3x_1 + x_2 - x_3 - x_4 \leq 18, \\ \frac{1}{10}x_1 + x_2 - x_4 = 2, \\ x \geq 0, \end{cases} \end{aligned}$$

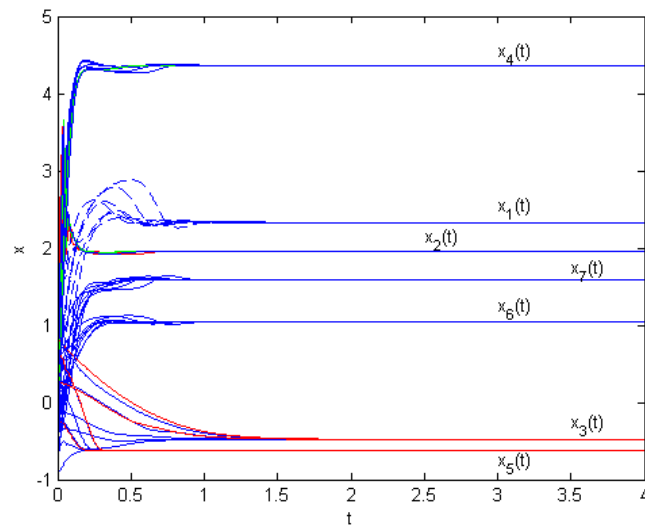
جواب بهینه $x^* = (0.982, 1.672, 0, 0)^T$ است که همگرایی مسیرها در شکل (۶.۲) نشان داده شده است.



شکل ۵.۲: مسیر جواب در مثال ۵.۳.۲



شکل ۶.۲: مسیر جواب در مثال ۶.۳.۲



شکل ۷.۲: مسیر جواب در مثال ۷.۳.۲

مثال ۷.۳.۲.

$$\min (x_1 - 10)^2 + 5(x_2 - 12)^2 + x_3^2 + 3(x_4 - 11)^2 + 10x_5^2 + 7x_6^2 + x_7^2 - 4x_6x_7 - 10x_6 - 8x_7$$

$$\text{s.t.} \begin{cases} 2x_1^2 + 3x_4^2 + x_3 + 4x_4^2 + 5x_5 - 127 \leq 0, \\ 7x_1 + 3x_2 + 10x_3^2 + x_4 - x_5 - 282 \leq 0, \\ 23x_1 + x_3^2 + 6x_6^2 - 8x_7 - 196 \leq 0, \\ 4x_1^2 + x_3^2 + 2x_4^2 - 3x_1x_2 + 5x_6 - 11x_7 \leq 0, \end{cases}$$

نقطه بهینه به صورت $x^* = (2/330, 1/951, -0/4775, 4/366, -0/625, 1/038, 1/594)^T$ می باشد. شکل (۷.۲) مسیرهای جواب را نمایش می دهد.

فصل ۳

یک روش منظم سازی آنتروپی برای حل مسائل مین-ماکس با برخی کاربردها

۱.۳ مقدمه

برخی مسائل علمی و مهندسی می توانند با مساله مین - ماکس^۱ بصورت زیر فرموله شوند [۱۰]:

$$\min_{x \in X} \max_{1 \leq i \leq m} \{f_i(x)\}, \quad (1.3)$$

که در آن X یک دامنه مشترک از مولفه توابع $f_i(x)$ ($i = 1, \dots, m$) است و f_i ها دوبار مشتق پذیر هستند. یک مشکل مهم مشتق پذیری تابع زیر است:

$$F(x) \equiv \max_{1 \leq i \leq m} \{f_i(x)\}, \quad x \in X. \quad (2.3)$$

توجه داشته باشید که مساله مین-ماکس (۱.۳) دارای تابع لاگرانژ^۲ زیر است:

$$L(x, \lambda) = \sum_{i=1}^m \lambda_i f_i(x), \quad (3.3)$$

برای هر $x \in X$ که در آن m بردار و λ نشان دهنده ضرایب لاگرانژ هستند که در شرایط نامنفی و نرمال بودن صدق می کنند، یعنی این ضرایب در مجموعه زیر قرار می گیرند:

$$\Lambda \equiv \{\lambda | \lambda \geq 0; \sum_{i=1}^m \lambda_i = 1\}. \quad (4.3)$$

با در نظر گرفتن برخی مفروضات منظم سازی، تابع ماکزیمم $F(x)$ را می توان به صورت زیر به دست آورد:

$$F(x) = \max_{\lambda \in \Lambda} L(x, \lambda), \quad x \in X. \quad (5.3)$$

^۱Min-max

^۲Lagrange

متاسفانه، $\max \sum_{i=1}^m \lambda_i f_i(x)$ روی $\lambda \in \Lambda$ ، به ندرت یک راه حل صریح و روشن برای بدست آوردن $\lambda(x)$ ارائه می دهد. در این فصل ما یک روش منظم سازی با جایگزین کردن تابع $L_p(x, \lambda)$ به جای $L(x, \lambda)$ به صورت زیر در نظر میگیریم:

$$L_p(x, \lambda) = \sum_{i=1}^m \lambda_i f_i(x) + \frac{1}{p} R(\lambda; \beta), \quad (6.3)$$

که در آن $p > 0$ پارامتر کنترل است، R یک تابع تنظیم است، و β بردار پارامتر اختیاری از R است. با انتخاب دقیق تابع R ، ماکزیمم تابع (۶.۳) را در نظر می گیریم:

$$F_p(x) \equiv \max_{\lambda \in \Lambda} L_p(x, \lambda). \quad (7.3)$$

مساله (۷.۳) یک تقریب برای مساله (۵.۳) است.

برخی از روش های منظم سازی شناخته شده از توابع زیر استفاده می کنند [۶]، [۱۸]:

$$R(\lambda; \beta) = -\frac{1}{2} \sum_{i=1}^m (\lambda_i - \beta_i)^2.$$

$$R(\lambda; \beta) = \sum_{i=1}^m l_n \lambda_i.$$

در قسمت بعد، یک روش منظم سازی با استفاده از توابع آنتروپی^۳ معرفی شده است. با این روش منظم سازی یک تقریب خوب و هموار $F_p(x)$ برای تابع $F(x)$ در (۲.۳) بدست می آید.

۲.۳ روش منظم سازی آنتروپی

۱.۲.۳ تابع آنتروپی شانون

فرض کنید $x \in X$ داده شده است، هنگامی که تابع آنتروپی شانون^۴ [۸] به عنوان تابع تنظیم به صورت

$$L_p(x, \lambda) \equiv \sum_{i=1}^m \lambda_i f_i(x) - \frac{1}{p} \sum_{i=1}^m (\ln \lambda_i) \lambda_i, \quad (8.3)$$

که در آن $\lambda \in \Lambda$. چون $L_p(x, \lambda)$ در λ اکیدا مقعر است، با ماکزیمم کردن آن روی Λ جواب منحصر به فرد بهینه زیر حاصل می شود:

$$\lambda_i^*(x, p) = \frac{\exp[p f_i(x)]}{Z}, \quad 1 \leq i \leq m, \quad (9.3)$$

که در آن

$$Z = \sum_{i=1}^m \exp[p f_i(x)]. \quad (10.3)$$

^۳Entropic

^۴Shannon's

با جایگزین کردن λ_i در (۸.۳) با $\lambda_i^*(x, p)$ داریم:

$$F_p(x) = \frac{1}{p} \ln \left\{ \sum_{i=1}^m \exp[p f_i(x)] \right\}, \quad (11.3)$$

که یک تابع هموار^۵ است. لازم به تاکید است که تابع $F_p(x)$ بدست آمده در اینجا معادل با ماکزیمم کردن $L_p(x, \lambda)$ روی Λ برای هر $x \in X$ است.

۲.۲.۳ خواص $F_p(x)$

لم ۱۰.۲.۳. $F_p(x)$ در X به صورت نقطه به نقطه به $F(x)$ همگرا^۶ است، زمانی که $p \rightarrow +\infty$.

