



دانشگاه صنعتی شاهرود
دانشکده ریاضی
گروه ریاضی کاربردی

یک کاربرد از مدل‌های شبکه‌های عصبی برای حل مسائل بهینه‌سازی سبد سرمایه

نگارش

نرگس طهماسبی

استاد راهنما

جناب آقای دکتر ناظمی

پایان‌نامه جهت اخذ درجه کارشناسی ارشد

شهریور ۱۳۹۱

چکیده

با پیشرفت فن‌آوری اطلاعات و ارتباطات و توسعه ارتباطات 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 یافته اند، موفقیت‌های چشم‌گیری در حل مؤثر و کارای مسائل بهینه‌سازی به‌دست آورده‌اند. روش‌هایی چون الگوریتم ژنتیک، جستجوی ممنوع، شبیه‌سازی تبریدی، شبکه عصبی و ... قابلیت‌های خود را در حل مسائل بزرگ عملی به‌خوبی 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 به‌صورت موازی در شبکه‌های عصبی امتیاز دیگری است که با توجه به گسترش سخت افزارهای موازی، امکان حل مسائل بسیار بزرگ را توسط این رویکرد ممکن می‌سازد.

در این پایان نامه دو مدل مختلف شبکه عصبی بازگشتی برای حل رده‌ای از مسائل بهینه‌سازی ارائه می‌شود. تحلیل وجود یکتایی، پایداری و همگرایی سراسری جواب‌ها مورد بررسی قرار می‌گیرند و عملکرد روش‌های ارائه شده با به‌کارگیری چند مثال از مسائل برنامه‌ریزی تصادفی، برنامه‌ریزی کسری، بهینه‌سازی مقاوم و بهینه‌سازی سبد سرمایه نشان داده می‌شود. در انتها نتایج کار و پیشنهاداتی برای کارهای آتی ارائه می‌دهیم.

واژه‌های کلیدی: شبکه‌های عصبی، پایداری، مدل میانگین-واریانس، برنامه‌ریزی مخروط مرتبه دوم، برنامه‌ریزی تصادفی، برنامه‌ریزی کسری، بهینه‌سازی مقاوم.

فهرست مطالب

ث	لیست جداول	
ج	لیست تصاویر	
۱	مقدمه و مفاهیم اولیه	۱
۲	۱.۱ مقدمه	۱.۱
۲	۲.۱ توابع محدب	۲.۱
۷	۳.۱ پایداری	۳.۱
۱۰	۴.۱ تابع انرژی	۴.۱
۱۲	۵.۱ برنامه‌ریزی مخروط مرتبه دوم	۵.۱
۱۴	۶.۱ مفاهیم آماری	۶.۱
۱۷	۷.۱ مفاهیم اقتصادی	۷.۱
۱۷	۱.۷.۱ بازده	۱.۷.۱
۱۹	۲.۷.۱ ریسک و مفهوم آن	۲.۷.۱
۲۲	۲ مقدمه‌ای بر شبکه‌های عصبی	۲
۲۳	۱.۲ مقدمه	۱.۲
۲۴	۲.۲ ساختار شبکه‌های عصبی مصنوعی	۲.۲
۲۶	۳.۲ مدل ریاضی یک سلول عصبی	۳.۲
۲۹	۴.۲ تاریخچه کاربرد شبکه‌های عصبی مصنوعی در حل مسائل بهینه‌سازی	۴.۲
۳۱	۵.۲ مدل‌های شبکه عصبی ارائه شده پیشین برای حل مسائل بهینه‌سازی	۵.۲
۳۹	۳ بهینه‌سازی سبد سرمایه	۳
۴۰	۱.۳ مقدمه	۱.۳
۴۲	۲.۳ مسأله بهینه‌سازی سبد سرمایه	۲.۳
۴۶	۳.۳ مدل‌های جایگزین	۳.۳
۴۸	۱.۳.۳ مدل‌های درجه دوم	۱.۳.۳
۵۰	۲.۳.۳ مدل‌های خطی	۲.۳.۳
۵۳	۴.۳ مدل شبکه عصبی	۴.۳

۵۴	 تحلیل پایداری و همگرایی	۵.۳
۶۰	 مثال‌های عددی	۶.۳

۷۲	مسائل برنامه‌ریزی تصادفی	۴
۷۳	 مقدمه	۱.۴
۷۳	 ۱.۱.۴ برنامه‌ریزی خطی تصادفی	۱.۴
۷۸	 مسائل بهینه‌سازی با قیود تصادفی	۲.۴
۸۶	 مدل شبکه عصبی	۳.۴
۸۸	 تحلیل پایداری و همگرایی	۴.۴
۹۵	 مثال‌های عددی	۵.۴

۱۰۴	مسائل برنامه‌ریزی کسری	۵
۱۰۵	 مقدمه	۱.۵
۱۰۵	 برنامه‌ریزی کسری	۲.۵
۱۰۷	 ۱.۲.۵ برنامه‌ریزی کسری با صورت و مخرج خطی	۱.۲.۵
۱۰۸	 ۲.۲.۵ برنامه‌ریزی کسری با تابع هدف درجه دوم	۲.۲.۵
۱۰۹	 ۳.۲.۵ برنامه‌ریزی کسری محدب با صورت درجه دوم	۳.۲.۵
	 ۴.۲.۵ برنامه‌ریزی با تابع هدف مختلط از کسری و درجه دوم و همراه با قیود	۴.۲.۵
۱۱۰	 درجه دوم	
۱۱۲	 مثال‌های عددی	۳.۵

۱۲۷	بهینه‌سازی مقاوم	۶
۱۲۸	 مقدمه	۱.۶
۱۲۸	 عدم قطعیت در پارامترهای مدل	۲.۶
۱۳۲	 انواع مجموعه‌های عدم قطعیت در بهینه‌سازی مقاوم	۳.۶
۱۳۵	 برنامه‌ریزی خطی مقاوم	۴.۶
۱۳۷	 اشکال گوناگون مقاوم بودن	۵.۶
۱۳۷	 ۱.۵.۶ مقاوم بودن محدودیت‌ها	۱.۵.۶
۱۳۸	 ۲.۵.۶ مقاوم بودن تابع هدف	۲.۵.۶
۱۴۰	 مدل‌های بهینه‌سازی مقاوم در مسائل مالی	۶.۶
۱۴۰	 ۱.۶.۶ انتخاب سبد سرمایه مقاوم	۱.۶.۶
۱۴۲	 مثال عددی	۲.۶.۶

۱۴۷	نتایج و پیشنهاداتی برای کارهای آینده
-----	--------------------------------------

۱۵۰	مراجع و مآخذ
-----	--------------

۱۶۸	واژه‌نامه انگلیسی به فارسی
-----	----------------------------

لیست جداول

۶۸		سبد سرمایه‌های کارا	۱.۳
۷۱		مقایسه مدل‌ها	۲.۳
۹۶		مقادیر مثال ۱.۵.۴	۱.۴
۹۷		مقادیر مثال ۲.۵.۴	۲.۴
۱۰۲		مقادیر مثال ۴.۵.۴	۳.۴

لیست تصاویر

۷	پایداری لیپانوف	۱.۱
۸	پایداری مجانبی	۲.۱
۱۳	مخروط مرتبه دوم در $\mathbb{R}^3$	۳.۱
۱۶	تابع چگالی نرمال استاندارد	۴.۱
۲۵	شبکه‌های عصبی مصنوعی	۱.۲
۲۶	مدل ساده یک سلول عصبی	۲.۲
۲۷	مدل بلوکی سلول عصبی	۳.۲
	انواع تابع تحریک: الف) تابع حد آستانه، ب) تابع قسمتی خطی، ج) تابع سیگموئید	۴.۲
۲۹	با شیب‌های متغیر	
۴۴	مرز کارایی و منحنی مطلوبیت	۱.۳
۵۵	بلوک دیاگرام شبکه عصبی (۲۸.۳)-(۲۹.۳)	۲.۳
۶۲	رفتار گذرا شبکه عصبی (۲۸.۳)-(۲۹.۳) با ۱۰ نقطه آغازین دلخواه در مثال ۱.۶.۳	۳.۳
۶۳	رفتار گذرا شبکه عصبی (۲۸.۳)-(۲۹.۳) با ۳۰ نقطه آغازین دلخواه در مثال ۲.۶.۳	۴.۳
	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$	۵.۳
۶۳	در مثال ۲.۶.۳	
	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$	۶.۳
۶۴	در مثال ۳.۶.۳	
۶۶	رفتار گذرا شبکه عصبی (۲۸.۳)-(۲۹.۳) با ۴۰ نقطه آغازین دلخواه در مثال ۴.۶.۳	۷.۳
۶۷	رفتار گذرا شبکه عصبی (۲۸.۳)-(۲۹.۳) با ۳۰ نقطه آغازین دلخواه در مثال ۵.۶.۳	۸.۳
۶۸	رفتار همگرایی از $\ x(t) - x^*\ ^2$ در مثال ۵.۶.۳	۹.۳
۶۹	مرز کارایی سبد سرمایه‌های بدست آمده در مثال ۶.۶.۳	۱۰.۳
۷۰	رفتار گذرا شبکه عصبی (۲۸.۳)-(۲۹.۳) با ۳۰ نقطه آغازین دلخواه در مثال ۶.۶.۳	۱۱.۳
۷۰	رفتار همگرایی از $\ x(t) - x^*\ ^2$ در مثال ۶.۶.۳	۱۲.۳
۸۹	بلوک دیاگرام شبکه عصبی (۵۶.۴)-(۵۷.۴)	۱.۴
۹۷	رفتار گذرا شبکه عصبی (۵۶.۴)-(۵۷.۴) با ۳۰ نقطه آغازین دلخواه در مثال ۱.۵.۴	۲.۴
۹۸	منحنی فاز شبکه عصبی (۵۶.۴)-(۵۷.۴) با ۱۴ نقطه آغازین در مثال ۱.۵.۴	۳.۴

۴.۴	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$
۹۸	در مثال ۱.۵.۴
۵.۴	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۲ نقطه آغازین در مثال ۲.۵.۴
۶.۴	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با ۲۰ نقطه اولیه دلخواه در مثال ۲.۵.۴
۷.۴	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۵۰ نقطه آغازین دلخواه در مثال ۳.۵.۴
۸.۴	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۰۰ نقطه آغازین دلخواه در مثال ۴.۵.۴
۹.۴	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۰ نقطه آغازین در مثال ۴.۵.۴
۱.۵	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۵۰ نقطه آغازین دلخواه در مثال ۱.۳.۵
۲.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$
۱۱۴	در مثال ۱.۳.۵
۳.۵	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۰۰ نقطه آغازین دلخواه در مثال ۲.۳.۵
۴.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با ۲۰ نقطه اولیه متفاوت در مثال ۲.۳.۵
۵.۵	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۳ نقطه آغازین در مثال ۳.۳.۵
۶.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با y_0 اولیه در مثال ۳.۳.۵
۷.۵	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۷ نقطه آغازین در مثال ۴.۳.۵
۸.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (0, -1, -2, -3, -4, -5, -6, -7, -8)^T$
۱۱۹	در مثال ۴.۳.۵
۹.۵	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۵۰ نقطه آغازین دلخواه در مثال ۵.۳.۵
۱۰.۵	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۷ نقطه آغازین در مثال ۵.۳.۵
۱۱.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با ۳۰ نقطه اولیه متفاوت در مثال ۵.۳.۵
۱۲.۵	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با $x_0 = (-1, 1)^T$ در مثال ۶.۳.۵
۱۳.۵	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۸ نقطه آغازین در مثال ۶.۳.۵
۱۴.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (-1, 1, -1, 1, -1, 1, -1, 1, -1)^T$
۱۲۴	در مثال ۶.۳.۵
۱۵.۵	رفتار گذرا شبکه عصبی با ۵۰ نقطه آغازین دلخواه در مثال ۷.۳.۵
۱۶.۵	منحنی فاز شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۱۸ نقطه آغازین در مثال ۷.۳.۵
۱۷.۵	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با y_0 در مثال ۷.۳.۵
۱.۶	مقاوم بودن محدودیتها
۲.۶	مقاوم بودن تابع هدف
۳.۶	ناحیه شدنی مسأله (۱۵.۶) - (۱۶.۶)
۴.۶	مجموعه جوابهایی با پیشیمانی کمتر از ۷۵٪ مسأله (۱۵.۶) - (۱۶.۶)
۵.۶	رفتار گذرا شبکه عصبی (۵۶.۴) - (۵۷.۴) با ۲۰ نقطه آغازین دلخواه
۶.۶	رفتار همگرایی از $\ x(t) - x^*\ ^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$

فصل ۱

مقدمه و مفاهیم اولیه

۱.۱ مقدمه

در این فصل تعاریف مورد نیاز در فصل‌های بعدی به‌طور مختصر آورده شده است. ابتدا تعاریف مرتبط با تابع محدب^۱ را بیان نموده و در خصوص مسائل بهینه‌سازی مقید و نامقید به ذکر شرایط لازم و کافی بهینگی^۲ می‌پردازیم. سپس مفاهیم پایداری^۳ و تابع انرژی در سیستم‌های دینامیکی را ارائه می‌دهیم. در ادامه مسائل برنامه‌ریزی مخروط مرتبه دوم را معرفی کرده، و در پایان نیز برخی مفاهیم آماری و اقتصادی را بیان خواهیم کرد.

۲.۱ توابع محدب

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

تعریف ۱.۲.۱. مجموعه $E \subset \mathbb{R}^n$ را محدب می‌گوییم اگر:

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

تعریف ۲.۲.۱. تابع $f(x) = f(x_1, x_2, \dots, x_n) \in \mathbb{R}$ بر روی یک مجموعه نقاط واقع در یک مجموعه

محدب $E \subset \mathbb{R}^n$ یک تابع محدب نامیده می‌شود 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 گوییم هر گاه $-f$ محدب باشد.

تعریف ۳.۲.۱. در مسأله (۱.۱)، 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 مسأله (۱.۱) است.

^۱Convex

^۲Optimality

^۳Stability

تعریف ۴.۲.۱. تابع $\|\cdot\| : \mathbb{R}^n \rightarrow \mathbb{R}$ را تابع نرم گوییم هرگاه

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

و برای بردار $x \in \mathbb{R}$ داریم:

$$\|x\|_1 = |x_1| + \dots + |x_n|$$

$$\|x\|_2 = \sqrt{x^T x} = \sqrt{x_1^2 + \dots + x_n^2}$$

$$\|x\|_r = \sqrt[r]{x_1^r + \dots + x_n^r}, \quad r \geq 1$$

تعریف ۵.۲.۱. 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.

شرایط لازم و کافی برای مسائل نامقید

قضیه ۶.۲.۱. (شرط لازم مرتبه اول [۱]): فرض کنید X یک مجموعه باز در $\mathbb{R}^n$ باشد. مسأله برنامه‌ریزی

نامقید زیر را در نظر بگیرید:

$$\begin{aligned} & \text{Minimize} && f(x) \\ & \text{subject to} && \\ & && x \in X, \end{aligned} \tag{۱.۱}$$

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

(مشتق تابع f در نقطه $\hat{x}$ برابر صفر).

قضیه ۷.۲.۱. (شرط لازم مرتبه دوم [۱]): 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 f'' در $\hat{x}$ نیمه معین مثبت است.

^۴Hessian

قضیه ۸.۲.۱. ([۱]): 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 ۹.۲.۱. ([۱]): اگر $f(x)$ بر مجموعه محدب E تابعی محدب باشد آنگاه $f(x)$ دارای یک کمینه موضعی^۵ است. اگر چنین کمینه‌ای یافت شود، یک کمینه سراسری^۶ است و بر روی مجموعه محدب به دست آورده می‌شود.

قضیه ۱۰.۲.۱. ([۱]): فرض کنید E مجموعه نقاطی از $\mathbb{R}^n$ باشد که در قیود $g_k(x) \leq 0, x \leq 0$ برای $k = (1, \dots, m)$ صدق کند. اگر $g_k(x)$ ها توابعی محدب باشند آنگاه E یک مجموعه محدب است.

دو قضیه بالا بیان می‌کند که برای یک مسأله برنامه ریزی کلی کمینه سازی، اگر $f(x)$ و همه $g_k(x)$ ها توابعی محدب باشند آنگاه یک کمینه موضعی از $f(x)$ تحت این مجموعه از قیود، یک کمینه سراسری است.

تعریف ۱۱.۲.۱. 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،

۱. نیمه معین مثبت^۷ روی E نامیده می‌شود، اگر

$$\forall V \in \mathbb{R}^n, V \neq 0, X \in E : V^T M(X) V \geq 0.$$

۲. معین مثبت^۸ روی E نامیده می‌شود، اگر

$$\forall V \in \mathbb{R}^n, V \neq 0, X \in E : V^T M(X) V > 0.$$

^۵Local minimum

^۶Global minimum

^۷Positive semi-definite

^۸Positive definite

۳. معین مثبت قوی روی E نامیم، اگر عددی مثبت مانند α وجود داشته باشد به طوری که

$$\forall V \in \mathbb{R}^n, X \in E : V^T M(X) V \geq \alpha \|V\|^2.$$

تعاریف فوق را می توان بر اساس مفهوم مقدار ویژه نیز ارائه داد.

اگر $\gamma(x)$ کوچکترین مقدار ویژه ماتریس $M(X)_{n \times n}$ باشد آنگاه:

۱. $M(X)_{n \times n}$ روی E نیمه معین مثبت اگر و فقط اگر به ازای هر $X \in E, \gamma(x) \geq 0$.

۲. $M(X)_{n \times n}$ روی E معین مثبت اگر و فقط اگر به ازای هر $X \in E, \gamma(x) > 0$.

۳. $M(X)_{n \times n}$ روی E معین مثبت قوی اگر و فقط اگر به ازای هر $X \in E, \gamma(x) \geq \alpha \geq 0$.