برهان. فرض کنید $x \in X$ و $V(x)$ را به صورت یک تابع برداری با مولفه

$$v_i(x) = \exp[f_i(x)], \quad 1 \leq i \leq m,$$

تعریف کنید در این صورت نرم- l_p از $V(x)$ به صورت زیر بدست می آید:

$$\|V(x)\|_p \equiv \left\{ \sum_{i=1}^m [v_i(x)]^p \right\}^{\frac{1}{p}} = \left\{ \sum_{i=1}^m \exp[p f_i(x)] \right\}^{\frac{1}{p}}. \quad (12.3)$$

از اینرو

$$F_p(x) = \ln \|V(x)\|_p \quad (13.3)$$

بنابراین

$$\lim_{p \rightarrow \infty} \|V(x)\|_p = \max_{1 \leq i \leq m} v_i(x) = \exp[F(x)], \quad (14.3)$$

و

$$\lim_{p \rightarrow \infty} F_p(x) = F(x). \quad (15.3)$$

□

لم ۲.۲.۳. $F_p(x)$ در X به صورت یکنواخت به $F(x)$ همگرا^۶ است، زمانی که $p \rightarrow +\infty$.

برهان. با توجه به خواص نرم l_p برای $x \in X$ ، یک تابع نزولی یکنواخت در p است، یعنی:

$$F_s(x) \leq F_r(x), \quad r \leq s. \quad (16.3)$$

^۵Smooth

^۶Converges

از این رو از لم (۱۰.۲.۳)، برای هر $p > 0$ داریم:

$$0 \leq F_p(x) - F(x) = \ln \|V(x)\|_p - \ln \{\exp F(x)\} \quad (۱۷.۳)$$

$$= \left(\frac{1}{p}\right) \ln \sum_{i=1}^m \exp\{p[f_i(x) - F(x)]\} \quad (۱۸.۳)$$

$$\leq \left(\frac{1}{p}\right) \ln(m). \quad (۱۹.۳)$$

به عبارت دیگر برای هر $p > 0$

$$F(x) \leq F_p(x) \leq F(x) + \frac{\ln(m)}{p}, \quad \forall x \in X. \quad (۲۰.۳)$$

از این رو $F_p(x)$ همگرا به $F(x)$ به طور یکنواخت در X است، وقتی $p \rightarrow +\infty$. $\square$

لم ۳.۲.۳. $F_p(x)$ محدب است، اگر به ازای $i = 1, 2, \dots, m$ یک تابع محدب، تعریف شده روی یک مجموعه محدب X باشد.

برهان. برای هر $x, y \in X$ و $\lambda \in (0, 1)$ داریم:

$$\begin{aligned} F_p(\lambda x + (1 - \lambda)y) &= \frac{1}{p} \ln \sum_{i=1}^m \exp[p f_i(\lambda x + (1 - \lambda)y)] \\ &\leq \frac{1}{p} \ln \sum_{i=1}^m \exp[\lambda p f_i(x) + (1 - \lambda)p f_i(y)] \\ &= \frac{1}{p} \ln \sum_{i=1}^m \{\exp[p f_i(x)]\}^\lambda \{\exp[p f_i(y)]\}^{1-\lambda}. \end{aligned} \quad (۲۱.۳)$$

نامساوی هلدر^۷ [۹] نشان می دهد که

$$\sum_{i=1}^m \{\exp[p f_i(x)]\}^\lambda \{\exp[p f_i(y)]\}^{1-\lambda} \leq \left\{ \sum_{i=1}^m \exp[p f_i(x)] \right\}^\lambda \left\{ \sum_{i=1}^m \exp[p f_i(y)] \right\}^{1-\lambda}. \quad (۲۲.۳)$$

در نتیجه

$$\begin{aligned} F_p(\lambda x + (1 - \lambda)y) &\leq \frac{\lambda}{p} \ln \sum_{i=1}^m \exp[p f_i(x)] + \frac{1 - \lambda}{p} \ln \sum_{i=1}^m \exp[p f_i(y)] \\ &= \lambda F_p(x) + (1 - \lambda) F_p(y). \end{aligned} \quad (۲۳.۳)$$

این نشان می دهد که $F_p(x)$ یک تابع محدب روی X است. $\square$

به طور خلاصه ما تمام نتایج را در قضیه اصلی زیر داریم:

قضیه ۴.۲.۳. فرض کنید $X = \mathbb{R}^n$ و $f_i(x)$ محدب است، برای $i = 1, 2, \dots, m$. یک جواب بهینه تقریبی از مساله مین - ماکس (۱.۳) را با حل یک مساله برنامه ریزی محدب نا مقید $\{\min_x F_p(x)\}$ با p به اندازه کافی بزرگ بدست می آید.

^۷Holder

۳.۳ کاربردهای مرتبط

در این قسمت برخی کاربردها از آنالیز عددی و مسائل بهینه سازی معادل مساله را بررسی میکنیم.

۱.۳.۳ سیستم های نامساوی و مساوی

یک سیستم نامساوی در $\mathbb{R}^n$ را به صورت زیر در نظر بگیرید:

$$f_i(x) \leq 0, \quad i = 1, 2, \dots, m. \quad (24.3)$$

تعریف می کنیم:

$$F(x) \equiv \max_{1 \leq i \leq m} \{f_i(x)\}. \quad (25.3)$$

می دانیم که سیستم (۲۴.۳) جواب دارد اگر و تنها اگر $F(x)$ یک مینیمم کننده x^* داشته باشد آنچنانکه $F(x^*) \leq 0$. با استفاده از $F_p(x)$ بجای $F(x)$ ، حل مساله (۲۴.۳) معادل با حل مساله مینیمم سازی نامقید با تابع هدف هموار $F_p(x)$ است. همچنین داریم:

$$\nabla F_p(x) = J(x)\lambda, \quad (26.3)$$

و

$$\nabla^2 F_p(x) = \sum_{i=1}^m \lambda_i \nabla^2 f_i(x) + pJ(x)(\Gamma - \lambda\lambda^T)J(x)^T, \quad (27.3)$$

که در آن $J(x)$ ماتریس ژاکوبین $[\partial f_i / \partial x_j]$ ، λ یک بردار با مولفه های $\lambda_i = \exp[p f_i(x)] / \{\sum_{i=1}^m \exp[p f_i(x)]\}^{-1}$ است و $\Gamma = \text{diag}(\lambda_i)$.

برای حل یک سیستم نامساوی خطی به صورت

$$f_i(x) \equiv A_i(x) - b_i \leq 0, \quad i = 1, \dots, m, \quad (28.3)$$

تابع آنتروپی متناظر بصورت زیر در نظر گرفته می شود:

$$F_p(x) = \frac{1}{p} \ln \left\{ \sum_{i=1}^m \exp[p(A_i x - b_i)] \right\}. \quad (29.3)$$

گرادیان^۸ و هسین^۹ بصورت زیر تعریف می شود:

$$\nabla F_p(x) = A^T \lambda, \quad (30.3)$$

و

$$\nabla^2 F_p(x) = pA^T(\Gamma - \lambda\lambda^T)A, \quad (31.3)$$

که در آن

$$\lambda_i = \frac{\exp[p(A_i x - b_i)]}{\sum_{i=1}^m \exp[p(A_i x - b_i)]}. \quad (32.3)$$

^۸Gradient

^۹Hessian

برای حل دستگاه معادلات تعریف شده توسط

$$f_i(x) = 0, \quad i = 1, 2, \dots, m, \quad (33.3)$$

ما می توانیم هر معادله را توسط دو نامعادله $f_i(x) \leq 0$ و $-f_i(x) \leq 0$ یا با یک نامعادله منحصر به فرد $f_i^2(x) \leq 0$ جایگزین کنیم. البته اولی افزایش ابعاد مساله و دومی پیچیده گی توابع غیر خطی ایجاد شده را در بر دارد.