شرایط لازم و کافی برای مسائل مقید

قضیه ۱۲.۲.۱. (شرایط لازم کروش-کان-تاکر)^۹ $(K, K.T)$: فرض کنید X یک مجموعه باز در $\mathbb{R}^n$

باشد. مسأله کمینه سازی مقید زیر را در نظر بگیرید:

$$\begin{aligned} & \text{Minimize} \quad f(x_1, x_2, \dots, x_n) \\ & \text{subject to} \\ & \begin{cases} g_k(x_1, x_2, \dots, x_n) \leq 0, & k = 1, \dots, m, \\ h_j(x_1, x_2, \dots, x_n) = 0, & j = 1, \dots, l, \\ x = (x_1, x_2, \dots, x_n)^T \in X. \end{cases} \end{aligned} \quad (2.1)$$

که در آن توابع $f: \mathbb{R}^n \rightarrow \mathbb{R}$ ، $g_k: \mathbb{R}^n \rightarrow \mathbb{R}$ و $h_j: \mathbb{R}^n \rightarrow \mathbb{R}$ 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

^۹Krvsh-Kuhn-Tucker

موضعی برای (۲.۱) باشد، آنگاه اسکالرهایی یکتای $u_k \geq 0$ برای $k \in K$ و $v_j \geq 0$ 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 $K.K.T$ می‌تواند به صورت زیر نوشته شود:

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

همچنین شرایط $K.K.T$ می‌تواند به فرم ماتریسی زیر نوشته شود:

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

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 ۱۳.۲.۱. (شرایط کافی کروش-کان-تاکر $(K.K.T)$): فرض کنید X یک مجموعه باز در $\mathbb{R}^n$

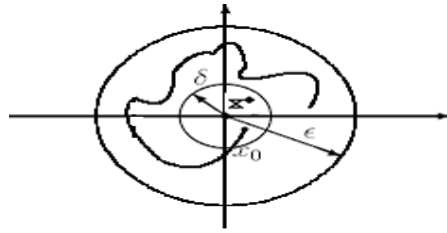
باشد، مسأله (۲.۱) را در نظر بگیرید. فرض کنید $\hat{x}$ یک جواب شدنی باشد و $K = \{k : g_k(\hat{x}) = 0\}$ ، اگر

شرایط $K.K.T$ در $\hat{x}$ برقرار باشد، یعنی اسکالرهایی $u_k \geq 0$ برای $k \in K$ و $v_j \geq 0$ 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 (۲.۱) خواهد بود.



شکل ۱.۱: پایداری لیاپانوف

۳.۱ پایداری

سیستم دینامیکی زیر را در نظر بگیرید:

$$\begin{cases} \dot{x} = f(x(t)), \\ x(t_0) = x_0 \in \mathbb{R}^n \end{cases} \quad (۳.۱)$$

که در آن $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ است. x^* یک نقطه تعادل^{۱۰} برای (۳.۱) نامیده می‌شود اگر $f(x^*) = 0$.

تعریف ۱.۳.۱. فرض کنید که $x(t)$ یک جواب (۳.۱) باشد، نقطه تعادل تنها x^* ، پایدار به مفهوم لیاپانوف

^{۱۱} است با توجه به شکل (۱.۱)، اگر برای هر $x(t_0) = x_0$ و هر $\epsilon > 0$ ، یک $\delta > 0$ وجود داشته باشد

به طوری که اگر $\|x(t) - x_0\| \leq \delta$ آنگاه

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

تعریف ۲.۳.۱. سیستم دینامیکی (۳.۱) همگرای سراسری^{۱۲} به مجموعه

$$\Omega^* = \{x \mid \text{سیستم (۳.۱) جواب دارد}\}$$

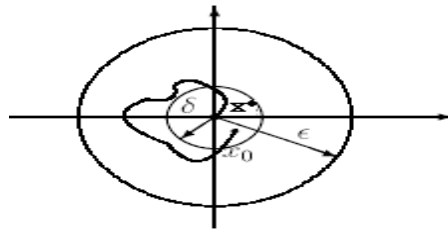
گفته می‌شود اگر هر جواب دلخواه $x(t)$ از سیستم در رابطه زیر صدق کند:

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

^{۱۰}Equilibrium point

^{۱۱}Lyapunov

^{۱۲}Global convergance



شکل ۲.۱: پایداری مجانبی

که در آن

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

تعریف ۳.۳.۱. سیستم دینامیکی در نقطه تعادل یکتا x^* پایدار مجانبی سراسری^{۱۳} نامیده می‌شود 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 ۴.۳.۱. مجموعه $M \subseteq \mathbb{R}^n$ ، یک مجموعه پایدار^{۱۴} نسبت به سیستم (۳.۱) گفته می‌شود اگر

$x(t_0) \in M$ و $t_0 \geq 0$ ، آنگاه به ازای هر $t \geq t_0$ $x(t) \in M$ باشد.

تعریف ۵.۳.۱. نگاشت F در شرط لیب شیتز^{۱۵} صدق می‌کند اگر عدد ثابت L وجود داشته باشد، به طوری

که برای هر دو نقطه $x, y \in \mathbb{R}^n$ داشته باشیم:

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

F ، را لیب شیتز محلی روی $\mathbb{R}^n$ نامیم اگر به ازای هر نقطه از $\mathbb{R}^n$ ، یک همسایگی مانند $D_0 \subset \mathbb{R}^n$ وجود

داشته باشد به طوری که نامساوی بالا برای هر دو نقطه $x, y \in D_0$ برقرار باشد.

^{۱۳}Globally asymptotically stable

^{۱۴}Invariant set

^{۱۵}Lipschitz

تعریف ۶.۳.۱. نگاشت $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ یکنوا^{۱۶} 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^{۱۷} 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 ۷.۳.۱. ([۳]): 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 ۸.۳.۱. ([۴]): در سیستم دینامیکی (۳.۱) فرض می‌کنیم که $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ یک تابع پیوسته، آنگاه برای هر $t_0 > 0$ و $x_0 \in \mathbb{R}^n$ ، یک جواب محلی^{۱۸} $x(t)$ به ازای $t \in [t_0, \tau)$ که $t_0 > \tau$ 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.

تذکر:

همیشه می‌توان نقطه تعادل را به صورت $x^* = 0$ در نظر گرفت. برای هر نقطه تعادل دیگر می‌توان با استفاده از تغییرمتغیر، یک سیستم جدید با نقطه تعادل y^* تعریف کرد. برای این منظور تعریف می‌کنیم:

$$y = x - x^* \Rightarrow \dot{y} = \dot{x} = f(x) \implies f(x) = f(y + x^*) = g(y).$$

^{۱۶}Monotone

^{۱۷}Strictly monotonic

^{۱۸}local solution

در نتیجه نقطه تعادل سیستم جدید $y^* = 0$ ، $\dot{y} = g(y)$ است زیرا

$$g(0) = f(0 + x^*) = f(x^*) = 0.$$

در نتیجه x^* نقطه تعادل برای سیستم $\dot{x} = f(x)$ است اگر و فقط اگر $y = 0$ نقطه تعادل سیستم $\dot{y} = g(y)$ باشد.

۴.۱ تابع انرژی

قضیه ۱.۴.۱. ([۵]): فرض کنید که $x = 0$ یک نقطه تعادل برای سیستم (۳.۱) و تابع $E : \Omega \subseteq \mathbb{R}^n \rightarrow \mathbb{R}$ به طور پیوسته مشتق پذیر باشد اگر

$$1. \quad E(0) = 0,$$

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

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

آنگاه $x = 0$ نقطه پایداری سیستم خواهد بود و $E(x)$ را «تابع لیپانوف» یا «تابع انرژی» برای سیستم (۳.۱) نامیم.

قضیه ۲.۴.۱. ([۵]): فرض کنید که $x = 0$ یک نقطه تعادل برای سیستم (۳.۱) باشد و $E : \Omega \subseteq \mathbb{R}^n \rightarrow \mathbb{R}$ یک تابع به طور پیوسته مشتق پذیر باشد اگر

$$1. \quad E(0) = 0,$$

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

$$3. \quad \text{به ازای هر } x \in \Omega \subseteq \mathbb{R}^n \setminus \{0\}, \quad \dot{E}(X) < 0,$$

آنگاه $x = 0$ پایدار مجانبی خواهد بود.

قضیه ۳.۴.۱ ([۵]): فرض کنید که $x = 0$ یک نقطه تعادل برای سیستم (۳.۱) و تابع $E : \Omega \subseteq \mathbb{R}^n \rightarrow \mathbb{R}$

به طور پیوسته مشتق پذیر باشد اگر

$$1. \quad E(0) = 0,$$

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

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

آنگاه نقطه $x = 0$ پایدار مجانبی سراسری است.

قضیه ۴.۴.۱ ([۵۴]): اصل تغییرناپذیری (ناوردای) لازال^{۱۹} : فرض می کنیم تابع $V : \mathbb{R}^n \rightarrow \mathbb{R}$ به طور

پیوسته مشتق پذیر باشد، همچنین فرض می کنیم

$$1. \quad M \subseteq \mathbb{R}^n \text{ یک مجموعه فشرده و پایدار نسبت به جواب سیستم (۳.۱) باشد،}$$

$$2. \quad \text{به ازای هر } x \in M, \dot{V}(x) \leq 0,$$

$$3. \quad E \text{ مجموعه همه نقاط } M \text{ باشد به طوری که } \dot{V}(x) = 0,$$

$$4. \quad E \text{ بزرگترین مجموعه پایدار در } N \text{ باشد،}$$

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

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

^{۱۹}Lasalle's invariance principle

۵.۱ برنامه‌ریزی مخروط مرتبه دوم

تعریف ۱.۵.۱. مجموعه $C \subset \mathbb{R}^n$ یک مخروط ^{۲۰} نامیده می‌شود اگر:

$$\forall x \in C, \quad \lambda \geq 0 \Rightarrow \lambda x \in C.$$

تعریف ۲.۵.۱. یک مخروط، $C \subset \mathbb{R}^n$ مخروط مناسب نامیده می‌شود اگر:

۱. C محدب و بسته باشد،

۲. C بسته،

۳. C توپر باشد (درون آن ناتهی باشد)،

۴. C نوک تیز باشد ($x \in C, -x \in C \Rightarrow x = 0$).

فرض کنید که $\mathcal{K}$ یک مخروط مناسب است، در این صورت نامساوی مخروط به صورت زیر تعریف می‌شود:

$$x \preceq_{\mathcal{K}} y \Leftrightarrow x - y \in \mathcal{K}.$$

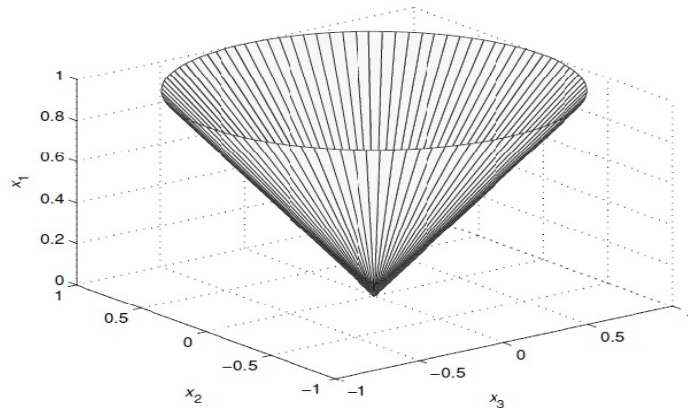
توجه کنید که نامساوی مخروط $\preceq_{\mathcal{K}}$ به طور کلی نسبت به مخروط مناسب $\mathcal{K}$ معنی می‌دهد.

به منظور تعریف کردن مخروط مرتبه دوم ^{۲۱} در $\mathbb{R}^n$ ($n \geq 1$)، افزایی از یک بردار $x \in \mathbb{R}^n$ را معرفی کرده که توسط $x = \begin{bmatrix} x_0 \\ x_1 \end{bmatrix}$ داده می‌شود که در آن $x_0 \in \mathbb{R}$ و $x_1 \in \mathbb{R}^{n-1}$. بنابراین مخروط مرتبه دوم به صورت مجموعه‌ای از $x \in \mathbb{R}^n$ به طوری که $|x_0| \geq \|x_1\|$ تعریف می‌شود، که این مخروط نیز شرایط مخروط مناسب را دارد. پس نامساوی ذکر شده همان نامساوی $\succeq_{\mathcal{K}}$ است. همچنین مخروط مرتبه دوم C_q را به صورت زیر تعریف می‌کنیم:

$$C_q = \{x \in \mathbb{R}^n \mid x_1^2 \geq \|x_{2:n}\|^2, x_1 \geq 0\},$$

^{۲۰} Cone

^{۲۱} Second Order Cone



شکل ۳.۱: مخروط مرتبه دوم در $\mathbb{R}^3$

که در آن $x_{2:n} = (x_2, \dots, x_n)$ است، به مخروط مرتبه دوم، مخروط لورنتس (مخروط بستنی) نیز گفته می‌شود. مخروط مرتبه دوم در فضای سه بعدی در شکل (۳.۱) نشان داده شده است. اکنون حالت خاصی از برنامه‌ریزی مخروط مرتبه دوم را به صورت زیر تعریف می‌کنیم که برای حل بعضی مسائل بهینه‌سازی استفاده می‌شود:

$$\begin{aligned} & \text{Minimize} && f^T x \\ & \text{subject to} && \\ & && \|A_i x + b_i\| \leq c_i^T x + d_i, \quad i = 1, \dots, N \end{aligned} \quad (4.1)$$

که در آن $x \in \mathbb{R}^n$ متغیر بهینه‌سازی و $f \in \mathbb{R}^n$ ، $A_i \in \mathbb{R}^{(n_i-1) \times n}$ ، $b_i \in \mathbb{R}^{n_i-1}$ ، $c_i \in \mathbb{R}^n$ و $d_i \in \mathbb{R}$ پارامترهای مسأله هستند. نرمی که در قیود به کار رفته است، نرم استاندارد اقلیدسی است، یعنی

$\|u\|_2 = (u^T u)^{\frac{1}{2}}$. یک مخروط مرتبه دوم استاندارد از بعد k به صورت زیر تعریف می‌شود:

$$c_k = \left\{ \begin{bmatrix} u \\ t \end{bmatrix} \mid u \in \mathbb{R}^{k-1}, t \in \mathbb{R}, \|u\| \leq t \right\}.$$

برای $k = 1$ مخروط مرتبه دوم بالا به صورت زیر

$$c_1 = \{ t \mid t \in \mathbb{R}, t \geq 0 \}.$$

مجموعه نقاطی که در یک قید مخروط مرتبه دوم صدق می کنند، نقش معکوس مخروط مرتبه دوم تحت یک نگاشت آفین هستند:

$$\|A_i x + b_i\| \leq c_i^T x + d_i \Leftrightarrow \begin{bmatrix} A_i \\ c_i^T \end{bmatrix} x + \begin{bmatrix} b_i \\ d_i \end{bmatrix} \in C_{n_i}.$$

که منظور از C_{n_i} مخروط مرتبه دوم از بعد n_i است که مجموعه‌ای محدب است. بنابراین برنامه‌ریزی مخروط مرتبه دوم (۴.۱) یک برنامه‌ریزی محدب است، برای مثال هنگامی که $n_i = 1$ ، برنامه‌ریزی مخروط مرتبه دوم به برنامه‌ریزی خطی^{۲۲} تبدیل می‌شود:

$$\begin{aligned} & \text{Minimize} && f^T x \\ & \text{subject to} && \\ & && c_i^T x + d_i \geq 0, \quad i = 1, \dots, N. \end{aligned} \quad (5.1)$$

حالت خاص و جالب دیگری که می‌شود به دست آورد زمانی که $c_i = 0$ است، که در این حالت i امین قید مخروط مرتبه دوم به $\|A_i x + b_i\| \leq d_i$ کاهش می‌یابد، که آن معادل (با فرض $d_i \geq 0$) قید درجه دو (محدب) $\|A_i x + b_i\|^2 \leq d_i^2$ است. بنابراین وقتی که c_i ها صفر هستند، برنامه‌ریزی مخروط مرتبه دوم به یک برنامه‌ریزی خطی با قیود درجه دوم^{۲۳} منجر می‌شود. همچنین می‌بینیم که برنامه‌ریزی درجه دوم محدب و برنامه‌ریزی درجه دوم مقید شده و تعداد زیاد دیگری از مسائل بهینه سازی را می‌توان به صورت برنامه‌ریزی مخروط مرتبه دوم فرمول بندی کرد.

۶.۱ مفاهیم آماری

تعریف ۱.۶.۱. تابع $f(x) = \begin{cases} > 0 & x \in X(s) \\ = 0 & x \notin X(s) \end{cases}$ را تابع جرم احتمال متغیر تصادفی گسسته X گویند اگر $\sum_{x \in X(s)} f(x) = 1$ و برای پیشامد A داریم، $P(A) = \sum_{x \in A} f(x)$.

^{۲۲}Linear programming(LP)

^{۲۳}Quadratic Constrained Linear Programming

تعریف ۲.۶.۱. امید ریاضی یا میانگین متغیر تصادفی گسسته X ، که با $E(X)$ یا μ_x نمایش داده شده، به صورت زیر تعریف می شود:

$$E(X) = \mu_x = \sum_x xP(X = x).$$

تعریف ۳.۶.۱. واریانس متغیر تصادفی X ، که با $\text{Var}(X)$ یا σ^2 نمایش داده می شود، به صورت زیر تعریف می شود:

$$\sigma^2 = \text{Var}(X) = E(X - \mu)^2,$$

همچنین σ را انحراف معیار متغیر تصادفی X گوئیم.

تعریف ۴.۶.۱. کوواریانس دو متغیر تصادفی X و Y عبارت است از:

$$\text{Cov}(X, Y) = E[(X - \mu_x)(Y - \mu_y)] = E(XY) - E(X)E(Y).$$

تعریف ۵.۶.۱. ضریب همبستگی^{۲۴} دو متغیر تصادفی X و Y عبارت است از:

$$\rho = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var}(X) \text{Var}(Y)}}.$$

تعریف ۶.۶.۱. تابع $f(x) = \begin{cases} > 0 & x \in X(s) \\ = 0 & x \notin X(s) \end{cases}$ را تابع چگالی احتمال متغیر تصادفی پیوسته X

گویند اگر $\int_{-\infty}^{\infty} f(x)dx = 1$ و برای پیشامد B داریم، $P(B) = \int_B f(x)dx$.

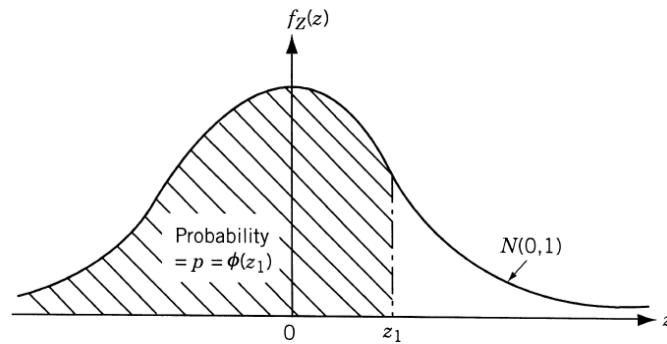
تعریف ۷.۶.۱. برای هر متغیر تصادفی X (اعم از گسسته یا پیوسته) تابع $[0, 1] \rightarrow \mathbb{R}$ ، F_X ، تابع توزیع تجمعی X است که به صورت $F_X(x) = P(X \leq x)$ تعریف می شود.