۲.۳.۳ حل مسائل برنامه ریزی خطی

مساله برنامه ریزی خطی زیر را در نظر بگیرید:

$$\min_x f(x) \equiv c^T x \quad (34.3)$$

$$\text{s.t.} \quad Ax = b, \quad (35.3)$$

$$e^T x = M, \quad (36.3)$$

$$x \geq 0, \quad (37.3)$$

که در آن e یک بردار یکه و M یک ثابت است. دوگان خطی (۳۴.۳) - (۳۷.۳) به صورت زیر تعریف می شود:

$$\min_y d(y) \equiv -b^T y - M y_{m+1} \quad (38.3)$$

$$\text{s.t.} \quad \sum_{i=1}^m a_{ij} y_i - c_j \leq -y_{m+1}, \quad j = 1, 2, \dots, n. \quad (39.3)$$

به وضوح، این مساله معادل با مساله مینیمم سازی بدون محدودیت زیر است:

$$\min_y \left\{ F(y) \equiv -b^T y + M \max_{1 \leq j \leq n} \left\{ \sum_{i=1}^m a_{ij} y_i - c_j \right\} \right\}. \quad (40.3)$$

یادآوری می کنیم که فرم استاندارد برنامه ریزی خطی کارماکار^{۱۰} [۴] با فرض $b = 1$ و $M = 1$ در نظر گرفته می شود. در این مورد، دوگان آن یک مساله مین - ماکس به صورت زیر است:

$$\min_y \left\{ F(y) \equiv \max_{1 \leq j \leq n} \left\{ \sum_{i=1}^m a_{ij} y_i - c_j \right\} \right\}. \quad (41.3)$$

روش منظم سازی آنتروپی برای حل مساله (۴۱.۳) منجر به مساله زیر می شود:

$$\min_y \left\{ F_p(y) \equiv \frac{1}{p} \sum_{j=1}^n \exp \left[p \left(\sum_{i=1}^m a_{ij} y_i - c_j \right) \right] \right\}. \quad (42.3)$$

^{۱۰} Karmarkar's

۳.۳.۳ حل مسائل مین - ماکس مقید

مسئله مین - ماکس مقید زیر را در نظر بگیرید:

$$\min_x \left\{ f(x) \equiv \max_{1 \leq i \leq m} \{f_i(x)\} \right\} \quad (43.3)$$

$$\text{s.t. } g_j(x) \leq 0, \quad j = 1, 2, \dots, l. \quad (44.3)$$

توجه داشته باشید که زمانی که $m = 1$ ، با یک مسئله برنامه ریزی غیر خطی با l محدودیت مواجه هستیم. مسئله (43.3) و (44.3) معادل با مسئله زیر است:

$$\min_{x, \alpha} \alpha \quad (45.3)$$

$$\text{s.t. } f_i(x) - \alpha \leq 0, \quad 1 \leq i \leq m, \quad (46.3)$$

$$g_j(x) \leq 0, \quad 1 \leq j \leq l. \quad (47.3)$$

هنگامی که از روش ذکر شده در [۷] استفاده می شود، کار اصلی در هر تکرار عبارت است از پیدا کردن مرکز یک مجموعه تراز بصورت زیر

$$G_\alpha(x) \equiv \{x | f_i(x) - \alpha \leq 0, \quad 1 \leq i \leq m; \quad g_j(x) \leq 0, \quad 1 \leq j \leq l\}, \quad (48.3)$$

با تعریف یک تابع فاصله

$$d(x) \equiv \min \{-f_i(x) + \alpha, \quad 1 \leq i \leq m; \quad -g_j(x), \quad 1 \leq j \leq l\}, \quad (49.3)$$

مرکز $G_\alpha(x)$ را می توان با ماکزیمم کردن $d(x)$ یا به طور معادل با مینیمم کردن $-d(x)$ از طریق تقریب هموار آن بصورت زیر پیدا کرد:

$$d_p(x) = \frac{1}{p} \ln \left\{ \sum_{i=1}^m \exp[p(f_i(x) - \alpha)] + \sum_{j=1}^l \exp[p g_j(x)] \right\}. \quad (50.3)$$

فصل ۴

یک مدل شبکه عصبی تصویری برای حل رده ای از مسائل بهینه سازی مشتق ناپذیر

۱.۴ مقدمه

در این فصل ما به تشریح مسائل بهینه سازی مشتق ناپذیر^۱ به صورت زیر میپردازیم

$$\min_{x \in \Omega} f(x) = \max_{1 \leq i \leq s} \{f_i(x)\} \quad (1.4)$$

$$\text{s.t.} \quad h_j(x) \leq 0, \quad j = 1, 2, \dots, t \quad (2.4)$$

که در آن

$$\Omega = \{x \in \mathbb{R}^n | l_i \leq x_i \leq h_i, \quad i = 1, 2, \dots, n\},$$

$$f_i(x) (i = 1, 2, \dots, s), h_j(x) (j = 1, 2, \dots, t)$$

توابع محدب دوبار مشتق پذیر هستند. هرچند توابع $f_i(x) (i = 1, 2, \dots, s)$ به طور پیوسته مشتق

$$\text{پذیر هستند ولی } f(x) = \max_{1 \leq i \leq s} \{f_i(x)\}$$

یک تابع ناهموار^۲ است. مساله (۱.۴) و (۲.۴)، بنابراین به عنوان یک مساله بهینه سازی ناهموار طبقه بندی میشود. این نوع مسائل بهینه سازی در بسیاری از زمینه های مختلف در ریاضیات مهندسی و مسائل مالی مطرح میشود ([۱۵]-[۱۱]) در بسیاری از مقالات قبلی این مسائل را با الگوریتم های عددی حل میکنند ([۲۲]-[۱۳]). از آنجا که زمان محاسبات تا حد زیادی به ابعاد و ساختار مسائل بستگی دارد الگوریتم های عددی معمولاً در مقیاس بزرگ و برای مسائل بهینه سازی زمان واقعی، کمتر موثر است. شبکه های عصبی با محاسبات موازی و همگرایی سریع می توانند به عنوان یک روش

^۱Non-differentiable

^۲non-smooth

کارآمد برای حل مسائل با مقیاس بزرگ و یا مسائل بهینه سازی زمان واقعی در نظر گرفته شوند. به تازگی، شبکه های عصبی برای حل مسائل بهینه سازی به طور گسترده مورد مطالعه قرار گرفته و برخی نتایج مهم نیز بدست آمده است. ([۳۲]-[۲۳]) به عنوان مثال، چن^۳ و همکاران، [۲۶]، شیا^۴ و همکاران [۲۸، ۲۹] و گائو^۵ و همکاران [۳۰]، شبکه های عصبی را برای حل مسائل برنامه ریزی محدب مشتق پذیر بر اساس روش اولیه - دوگان و عملگر تصویر بکار برده اند. در این فصل ابتدا مساله (۱.۴) و (۲.۴) به یک برنامه ریزی غیر خطی محدب با محدودیت های نامساوی تبدیل می شود و بر اساس شرایط بهینگی کروش-کان-تاکر (KKT) برای بهینگی برنامه ریزی محدب، یک مدل شبکه عصبی برای حل مساله بدست آمده پیشنهاد خواهد شد. با ساختن یک تابع لیاپانف مناسب، پایداری مجانبی و همگرایی سراسری شبکه اثبات می شود.

۲.۴ مدل شبکه عصبی

یک روش مناسب برای تبدیل مساله (۱.۴) و (۲.۴) بصورت زیر است:

$$\begin{aligned} & \min_{x \in \Omega} z \\ \text{s.t. } & \begin{cases} f_i(x) \leq z, & i = 1, 2, \dots, s, \\ h_j(x) \leq 0 & j = 1, 2, \dots, m. \end{cases} \end{aligned} \quad (3.4)$$

واضح است که این تبدیل منجر به برنامه ریزی محدب مشتق پذیر با تعداد محدودیت ها و پیچیدگی محاسبات بالا میشود.

به تازگی، یک روش تقریبی برای حل (۱.۴) و (۲.۴) ارائه شده است. با استفاده از تابع آنتروپی معرفی شده در فصل ۳ مساله زیر را در نظر میگیریم:

$$\begin{aligned} \min_{x \in \Omega} F_p(x) &= \frac{1}{p} \ln \sum_{i=1}^s \exp(pf_i(x)) \\ \text{s.t. } & h_j(x) \leq 0 \quad j = 1, 2, \dots, m, \quad p \rightarrow +\infty. \end{aligned} \quad (4.4)$$

از تجزیه و تحلیل بالا، واضح است که مساله مشتق ناپذیر (۱.۴) و (۲.۴) معادل یک مساله برنامه ریزی محدب است. بنابراین مساله برنامه ریزی غیرخطی محدب زیر را با محدودیت های نامساوی غیر خطی بصورت زیر در نظر می گیریم:

$$\begin{aligned} & \min_{x \in \Omega} F(x) \\ \text{s.t. } & g_j(x) \leq 0, \quad j = 1, 2, \dots, m, \end{aligned} \quad (5.4)$$

^۳Chen

^۴Xia

^۵Gao

که $(j = 1, 2, \dots, m), g_j(x), F(x)$ محدب و به طور پیوسته مشتق پذیر هستند. فرض کنید

$$S = \{x \in \mathbb{R}^n : g_j(x) \leq 0; j = 1, 2, \dots, m\}$$

به خوبی می دانیم $x^* \in S$ جواب بهینه (۵.۴) است اگر و تنها اگر $\lambda^* = (\lambda_1^*, \lambda_2^*, \dots, \lambda_m^*)^T \in \mathbb{R}^m$ وجود داشته باشد به طوریکه $(x^{*T}, \lambda^{*T})^T$ در شرایط کروش-کان-تاکر زیر صدق کند:

$$\begin{cases} (x - x^*)^T (\nabla F(x^*) + \nabla g(x^*)^T \lambda^*) \geq 0, & \forall x \in \Omega, \\ \lambda_j^* g_j(x^*) = 0, & \lambda_j^* \geq 0, \quad g_j(x^*) \leq 0, \quad j = 1, \dots, m. \end{cases} \quad (6.4)$$