توزیع نرمال

فرض کنید تابع چگالی متغیر پیوسته X به صورت

$$f_X(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad -\infty < x < +\infty,$$

^{۲۴}Correlation Coefficient



شکل ۴.۱: تابع چگالی نرمال استاندارد

باشد، گوئیم X توزیع نرمال با میانگین μ و واریانس σ^2 دارد و می‌نویسیم $X \sim N(\mu, \sigma^2)$. همچنین اگر $\sigma^2 = 1$, $\mu = 0$ باشد، گوئیم Z توزیع نرمال استاندارد دارد یا $Z \sim N(0, 1)$. تابع توزیع تجمعی نرمال استاندارد را با

$$\Phi(x) = P(Z \leq x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{z^2}{2}} dz,$$

نشان می‌دهند و مقادیر تقریبی آن را از جدول نرمال استاندارد حساب می‌کنند. با داشتن این جدول و با توجه به اینکه

$$X \sim N(\mu, \sigma^2) \iff z = \frac{X - \mu}{\sigma} \sim N(0, 1),$$

می‌توان مقادیر تقریبی احتمالات را برای هر توزیع نرمال یافت. طبق شکل (۴.۱) داریم:

$$\Phi(z_1) = p, \quad z_1 = \Phi^{-1}(p),$$

و همچنین مقادیر z_1 مربوط به $p < 0.5$ را می‌توان به صورت زیر به دست آورد:

$$z_1 = \Phi^{-1}(p) = -\Phi^{-1}(1 - p)$$

۷.۱ مفاهیم اقتصادی

با توجه به وجود فضای نامطمئن^{۲۵}، مهمترین مفاهیم در تصمیم‌گیری برای سرمایه‌گذاری، ریسک و بازده^{۲۶} هستند.

۱.۷.۱ بازده

بازده در فرایند سرمایه‌گذاری، نیروی محرکی است که ایجاد انگیزه کرده و به عبارت بهتر پاداش سرمایه‌گذاران محسوب می‌گردد، چرا که تمام هدف سرمایه‌گذاری کسب بازده مطلوب است. از طرف دیگر تعیین تفاوت میان بازده تحقق یافته^{۲۷} و بازده مورد انتظار^{۲۸} از اهمیت بالایی برخوردار است. بازده تحقق یافته، بازدهی است که مربوط به گذشته بوده، حال آنکه بازده مورد انتظار، بازده تخمینی برای یک دارایی^{۲۹} در دوره‌ای از آینده است که این بازده با عدم اطمینان^{۳۰} همراه بوده و احتمال برآورده نشدن آن نیز وجود دارد. در مورد دارایی‌ها و مخصوصاً سهام، بازده از دو جزء تشکیل شده است:

۱. بازده ناشی از تغییر قیمت^{۳۱}، این جزء از اختلاف بین قیمت خرید دارایی و فروش دارایی حاصل می‌شود.

۲. بازده ناشی از دریافت سود^{۳۲}، این بازده ناشی از صورت جریان‌ات نقدی دریافتی به دلیل تملک

دارایی در طول سرمایه‌گذاری است (همانند سود تقسیمی).

^{۲۵}Uncertainty Space

^{۲۶}Return

^{۲۷}Realized Return

^{۲۸}Expected Return

^{۲۹}Asset

^{۳۰}Uncertainty

^{۳۱}Capital Gain

^{۳۲}Yeild

بازده مورد انتظار هر سهم

با معلوم بودن توزیع احتمال بازده بالقوه هر سهم داریم

$$E(R_i) = \sum_{k=1}^m (p_k) PR_k,$$

که در آن، p_k احتمال وقوع هر نرخ بازده بالقوه و PR_k بازده بالقوه یک سهم است.

سبد سرمایه

معنی ساده واژه سبد سرمایه (پورتفوی)^{۳۳}، سبد سرمایه‌گذاری به‌طور عام و سبد سهام به‌طور خاص، عبارت از ترکیب دارایی‌های سرمایه‌گذاری شده توسط یک سرمایه‌گذار اعم از فرد یا مؤسسه است. به لحاظ فنی، یک سبد سرمایه، مجموعه کامل دارایی‌های حقیقی و مالی سرمایه‌گذار را در بر می‌گیرد.

بازده مورد انتظار سبد سرمایه

در صورتی که سبد سرمایه با N دارایی داشته باشیم که وزن دارایی i ام در آن x_i باشد، بازده مورد انتظار سبد سرمایه به‌صورت زیر محاسبه می‌شود:

$$E(R_p) = \sum_{i=1}^N x_i E(R_i). \quad (۶.۱)$$

فروش استقراضی

فروش استقراضی^{۳۴} فروش اوراق بهادار به‌وسیله کسی که چنین اوراقی را ندارد. بنابراین باید آنها را قرض بگیرد و به خریدار تحویل دهد. فروش به این فرم باعث می‌شود که فروشنده از کاهش قیمت اوراق در آینده سود ببرد. فروشنده باید در زمان آینده اوراق را خریداری و بدهی خود را پرداخت نماید.

^{۳۳}Portfolio

^{۳۴}Short-Selling

۲.۷.۱ ریسک و مفهوم آن

به طور ساده می‌توان ریسک را عدم اطمینان به نتایج آتی تعریف کرد. به عبارت دیگر ریسک پدیده‌ای مربوط به آینده است که نمی‌توان آن را به‌طور دقیق پیش بینی نمود، زیرا با عدم اطمینان همراه است. هر چه این عدم اطمینان بیشتر باشد، ریسک هم بیشتر است. در بازار های مالی به همراه هر فرصتی ریسک هم وجود دارد و اصولاً نمی‌توان کلیه ریسکها را از بین برد زیرا کلیه فرصتها هم از بین خواهند رفت. سرمایه گذار موفق کسی است که سطح قابل قبولی از ریسک را بپذیرد.

تعاریف مختلفی از ریسک بیان شده از جمله، ریسک را معادل نوسان بازدهی دارایی در نظر می‌گیرند. استفاده از واریانس بازده (نوسان حول میانگین) به عنوان معیار ریسک بر مبنای همین نظر است. همچنین ریسک عبارت است از میزان اختلاف بازده واقعی یک سرمایه‌گذاری از بازده مورد انتظار آن، معمولاً در اقتصاد فرض بر اینکه سرمایه‌گذاران منطقی عمل می‌کنند و اطمینان را بر عدم اطمینان ترجیح می‌دهند و سرمایه‌گذاران ریسک‌گریزند^{۳۵}.

ریسک سبد سرمایه

ریسک سبد سهام، تابعی از سهام منفرد و کواریانس بین بازده‌های سهام منفرد است که به صورت زیر بیان می‌شود:

$$\text{Var}(R_p) = \sum_{i=1}^N x_i \text{Var}(R_i) + \sum_{i=1}^N \sum_{j=1, j \neq i}^N x_i x_j \text{Cov}(R_i, R_j). \quad (7.1)$$

که در آن $\text{Cov}(R_i, R_j) = \rho_{ij} \sigma_i \sigma_j$ کواریانس بین بازده‌های سهام‌های i و j است و ρ_{ij} ضریب همبستگی بازده‌های i و j است.

^{۳۵}Risk-averse Investor

اندازه‌گیری ریسک

تاکنون معیارهای مختلفی برای تعیین ریسک معرفی شده، شاخص‌های اندازه‌گیری ریسک اولین بار از طریق مطالعات شاخص‌های پراکندگی آماری محاسبه گردیدند و از آن به بعد روش‌های جدیدتری از جمله ریسک نامطلوب^{۳۶}، استفاده از دیرش^{۳۷} برای محاسبه حساسیت ارزش اوراق قرضه و در نهایت ارزش در معرض ریسک^{۳۸} معرفی گردیدند که همگی از روش‌های آماری استفاده می‌کنند.

در سال ۱۹۵۲، هری مارکوویتز با ارائه مدلی کمی به اندازه‌گیری ریسک پرداخت و با معرفی مدلی مبتنی بر ریسک و بازده و ارائه مجموعه کارا^{۳۹} برای اولین بار ریسک را در کنار بازده به‌عنوان متغیری دیگر جهت انتخاب دارایی برای سرمایه‌گذاری قرار داد. وی انحراف معیار را به‌عنوان شاخص پراکندگی، معیار عددی ریسک خواند. شاگرد او ویلیام شارپ، شاخص بتا (ضریب حساسیت) را برای تغییرات نسبی ارزش یک سهم در قبال تغییر ارزش بازار با معرفی خط مشخصات ارائه کرد و با معرفی مدل قیمت‌گذاری دارایی‌های سرمایه‌ای، مدیریت علمی سبد سرمایه را پایه‌گذاری کرد. مک کالی^{۴۰} معیار دیرش را به‌عنوان ملاک اندازه‌گیری ریسک اوراق بهادار با درآمد ثابت معرفی کرد، در ادامه کارش معیار تحذب را به‌عنوان شاخصی دقیق‌تر معرفی کرد. در سال ۱۹۹۶ وترستون^{۴۱} مدیر موسسه (J.P.Morgan) مدل ارزش در معرض ریسک را معرفی کرد. می‌توان طبقه‌بندی زیر را در مورد معیارهای اندازه‌گیری ریسک ارائه داد:

۱. حساسیت: تغییر یک متغیر وابسته بر اثر تغییر یک متغیر مستقل، مثل تغییرات قیمت در قبال تغییر یک واحد نرخ سود (دیرش، تحذب و بتا).

۲. نوسان: عبارت است از نوسان یک متغیر در اطراف میانگین و یا یک پارامتر تصادفی دیگر مثل

^{۳۶}Downside Risk

^{۳۷}Duration

^{۳۸}Value at Risk

^{۳۹}Efficient Frontier Line

^{۴۰}McCauley

^{۴۱}Westerstone

(واریانس و انحراف معیار).

۳. معیارهای ریسک نامطلوب: این معیارها که برعکس معیارهای نوسان، تنها بر بخش مخرب ریسک تمرکز دارند و در حقیقت نوسانات زیر سطح میانگین و یا متغیر هدف را مورد محاسبه قرار می‌دهند (نیم واریانس^{۴۲}، نیم انحراف معیار و ارزش در معرض ریسک).

تعریف ۱.۷.۱. نیم واریانس: شاخصی برای محاسبه پراکندگی صفت متغیر و برای انحرافات نامطلوب به کار برده می‌شود. به عبارت دیگر اگر ریسک را میزان ضرر تعریف کنیم، آنگاه تغییرات مطلوب (افزایش نرخ بازدهی دارایی مالی) به عنوان ریسک محسوب نمی‌شود و فقط آن دسته از مشاهداتی که کمتر از میانگین نرخ بازدهی هستند، به عنوان ریسک به حساب می‌آید:

$$SV = \frac{1}{K} \sum_{i=1}^K \{\max(0, E - r_T)\}^2,$$

که در آن، r_T بازدهی دارایی در طول دوره، K تعداد مشاهدات، E امید ریاضی نرخ بازدهی است.

نظریه مدرن سبد سرمایه MPT براساس رابطه بازدهی و ریسک محاسبه شده از طریق واریانس و انحراف معیار بازدهی تبیین می‌شود. نظریه فرامدرن سبد سرمایه $PMPT$ ^{۴۳} براساس رابطه ریسک نامطلوب به تبیین رفتار سرمایه‌گذار و معیار انتخاب سبد سرمایه بهینه می‌پردازد (برای توضیحات بیشتر، مراجع [۶، ۷] را ببینید).

^{۴۲}Semi Variance

^{۴۳}Post-Modern Portfolio Theory

فصل ۲

مقدمه‌ای بر شبکه‌های عصبی

۱.۲ مقدمه

در این فصل به بیان مقدمه‌ای از مدل‌های شبکه‌های عصبی می‌پردازیم. در ابتدا ساختار شبکه عصبی و مدل ریاضی یک سلول عصبی را بیان می‌کنیم و در ادامه به بیان تاریخچه‌ای از شبکه‌های عصبی در حل مسائل بهینه سازی می‌پردازیم. در انتها به بیان مدل‌های ارائه شده پیشین برای حل مسائل بهینه سازی خواهیم پرداخت.

کار بر روی شبکه‌های عصبی مصنوعی، با الهام از عملکرد مغز انسان از آنجا شروع شد که دانشمندان دریافتند مغز بشر به‌گونه‌ای کاملاً متفاوت از کامپیوترهای دیجیتالی مرسوم محاسبات را انجام می‌دهد. مغز یک سیستم پردازش داده بسیار پیچیده است که از واحدهای ساختاری به نام سلول عصبی یا نرون^۱ تشکیل شده است. شبکه عصبی سیستمی پویا و غیرخطی است که از تعداد زیادی واحد پردازش (نرون ها) و اتصالات بین این واحدهای پردازش تشکیل می‌شود.

یک شبکه عصبی بر خلاف کامپیوتر که نیازمند دستورهای کاملاً صریح و مشخص است، به مدل‌های ریاضی محض نیازی ندارد، بلکه مانند انسان تجربه کسب کرده و سپس نتیجه این تجربیات را تعمیم می‌دهد. امروزه این شبکه‌ها تبدیل به ابزاری قدرتمند و عمومی شده‌اند که برای حل مسائل از قبیل تخمین (تقریب)، تشخیص الگو^۲، رباتیک، کنترل، و به‌طور کلی در هر جا که نیاز به یادگیری نگاشت خطی و یا غیرخطی باشد، به کار می‌روند.[۸]

برای حل هر مسأله‌ای، شبکه‌های عصبی سه مرحله را طی می‌کنند: آموزش^۳، تعمیم^۴ و اجرا^۵. آموزش، فرایندی است که طی آن شبکه می‌آموزد تا الگوی موجود در ورودی‌ها را (که به‌صورت مجموعه داده‌های آموزشی است) بشناسد. برای این منظور هر شبکه عصبی از مجموعه‌ای از قوانین یادگیری

^۱Neuron

^۲Pattern recognition

^۳Traning

^۴generalization

^۵Operation

که نحوه یادگیری را تعریف می‌کنند استفاده می‌کند. تعمیم، توانایی شبکه برای ارائه جواب قابل قبول در قبال ورودی‌هایی است که در مجموعه آموزشی نبوده‌اند. یعنی پس از آنکه مثال‌های اولیه به شبکه آموزش داده شد، شبکه می‌تواند در قبال یک ورودی آموزش داده نشده قرار گیرد و یک خروجی مناسب را ارائه نماید. این خروجی براساس مکانیسم تعمیم که همانا چیزی جز فرایند درونیابی نیست به‌دست می‌آید. استفاده از شبکه برای انجام عملکردی که به آن منظور طراحی شده است را اجرا می‌گویند.

در اثر آموزش دیدن شبکه، وزن‌های داخلی^۶ (که بر روی ورودی‌های هر سلول اعمال می‌شود) تغییر می‌کند و به وضعیت مناسب می‌رسند. یکی از نقاط ضعف شبکه‌های عصبی این است که نتایج آموزش، یعنی وزن‌های داخلی، هیچگونه تصویر روشنی از اعتبار جواب‌های مسأله به‌دست نمی‌دهد. این وزن‌ها کاملاً قابل درک نیستند. با این وجود جواب‌های تولید شده توسط شبکه، اغلب صحیح است و با شرایط کمی حاکم بر محیط سازگاری دارد. این صحت جواب‌ها و صدق شرایط کمی حاکم بر محیط، گاهی مهم‌تر از توضیح‌پذیر بودن آن است. تاکنون اسامی متفاوتی برای شبکه‌های عصبی به‌کار رفته است. از جمله آنها می‌توان به مدل‌های جعبه سیاه^۷، مدل‌های تجربی^۸، تقریب زن‌های عام^۹ و مدل‌های موازی^{۱۰} اشاره نمود.[۹]

۲.۲ ساختار شبکه‌های عصبی مصنوعی

شبکه‌های عصبی مصنوعی از یک سری واحدهای ساختمانی اولیه تشکیل می‌شوند. هر واحد دارای چندین ورودی است که با هم ترکیب شده و پس از انجام عملیات پردازش یک خروجی را به‌دست می‌دهند. این واحدهای اولیه مانند شکل (۱.۲) به هم متصل هستند به‌طوری‌که خروجی هر واحد به عنوان ورودی واحد های دیگر مورد استفاده قرار می‌گیرد. واحدهای ساختمانی را سلول عصبی، نرون یا گره^{۱۰} نیز می‌نامند.

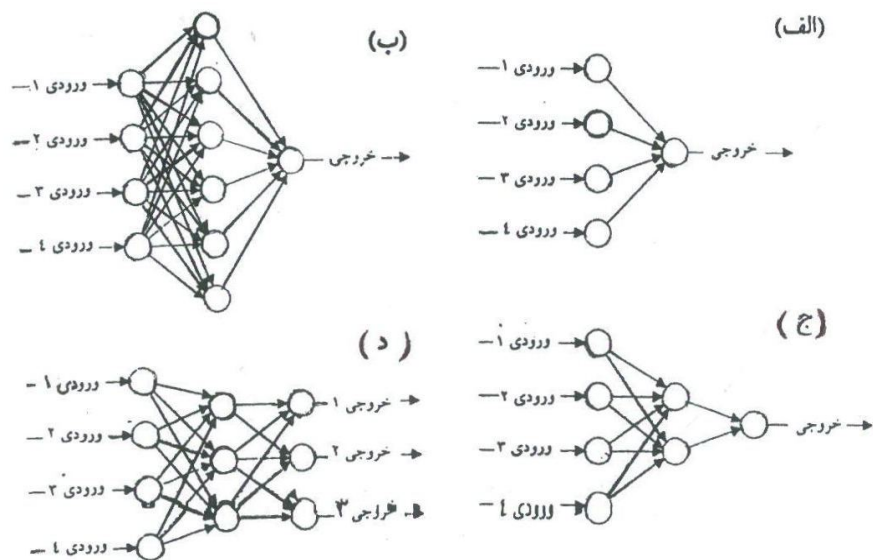
^۶Internal Weight

^۷Black-Box models

^۸Universal Approximators

^۹Parallel models

^{۱۰}Node



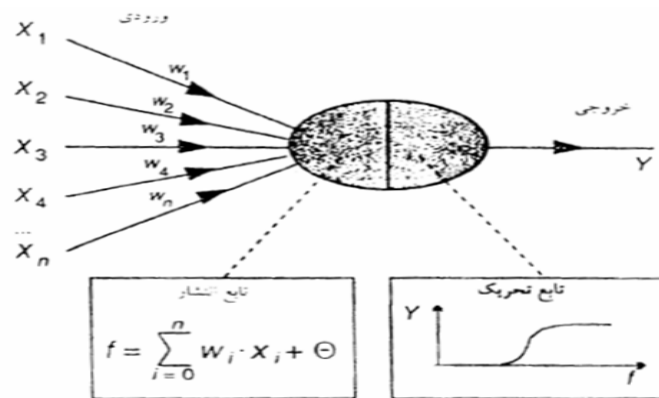
شکل ۱.۲: شبکه‌های عصبی مصنوعی

در شکل (۱.۲)، (الف)-یک شبکه بسیار ساده شامل چهار ورودی و یک خروجی، نتیجه آموزش این شبکه تکنیک رگرسیون لجستیک^{۱۱} در آمار را نشان می‌دهد. (ب)-شبکه دارای دو لایه میانی است که قدرت بیشتری در تشخیص الگو دارد. (ج)-افزایش تعداد لایه پنهان شبکه را قوی‌تر می‌کند ولی خطر بیش‌برازشی^{۱۲} در آن وجود دارد. (د)-شبکه عصبی با چندین خروجی.