لم ۱۰.۲.۴. [۲۸] جواب بهینه (۵.۴) است اگر و تنها اگر وجود داشته باشد $\lambda^* = (\lambda_1^*, \lambda_2^*, \dots, \lambda_m^*)^T \in \mathbb{R}^m$ به طوریکه:

$$\begin{cases} x^* - P_\Omega[x^* - \nabla F(x^*) - \nabla g(x^*)^T \lambda^*] = 0, \\ \lambda^* - [\lambda^* + g(x^*)]^+ = 0. \end{cases} \quad (7.4)$$

که

$$P_\Omega(x_i) = \begin{cases} l_i, & x_i < l_i, \\ x_i, & l_i \leq x_i \leq h_i, \\ h_i, & x_i > h_i. \end{cases} \quad (x_i)^+ = \begin{cases} x_i, & x_i \geq 0, \\ 0, & x_i < 0. \end{cases}$$

لم ۱۰.۲.۴ نشان می دهد که جواب بهینه (۵.۴) با حل (۷.۴) به دست می آید.

لم ۲۰.۲.۴. [۳۳] فرض کنید مجموعه $\Omega \subset \mathbb{R}^n$ یک مجموعه محدب بسته است. آنگاه

$$(v - P_\Omega(v))^T (P_\Omega(v) - u) \geq 0, \quad v \in \mathbb{R}^n, u \in \Omega, \quad (8.4)$$

$$\|P_\Omega(u) - P_\Omega(v)\| \leq \|u - v\|, \quad u, v \in \mathbb{R}^n. \quad (9.4)$$

شبکه عصبی تصویر فریز^۶ برای حل (۵.۴) به صورت زیر بیان می شود:

$$\begin{cases} \frac{dx}{dt} = -(x - P_\Omega[x - \nabla F(x) - \nabla g(x)^T \lambda]), \\ \frac{d\lambda}{dt} = -(\lambda - [\lambda + g(x)]^+). \end{cases} \quad (10.4)$$

متأسفانه شبکه عصبی (۱۰.۴) برای برخی مسائل برنامه ریزی محدب ناپایدار است [۲۸]. برای سادگی

^۶ Friesz

فرض می کنیم:

$$\begin{aligned} u(t) &= (x^T, \lambda^T)^T \in \mathbb{R}^{n+m}, \\ D &= \Omega \times \mathbb{R}_+^m, \\ \bar{\lambda} &= [\lambda + g(x)]^+, \\ \bar{x} &= P_\Omega[x - \nabla F(x) - \nabla g(x)^T \lambda], \end{aligned}$$

و D^* را مجموعه نقاط بهینه (۵.۴) در نظر می گیریم.

براساس نتایج فوق مدل شبکه عصبی جدید زیر را معرفی می کنیم:

$$\begin{cases} \frac{dx}{dt} = -(x - P_\Omega[x - \nabla F(x) - \nabla g(x)^T \bar{\lambda}]), \\ \frac{d\lambda}{dt} = -(\lambda - [\lambda + g(x)]^+). \end{cases} \quad (۱۱.۴)$$

راساس نتایج فوق مدل شبکه عصبی جدید زیر را معرفی می کنیم:

$$\begin{cases} \frac{dx}{dt} = -(x - P_\Omega[x - \nabla F(x) - \nabla g(x)^T \bar{\lambda}]), \\ \frac{d\lambda}{dt} = -(\lambda - [\lambda + g(x)]^+). \end{cases} \quad (۱۲.۴)$$

فرض کنید D^e مجموعه نقاط تعادل شبکه عصبی (۱۲.۴) باشد. طبق لم ۱۰.۲.۴، $x \in D^e$ یک نقطه تعادل شبکه عصبی (۱۲.۴) است اگر و تنها اگر $x \in D^*$ یک جواب بهینه (۵.۴) باشد. در این صورت $D^* = D^e$.

۳.۴ تجزیه و تحلیل پایداری

در این بخش پایداری نقطه تعادل و همگرایی جواب بهینه شبکه عصبی (۱۲.۴) را مطالعه می کنیم.

قضیه ۱۰.۳.۴. برای هر نقطه اولیه $u(t_0) = (x(t_0)^T, \lambda(t_0)^T)^T \in \mathbb{R}^{n+m}$ یک جواب پیوسته منحصر به فرد $u(t) = (x(t)^T, \lambda(t)^T)^T$ برای سیستم (۱۲.۴) وجود دارد. علاوه بر این $x(t) \in \Omega$ و $\lambda(t_0) \geq 0$.

برهان. $\nabla F(x)$ و $\nabla g_j(x)$ روی مجموعه محدب باز $D \subseteq \mathbb{R}^{n+m}$ شامل $\Omega \times \mathbb{R}_+^m$ به طور پیوسته مشتق پذیر هستند. پس $x - P_\Omega[x - \nabla F(x) - \nabla g(x)^T \lambda]$ و $\lambda - [\lambda + g(x)]^+$ به طور موضعی پیوسته لیپ شیتز هستند. طبق قضیه وجود و یکتایی جواب برای معادلات دیفرانسیل معمولی مسئله مقدار اولیه سیستم (۱۲.۴) یک جواب منحصر به فرد دارد.

فرض کنید نقطه اولیه $x_0 = x(t_0) \in \Omega$ و $\lambda_0 = \lambda(t_0) \geq 0$ داده شده است. چون:

$$\begin{cases} \frac{dx}{dt} + x = P_\Omega[x - \nabla F(x) - \nabla g(x)^T \bar{\lambda}], \\ \frac{d\lambda}{dt} + \lambda = [\lambda + g(x)]^+. \end{cases} \quad (۱۳.۴)$$

در نتیجه

$$\begin{cases} \int_{t_0}^t (\frac{dx}{dt} + x) e^t dt = \int_{t_0}^t P_{\Omega}[x - \nabla F(x) - \nabla g(x)^T \bar{\lambda}] e^t dt, \\ \int_{t_0}^t (\frac{d\lambda}{dt} + \lambda) e^t dt = \int_{t_0}^t [\lambda + g(x)]^+ e^t dt. \end{cases} \quad (۱۴.۴)$$

آنگاه

$$\begin{cases} x(t) = e^{-(t-t_0)} x_0 + e^{-t} \int_{t_0}^t e^t P_{\Omega}[x - \nabla F(x) - \nabla g(x)^T \bar{\lambda}] dt, \\ \lambda(t) = e^{-(t-t_0)} \lambda_0 + e^{-t} \int_{t_0}^t e^t [\lambda + g(x)]^+ dt. \end{cases} \quad (۱۵.۴)$$

با استفاده از قضیه مقدار میانگین برای انتگرال ها داریم:

$$\begin{cases} x(t) = e^{-(t-t_0)} x_0 + (1 - e^{-(t-t_0)}) P_{\Omega}[\hat{x} - \nabla F(\hat{x}) - \nabla g(\hat{x})^T \hat{\lambda}], \\ \lambda(t) = e^{-(t-t_0)} \lambda_0 + (1 - e^{-(t-t_0)}) [\hat{\lambda} + g(\hat{x})]^+. \end{cases} \quad (۱۶.۴)$$

□ چون $x(t_0) \in \Omega$ و $\lambda(t_0) \geq 0$ در نتیجه $x(t) \in \Omega$ و $\lambda(t) \geq 0$ و این اثبات را کامل می کند.

ملاحظه ۲.۳.۴. نتایج قضیه ۱.۳.۴ نشان می دهد که شبکه عصبی (۱۲.۴) خوش تعریف است. در اثبات قضیه ۱.۳.۴ فرض کردیم $(x_0^T, \lambda_0^T)^T \in \Omega \times \mathbb{R}_+^m$ یعنی نقطه اولیه باید شدنی باشد. اما اگر $(x_0^T, \lambda_0^T)^T \notin \Omega \times \mathbb{R}_+^m$ وقتی $t \rightarrow +\infty$ $(x(t, x_0)^T, \lambda(t, x_0)^T)^T \in \Omega \times \mathbb{R}_+^m$ می شود. بنابراین یک لحظه t_0 وجود دارد به طوریکه به ازای $t > t_0$ $(x(t, x_0)^T, \lambda(t, x_0)^T)^T$ در ناحیه شدنی قرار می گیرد و $(x(t, x_0)^T, \lambda(t, x_0)^T)^T$ در $\Omega \times \mathbb{R}_+^m$ می ماند.