شبکه‌های عصبی اولیه تنها دو لایه (ورودی و خروجی) داشتند و به‌همین دلیل تنها در مسائل خطی مفید بودند. ساختارهای توسعه یافته امروزی شامل سه و یا تعداد بیشتری لایه هستند. رایج‌ترین ساختار شبکه، سه لایه‌ای است (لایه ورودی، یک لایه پنهان و لایه خروجی) و عموماً قادر به حل اکثر مسائل پیچیده است.

^{۱۱}Logistic regression

^{۱۲}Over fitting



شکل ۲.۲: مدل ساده یک سلول عصبی

۳.۲ مدل ریاضی یک سلول عصبی

شکل (۲.۲) یک مدل سلول عصبی را نشان می‌دهد. بدنه این سلول از دو بخش تشکیل شده است اولین بخش را تابع ترکیب^{۱۳} می‌گویند. وظیفه تابع ترکیب این است که تمام ورودی‌ها را ترکیب و یک عدد تولید کند. مطابق شکل (۲.۲) هر ورودی دارای وزن مختص به خود است. ورودی‌ها در وزن‌های مربوطه ضرب و سپس با هم جمع می‌شوند. مجموع حاصل را مجموع موزون^{۱۴} گویند که رایج‌ترین تابع ترکیب است.

بخش دوم سلول عصبی، تابع انتقال^{۱۵} نام دارد که تابع ترکیب را به خروجی سلول تبدیل و سپس انتقال می‌دهد. تابع انتقال را تابع تحریک^{۱۶} نیز می‌نامند.

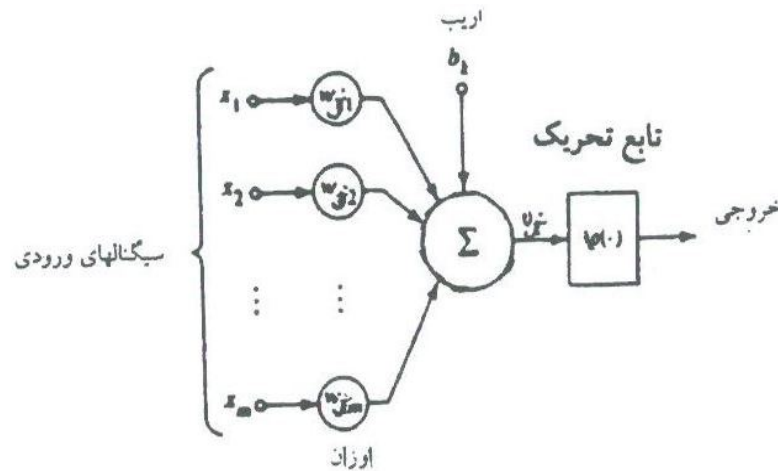
رایج‌ترین توابع تحریک، بر پایه مدل‌های بیولوژیک استوار گردیده است (برای توضیحات بیشتر مراجع [۹، ۱۱، ۱۰] را ببینید). اگر ورودی سلول i ام واقع در لایه ورودی را x_i در نظر بگیریم، این ورودی برای اتصال به سلول j ام لایه بعد در وزن w_{ji} ضرب می‌شود. حرف j در w_{ji} نشان دهنده شماره سلول در لایه

^{۱۳}Combination function

^{۱۴}Weighted sum

^{۱۵}Transfer function

^{۱۶}Activation function



شکل ۳.۲: مدل بلوکی سلول عصبی

بعد و حرف دوم اندیس معرف شماره سلول لایه قبلی است. شکل (۳.۲) مدل بلوکی یک سلول عصبی را نشان می‌دهد سلول مذکور دارای یک ورودی اضافی است که به آن اریبی^{۱۷} می‌گویند و با b_j نشان داده می‌شود. نقش اریبی افزایش یا کاهش مجموع موزون است.

به بیان ریاضی، سلول j بوسیله دو رابطه زیر تعریف می‌شود:

$$u_j = \sum_{i=1}^m w_{ji} x_i,$$

$$y_j = \phi(u_j + b_j),$$

که $x_1, x_2, \dots, x_m$ داده‌های ورودی، $w_{j1}, w_{j2}, \dots$ وزن‌های اتصال ورودی‌های ۱، ۲، ... به سلول j ، u_j ترکیب خطی ورودی‌ها، b_j اریبی، ϕ تابع تحریک و y_j خروجی سلول است. می‌توان مجموع u_j و b_j را به صورت زیر نمایش داد:

$$v_j = (u_j + b_j).$$

تابع تحریک یا تابع انتقال که با ϕ نشان داده می‌شود بر روی مجموع وزن‌ها (شامل اریبی) عمل می‌کند.

^{۱۷}Bias

این تابع می‌تواند خطی یا غیرخطی باشد و بر اساس نیاز خاص حل یک مسئله انتخاب می‌شود. انواع مختلفی از توابع تحریک وجود دارند که مهم‌ترین آنها بشرح زیر می‌باشند:

۱. تابع حد آستانه^{۱۸} : این تابع در شکل (۴.۲) قسمت (a) نشان داده شده و به‌صورت زیر بیان می‌شود

$$\phi(\nu) = \begin{cases} 1 & \nu \geq 0, \\ 0 & \nu < 0. \end{cases}$$

۲. تابع قسمتی خطی^{۱۹} این تابع که در شکل (۴.۲) قسمت (b) نشان داده شده است به‌صورت زیر تعریف می‌شود:

$$\phi(\nu) = \begin{cases} 1 & \nu \geq 0.5, \\ \nu & -0.5 < \nu < 0.5, \\ 0 & \nu < -0.5. \end{cases}$$

۳. تابع سیگموئید^{۲۰} نام این تابع به علت شکل S مانند آن است و پرکاربردترین تابع تحریک محسوب می‌شود. این تابع بین رفتار خطی و غیرخطی، به آرامی تغییر می‌کند. مثالی از تابع سیگموئید تابع لجستیک^{۲۱} است که به‌صورت زیر بیان می‌شود:

$$\phi(\nu) = \frac{1}{1 + \exp(-a\nu)},$$

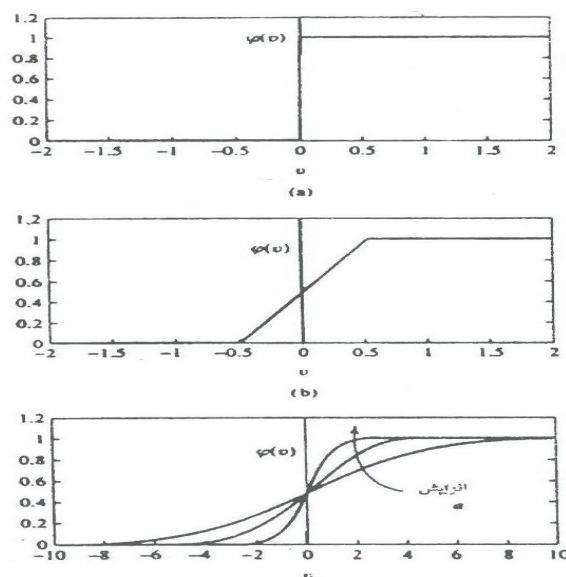
که در آن a پارامتر شیب تابع سیگموئید است، با تغییر پارامتر a ، توابع مختلفی با شیب‌های متفاوت مانند آنچه که در شکل (۴.۲) قسمت (c) نشان داده شده بدست می‌آید.

^{۱۸}Threshold function

^{۱۹}Piecewise-Linear function

^{۲۰}Sigmoid function

^{۲۱}Logistic function



شکل ۴.۲: انواع تابع تحریک: الف) تابع حد آستانه، ب) تابع قسمتی خطی، ج) تابع سیگموئید با شیب‌های متغیر

۴.۲ تاریخچه کاربرد شبکه‌های عصبی مصنوعی در حل مسائل بهینه‌سازی

مسائل بهینه‌سازی در طیف وسیعی از کاربردهای علوم و مهندسی تحلیل پردازش تصویر^{۲۲}، شناسایی سیستم‌ها^{۲۳}، حل معادلات، طراحی فیلتر^{۲۴}، تقریب توابع^{۲۵}، آنالیز رگرسیون^{۲۶} و... پدیدار می‌شود. در بسیاری از کاربردهای علوم و مهندسی جواب زمان واقعی^{۲۷} مسائل بهینه‌سازی بسیار مورد نیاز می‌شود. الگوریتم‌های قدیمی برای کامپیوترهای دیجیتال چندان موثر نیستند زیرا زمان مورد نیاز برای بدست آوردن جواب وابستگی زیادی به بُعد^{۲۸} و ساختار^{۲۹} مسأله دارد، این یک رهیافت ممکن و بسیار امیدوارکننده

^{۲۲}Signal processing

^{۲۳}System identification

^{۲۴}Filter design

^{۲۵}Function approximation

^{۲۶}Regression analysis

^{۲۷}Real time

^{۲۸}Dimension

^{۲۹}Structure

کاربرد شبکه‌های عصبی مصنوعی است.

به واسطه ذات شبکه‌های عصبی که می‌توانند حجم عظیمی از کارهای موازی را با هم انجام دهند رهیافت شبکه عصبی می‌تواند مسائل بهینه‌سازی در زمان واقعی را با سرعت بیشتری نسبت به سایر الگوریتم‌های بهینه‌سازی موجود حل کند. هاپفیلد^{۳۰} و تنک^{۳۱} در [۱۲] ابتدا شبکه عصبی را برای حل مسائل برنامه‌ریزی خطی پیشنهاد دادند، کار بنیادین آنها الهام بخش تحقیقات زیادی برای ارائه شبکه‌های عصبی دیگر برای حل مسائل برنامه‌ریزی خطی و غیرخطی شد. بدین ترتیب شبکه‌های عصبی بهینه‌سازی زیادی گسترش یافتند. به عنوان مثال کندی^{۳۲} و چاو^{۳۳} در [۱۳] یک شبکه عصبی برای حل مسائل برنامه‌ریزی غیرخطی پیشنهاد کردند که از روش گرادیان و روش تابع جریمه استفاده می‌نمودند و نقاط تعادلشان متناظر با تقریبی از جواب‌های بهینه بود. با استفاده از روش گرادیان و تصویر، بازردوم^{۳۴} و پاتیسن^{۳۵} [۱۴] یک شبکه عصبی برای حل مسائل بهینه‌سازی درجه دوم تنها با متغیرهای کراندار معرفی کردند. این شبکه عصبی تحت شرایط خاصی به تنها جواب بهینه همگرا بود.

بر پایه روش‌های گرادیان و تصویر، ایکسیا^{۳۶} و همکارانش [۲۰]-[۱۵] چندین شبکه عصبی برای حل مسائل برنامه‌ریزی درجه دوم و خطی با جواب‌های بهینه غیریکتا معرفی کردند و همگرایی سراسری آنها را به جواب‌های دقیق ثابت نمودند.

بر پایه بعضی از ملاحظات در اجرای سخت افزاری، یک شبکه عصبی با عملکرد محاسباتی خوب باید سه ویژگی اساسی داشته باشد:

۱. همگرایی شبکه عصبی به ازای هر وضعیت دلخواه داده شده تضمین شود.

^{۳۰}Hopfield

^{۳۱}Tank

^{۳۲}Kennedy

^{۳۳}Chua

^{۳۴}Bouzerdoum

^{۳۵}Pattison

^{۳۶}Xia

۲. در طراحی شبکه ترجیحاً هیچ پارامتر مجهولی وجود نداشته باشد.

۳. نقاط تعادل شبکه متناظر با جواب‌های دقیق یا تقریبی باشد.

از دیدگاه ریاضی این خصوصیات مربوط به تکنیک‌های بهینه‌سازی هستند که برای طراحی مدل‌های شبکه‌های عصبی بهینه‌سازی به کار می‌روند.

در مجموع برای فرمول‌بندی یک مسئله بهینه‌سازی برحسب یک شبکه عصبی دو نوع روش وجود دارد. یک رهیافت که عموماً در ایجاد یک شبکه عصبی بهینه‌سازی استفاده می‌شود این است که در ابتدا مسئله بهینه‌سازی را به یک مسئله بهینه‌سازی نامقید تبدیل کرده و سپس یک شبکه عصبی طراحی می‌کنند که مسئله نامقید را با استفاده از روش گرادیان حل کند. روش گرادیان یک مزیت برای آن دسته از شبکه‌های عصبی است که مستقیماً از مشتق تابع انرژی استفاده می‌کنند اما عیب آنها این است که همگرایی در این موارد تضمین نمی‌شود بویژه درحالتی که مجموعه‌های جواب کراندار نیستند.

رهیافت دیگر این است که یک مجموعه از معادلات دیفرانسیل ایجاد می‌کنند به‌طوری‌که نقاط تعادل آنها متناظر با جواب مطلوب باشد و بعد یک تابع لیاپانوف مناسب پیدا کرده، به‌طوری‌که همه مسیرهای سیستم به نقاط تعادل همگرا باشند.

۵.۲ مدل‌های شبکه عصبی ارائه شده پیشین برای حل مسائل بهینه‌سازی

بحث و بررسی در مورد کاربرد شبکه‌های عصبی (شبکه‌های عصبی مصنوعی) در بهینه‌سازی از اوایل سال ۱۹۸۰ میلادی آغاز شده است. نتایج پژوهش‌های انجام شده از آن زمان تاکنون، مواردی چون برنامه‌ریزی خطی، برنامه‌ریزی درجه دوم، برنامه‌ریزی هندسی و برنامه‌ریزی غیرخطی را در برمی‌گیرد. ایده اصلی در استفاده از شبکه‌های عصبی مصنوعی برای مسائل بهینه‌سازی استفاده از یک تابع انرژی (نامنفی) و یک سیستم دینامیکی است که این دو بیان‌کننده مدل‌های شبکه‌های عصبی مصنوعی متناظر مسائل

بهینه‌سازی هستند. سیستم دینامیکی بیان شده معمولاً یک دستگاه معادلات دیفرانسیل غیرخطی مرتبه اول است. انتظار می‌رود که برای یک نقطه آغازین نقطه پایداری دستگاه معادلات دیفرانسیل غیرخطی به دست آمده، جواب بهینه مسأله بهینه‌سازی اصلی باشد. یک اصل اساسی در استفاده از تابع انرژی این است که برای همگرایی دستگاه معادلات دیفرانسیل به نقطه پایداری آن، تابع انرژی متناظر با آن حتماً باید نامنفی باشد. البته می‌توان مدل شبکه‌های عصبی نظیر مسائل بهینه‌سازی را حتی بدون در نظر گرفتن تابع انرژی نیز بیان کرد. بنابراین برای یک مدل متناظر با مسائل بهینه‌سازی، اصل اساسی استفاده از شبکه‌های عصبی در اینگونه مسائل به صورت زیر بیان می‌شود:

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

مدل‌های مطرح شده متناظر مسائل مختلف بهینه‌سازی را می‌توان به دو قسمت مدل‌های دوگانی و مدل‌های جریمه‌ای تقسیم نمود. نظریه دوگانی و توابع جریمه‌ای دو نظریه بسیار مهم در مسائل بهینه‌سازی هستند که اکثر این مسائل بر مبنای این دو روش کلاسیک قابل حل هستند. در نظریه توابع جریمه‌ای معمولاً از مدل‌های گرادیانی برای معرفی مدل مورد نظر استفاده می‌شود، ولی در نظریه دوگانی از مدل‌های اولیه -دوگان برای حل این مسائل استفاده می‌شود. اگر بتوان متناظر با هر مسأله بهینه‌سازی، روش مشخصی را ارائه نمود که آن روش برای حل آن مسأله شرایط لازم و کافی را برآورده سازد، آنگاه می‌توان متناظر با آن روش، یک مدل شبکه عصبی برای مسأله مورد نظر بسازیم. در زیر به بیان چند مدل که تاکنون برای مسائل برنامه‌ریزی خطی و درجه دوم و غیرخطی ارائه شده‌اند می‌پردازیم.

مدل اول ([۲۱])

مسئله برنامه‌ریزی خطی زیر را در نظر بگیرید:

$$\begin{cases} \text{Minimize} & f(x) = a^T x \\ \text{s.t.} & \\ & g(x) = Dx - b = \circ, \\ & x \geq \circ, \end{cases} \quad (۱.۲)$$

که در آن D یک ماتریس $m \times n$ و $Rank(D) = m$ و $a \in \mathbb{R}^m$ ، $x, b \in \mathbb{R}^n$ است. دوگان (۱.۲) به صورت

زیر بیان می‌شود:

$$\begin{cases} \text{Maximize} & \bar{f}(y) = b^T y \\ \text{s.t.} & \\ & \bar{g}(y) = \bar{D}^T y - a \leq \circ, \\ & y \in \mathbb{R}^n. \end{cases} \quad (۲.۲)$$

در مرجع [۲۱] مدل متناظر با (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 $\beta = \| (x + D^T y - a)^+ - x \|_2^2$ و $(x, y) \in \{(x, y) \mid x \geq \circ\}$.

مدل دوم ([۲۱])

مسئله برنامه‌ریزی درجه دوم زیر را در نظر بگیرید:

$$\begin{cases} \text{Minimize} & f(x) = \frac{1}{2} x^T A x + a^T x \\ \text{s.t.} & \\ & Dx - b = \circ, \\ & x \geq \circ, \end{cases} \quad (۳.۲)$$

که در آن A یک ماتریس متقارن معین مثبت است. دوگان (۳.۲) به صورت زیر است:

$$\begin{cases} \text{Maximize} & \bar{f}(y) = b^T y - \frac{1}{2} x^T A x \\ \text{s.t.} & \\ & \bar{D} y - A x - a \leq \circ. \end{cases} \quad (۴.۲)$$

در مرجع [۲۱] مدل متناظر با (۳.۲) و (۴.۲) به صورت زیر نمایش داده شده است:

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

که در آن $\gamma = \| (x + D^T y - A x - a)^+ - x \|_2^2$ و $(x, y) \in \{(x, y) \mid x \geq \circ\}$.

مدل سوم ([۱۲])

مسئله برنامه‌ریزی غیرخطی زیر را در نظر بگیرید:

$$\begin{cases} \text{Minimize} & f(x) \\ \text{s.t.} & \\ & g(x) = [g_1(x), g_2(x), \dots, g_m(x)]^T \leq 0, \end{cases} \quad (5.2)$$

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 (۵.۲) هافیلد و تنک در سال

۱۹۸۶، یک مدل شبکه عصبی به صورت زیر ارائه دادند:

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

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 می‌شوند:

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

که $\frac{1}{\rho_i}$ ضریب هدایت گرمایی هر نرون، d_{ji} وزن‌های تخصیص داده شده به نرونهای 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 ([۱۳])

برای حل (۵.۲)، کندی و چوآ [۱۳] مدلی دیگر به صورت زیر ارائه داده‌اند:

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

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

$$E_2(x) = f(x) + \frac{s}{2} \sum_{j=1}^m [g_j^+]^2.$$

مدل پنجم ([۲۲])

مجدداً برای حل (۵.۲) توسط رودریگز-واسکوئز مدلی به‌صورت زیر ارائه شده است:

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

که در آن u_x متناظر اندیس‌های شدنی x است چنان‌که

$$u_x = \begin{cases} 1 & g(x) \leq 0, \\ 0 & 0.w. \end{cases}$$

تابع انرژی متناظر این مدل نیز به‌صورت زیر انتخاب شده است:

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

همان‌طور که ملاحظه می‌کنید، مدل‌های بیان شده به‌طور متوالی اصلاح شده‌اند و این به‌دلیل مشکلاتی

بوده است که هرکدام از این مدل‌ها داشته‌اند. ([۲۱])

مدل ششم ([۴])

برنامه‌ریزی محدب درجه دوم ${}^{\mathfrak{V}}(CQP)$ زیر را در نظر بگیرید:

$$\begin{cases} \text{Minimize} & f(x) = \frac{1}{\mathfrak{P}} x^T Q x + D^T x \\ \text{s.t.} & \\ & g(x) = Ax - b \leq 0, \\ & h(x) = Ex - f = 0. \end{cases} \quad (6.2)$$

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. دوگان برنامه‌ریزی محدب درجه دوم (۶.۲) به‌صورت زیر نوشته می‌شود:

$$\begin{cases} \text{Maximize} & \bar{f}(x, u, v) = f(x) + u^T (Ax - b) + v^T (Ex - f) \\ \text{s.t.} & \\ & \bar{g}(x, u, v) = Qx + D + A^T u + E^T v = 0, \\ & u \geq 0. \end{cases} \quad (7.2)$$

^{۳۷}Convex quadratic programming

بر طبق شرایط $K.K.T$ برای (۶.۲) و (۷.۲)، $x^* \in \mathbb{R}^n$ جواب بهینه (۶.۲) است اگر و تنها اگر بردارهای

$u^* \in \mathbb{R}^m$ و $v^* \in \mathbb{R}^l$ وجود داشته باشد به طوری که در رابطه زیر صدق کند:

$$\begin{cases} u^* \geq 0, Ax^* - b \leq 0, u^{*T}(Ax^* - b) = 0, \\ Qx^* + D + A^T u^* + E^T v^* = 0, \\ Ex^* - f = 0. \end{cases} \quad (۸.۲)$$

برای حل مسأله (۶.۲) تابع زیر را که معروف به تابع فیشر-برمیستر^{۳۸} است، معرفی می‌کنیم:

$$\begin{cases} \phi : R^r \rightarrow R^1 \\ \phi(a, b) = \sqrt{a^r + b^r} - a - b. \end{cases}$$

قضیه ۱.۵.۲. (۲۳):

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^{۳۹} هموار

است.

تابع ϕ که در شرط (۱) قضیه (۱.۵.۲) صدق کند، تابع NCP ^{۴۰} نامیده می‌شود.

متناظر با دستگاه (۷.۲) و (۶.۲) دستگاه زیر را در نظر می‌گیریم:

$$\phi(x, u, v) = \begin{bmatrix} \phi(\alpha, Qx + D + A^T u + E^T v) \\ \phi(u, Ax - b) \\ \phi(\beta, Ex - f) \end{bmatrix}, \quad (۹.۲)$$

که در آن $\alpha = (\alpha_1, \dots, \alpha_m)^T > 0$ ، $u = (u_1, \dots, u_m)$ و $\beta = (\beta_1, \dots, \beta_l)^T$ است. مدل شبکه عصبی

متناظر با (۹.۲) به صورت زیر بیان شده است:

$$\begin{cases} \frac{dy}{dt} = -\tau \nabla E(y(t)), \quad \tau > 0 \\ y(t_0) = y_0 \in \mathbb{R}^{n+m+l}. \end{cases} \quad (۱۰.۲)$$

^{۳۸}Fisher-Bermister

^{۳۹}Stringly smooth

^{۴۰}Nonlinear Complementarity Problem

که در آن $y = (x^T, u^T, v^T)^T$ و τ ضریبی برای افزایش سرعت همگرایی و کاهش تعداد تکرار ها در روش عددی انتخاب شده برای حل دستگاه (۹.۲) است. تابع انرژی متناظر این مدل به صورت زیر انتخاب شده است:

$$E(y) = \frac{1}{4} \| \phi(y) \|^2.$$

مدل هفتم ([۴])

برای طراحی مدل هفتم، مجدداً شرایط $K.K.T$ را برای حل مسائل (۷.۲) و (۶.۲) به صورت زیر در نظر می گیریم:

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

آنگاه متناظر (۶.۲) و (۷.۲) دستگاه زیر تعریف می شود:

$$U(x, u, v) = \begin{bmatrix} -(Qx + D + A^T u + E^T v) \\ (u + Ax - b)^+ - u \\ (Ex - f) \end{bmatrix}. \quad (۱۱.۲)$$

مدل شبکه عصبی متناظر برای حل (۶.۲) و (۷.۲) به صورت زیر بیان شده است:

$$\begin{cases} \frac{dy}{dt} = \kappa U(y), \\ y(t_0) = y_0. \end{cases} \quad (۱۲.۲)$$

مدل هشتم ([۲۴])

مسأله برنامه ریزی درجه دوم (۶.۲) را در نظر بگیرید. دوگان و شرایط بهینگی برای آن به ترتیب در فرم (۷.۲) و (۸.۲) بیان شده اند که در آن $Q \in \mathbb{R}^{n \times n}$ یک ماتریس نیمه معین مثبت است. برای مسائل (۶.۲) و (۷.۲) یک مدل شبکه عصبی به صورت زیر ارائه می شود:

$$\frac{dy}{dt} = (I + M^T) \begin{bmatrix} -(Qx + D + A^T u + E^T v) \\ (u + Ax - b)^+ - u \\ (Ex - f) \end{bmatrix}, \quad (۱۳.۲)$$

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 ([۲۵])

مسأله برنامه‌ریزی درجه دوم زیر را در نظر بگیرید:

$$\begin{cases} \text{Minimize} & f(x) = \frac{1}{2}x^T Ax + a^T x \\ \text{s.t.} & \\ & Dx \leq b, \\ & x \geq \circ. \end{cases} \quad (۱۴.۲)$$

دوگان آن عبارت است از:

$$\begin{cases} \text{Maximize} & \bar{f}(w) = b^T w - \frac{1}{2}x^T Ax \\ \text{s.t.} & \\ & \bar{D}^T w - Ax \leq a, \\ & w \geq \circ. \end{cases} \quad (۱۵.۲)$$

که در آن $A \in \mathbb{R}^{n \times n}$ متقارن و نیمه معین مثبت است، $a \in \mathbb{R}^n$ و $b \in \mathbb{R}^m$ ، $D \in \mathbb{R}^{m \times n}$ است. با توجه

به شرایط KKT برای مسائل برنامه‌ریزی محدب، $(x^{*T}, w^{*T})^T$ به ترتیب یک جواب بهینه برای مسائل

(۱۴.۲) و (۱۵.۲) است اگر و فقط اگر $(x^{*T}, w^{*T})^T$ در شرایط $K.K.T$ 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 (۱۴.۲) و (۱۵.۲) در [۲۵] به صورت زیر بیان شده است:

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

که در آن

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

M یک ماتریس نیمه معین مثبت است زیرا $u^T Mu = x^T Ax \geq \circ$. همچنین $B = \lambda I$ است که

$\circ < \lambda$ و با افزایش λ سرعت همگرایی افزایش می‌یابد.

فصل ۳

بهینه‌سازی سبد سرمایه

۱.۳ مقدمه

در این فصل ابتدا مقدمه‌ای بر سرمایه‌گذاری سبد سرمایه داریم سپس مسأله بهینه‌سازی سبد سرمایه را بیان کرده و مروری بر کارهای انجام شده در این زمینه داریم. در ادامه چندین مدل جایگزین برای مدل مبنا را معرفی می‌کنیم. برای مسائل درجه دوم محدب مدل شبکه عصبی را معرفی کرده و به اثبات پایداری و همگرایی آن می‌پردازیم. در پایان کارایی مدل شبکه عصبی را با چندین مثال عددی از مسائل بهینه‌سازی سبد سرمایه نشان می‌دهیم.

بحث سرمایه‌گذاری در اوراق بهادار و تحلیل آن را می‌توان به دو چارچوب کلی و متفاوت تقسیم کرد ([۷]):

۱. تحلیل و انتخاب دارایی به‌طور جداگانه با استفاده از ابزارها و روشهای تحلیل تکنیکی^۱ و تحلیل بنیادی^۲.

۲. تشکیل سبد دارایی‌ها با استفاده از نظریه نوین سبد سرمایه^۳ (MPT) و سایر تئوریهای بازار سرمایه.

مسأله انتخاب مجموعه‌ی بهینه‌ای از دارایی‌ها، یکی از مسائل بازار سرمایه است که اهمیت زیادی برخوردار است. از حوزه اقتصاد خرد، اهمیت تصمیمات سرمایه‌گذاری ناشی از این مسأله است که در واقع فرد سرمایه‌گذار، مصرف امروز را با امید مصرف بیشتر به زمانی در آینده موکول می‌کند. در واقع تصمیم بهینه سرمایه‌گذاری، میزان مطلوبیت هر فرد است که، با توجه به ترجیحات شخصی وی تعیین می‌گردد و لزوماً با سایر افراد یکسان نخواهد بود، ولی می‌توان با معیارهایی مطلوبیت فرد را از انتخاب مجموعه‌ای از دارایی‌های سرمایه‌ای، مشخص کرد. از جمله مهمترین این معیارها، ریسک و بازده کل مجموعه انتخاب

^۱Technical Analysis

^۲Fundamental Analysis

^۳Portfolio

شده است. هدف مدیریت مجموعه دارایی‌ها به‌طور عام و مجموعه سهام (سبد سهام) به‌طور خاص، تعیین دارایی‌های درون سبد بگونه‌ای که ریسک حداقل و بازده حداکثر شود. هر چه ریسک سبد سهام بیشتر باشد، احتمال دریافت بازده بالاتر هم بیشتر است.

در دنیای واقعی با توجه به‌اینکه درجه ریسک‌پذیری افراد با یکدیگر متفاوت است، بازده دارایی‌ها نیز به‌دلیل وجود عوامل متعدد مؤثر بر آن، غیر قابل پیش‌بینی خواهد بود. به دلیل وجود این عدم اطمینان در بازار، مسأله تنوع بخشی در دارایی‌ها از اهمیت ویژه‌ای برخوردار می‌گردد. طبقات مختلف دارایی شامل انواع اوراق بهادار، پول نقد، زمین، طلا و... هستند که ریسک و بازده هر یک از آنها با هم متفاوت است.

تئوریهای سرمایه‌گذاری در چند دهه اخیر از پیشرفتهای زیادی برخوردار بودند و در سیر تاریخی تکاملی خود، به فرمولهای کاربردی زیادی دست یافتند. حجم تجارت و سرمایه‌گذاری طی قرن بیستم و اوایل بیست و یکم گسترش فراوانی داشته است. این گسترش و تغییرات حاصل از آن، موجب شده است که معیارهای متفاوتی برای اخذ تصمیم سرمایه‌گذار در مقایسه با دوره‌های گذشته به‌کار گرفته شود. از عوامل مؤثر در انتخاب و انجام سرمایه‌گذاری، توجه سرمایه‌گذار به مفهوم ریسک و بازده و رابطه این دو با یکدیگر است. در قرن ۱۸ میلادی برنولی و کرامر^۴ به این نتیجه رسیدند که تصمیمات تحت شرایط عدم اطمینان نباید صرفاً براساس بازده موردانتظار صورت پذیرد.

تا اوایل قرن بیستم میلادی، سرمایه‌گذاران جهت اخذ تصمیم در فرآیند سرمایه‌گذاری از نسبتهای بازده سرمایه‌گذاری استفاده می‌کردند و توجهی به مفاهیم ارزش زمانی پول و ریسک سرمایه‌گذاری نداشتند. از دهه ۱۹۲۰ مفهوم ارزش زمانی پول با روشهای تنزیلی^۵ وارد حوزه ادبیات مالی و سرمایه‌گذاری شد. این روشها، تحولی قابل توجه در انتخاب طرحهای سرمایه‌گذاری به‌وجود آوردند. لیکن همچنان رفتار متفاوت سرمایه‌گذاران در مواجهه با ریسک نادیده گرفته می‌شد. در واقع با وجود اینکه نظریه مطلوبیت پول تا

^۴Bernoulli and Cramer

^۵Discounting Methods

حدودی به تکامل معیارهای انتخاب کمک نموده بود، لیکن هنوز از جامعیت کافی برخوردار نبودند. تا دهه ۱۹۵۰، ریسک یک عامل کیفی به‌شمار می‌رفت تا اینکه هری مارکوویتز برای نخستین بار ریسک را کمیت‌پذیر نمود و انحراف معیار جریانات نقدی طرحهای سرمایه‌گذاری را به‌عنوان کمیت سنجش ریسک معرفی کرد. چندی بعد ویلیام شارپ در [۲۶] با تبیین ضریب حساسیت بتا (β) به‌عنوان معیار ریسک، مدل ساده‌تر و کاربردی‌تری را به دنیای تئوریهای سرمایه‌گذاری عرضه کرد. این روش امروز به‌عنوان مدل تک شاخصی معروف است. در ادامه این روند و در اواسط دهه ۶۰ شارپ و لینتر [۲۷] بر پایه تئوری بازار سرمایه، مدلی را توسعه دادند که امروزه تحت عنوان مدل قیمت‌گذاری دارایی‌های سرمایه‌ای^۶ شناخته می‌شود. این مدل ریسک سیستماتیک و غیرسیستماتیک را به‌عنوان اجزای اصلی ریسک از یکدیگر تفکیک می‌کند. در دهه ۷۰، پروفیسور راس، مدل آربیتراژ^۷ را در [۲۸] پایه‌گذاری کرد. (برای توضیحات بیشتر به [۳۰]، [۲۹] مراجعه شود).

۲.۳ مسأله بهینه‌سازی سبد سرمایه

در سال ۱۹۵۲ میلادی، هری مارکوویتز مدل بنیادی سبد سرمایه را ارائه داد که مبنایی بر تئوری مدرن سبد سرمایه گردید. پیش از ارائه این تئوری سرمایه‌گذاران علیرغم آشنایی با مفاهیم ریسک و بازده، قادر به اندازه‌گیری آن به روش کمی نبودند. سرمایه‌گذاران ضمن ایجاد تنوع، قرار ندادن همه تخم مرغها را در یک سبد را نیز قبول داشتند. مارکوویتز نخستین کسی بود که مفهوم تنوع بخشی در سبد سرمایه‌گذاری به‌طور عام و سبد سهام به‌طور خاص را بسط و گسترش داد. او این موضوع را که چگونه متنوع سازی سبد سرمایه می‌تواند منجر به کاهش ریسک سرمایه‌گذاری یک سرمایه‌گذار شود را به‌صورت کمی بیان نمود.

تئوری مدرن سبد سرمایه، یک نگرش کل‌گرا به بازار سهام است. این نظریه بر خلاف روشهای نموداری

^۶Capital Asset Pricing Model(CAPM)

^۷Arbitrage Pricing Theory(APT)

(تکنیکال) و بنیادی، به مجموعه‌ی سهام در سبد توجه دارد. مارکوویتز در فرمول‌بندی مدل میانگین-واریانس^۸ خود، به هدف سرمایه‌گذاری توجه خاصی داشت. به نظر وی، سرمایه‌گذار عاقل به دنبال سرمایه‌گذاری در طرح‌هایی است که بازدهی بیشتر و ریسک کمتر داشته باشد. وی ریسک یک سرمایه‌گذاری را تنها در انحراف معیار آن طرح جستجو نمی‌کرد، بلکه به رابطه بین دارایی‌های مختلف در سبد سرمایه و تأثیر این رابطه بر ریسک کل هم توجه داشت.

یکی دیگر از مفاهیمی که مارکوویتز مطرح نمود، سبد سرمایه کارا (سبد سهام کارا)^۹ بود، که به معنای ترکیب مطلوب اوراق بهادار است به نحوی که ریسک آن سبد سرمایه در ازای بازده معینی حداقل شود یا بازده سبد سرمایه در ازای ریسک معین حداکثر شود. به این ترتیب، مجموعه این سبد سرمایه‌ها با عنوان مجموعه کارا شناخته شده و توسط حل مسأله پارامتریک غیرخطی قابل شناسایی است. در نمودار ریسک-بازدهی مجموعه کارا، مرز کارا نیز نامیده می‌شود.