قضیه ۳.۳.۴. فرض کنید $f(x)$ و $g_j(x)$ ، $j = 1, 2, \dots, m$ ، روی مجموعه محدب باز $N \subseteq \mathbb{R}^{n+m}$ شامل $\Omega \times \mathbb{R}_+^m$ محدب و مشتق پذیر باشند. آنگاه شبکه عصبی (۱۲.۴) به مفهوم لیاپانوف پایدار است و برای هر نقطه اولیه $(x(t_0)^T, \lambda(t_0)^T)^T \in \mathbb{R}^{n+m}$ مسیر جواب (۱۲.۴) به یک نقطه در D^* همگرا خواهد بود. به خصوص شبکه عصبی (۱۲.۴) هنگامی که D^* فقط یک نقطه دارد پایدار مجانبی سراسری است.

برهان. با استفاده از قضیه ۱.۳.۴، $\forall (x_0^T, \lambda_0^T)^T \in \Omega \times \mathbb{R}_+^m$ یک جواب پیوسته منحصر به فرد $(x(t)^T, \lambda(t)^T)^T \subseteq \Omega \times \mathbb{R}_+^m$ برای سیستم (۱۲.۴) وجود دارد. یک تابع لیاپانوف به صورت زیر تعریف می کنیم:

$$\begin{aligned} V(x, \lambda) = & F(x) - F(x^*) + \frac{1}{\gamma} [(\bar{\lambda})^2 - (\lambda^*)^2] - (x - x^*)^T (\nabla F(x^*) + \nabla g(x^*)^T \lambda^*) \\ & - (\lambda - \lambda^*)^T \lambda^* + \frac{1}{\gamma} \|x - x^*\|^2 + \frac{1}{\gamma} \|\lambda - \lambda^*\|^2. \end{aligned} \quad (۱۷.۴)$$

باتوجه به اینکه $\|\bar{\lambda}\|^2 = \sum_{j=1}^m [(\lambda_j + g_j(x))^+]^2$ و

$$[(\lambda_j + g_j(x))^+]^2 = \begin{cases} [(\lambda_j + g_j(x))]^2, & \lambda_j + g_j(x) \geq 0, \\ 0, & o.w. \end{cases} \quad (۱۸.۴)$$

داریم:

$$\nabla \|\bar{\lambda}\|^2 = \nabla \left(\sum_{j=1}^m [(\lambda_j + g_j(x))^+]^2 \right) = \begin{pmatrix} 2 \nabla g(x)^T \bar{\lambda} \\ 2 \bar{\lambda} \end{pmatrix}. \quad (۱۹.۴)$$

با محاسبه مشتق $V(t)$ در طول مسیر سیستم (۱۲.۴) با $-\lambda + \bar{\lambda} = g(x) - (\lambda + g(x))^-$ داریم:

$$\begin{aligned} \frac{dV(x, \lambda)}{dt} &= [\nabla F(x) + \nabla g(x)^T \bar{\lambda} - \nabla F(x^*) - \nabla g(x^*)^T \lambda^* + x - x^*]^T (-x + \bar{x}) \\ &\quad + \frac{(\bar{\lambda} + \lambda - 2\lambda^*)^T (-\lambda + \bar{\lambda})}{2} \\ &= -\|x - \bar{x}\|^2 \\ &\quad + [\nabla F(x) + \nabla g(x)^T \bar{\lambda} - \nabla F(x^*) - \nabla g(x^*)^T \lambda^* + \bar{x} - x^*]^T (-x + \bar{x}) \\ &\quad - \frac{1}{2} \|\lambda - \bar{\lambda}\|^2 + (\bar{\lambda} - \lambda^*)^T (-\lambda + \bar{\lambda}) \\ &= -\|x - \bar{x}\|^2 - [\nabla F(x) - \nabla F(x^*)]^T (x - x^*) - (\nabla F(x^*) \\ &\quad + \nabla g(x^*)^T \lambda^*)^T (\bar{x} - x^*) \\ &\quad - [x - \nabla F(x) - \nabla g(x)^T \bar{\lambda} - \bar{x}]^T (\bar{x} - x^*) \\ &\quad - (\nabla g(x)^T \bar{\lambda} - \nabla g(x^*)^T \lambda^*)^T (\bar{x} - x^*) \\ &\quad - \frac{1}{2} \|\lambda - \bar{\lambda}\|^2 + (\bar{\lambda} - \lambda^*)^T [g(x) - (\lambda + g(x))^-] \\ &= -\|x - \bar{x}\|^2 - \frac{1}{2} \|\lambda - \bar{\lambda}\|^2 - [\nabla F(x) - \nabla F(x^*)]^T (x - x^*) \\ &\quad - [\nabla F(x^*) + \nabla g(x^*)^T \lambda^*]^T (\bar{x} - x^*) - [x - \nabla F(x) - \nabla g(x)^T \bar{\lambda} - \bar{x}]^T (\bar{x} - x^*) \\ &\quad - \bar{\lambda}^T [-\nabla g(x)(x^* - x) - g(x) + g(x^*)] + \bar{\lambda}^T g(x^*) \\ &\quad - \bar{\lambda}(\lambda + g(x))^- + (\lambda^*)^T [\nabla g(x^*)(x - x^*) - g(x) + g(x^*)] \\ &\quad - (\lambda^*)^T g(x^*) + (\lambda^*)^T (\lambda + g(x))^- \\ &= -\|x - \bar{x}\|^2 - \frac{1}{2} \|\lambda - \bar{\lambda}\|^2 - [\nabla F(x) - \nabla F(x^*)]^T (x - x^*) \\ &\quad - [\nabla F(x^*) + \nabla g(x^*)^T \lambda^*]^T (\bar{x} - x^*) - [x - \nabla F(x) - \nabla g(x)^T \bar{\lambda} - \bar{x}]^T (\bar{x} - x^*) \\ &\quad - \bar{\lambda}^T [g(x^*) - g(x) - \nabla g(x)(x^* - x)] + \bar{\lambda}^T g(x^*) \\ &\quad - (\lambda^*)^T [g(x) - g(x^*) - \nabla g(x^*)(x - x^*)] + (\lambda^*)^T (\lambda + g(x))^- \\ &\leq -\|x - \bar{x}\|^2 - \frac{1}{2} \|\lambda - \bar{\lambda}\|^2 - [\nabla F(x) - \nabla F(x^*)]^T (x - x^*) \\ &\quad - [\nabla F(x^*) + \nabla g(x^*)^T \lambda^*]^T (\bar{x} - x^*) - [x - \nabla F(x) - \nabla g(x)^T \bar{\lambda} - \bar{x}]^T (\bar{x} - x^*) \\ &\quad - \bar{\lambda}^T [g(x^*) - g(x) - \nabla g(x)(x^* - x)] \\ &\quad - (\lambda^*)^T [g(x) - g(x^*) - \nabla g(x^*)(x - x^*)]. \end{aligned} \quad (۲۰.۴)$$

در نامساوی لم ۲.۲.۴ فرض کنید $\bar{\lambda} = \nabla g(x)^T$ و $v = x - \nabla f(x) - \nabla g(x)^T \bar{\lambda}$ و $u = x^*$ لذا داریم:

$$(x - \nabla F(x) - \nabla g(x)^T \bar{\lambda} - \bar{x})^T (\bar{x} - x^*) \geq 0. \quad (21.4)$$

به واسطه مشتق پذیری و تحدب $F(x)$ و $g(x)$ ، $\forall x \in \Omega$ داریم:

$$\begin{cases} [\nabla F(x) - \nabla F(x^*)]^T (x - x^*) \geq 0, \\ g(x^*) - g(x) - \nabla g(x)^T (x^* - x) \geq 0, \\ g(x) - g(x^*) - \nabla g(x^*)^T (x - x^*) \geq 0. \end{cases} \quad (22.4)$$

با جایگذاری (۶.۴)، (۲۱.۴) و (۲۲.۴) در (۲۰.۴) داریم:

$$\frac{dV(x, \lambda)}{dt} \leq -\|x - \bar{x}\|^2 - \frac{1}{4}\|\lambda - \bar{\lambda}\|^2 \leq 0. \quad (23.4)$$

این به این معناست که شبکه عصبی (۱۲.۴) به مفهوم لیاپانوف پایدار است.

چون

$$V(x, \lambda) \geq \frac{1}{4}(\|x - x^*\|^2 + \|\lambda - \lambda^*\|^2).$$

بنابراین $V(x, \lambda)$ معین مثبت و شعاعی بی کران است. بنابراین یک زیر دنباله همگرا به صورت زیر وجود دارد.

$$\{(x(t_k)^T, \lambda(t_k)^T)^T \mid t_0 < t_1 < \dots < t_k < t_{k+1} < \dots\}, \quad t_k \rightarrow \infty, \quad k \rightarrow \infty,$$

به طوریکه

$$\lim_{k \rightarrow \infty} (x(t_k)^T, \lambda(t_k)^T)^T = (\hat{x}^T, \hat{\lambda}^T)^T.$$

ازقضیه لسل داریم:

$\{(x(t)^T, \lambda(t)^T)^T \mid t \rightarrow \infty\}$ بزرگترین مجموعه تغییر ناپذیر در

$$K = \{(x(t)^T, \lambda(t)^T)^T \mid \frac{dV(x, \lambda)}{dt} = 0\},$$

است.

از طرفی (۱۲.۴) و (۲۳.۴) نشان می دهند که

$$\frac{d\lambda}{dt} = 0, \frac{dx}{dt} = 0 \iff \frac{dV(x, \lambda)}{dt} = 0.$$

بنابراین $(\hat{x}^T, \hat{\lambda}^T)^T \in D^*$ با $M \subseteq k \subseteq D^*$. با جایگذاری $x^* = \hat{x}$ و $\lambda^* = \hat{\lambda}$ در (۱۷.۴) تابع لیاپانوف دیگری را به صورت زیر تعریف می کنیم:

$$\begin{aligned} \hat{V}(x, \lambda) = & F(x) - F(\hat{x}) + \frac{1}{4}[(\bar{\lambda})^2 - (\hat{\lambda})^2] - (x - \hat{x})^T (\nabla F(\hat{x}) + \nabla g(\hat{x})\hat{\lambda}) \\ & - (\lambda - \hat{\lambda})^T \hat{\lambda} + \frac{1}{4}\|x - \hat{x}\|^2 + \frac{1}{4}\|\lambda - \hat{\lambda}\|^2. \end{aligned} \quad (24.4)$$

پس $\hat{V}(x, \lambda)$ به طور پیوسته مشتق پذیر است و $\hat{V}(\hat{x}, \hat{\lambda}) = 0$. با توجه به اینکه

$$\lim_{k \rightarrow \infty} (x(t_k)^T, \lambda(t_k)^T)^T = (\hat{x}^T, \hat{\lambda}^T)^T,$$

بنابراین

$$\lim_{k \rightarrow \infty} \hat{V}(x(t_k)^T, \lambda(t_k)^T)^T = \hat{V}(\hat{x}, \hat{\lambda}) = 0.$$

بنابراین $\forall \epsilon > 0$ وجود دارد $q > 0$ به طوریکه برای هر $t > t_q$ داریم $\hat{V}(x, \lambda) < \epsilon$. به طور مشابه با تجزیه و تحلیل بالا می توانیم ثابت کنیم که تابع $\frac{d\hat{V}(x, \lambda)}{dt} \leq 0$ در نتیجه برای $t \geq t_q$ داریم:

$$\|x(t) - \hat{x}\|^2/2 + \|\lambda(t) - \hat{\lambda}\|^2/2 \leq \hat{V}(x, \lambda) \leq \epsilon.$$

این یعنی $\lim_{t \rightarrow \infty} x(t) = \hat{x}$ و $\lim_{t \rightarrow \infty} \lambda(t) = \hat{\lambda}$ است. بنابراین مسیر جواب شبکه عصبی (۱۲.۴) به نقطه تعادل $(\hat{x}^T, \hat{\lambda}^T)^T$ همگرای سراسری است. همچنین $(\hat{x}^T, \hat{\lambda}^T)^T$ یک جواب بهینه (۵.۴) است. به خصوص اگر $D^* = \{((x^*)^T, (\lambda^*)^T)^T\}$ آنگاه برای هر $x_* \in \Omega$ و $\lambda_* \geq 0$ جواب $(x^T, \lambda^T)^T$ با نقطه اولیه $(x_*^T, \lambda_*^T)^T$ با تجزیه و تحلیل بالا به $((x^*)^T, (\lambda^*)^T)^T$ همگرا می شود. این یعنی شبکه عصبی (۱۲.۴) پایدار مجانبی سراسری است و این اثبات را کامل می کند. $\square$

۴.۴ مثال شبیه سازی شده

مثال ۱.۴.۴. برنامه ریزی مین ماکس زیر را در نظر بگیرید:

$$\begin{aligned} \min \quad & f(x) = \max\{x_1^2 + x_2^2; (2 - x_1)^2 + (2 - x_2)^2; 2e^{-x_1 + x_2}\} \\ \text{s.t.} \quad & \begin{cases} x_1^2 + x_2^2 - 2x_1 + x_2 \leq 4, \\ |x_1| \leq 2, \quad |x_2| \leq 2. \end{cases} \end{aligned} \quad (25.4)$$

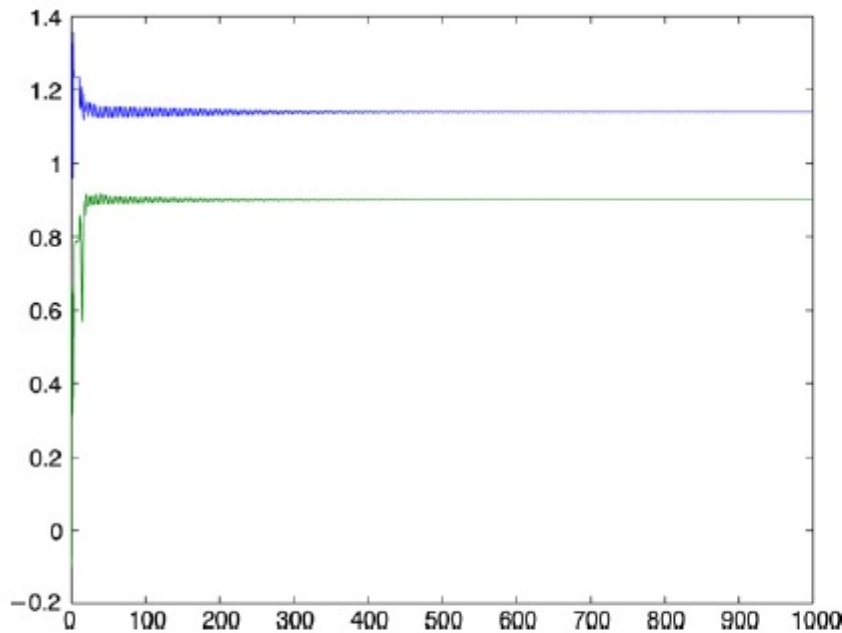
مساله (۲۵.۴) متعلق به یک رده از مسائل بهینه سازی مشتق ناپذیر است.

حالت ۱: فرض کنید $z = \max\{x_1^2 + x_2^2; (2 - x_1)^2 + (2 - x_2)^2; 2e^{-x_1 + x_2}\}$ (۲۵.۴) را می توان به برنامه ریزی محدب معادل زیر تبدیل کرد:

$$\begin{aligned} \min \quad & z \\ \text{s.t.} \quad & \begin{cases} x_1^2 + x_2^2 - z \leq 0 \\ (2 - x_1)^2 + (2 - x_2)^2 - z \leq 0 \\ 2e^{-x_1 + x_2} - z \leq 0 \\ x_1^2 + x_2^2 - 2x_1 + x_2 \leq 4 \\ |x_1| \leq 2; \quad |x_2| \leq 2 \end{cases} \end{aligned} \quad (26.4)$$

نتیجه شبیه سازی نشان می دهد که مسیر جواب همگرا به نقطه تعادل $x^e = (1/1392, 0/8996)^T$ است.

شکل ۱.۴ نشان می دهد که مسیر شبکه به طور سراسری همگرا است.