مرز کارا مجموعه تعداد زیادی سبد سرمایه سهام ارائه می‌دهد که انتخاب هر کدام از این سبد سرمایه‌ها به منحنی مطلوبیت^{۱۰} سرمایه‌گذار بستگی دارد. هر سرمایه‌گذار پورتفوئی که در نقطه مماس شدن منحنی مرز کارای مارکوویتز و منحنی مطلوبیت وی قرار دارد را به عنوان بهترین سبد سرمایه انتخاب خواهد کرد. شکل (۱.۳) مرز کارای مارکوویتز و منحنی مطلوبیت سرمایه‌گذار را نشان می‌دهد.

مدل مارکوویتز براساس مفروضاتی شکل گرفته است. این مفروضات عبارتند از:

۱. سرمایه‌گذار دارای رفتار عقلایی است.

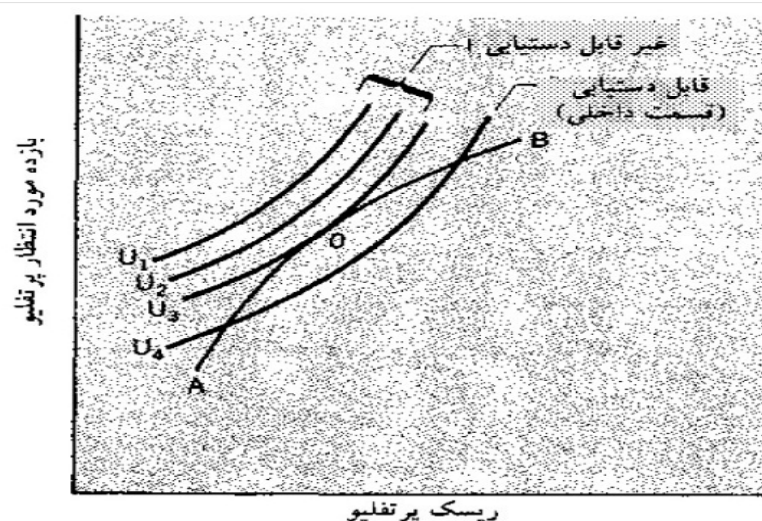
۲. سرمایه‌گذاران، ریسک‌گریزند و دارای مطلوبیت مورد انتظار افزایشی هستند.

۳. سرمایه‌گذاران سبد سرمایه خود را بر مبنای میانگین و واریانس مورد انتظار بازدهی انتخاب می‌کنند.

^۸Mean-Variance

^۹Efficient Portfolio

^{۱۰}Utility Curve



شکل ۱.۳: مرز کارایی و منحنی مطلوبیت

بنابراین منحنیهای بی تفاوتی آنها تابعی از نرخ بازده و واریانس مورد انتظار است.

۴. هر گزینه سرمایه‌گذاری تا بی نهایت قابل تقسیم است.

۵. سرمایه‌گذاران افق زمانی یک دوره‌ای داشته و این برای همه مشابه است.

۶. سرمایه‌گذاران در یک سطح مشخص از ریسک، بازده بالاتری را ترجیح می‌دهند و بالعکس برای

سطح معینی از بازدهی، خواهان کمترین ریسک هستند.

۷. بازارها کامل هستند (هزینه مالیات و معاملات وجود ندارد).

در رویکرد مارکوفیتز بازده دارایی‌ها به صورت متغیر تصادفی در نظر گرفته می‌شود که احتمال آنها از تابع

توزیع نرمال به دست می‌آید. میانگین به عنوان معیار کارایی سبد سرمایه و انحراف معیار^{۱۱} هم به عنوان

شاخص ریسک سبد سرمایه در نظر گرفته می‌شود. ورودیهای اصلی مدل مارکوفیتز عبارتند از بازده مورد

انتظار، ماتریس واریانس-کواریانس بازدهی دارایی‌ها و ضریب همبستگی بین بازدهی‌ها. مدل کلاسیک

^{۱۱} Standard Deviation

MV به صورت زیر بیان می‌شود ([۳۱]):

$$\text{Minimize } Z_{Qp} = \frac{1}{2} X^T \Sigma X \quad (۱.۳)$$

subject to

$$\begin{cases} \mu^T X = R, \\ \sum_{i=1}^n x_i = 1, \\ x_i \geq 0, \quad i = 1, 2, \dots, n, \end{cases} \quad (۲.۳)$$

که در آن

$$E[X] = x_1 \mu_1 + \dots + x_n \mu_n = \mu^T X,$$

$$V[X] = \sum_{i,j} \rho_{ij} \sigma_i \sigma_j x_i x_j = X^T \Sigma X,$$

وقتی که ρ_{ij} ضریب همبستگی بین سهم i ام و j ام، μ_i بازده مورد انتظار سهم i ام است و قید اول نشان دهنده سطح مشخصی از بازدهی مورد انتظار یعنی R ، قید دوم این اطمینان را می‌دهد که تمام بودجه سرمایه‌گذاری شود و قید آخر هم نشان دهنده این است که سرمایه‌گذار اجازه فروش استقراضی^{۱۲} را ندارد. متأسفانه مدل مارکوویتز یک سری معایب و فرضیات غیرواقعی دارد که از جمله می‌توان به موارد زیر اشاره کرد:

۱. با افزایش تعداد دارایی، حجم محاسبات ماتریس کواریانس بیش از اندازه بزرگ می‌شود.

۲. محدودیتهایی که در دنیای واقعی وجود دارد، در مدل مارکوویتز در نظر گرفته نشده است، همانند

محدودیت سهام ثابت در سبد سرمایه^{۱۳}، هزینه‌های معاملاتی^{۱۴}، خرید آستانه‌ای^{۱۵}، معاملات

بلوکی^{۱۶}، معادلات سهام در دسته‌های ثابت^{۱۷} و ...، اکثر این محدودیتها از توابع غیرخطی

پیروی می‌کنند که مدل‌ها را برای حل دچار پیچیدگی می‌نماید.

^{۱۲} Short Selling

^{۱۳} Cardinality Constraint

^{۱۴} Transaction Cost

^{۱۵} Buy-in Thresholdy

^{۱۶} Roundlots

^{۱۷} Transaction Lot

۳. فرض نرمال بودن تابع توزیع بازده دارایی‌ها فرض قابل قبولی نیست.

۴. معیار عمومی ریسک، واریانس و یا انحراف معیار است که این معیار برای یک دارایی با توزیع نرمال و در بازار کارا قابل قبول است. اگر این دو خصوصیت برای دارایی وجود نداشته باشد، واریانس معیار مناسبی نیست.

۳.۳ مدل‌های جایگزین

در سال ۱۹۵۹ مارکوویتز^{۱۸} به مزایای ریسک نامطلوب اشاره کرد، وی در [۳۲] به این نتیجه رسید که سرمایه‌گذاران به دو علت مایل به حداقل ساختن ریسک نامطلوب هستند، نخست اینکه سرمایه‌گذاران ابتدا به امنیت اصل سرمایه می‌اندیشند و دوم اینکه اگر توزیع متغیرهای تصادفی (نرخ بازدهی) از نوع نرمال نباشد، آنگاه معیار ریسک نامطلوب مفید خواهد بود. بعدها در سال ۱۹۹۱ وی هنگام دریافت جایزه نوبل اقتصاد برای ارائه نظریه مدرن سبد سرمایه و مدل (MV) نیز، به این مسأله اشاره داشت.

سال ۱۹۷۳ لی و لرو^{۱۹} در [۳۳]، مدلی به‌منظور بهینه‌سازی سبدسرمایه با به کارگیری برنامه‌ریزی آرمانی ارائه نمودند. چندی بعد در سال ۱۹۸۰ لی و چیسیر^{۲۰} در [۳۴]، نگرش برنامه‌ریزی آرمانی چند معیاره^{۲۱} را مطرح کردند. کونو و یامازاکی در سال ۱۹۹۱، مدلی از طریق برنامه‌ریزی خطی برای بهینه‌سازی سبد سرمایه‌گذاری ارائه کردند. پس از آن، در سال ۱۹۹۷، اسپرانزا و مانسینی^{۲۲} در [۳۵]، مدلی از برنامه‌ریزی صحیح آمیخته^{۲۳} را با خصوصیات واقعی مثل هزینه‌های معاملات و حداقل واحدهای معاملات ارائه داد و برای حل آن از الگوریتم‌های ابتکاری کمک گرفتند و مدل را برای بازار سهام میلان آزمایش کردند. سال ۱۹۹۸ یونگ^{۲۴} در [۳۶]، مدل بهینه‌سازی سبد سرمایه قابل حل توسط برنامه‌ریزی خطی را که در آن

^{۱۸}Markowitz

^{۱۹}Lee and Lerro

^{۲۰}Lee and Chesser

^{۲۱}Multi-criteria goal programming

^{۲۲}Mansini and Speranza

^{۲۳}Mixed integer programming

^{۲۴}Young

ریسک بر مبنای بدترین حالت تعریف شده و آن را رویکرد "حداقل نمودن حداکثر" ^{۲۵} نامید. سال ۱۹۹۹ اکسیا و همکاران در [۳۷] مدلی خاص را معرفی کردند. در فوریه همین سال چانگ ^{۲۶} و همکاران، در [۳۸] به ارائه مدلی پارامتریک و غیرخطی جهت بهینه‌سازی سهام پرداختتند. در ادامه در نوامبر همین سال، جوناس پالمکوئیست ^{۲۷} و همکارانش در [۳۹] به معرفی بهینه‌سازی سبد سرمایه با محدودیت‌ها و تابع هدف ارزش در معرض ریسک شرطی ^{۲۸}، مدلی از نوع برنامه‌ریزی خطی ارائه دادند. در سال ۲۰۰۲ لوبو ^{۲۹} و همکاران در [۴۰] به بهینه‌سازی سبد با اعمال هزینه‌های معاملات خطی و ثابت پرداختتند. در همین سال، کریستوس پاپاریستودولو ^{۳۰} در [۴۱] به ارائه مدل‌های موجود در این زمینه پرداخت و سپس با طرح مسأله‌ای متشکل از ۵ سهام در دوره زمانی ۱۲ ماهه به مقایسه سبدهای به دست آمده در هر مدل پرداخت. در ژانویه ۲۰۰۳ آلکسی کخلو ^{۳۱} و همکاران در مقاله‌ای تحت عنوان "بهینه‌سازی سبد سرمایه با محدودیت‌های کاهش دهنده" ^{۳۲} به ارائه مدلی در این زمینه پرداخت. در ژوئای ۲۰۰۳ مانسینی ^{۳۳} و همکارانش در [۴۲] به معرفی مدل‌های *LP* پرداخته‌اند. در سال ۲۰۰۴ گاندریو ^{۳۴} و همکاران برای حل مسائل بهینه‌سازی غیرخطی سبد سرمایه، روش نقطه درونی ^{۳۵} اولیه-دوگان را معرفی کردند. در ادامه به بیان چندین مدل از مدل‌های نامبرده می‌پردازیم.

^{۲۵}Minimax^{۲۶}Chang^{۲۷}Palmquist^{۲۸}Conditional Value at Risk^{۲۹}Lobo^{۳۰}Papahristodoulou^{۳۱}Alexi Chekhlov^{۳۲}Draw Down^{۳۳}Mansini^{۳۴}Gandryv^{۳۵}Interior-point method

۱.۳.۳ مدل‌های درجه دوم

فرناندز و گومز^{۳۶} در [۴۳]، مدل اصلاح یافته مارکویتز را تحت عنوان "مدل میانگین-واریانس با مؤلفه‌های محدود^{۳۷}" معرفی کردند که شکل عمومی آن به صورت زیر است:

$$\text{Minimize } Z = \frac{1}{2} X^T \Sigma X \quad (۳.۳)$$

subject to

$$\begin{cases} \mu^T X = R, \\ \sum_{i=1}^n x_i = 1, \\ \varepsilon_i \leq x_i \leq \delta_i, \quad i = 1, 2, \dots, n. \end{cases} \quad (۴.۳)$$

در صورتی که محدودیت مربوط به دارایی‌های منتخب به مسأله فوق اضافه شود، مدل مربوطه به صورت زیر در می‌آید:

$$\text{Minimize } \lambda \left[\sum_{i=1}^n \sum_{j=1}^n z_i x_i z_j x_j \rho_{ij} \right] - (1 - \lambda) \left[\sum_{i=1}^n z_i x_i \mu_i \right] \quad (۵.۳)$$

subject to

$$\begin{cases} \sum_{i=1}^n x_i = 1, \\ \sum_{i=1}^n z_i = k, \\ \varepsilon_i z_i \leq x_i \leq \delta_i z_i, \\ z_i \in [0, 1], \\ x_i \geq 0, \quad i = 1, 2, \dots, n. \end{cases} \quad (۶.۳)$$

که در آن ضریب همبستگی بین سهم i ام و j ام، μ_i بازده مورد انتظار سهم i ام است و $\lambda \in [0, 1]$ و z_i متغیر تصمیم در مورد سرمایه‌گذاری در سهم i ام است. مجموع تعداد سهامی که در سبد خواهد بود بنا به محدودیت دوم، k تا خواهد بود. برای حل این مسأله از الگوریتم‌های ابتکاری مثل تکنیک بهینه‌سازی جمعی ذرات^{۳۸} کمک گرفته شده است.

^{۳۶}Fernandez and Gomez

^{۳۷}Cardinality Constrained Mean-Variance(CCMV)

^{۳۸}Particle Swarm Optimization(PSO)

مدلهایی با فرم درجه دوم قطری

تجزیه چولسکی [۴۴] ماتریس کواریانس V به صورت زیر بیان می‌شود:

$$V = L^T L$$

که در آن $L_{n \times n}$ یک ماتریس پایین مثلثی است، با تعریف بردار جدید $Y_{n \times 1}$ به صورت $Y = LX$ با عناصر،

$$y_i = \sum_{j=1}^i l_{ij} x_j \quad i = 1, \dots, n,$$

فرمول‌بندی یکسان از مسأله مارکوویتز به دست می‌آید. بنابراین مدل قطری به صورت زیر بیان می‌شود:

$$\text{Minimize} \quad Z_{diag1} = \sum_{i=1}^n y_i^2 \quad (7.3)$$

subject to

$$\begin{cases} y_i = \sum_{j=1}^i l_{ij} x_j & i = 1, 2, \dots, n, \\ \sum_{i=1}^n x_i \mu_i = R, \\ \sum_{i=1}^n x_i = 1, \\ x_i \geq 0, & i = 1, 2, \dots, n, \\ g_i^y \leq y_i \leq h_i^y & i = 1, 2, \dots, n. \end{cases} \quad (8.3)$$

برای فرم درجه دوم کلی، y_i متغیرهای آزاد $(-\infty < y_i < +\infty)$ هستند.

رویکرد مشابه در تجزیه ماتریس کواریانس V ، بازدهی‌های R مشاهده شده در کل دوره T را نیز در این

تجزیه اعمال کنیم. اگر $\bar{R}$ را ماتریس میانگین بازدهی‌ها در نظر بگیریم آنگاه ماتریس V به صورت زیر

محاسبه می‌شود:

$$V = \frac{1}{n-1} (R - \bar{R})^T (R - \bar{R}).$$

با تعریف $S_{T \times n} = \frac{1}{\sqrt{n-1}} (R - \bar{R})$ ، ماتریس V به صورت $V = S^T S$ تعریف می‌شود، بنابراین مدل قطری

دوم به صورت زیر بیان می‌شود:

$$\text{Minimize} \quad Z_{diag2} = \sum_{t=1}^T y_t^2 \quad (9.3)$$

subject to

$$\begin{cases} y_t = \sum_{i=1}^n s_{ti} x_i & t = 1, 2, \dots, T, \\ \sum_{i=1}^n x_i \mu_i = R, \\ \sum_{i=1}^n x_i = 1, \\ x_i \geq 0, & i = 1, 2, \dots, n, \\ g_t^y \leq y_t \leq h_t^y & t = 1, 2, \dots, T. \end{cases} \quad (10.3)$$

۲.۳.۳ مدل‌های خطی

مدل میانگین قدرمطلق انحرافات^{۳۹}

بعد از مارکویتز، شارپ در سال ۱۹۷۱ تلاش فراوانی کرد تا مدل خطی از بهینه‌سازی سبد سرمایه ارائه دهد، بعد از آن کونو^{۴۰} در [۴۵]، مدلی را با استفاده از توابع قطعه‌ای خطی پیشنهاد کرد. در مدل MAD ریسک به جای واریانس توسط میانگین قدرمطلق انحرافات اندازه‌گیری می‌شود. هم‌چنین آنها در [۴۶] نشان دادند که مدل مارکویتز تحت این فرض که بازدهی‌ها دارای توزیع نرمال چند متغیره باشند، معادل مدل آنهاست. فرض کنید m_t انحرافات مطلق بازدهی سبد سرمایه از میانگین آن در زمان t باشد، بنابراین مدل را به‌صورت زیر بیان می‌کنیم:

$$\text{Minimize} \quad Z_{mad} = \frac{1}{T} \sum_{t=1}^T m_t \quad (11.3)$$

subject to

$$\begin{cases} \sum_{i=1}^n (r_{it} - \mu_i) x_i \leq m_t & t = 1, 2, \dots, T, \\ \sum_{i=1}^n (r_{it} - \mu_i) x_i \geq -m_t & t = 1, 2, \dots, T, \\ \sum_{i=1}^n x_i \mu_i = R, \\ \sum_{i=1}^n x_i = 1, \\ m_t \geq 0, & t = 1, 2, \dots, T, \\ x_i \geq 0, & i = 1, 2, \dots, n. \end{cases} \quad (12.3)$$

سیمن^{۴۱} در [۴۷]، درباره معایب و مزایای دو مدل MV و MAD بحث کرده است.