شکل ۱.۴: مسیرهای شبکه های عصبی در حالت (۱)

بنابراین، جواب بهینه بدست آمده عبارتست از $x^* = (۱/۱۳۹۲, ۰/۸۹۹۶)^T$. همچنین

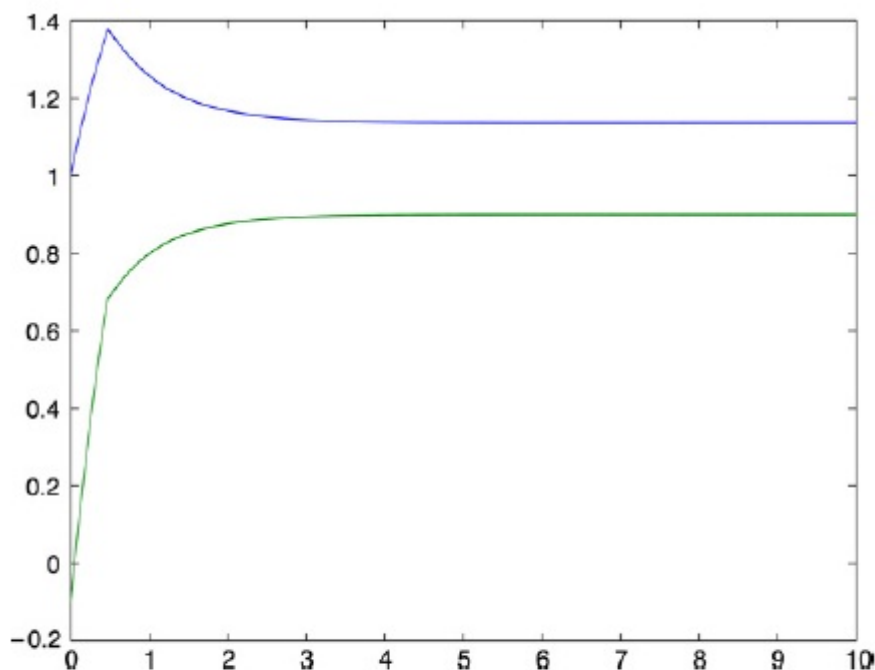
$$f_1(x^*) = ۱/۹۵۲۷, f_2(x^*) = ۱/۹۵۱۹, f_3(x^*) = ۱/۵۷۴۰$$

حالت ۲: فرض کنید $F(x) = \frac{1}{p} \ln \sum_{i=1}^3 \exp(pf_i(x))$. واضح است که، $F(x)$ یک تابع محدب هموار است برای هر $p > 0$. مساله برنامه ریزی محدب زیر را در نظر بگیرید:

$$\begin{aligned} \min F(x) &= \frac{1}{p} \ln \sum_{i=1}^3 \exp(pf_i(x)) \\ \text{s.t. } &\begin{cases} x_1^2 + x_2^2 - 2x_1 + x_2 \leq 4, \\ |x_1| \leq 2; \quad |x_2| \leq 2. \end{cases} \end{aligned} \quad (۲۷.۴)$$

با $p = ۱۰^5$ نتیجه شبیه سازی نشان می دهد مسیر جواب همگرا به نقطه تعادل $x^e = (۱/۱۳۷۷, ۰/۹۰۰۶)^T$ است. شکل ۲.۴ نشان می دهد که مسیر شبکه عصبی به طور سراسری همگراست. بنابراین، جواب بهینه بدست آمده عبارتست از $x^* = (۱/۱۳۹۲, ۰/۰/۸۹۹۶)^T$. همچنین داریم:

$$f_1(x^*) = ۱/۹۵۲۲, f_2(x^*) = ۱/۹۵۲۳, f_3(x^*) = ۱/۵۷۷۹$$



شکل ۲.۴: مسیرهای شبکه عصبی در حالت (۲)

۵.۴ نتیجه گیری و پیشنهادات برای کارهای آتی

در این فصل، ما روش های بهینه سازی برای حل رده ای از مسائل بهینه سازی مشتق ناپذیر را بررسی کردیم. ابتدا مساله برنامه ریزی مشتق ناپذیر به یک مساله برنامه ریزی محدب مشتق پذیر تبدیل می شود. با ساختن یک مدل شبکه عصبی برای حل برنامه ریزی محدب، پایداری نقطه تعادل و همگرایی مسیرها به جواب بهینه بررسی می شود. کارایی مدل پیشنهادی با یک مثال داده می شود. از جمله مزیت های مدل ارائه شده می توان به موارد زیر اشاره کرد:

- حل مساله ریاضیات چند انگشتی
- کاربرد مدل ارائه شده در مسائل پردازش تصویر
- حل مساله رگرسیون و طبقه بندی داده ها
- حل مسائل بهینه سازی سبد سهام مشتق ناپذیر
- این مدل نیاز به پارامتر جریمه ندارد

- با وجود سادگی این مدل، در رده بسیاری از مسائل بهینه‌سازی محدب قابل استفاده است
- این مدل هیچ‌گونه وابستگی به نقطه شروع نداشته و این نقطه می‌تواند خارج از ناحیه شدنی و نقطه‌ای تصادفی باشد
- این مدل می‌تواند قادر به حل مسائل وابسته به زمان و مسائل با ابعاد بزرگ باشد

۱.۵.۴ پیشنهاداتی برای کارهای آتی

- در زمینه آنالیز عددی و محاسبات عددی، مینیم کردن ماکزیمم خطای چند جمله ای درونیاب با استفاده از تابع آنتروپی و چند جمله ای های چبیشف بر مبنای مسائل مین ماکس
- در زمینه مهندسی پزشکی، مینیم کردن خطای دستگاه رادیو تراپی با ماکزیمم تابش پرتو بر اساس شبکه عصبی و مسائل مین ماکس
- در زمینه نظامی مینیم کردن خطای موشک با ماکزیمم برد پرتاب، بر اساس مسائل مین ماکس

مراجع

- [1] A. R. Nazemi, A dynamic system model for solving convex nonlinear optimization problems, *comman Nonlinear Sci Number Simulat* 17, 2012, 1696-1705
- [2] A. R. Nazemi, Solving general convex nonlinear optimization problems by an efficient neurodynamic model, *Engineering Application of Artificial Intelligence* 26, 2013, 685-696
- [3] Bazaraa. M. S. Sherali, H. D. and Shetty, C. M., 1979. *Nonlinear Programming, Theory and Algorithms*, John Wiley Sons, New York.
- [4] Fang, S. C, Puthenpura. S, (1993) *Linear optimization and extensions: Theory and algorithms*. Prentice Hall, Englewood Cliffs, New Jersey
- [5] Gigola. C, Gomez. S, (1990) A regularization method for solving the finite convex min-max problem. *SIAM Journal on Numerical Analysis* 27:1621-1634
- [6] Hiriart-Urruty J-B, Lemarechal. C, (1993) *Convex analysis and minimization algorithm*. Springer-Verlag, Berlin
- [7] Huard. P, (1967) Resolution of mathematical programming with nonlinear constraints by the method of centers. In: Abadie J (ed) *Nonlinear programming*, North-Holland, Amsterdam: 207-219
- [8] Jaynes ET (1957) Information theory and statistical mechanics. *Physics Review* 106:620-630
- [9] Kazarinoff ND (1961) *Analytic inequalities*. Holt, Rinehart and Winston, New York
- [10] Polak E, Mayne DQ, Higgins JE (1991) Superlinearly convergent algorithm for min-max problems. *Journal of Optimization Theory and Applications* 69:407-439
- [11] E. Cherkaev, A. Cherkaev, Minimax optimization problem of structural design, *Computers and Structures* 86 (2008) 1426-1435.
- [12] K.L. Teo, X.Q. Yang, Portfolio selection problem with minimax type risk function, *Journal of Annals of Operations Research* 101 (2001) 333-349.

- [13] Y. H. Yu, L. Gao, Nonmonotone linear search algorithm for constrained minimax problem, *Journal of Optimization Theory and Application* 115 (2002) 419-446.
- [14] X. Li, S. Fang, On the entropy regularization method for solving min-max problems with applications, *Mathematics Methods of Operations Research* 46 (1997) 119-130.
- [15] X. Li, An efficient approach to a class of non-smooth optimization problems, *Science in China (Series A)* 37 (1994) 323-330.
- [16] G. Di Pillo, L. Grippo, S. Lucidi, Smooth transformation of the generalized minimax problem, *Journal of Optimization Theory and Application* 95 (1997) 1-24.
- [17] C. Charalambous, A.R. Conn, An efficient method to solve the minimax problem directly, *SIAM Journal on Numerical Analysis* 15 (1978) 162-187.
- [18] C. Gigola, S. Gomez, A regularization method for solving the finite convex min-max problem, *SIAM Journal on Numerical Analysis* 27 (1990) 1621-1634.
- [19] A. Vardi, New minimax algorithm, *Journal of Optimization Theory and Application* 75 (1992) 613-634.
- [20] F. Ye, H. Liu, S. Zhou, S. Liu, A smoothing trust-region Newton-CG method for minimax problem, *Applied Mathematics and Computation* 199 (2008) 581-589.
- [21] M. Chao, Z. Wang, Y. Liang, Q. Hu, Quadratically constraint quadratical algorithm model for nonlinear minimax problems, *Applied Mathematics and Computation* 205 (2008) 247-262.
- [22] Z. Zhu, X. Cai, J. Jian, An improved SQP algorithm for solving minimax problems, *Applied Mathematics Letters* (in press).
- [23] D.W. Tank, J.J. Hopfield, Simple 'neural' optimization network: An A/D converter signal decision circuit and a linear programming circuit, *IEEE Transactions on Circuits and System* 33 (1986) 533-541.
- [24] M.P. Kennedy, L.O. Chua, Neural networks for nonlinear programming, *IEEE Transactions on Circuits and System* 35 (1988) 554-562.
- [25] S. Effati, A.R. Nazemi, Neural network models and its application for solving linear and quadratic programming problems, *Applied Mathematics and Computation* 172 (2006) 305-331.
- [26] K.Z. Chen, Y. Leung, K.S. Leung, X.B. Gao, A neural network for nonlinear programming problems, *Neural Computing Application* 11 (2002) 103-111.

- [27] Q. Tao, J. Cao, M. Xue, H. Qiao, A high performance neural network for solving nonlinear programming problems with hybrid constraints, *Physics Letters A* 288 (2001) 88-94.
- [28] Y. Xia, H. Leung, J. Wang, A projection neural network and its application to constrained optimization problems, *IEEE Transactions on Circuits and System-I* 49 (2002) 447-458.
- [29] Y. Xia, J. Wang, A recurrent neural networks for nonlinear convex optimization subject to nonlinear inequality constraints, *IEEE Transactions on Circuits and System-I* 51 (2004) 1385-1394.
- [30] X.B. Gao, A novel neural network for nonlinear convex programming, *IEEE Transactions on Neural Networks* 15 (2004) 613-621.
- [31] M. Forti, P. Nistri, M. Quincampoix, Generalized neural network for nonsmooth nonlinear programming problems, *IEEE Transactions on Circuits and System-I* 51 (2004) 1741-1754.
- [32] X. Xue, W. Bian, Subgradient-based neural network for nonsmooth convex optimization problems, *IEEE Transactions on Circuits and System-I* 55 (2008) 2378-2391.
- [33] D. Kinderlehrer, G. Stampcchia, *An Introduction to Variational Inequalities and their Applications*, Academic, New York, 1980.

واژه‌نامه فارسی به انگلیسی

LaSalle Principle of Invariance	اصل تغییرناپذیری لسال
Strictly convex	اکیداً محدب
Strictly Monotone	اکیداً یکنوا
Convex Optimization	بهینه‌سازی محدب
Nonconvex Optimization	بهینه‌سازی نامحدب
Globally Asymptotically Stable	پایدار مجانبی سراسری
Stability in the Sense of Lyapunov	پایداری به مفهوم لیاپانوف
Globally Exponential Stability	پایداری نمایی سراسری
Lipschitz Continuous	پیوسته لیپ شیتز
Locally Lipschitz Continuous	پیوسته لیپ‌شیتز محلی
Energy Function	تابع انرژی
Regression Function	تابع رگرسیون
Hubber Function	تابع هوبر
Trancpose Jacobian	ترانپازه ژاکوبین
Activation Function	تابع فعال‌سازی
Firing	تحریک شدن
Feasible Solution	جواب شدنی
Least Squares	حداقل مربعات
State Trajector	خط سیر
Local Duality	دوگانی موضعی
Trasient Behavior	رفتار ناپایدار
Euler Method	روش اویلر
Discrete time	زمان گسسته
Neural Networks	شبکه‌های عصبی
Feedback Neural Network	شبکه‌های عصبی بازگشتی

Artificial Neural Networks	شبکه‌های عصبی مصنوعی
Lagrangian Multiplier	ضرب‌گر لاگرانژ
Projection Operator	عملگر تصویر
Global Minimum	کمینه سراسری
Local Minimum	کمینه موضعی
Hessian Matrix	ماتریس هسین
Jacobian Matrix	ماتریس ژاکوبین
Support Vector Machine	ماشین بردار پشتیبانی
Dual Variable	متغیر دوگان
Invariant Set	مجموعه تغییرناپذیر
Level Set	مجموعه سطح
Quadratic Programming	مسئله برنامه ریزی درجه دوم
Lagrangian Dual Problem	مسئله حداقل مربعات
Dual Program	مسئله دوگان
Linear Complementarity Problem	مسئله مکمل خطی
Nonlinear Complementarity Problem	مسئله مکمل غیرخطی
Positive Definite	معین مثبت
Strongly Positive Definite	معین مثبت قوی
Negative Definite	معین منفی
Nonnegative Region	ناحیه نامنفی
Initial Point	نقطه اولیه
Equilibrium Point	نقطه تعادل
Differentiable Mapping	نگاشت مشتق پذیر
Positive Semidefinite	نیمه معین مثبت
Globally Convergent	همگرای سراسری

واژه‌نامه انگلیسی به فارسی

Activation Function	تابع فعال سازی
Artificial Neural Networks	شبکه های عصبی مصنوعی
Convex Optimization	بهینه سازی محدب
Differentiable Mapping	نگاشت مشتق پذیر
Discrete time	زمان گسسته
Dual Program	مسأله دوگان
Dual Variable	متغیر دوگان
Equilibrium Point	نقطه تعادل
Energy Function	تابع انرژی
Euler Method	روش اویلر
Feasible Solution	جواب شدنی
Feedback Neural Network	شبکه های عصبی بازگشتی
Firing	تحریک شدن
Globally Asymptotically Stable	پایدار مجانبی سراسری
Globally Convergent	همگرای سراسری
Globally Exponential Stability	پایدار نمایی سراسری
Global Minimum	کمینه سراسری
Hessian Matrix	ماتریس هسین
Hubber Function	تابع هوبر
Initial Point	نقطه اولیه
Invariant Set	مجموعه تغییر ناپذیر
Jacobian Matrix	ماتریس ژاکوبین
Lagrangian Dual Problem	مسأله دوگان لاگرانژ
Lagrangian Multiplier	ضربگر لاگرانژ
Least Squares	حداقل مربعات

Level Set.....	مجموعه سطح
Linear Complementarity Problem.....	مسأله مکمل خطی
Lipschitz Continuous.....	پیوسته لیپ شیتز
Local Duality.....	دوگانی موضعی
Locall Minimum.....	کمینه موضعی
Locally Lipschitz Continuous.....	پیوسته لیپ شیتز محلی
Negative Definite.....	معین منفی
Neural Networks.....	شبکه های عصبی
Nonconvex Optimization.....	بهینه سازی نا محدب
Nonlinear Complementarity Problem.....	مسأله مکمل غیرخطی
Nonnegative Region.....	ناحیه نامنفی
Positive Definite.....	معین مثبت
Projection Operator.....	عملگر تصویر
Positive Semidefinite.....	نیمه معین مثبت
Quadratic Programming.....	مسأله برنامه ریزی درجه دوم
Regression Function.....	تابع رگرسیون
Stability in the Sense of Lyapunov.....	پایداری به مفهوم لیاپانوف
State Trajector.....	خط سیر
Strictly convex.....	اکیداً محدب
Strictly Monotone.....	اکیداً یکنوا
Support Vector Machine.....	ماشین بردار پشتیبانی
Strongly Positive Definite.....	معین مثبت قوی
Trancpose Jacobian.....	ترانهاده ژاکوبین
Trasient Behavior.....	رفتار ناپایدار

Aabstract

In this thesis, we describe a neural network model based on a dynamic optimization technique for solving a class of non-smooth optimization problems with min-max objective function. The basic idea is to replace the min-max function by a smooth one using an entropy function. With this smoothing technique, the non-smooth problem is converted into an equivalent differentiable convex programming problem. A neural network model is then constructed based on Karush-Kuhn-Tucker optimality conditions. It is investigated that the proposed neural network is stable in the sense of Lyapunov and can converge to an exact optimal solution of the original problem. The effectiveness of the method are demonstrated by several numerical simulations.

keywords: Stability, Convergence, Non-differentiable Programming, Optimization, Neural network, Min-Max Problem, Linear and Nonlinear Programming, Entropy



Shahrood University

Faculty of Mathematical Sciences

**An application of dynamic systems for
solving a class of non-smooth optimization
problems**

Abdolsakhi Nazari

Supervisor

Dr. Alireza Nazemi

2015