^{۳۹} Mean-absolute Deviation (MAD)

^{۴۰} Konno

^{۴۱} Simaan

مدل Maxmin

یانگ در سال ۱۹۹۸ در [۳۶] برای مسأله انتخاب سبد سرمایه، از حداقل بازدهی که در دوره‌های گذشته حاصل شده است به عنوان تخمین ریسک استفاده می‌کنند. به عبارتی دیگر به دنبال حداقل کردن، حداکثر زیانی که در طول دوره ممکن است رخ دهد، هستیم. اکنون متغیر M_p را که نشان‌دهنده حداقل بازدهی بدست آمده از سبد سرمایه را به فرم $M_p = \min_t \sum_{i=1}^n r_{it}x_i$ معرفی می‌کنیم، بنابراین مدل به صورت زیر قابل بیان است:

$$\text{Maximize } Z_{mm} = M_p \quad (۱۳.۳)$$

subject to

$$\begin{cases} \sum_{i=1}^n r_{it}x_i \geq M_p & t = 1, 2, \dots, T, \\ \sum_{i=1}^n x_i \mu_i = R, \\ \sum_{i=1}^n x_i = 1, \\ x_i \geq 0, & i = 1, 2, \dots, n, \\ l^{M_p} \leq M_p \leq u^{M_p}. \end{cases} \quad (۱۴.۳)$$

مدل Minimax

کای^{۴۲} و همکاران در [۴۸]، مدلی جدید ریسک را بر پایه MAD ارائه کردند، قبل از معرفی مدل یک تعریف ارائه می‌دهیم:

تعریف ۱.۳.۳. مدل L_∞ ، ماکسیمم انحراف مطلق ریسک به صورت زیر است:

$$L_\infty(x) = \max_{1 \leq j \leq n} E(|r_j x_j - \mu_j x_j|).$$

بنابراین مدل $Minimax$ به صورت زیر بیان می‌شود:

$$\text{Minimize } Z_{mm} = l_\infty(x) = \max_{1 \leq j \leq n} E(|r_j x_j - \mu_j x_j|) \quad (۱۵.۳)$$

^{۴۲}Cai

subject to

$$\begin{cases} \sum_{j=1}^n x_j \mu_j = R, \\ \sum_{j=1}^n x_j = 1, \\ x_j \geq 0, \quad j = 1, 2, \dots, n. \end{cases} \quad (16.3)$$

این مدل را می‌توان به‌فرم خطی زیر تبدیل کرد: [۴۹]

$$\text{Minimize} \quad y \quad (17.3)$$

subject to

$$\begin{cases} q_j x_j \leq y, \\ \sum_{j=1}^n x_j \mu_j = R, \\ \sum_{j=1}^n x_j = 1, \\ x_j \geq 0, \quad j = 1, 2, \dots, n. \end{cases} \quad (18.3)$$

که در آن $(j = 1, 2, \dots, n)$ برای $y = \max q_j$ و $q_j = E(|r_j - \mu_j|)$ است.

همچنین، جایگزین دیگر تابع ریسک L_∞ را که توسط تنو^{۴۳} و یانگ در [۵۰] پیشنهاد شد، به‌صورت

زیر است: که در آن متغیرهای تصادفی و μ_{jt} مقادیر مورد انتظار از r_{jt} برای $(i = 1, 2, \dots, n)$ و

$(t = 1, 2, \dots, T)$ هستند. این تابع توسعه یافته L_∞ است. بنابراین مدل دیگر از Minimax به‌صورت

زیر ارائه شده است:

$$\text{Minimize} \quad Z_{mm} = H_\infty^T(x) = \frac{1}{T} \sum_{t=1}^T \max_{1 \leq j \leq n} E(|r_{jt}x_j - \mu_{jt}x_j|) \quad (19.3)$$

subject to

$$\begin{cases} \sum_{j=1}^n x_j \mu_j = R, \\ \sum_{j=1}^n x_j = 1, \\ x_j \geq 0, \quad j = 1, 2, \dots, n. \end{cases} \quad (20.3)$$

این مدل را نیز، می‌توان به‌فرم خطی زیر تبدیل کرد: [۴۹]

$$\text{Minimize} \quad \frac{1}{T} \sum_{t=1}^T y_t \quad (21.3)$$

subject to

^{۴۳}Teo

$$\begin{cases} a_{jt}x_j \leq y_t, & i = 1, 2, \dots, n, t = 1, 2, \dots, T, \\ \sum_{j=1}^n x_j \mu_j = R, \\ \sum_{j=1}^n x_j = 1, \\ x_j \geq 0, & j = 1, 2, \dots, n. \end{cases} \quad (22.3)$$

که در آن $a_{jt} = E|r_{jt} - \mu_{jt}|$ برای $(i = 1, 2, \dots, n)(t = 1, 2, \dots, T)$ هستند. برای یافتن مدل‌های بیشتر مراجع [۳۲، ۵۱، ۵۲] را ببینید.

۴.۳ مدل شبکه عصبی

با توجه به تحلیل‌های بالا، واضح است که بیشتر مسائل بهینه‌سازی سبد سرمایه به صورت برنامه‌ریزی خطی و درجه دوم هستند، بنابراین ما فرم کلی مسائل برنامه‌ریزی درجه دوم را به صورت زیر در نظر می‌گیریم:

$$\text{Minimize } \frac{1}{2}x^T Qx + D^T x \quad (23.3)$$

subject to

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

که در آن $Q \in \mathbb{R}^{n \times n}$ ماتریس متقارن و نیمه معین مثبت، $x \in \mathbb{R}^n$ ، $E \in \mathbb{R}^{l \times n}$ ، $b \in \mathbb{R}^n$ ، $A \in \mathbb{R}^{m \times n}$ و $rank(A, E) = m + l$. با استفاده از تکنیک‌های استاندارد بهینه‌سازی، مسئله (۲۳.۳)–(۲۴.۳) را به سیستم دینامیکی غیرخطی تبدیل می‌کنیم.

با توجه به قضایای ۱۲.۲.۱ و ۱۳.۲.۱، جواب بهینه (۲۳.۳)–(۲۴.۳) است اگر و تنها اگر بردارهای $u^* \in \mathbb{R}^m$ و $v^* \in \mathbb{R}^l$ موجود باشند به طوری که $(x^{*T}, u^{*T}, v^{*T})^T$ در سیستم $K.K.T$ زیر صدق کند:

$$\begin{cases} Qx^* + D^T + A^T u^* + E^T v^* = 0 \\ u^* \geq 0, Ax^* - b \leq 0, u^{*T}(Ax^* - b) = 0, \\ Ex^* - f = 0. \end{cases} \quad (25.3)$$

x^* را یک نقطه $K.K.T$ برای (۲۳.۳) و (۲۴.۳) می‌گویند. بنابراین x^* جواب بهینه (۲۳.۳)–(۲۴.۳) است

اگر و تنها اگر x^* یک نقطه $K.K.T$ ، (۲۳.۳)–(۲۴.۳) باشد.

اکنون فرض کنید که $x(\cdot)$ ، $u(\cdot)$ و $v(\cdot)$ متغیرهای وابسته به زمان هستند. هدف بیان یک ساختار دینامیکی پیوسته‌ای است که به نقطه $K.K.T$ مسأله (۲۳.۳)–(۲۴.۳) هدایت شود. ما مدل شبکه عصبی بازگشتی را با معادله دینامیکی به صورت زیر بیان می‌کنیم:

$$\begin{cases} \frac{dx}{dt} = -(Qx + D + A^T(u + Ax - b)^+ + E^T v), \\ \frac{du}{dt} = (u + Ax - b)^+ - u, \\ \frac{dv}{dt} = Ex - f. \end{cases} \quad (26.3)$$

با نقطه آغازین $(x_0^T, u_0^T, v_0^T)^T$ ، که در آن

$$(u + Ax - b)^+ = ([(u + Ax - b)_1]^+, [(u + Ax - b)_2]^+, \dots, [(u + Ax - b)_m]^+),$$

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

برای سادگی، تعریف می‌کنیم $y = (x^T, u^T, v^T)^T \in \mathbb{R}^{n+m+l}$ ، D^* مجموعه نقاط بهینه از (۲۳.۳)–

(۲۴.۳) و دوگانش و

$$\zeta(y) = \begin{bmatrix} -(Qx + D + A^T(u + Ax - b)^+ + E^T v) \\ (u + Ax - b)^+ - u \\ Ex - f \end{bmatrix} \quad (27.3)$$

بنابراین شبکه عصبی (۲۶.۳) را می‌توان به صورت زیر نوشت:

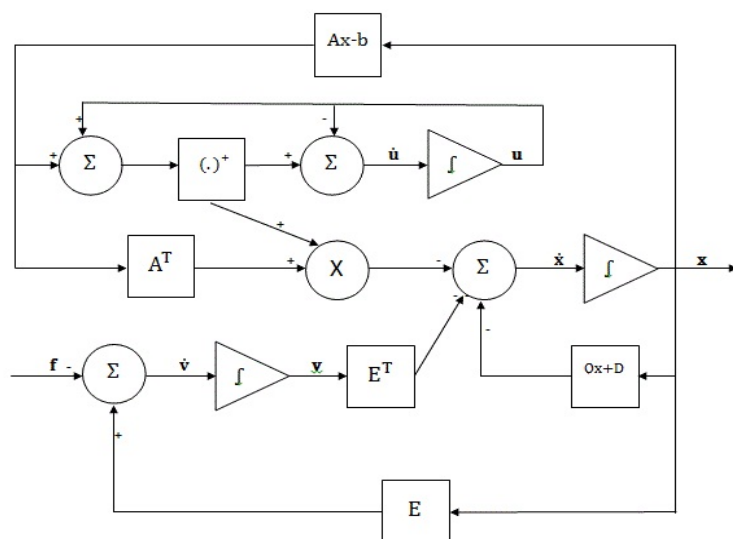
$$\frac{dy}{dt} = \kappa \zeta(y), \quad (28.3)$$

$$y(t_0) = y_0. \quad (29.3)$$

وقتی κ پارامتر اسکالر و بیان کننده نرخ همگرایی شبکه عصبی (۲۸.۳)–(۲۹.۳) است که برای سادگی تحلیل‌مان، κ را برابر ۱ در نظر می‌گیریم. شکل (۲.۳) بیانگر این است که مدل فوق روی سخت افزار چگونه اجرا می‌شود.

۵.۳ تحلیل پایداری و همگرایی

در این بخش پایداری و همگرایی شبکه عصبی پیشنهاد شده (۲۹.۳)–(۲۸.۳) بررسی می‌شود.



شکل ۲.۳: بلوک دیاگرام شبکه عصبی (۲۹.۳)–(۲۸.۳)

قضیه ۱.۵.۳. فرض کنید $y^* = (x^{*T}, u^{*T}, v^{*T})^T$ ، نقطه تعادل از شبکه عصبی (۲۹.۳)–(۲۸.۳) باشد.

آنگاه $x^* \in \mathbb{R}^n$ از مسأله $K.K.T$ (۲۳.۳)–(۲۴.۳) است و برعکس، اگر $x^* \in \mathbb{R}^n$ جواب بهینه مسأله (۲۳.۳)–

(۲۴.۳) باشد، آنگاه بردارهای $u^* \in \mathbb{R}^m$ و $v^* \in \mathbb{R}^l$ وجود دارد به طوری که $y^* = (x^{*T}, u^{*T}, v^{*T})^T$ یک

نقطه تعادل از شبکه عصبی پیشنهاد شده (۲۹.۳)–(۲۸.۳) است.

اثبات. فرض کنید $y^* = (x^{*T}, u^{*T}, v^{*T})^T$ نقطه تعادل شبکه عصبی فوق باشد، آنگاه $\frac{dx^*}{dt} = 0, \frac{du^*}{dt} = 0$

و $\frac{dv^*}{dt} = 0$ بنابراین

$$\begin{cases} Qx^* + D + A^T(u^* + Ax^* - b)^+ + E^T v^* = 0, \\ (u^* + Ax^* - b)^+ - u^* = 0, \\ Ex^* - f = 0. \end{cases} \quad (30.3)$$

از طرفی، داریم $(u^* + Ax^* - b)^+ = u^*$ اگر و تنها اگر

$$u^* \geq 0, \quad Ax^* - b \leq 0, \quad u^{*T}(Ax^* - b) = 0. \quad (31.3)$$

با جایگذاری (۳۱.۳) در (۳۰.۳) داریم:

$$Qx^* + D + A^T u^* + E^T v^* = 0 \quad (32.3)$$

از (۳۴.۳) و (۳۰.۳)، به نظر می‌رسد که $y^* = (x^{*T}, u^{*T}, v^{*T})^T$ در شرایط $K.K.T$ صدق می‌کند و این اثبات را کامل می‌کند. $\square$

لم ۲.۵.۳. برای هر نقطه آغازین $y(t_0) = (x(t_0)^T, u(t_0)^T, v(t_0)^T)^T$ ، سیستم (۲۸.۳)–(۲۹.۳) دارای جواب یکتای پیوسته $y(t) = (x(t)^T, u(t)^T, v(t)^T)^T$ است.

اثبات. واضح است که $Qx + D + A^T(u + Ax - b)^+ + E^T v$ ، $(u + Ax - b)^+ - u$ و $Ex - f$ لیپ شیتز محلی پیوسته روی مجموعه باز $D \subseteq \mathbb{R}^{n+m+l}$ هستند، با توجه به قضیه (۸.۳.۱)، شبکه عصبی (۲۸.۳)–(۲۹.۳) دارای جواب یکتای پیوسته $y(t)$ برای $t \in [t_0, \tau)$ وقتی $\tau \rightarrow \infty$ است. $\square$

لم ۳.۵.۳. ماتریس ژاکوبی $\nabla \zeta(y)$ از نگاشت ζ تعریف شده در (۲۷.۳)، ماتریس نیمه معین منفی است. اثبات. بدون کاستن از کلیت مسأله، فرض می‌کنیم $0 < p < m$ موجود است به طوری که:

$$(u + Ax - b)^+ = \left((u + Ax - b)_1, (u + Ax - b)_2, \dots, (u + Ax - b)_p, \underbrace{0, 0, \dots, 0}_{m-p} \right).$$

با محاسبات ساده می‌توان نشان داد:

$$\nabla \zeta(y) = \begin{bmatrix} -Q_{n \times n} - (A^p)^T A^p & -(A^p)^T & -E^T \\ A^p & \mathcal{T}_{m \times m} & O_{m \times l} \\ E & O_{l \times m} & O_{l \times l} \end{bmatrix},$$

که در آن

$$A^p = \begin{bmatrix} U_{p \times n} \\ O_{(m-p) \times n} \end{bmatrix} = \begin{bmatrix} A_1 \\ A_2 \\ \dots \\ \dots \\ A_p \\ O_{1 \times n} \\ O_{1 \times n} \\ \dots \\ \dots \\ O_{1 \times n} \end{bmatrix},$$

۹

$$\mathcal{T}_{m \times m} = \begin{bmatrix} O_{p \times p} & O_{p \times (m-p)} \\ O_{(m-p) \times p} & -I_{(m-p) \times (m-p)} \end{bmatrix}.$$

که O ماتریس صفر است و $A^p (A^p)^T$ ماتریس نیمه معین مثبت است. و از طرفی ماتریس $\mathcal{T}_{m \times m}$ ، نیمه معین منفی است. بنابراین ماتریس ژاکوبی $\nabla \zeta(y)$ ، نیمه معین منفی است.

اگر $p = m$ یعنی $(u + Ax - b)^+ = ((u + Ax - b)_1, (u + Ax - b)_2, \dots, (u + Ax - b)_m)$ ، آنگاه

$$\nabla \zeta(y) = \begin{bmatrix} -Q_{n \times n} - A^T A & -A^T & -E^T \\ A & O_{m \times m} & O_{m \times l} \\ E & O_{l \times m} & O_{l \times l} \end{bmatrix}.$$

مشابه حالت قبلی، می‌توان نشان داد که $\nabla \zeta(y)$ ماتریس نیمه معین منفی است. سرانجام، اگر $p = 0$

یعنی $(u + Ax - b)^+ = (\underbrace{0, 0, \dots, 0}_m)$ ، آنگاه داریم:

$$\nabla \zeta(y) = \begin{bmatrix} O_{n \times n} & O_{n \times m} & -E^T \\ O_{m \times n} & -I_{m \times m} & O_{m \times l} \\ E & O_{l \times m} & O_{l \times l} \end{bmatrix}.$$

در این مورد هم، به‌سادگی می‌توان نشان داد که $\nabla \zeta(y)$ ماتریس نیمه معین منفی است، این اثبات را

□

کامل می‌کند.

قضیه ۴.۵.۳. اگر فرض لم (۳.۵.۳) برقرار باشد، آنگاه شبکه عصبی پیشنهاد شده در (۲۹.۳)–(۲۸.۳)،

پایدار مجانبی به مفهوم لیاپانوف و همگرایی سراسری به $y^* = (x^{*T}, u^{*T}, v^{*T})^T$ است که در آن x^*

جواب بهینه (۲۳.۳)–(۲۴.۳) است.

اثبات. تابع لیاپانوف $(E : \mathbb{R}^{m+n+l} \rightarrow \mathbb{R})$ را به‌صورت زیر را در نظر بگیرید:

$$E(y) = \|\zeta(y)\|^2 + \frac{1}{2} \|y - y^*\|^2. \quad (33.3)$$

از (۲۷.۳) داریم:

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

با محاسبه مشتق $E(y(t))$ از شبکه عصبی (۲۹.۳)–(۲۸.۳) ، داریم:

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

با به کار بردن لم (۳.۵.۳) ، داریم

$$\zeta^T(y)(\nabla \zeta(y)^T + \nabla \zeta(y)) \zeta(y) \leq 0, \quad \forall y \neq y^*. \quad (۳۴.۳)$$

همچنین با استفاده از تعریف (۶.۳.۱) و قضیه (۷.۳.۱) داریم:

$$(y - y^*)^T (\zeta(y) - \zeta(y^*)) = (y - y^*)^T \zeta(y) \leq 0, \quad \forall y \neq y^*.$$

بنابراین

$$\frac{dE(y(t))}{dt} \leq 0. \quad (۳۵.۳)$$

و این بدان معنی است که شبکه عصبی (۲۹.۳)–(۲۸.۳) ، پایدار مجانبی به مفهوم لیاپانوف است. در ادامه

چون

$$E(y) \geq \frac{1}{4} \|y - y^*\|^2 \quad (۳۶.۳)$$

زیر دنباله‌های همگرای

$$\{(x(t_k)^T, u(t_k)^T, v(t_k)^T)^T | t_0 < t_1 < \dots < t_k < t_{k+1}\}$$

وقتی $t_k \rightarrow \infty$ از $k \rightarrow \infty$ وجود دارد، به‌طوری‌که

$$\lim_{k \rightarrow \infty} (x(t_k)^T, u(t_k)^T, v(t_k)^T)^T = (\bar{x}^T, \bar{u}^T, \bar{v}^T)^T,$$

وقتی $(\bar{x}^T, \bar{u}^T, \bar{v}^T)^T$ در رابطه زیر صدق می‌کند:

$$\frac{dE(y(t))}{dt} = 0.$$

این نشان می‌دهد که $(\bar{x}^T, \bar{u}^T, \bar{v}^T)^T$ ، یک نقطه ω -حدی از $\{(x(t)^T, u(t)^T, v(t)^T)^T | t \geq t_0\}$ است (مرجع [۵۴]، صفحه ۲۲۵ و پیوست [۵۵] را ببینید). با استفاده از قضیه مجموعه ناوردای لازال در (۴.۴.۱)، در می‌یابیم $\{(x(t)^T, u(t)^T, v(t)^T)^T \rightarrow M\}$ ، وقتی که $t \rightarrow \infty$ و M ، بزرگترین مجموعه ناوردا در $\{(x(t)^T, u(t)^T, v(t)^T)^T | \frac{dE(y(t))}{dt} = 0\}$ است. از (۲۹.۳)–(۲۸.۳) و (۳۵.۳) داریم،

$$M \subseteq K \subseteq D^* \text{ با } (\bar{x}^T, \bar{u}^T, \bar{v}^T)^T \in D^* \text{ بنابراین } \frac{dx}{dt} = 0, \frac{du}{dt} = 0, \frac{dv}{dt} = 0 \Leftrightarrow \frac{dE(y(t))}{dt} = 0.$$

با جایگذاری $x^* = \bar{x}$ ، $u^* = \bar{u}$ و $v^* = \bar{v}$ در (۳۳.۳)، ما تابع لیاپانوف دیگری، تعریف می‌کنیم:

$$\bar{E}(y) = \|\zeta(y)\|^2 + \frac{1}{\gamma} \|y - \bar{y}\|^2. \quad (۳۷.۳)$$

آنگاه $\bar{E}(y)$ به‌طور پیوسته مشتق پذیر و $\bar{E}(\bar{y}) = 0$. توجه داریم که

$$\lim_{k \rightarrow \infty} (x(t_k)^T, u(t_k)^T, v(t_k)^T)^T = (\bar{x}^T, \bar{u}^T, \bar{v}^T)^T,$$

از اینرو داریم $\lim_{k \rightarrow \infty} \bar{E}(x(t_k)^T, u(t_k)^T, v(t_k)^T)^T = \bar{E}(\bar{x}, \bar{u}, \bar{v})$. پس، $\forall \epsilon > 0$ ، $q > 0$ موجود است به‌طوری‌که برای هر $t \geq t_q$ داریم $\bar{E}(y(t)) < \epsilon$. مشابه ما می‌توانیم بدست آوریم $\frac{d\bar{E}(y(t))}{dt} \leq 0$. برای $t \geq t_q$ داریم،

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

در نتیجه، $\lim_{t \rightarrow \infty} \|y(t) - \bar{y}\| = 0$ و $\lim_{t \rightarrow \infty} y(t) = \bar{y}$. بنابراین، شبکه عصبی پیشنهاد شده (۲۹.۳)–(۲۸.۳) همگرای سراسری به نقطه تعادل $\bar{y} = (\bar{x}^T, \bar{u}^T, \bar{v}^T)^T$ است که در آن، $\bar{x}$ جواب بهینه (۲۳.۳)–(۲۴.۳) است. $\square$

از قضیه (۴.۵.۳) می‌توانیم گزاره زیر را نتیجه بگیریم.

گزاره ۵.۵.۳. اگر فرضیات لم (۳.۵.۳) برقرار باشد و $D^* = \{(x^{*T}, u^{*T}, v^{*T})^T\}$ ، آنگاه شبکه عصبی (۲۹.۳)–(۲۸.۳) برای حل کردن (۲۳.۳)–(۲۴.۳)، پایدار مجانبی سراسری به نقطه تعادل یکتای $y^* = (x^{*T}, u^{*T}, v^{*T})^T$ است.

۶.۳ مثال‌های عددی

به منظور نشان دادن کارایی و تأثیر شبکه عصبی پیشنهاد شده، در این بخش ما به ذکر ۷ مثال از ۲ مسأله مختلف می‌پردازیم. برای برخی مثال‌ها، عملکرد شبکه عصبی را برای مقادیر مختلف $y(0)$ بررسی می‌کنیم، همچنین ما مقایسه عملکرد عددی از شبکه عصبی را با مقادیر گوناگون κ و مقادیر مختلف $\|x(t) - x^*\|^2$ با نقطه آغازین مشخص y داریم. نتایج شبیه‌سازی در پکیج *ode45s* نرم افزار *Matlab* انجام شده و جواب واقعی مسائل نیز توسط نرم افزار *Lingo ۱۱* محاسبه شده است.

مسأله I

مدیر سبد سرمایه مؤسسه سیگما، قصد دارد یک سبد سرمایه بهینه برای یک مشتری انتخاب کند. به جای در نظر گرفتن تعداد زیاد دارایی‌ها در سبد، برای سادگی فرض می‌کنیم ۵ سهم (A, B, C, D, E) از بازار بورس استکهلم در اختیار باشد. داده‌های تاریخی ۱۲ ماهه، ماتریس واریانس و انحرافات مطلق هر ماه در [۴۱] داده شده است. مشتری خواهان سرمایه گذاری $SEK ۱۰۰۰۰۰$ ، حداقل بازدهی دریافتی $(SEK ۳۰۰۰)$ ۳٪ است و در نظر دارد، در هر سهم حداکثر ۷۵٪ سرمایه گذاری داشته باشد. بنابراین مسأله فوق، حداقل کردن ریسک سبد، با در نظر گرفتن خواسته‌های مشتری و اجازه نداشتن فروش استقراضی در هر سهام است.

مثال ۱.۶.۳. مسأله برنامه‌ریزی درجه دوم به فرم (۱.۳)–(۲.۳) را در نظر بگیرید:

$$\left\{ \begin{array}{l} \text{Minimize} \\ 0.00146739x_A^2 + 2[0.00026179x_Ax_B + 0.00052103x_Ax_C + 0.00017382x_Ax_D \\ + 0.00043272x_Ax_E - 0.00001133x_Bx_C + 0.00059010x_Bx_D + 0.00048699x_Bx_E \\ - 0.00002990x_Cx_D - 0.00017064x_Cx_E + 0.00055264x_Dx_E] + 0.00111640x_B^2 \\ + 0.00105790x_C^2 + 0.00108543x_D^2 + 0.00100371x_E^2 \\ \text{s.t.} \\ 0.0207x_A + 0.0316x_B + 0.0323x_C + 0.0337x_D + 0.0376x_E \geq 0.03, \\ x_A + x_B + x_C + x_D + x_E = 1, \\ 0 \leq x_A, x_B, x_C, x_D, x_E \leq 0.75. \end{array} \right.$$

جواب بهینه این مسأله $x^* = (0, 0/144, 0/417, 0/127, 0/312)^T$. ما برای حل این مسأله شبکه

عصبی پیشنهاد شده در (۲۸.۳)-(۲۹.۳) را به کار می‌بریم. شکل (۳.۳) نشان می‌دهد که مسیرهای شبکه

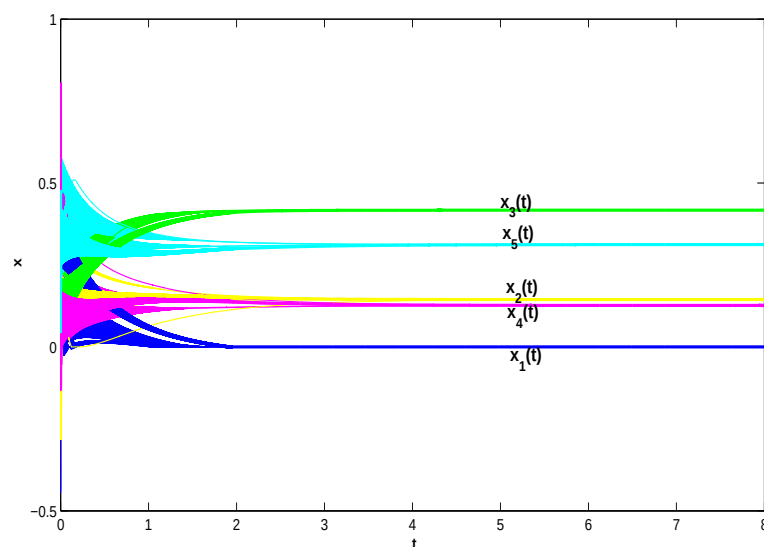
عصبی پیشنهادی با ۱۰ نقطه آغازین متفاوت، همگرا به جواب بهینه مسأله است.

مثال ۲.۶.۳. مسأله برنامه‌ریزی خطی به فرم (۱۱.۳)-(۱۲.۳) را در نظر بگیرید:

$$\left\{ \begin{array}{l} \text{Minimize} \quad \frac{1}{2} \sum_{t=1}^{12} y_t \\ \text{s.t.} \\ y_1 + 0/0333x_A + 0/0004x_B + 0/0083x_C + 0/0043x_D + 0/0114x_E \geq 0, \\ y_1 - 0/0333x_A - 0/0004x_B - 0/0083x_C - 0/0043x_D - 0/0114x_E \geq 0, \\ y_2 + 0/0243x_A + 0/0234x_B + 0/0203x_C + 0/0283x_D + 0/0294x_E \geq 0, \\ y_2 - 0/0243x_A - 0/0234x_B - 0/0203x_C - 0/0283x_D - 0/0294x_E \geq 0, \\ y_3 + 0/0507x_A + 0/0676x_B + 0/0123x_C + 0/0707x_D + 0/0766x_E \geq 0, \\ y_3 - 0/0507x_A - 0/0676x_B - 0/0123x_C - 0/0707x_D - 0/0766x_E \geq 0, \\ y_4 + 0/0387x_A + 0/0204x_B + 0/0087x_C + 0/0163x_D + 0/0134x_E \geq 0, \\ y_4 - 0/0387x_A - 0/0204x_B - 0/0087x_C - 0/0163x_D - 0/0134x_E \geq 0, \\ y_5 + 0/0223x_A + 0/0154x_B + 0/0173x_C + 0/0313x_D + 0/0114x_E \geq 0, \\ y_5 - 0/0223x_A - 0/0154x_B - 0/0173x_C - 0/0313x_D - 0/0114x_E \geq 0, \\ y_6 + 0/0263x_A + 0/0024x_B + 0/0003x_C + 0/0767x_D + 0/0006x_E \geq 0, \\ y_6 - 0/0263x_A - 0/0024x_B - 0/0003x_C - 0/0767x_D - 0/0006x_E \geq 0, \\ y_7 + 0/0334x_A + 0/0314x_B + 0/0013x_C + 0/0283x_D + 0/0174x_E \geq 0, \\ y_7 - 0/0334x_A - 0/0314x_B - 0/0013x_C - 0/0283x_D - 0/0174x_E \geq 0, \\ y_8 + 0/0153x_A + 0/0164x_B + 0/0113x_C + 0/0003x_D + 0/0126x_E \geq 0, \\ y_8 - 0/0153x_A - 0/0164x_B - 0/0113x_C - 0/0003x_D - 0/0126x_E \geq 0, \\ y_9 + 0/0597x_A + 0/0066x_B + 0/0747x_C + 0/0013x_D + 0/0144x_E \geq 0, \\ y_9 - 0/0597x_A - 0/0066x_B - 0/0747x_C - 0/0013x_D - 0/0144x_E \geq 0, \\ y_{10} + 0/0637x_A + 0/0084x_B + 0/0113x_C + 0/0223x_D + 0/0176x_E \geq 0, \\ y_{10} - 0/0637x_A - 0/0084x_B - 0/0113x_C - 0/0223x_D - 0/0176x_E \geq 0, \\ y_{11} + 0/0253x_A + 0/0044x_B + 0/0043x_C + 0/0233x_D + 0/0074x_E \geq 0, \\ y_{11} - 0/0253x_A - 0/0044x_B - 0/0043x_C - 0/0233x_D - 0/0074x_E \geq 0, \\ y_{12} + 0/0313x_A + 0/0486x_B + 0/0037x_C + 0/0087x_D + 0/0024x_E \geq 0, \\ y_{12} - 0/0313x_A - 0/0486x_B - 0/0037x_C - 0/0087x_D - 0/0024x_E \geq 0, \\ 0/0207x_A + 0/0316x_B + 0/0323x_C + 0/0337x_D + 0/0376x_E \geq 0/03, \\ x_A + x_B + x_C + x_D + x_E = 1, \\ 0 \leq x_A, x_B, x_C, x_D, x_E \leq 0/75, \\ y_t \geq 0, t = 1, \dots, 12. \end{array} \right.$$

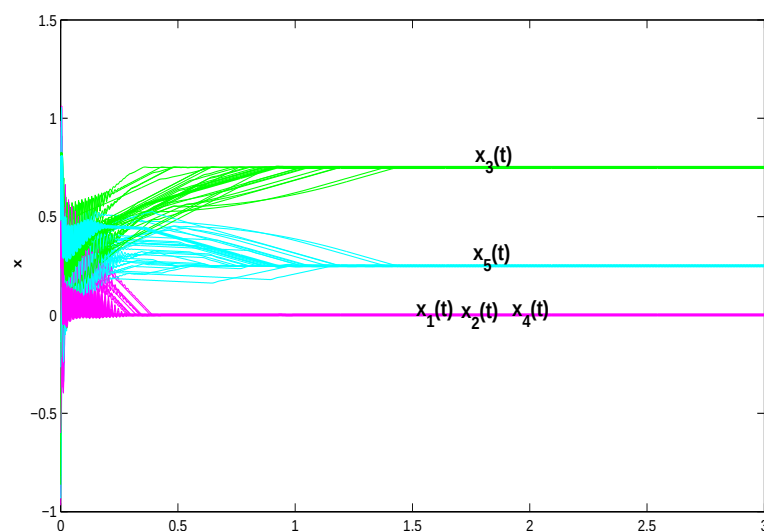
جواب بهینه این مسأله $x^* = (0, 0, 0/75, 0, 0/25)^T$. ما برای حل این مسأله شبکه عصبی پیشنهاد

شده در (۲۸.۳)-(۲۹.۳) را به کار می‌بریم. شکل (۴.۳) نشان می‌دهد که مسیرهای شبکه عصبی پیشنهادی

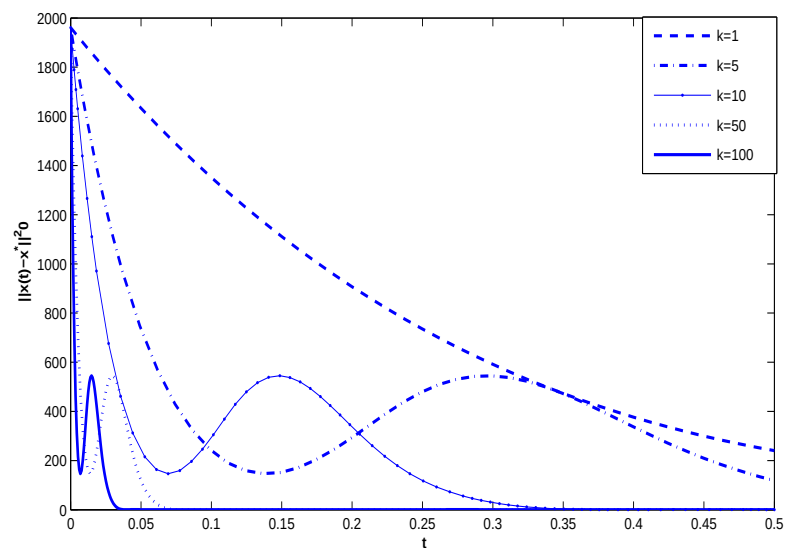


شکل ۳.۳: رفتار گذرا شبکه عصبی (۲۸.۳)–(۲۹.۳) با ۱۰ نقطه آغازین دلخواه در مثال ۱.۶.۳.

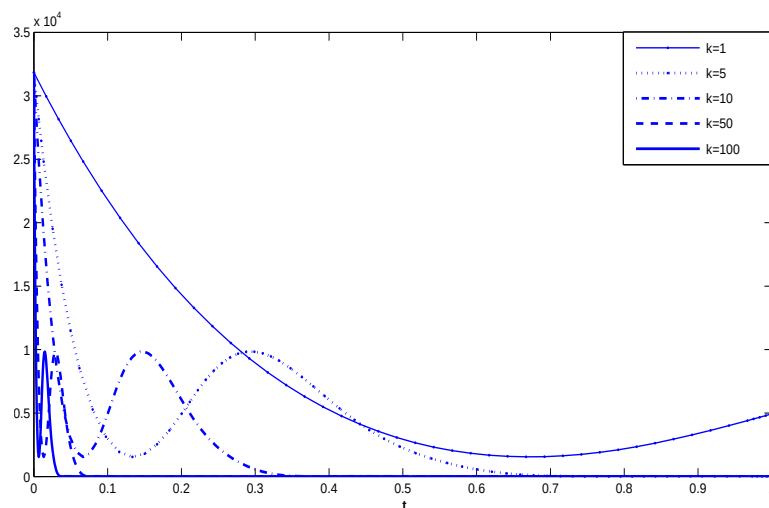
با ۳۰ نقطه آغازین متفاوت، همگرا به جواب بهینه مسأله است. ما همچنین تأثیر پارامتر κ از شبکه عصبی (۲۸.۳)–(۲۹.۳) روی مقدار $\|x(t) - x^*\|^2$ را در شکل (۵.۳) بررسی کردیم. واضح است که κ بزرگتر، نرخ همگرایی بهتری از $\|x(t) - x^*\|^2$ نتیجه می‌دهد.



شکل ۴.۳: رفتار گذرا شبکه عصبی $(28.3)-(29.3)$ با ۳۰ نقطه آغازین دلخواه در مثال ۲.۶.۳.



شکل ۵.۳: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$ در مثال ۲.۶.۳.



شکل ۶.۳: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$ در مثال ۳.۶.۳.

مثال ۳.۶.۳. مسأله برنامه‌ریزی خطی به فرم (۱۳.۳)-(۱۴.۳) را در نظر بگیرید:

$$\left\{ \begin{array}{ll} \text{Maximize} & z \\ \text{s.t.} & \\ & 0.054x_A + 0.032x_B + 0.064x_C + 0.038x_D + 0.049x_E \geq z, \\ & 0.045x_A + 0.055x_B + 0.056x_C + 0.062x_D + 0.067x_E \geq z, \\ & -0.030x_A - 0.036x_B - 0.048x_C - 0.037x_D - 0.039x_E \geq z, \\ & -0.018x_A + 0.052x_B + 0.007x_C + 0.050x_D + 0.051x_E \geq z, \\ & 0.043x_A + 0.047x_B + 0.053x_C + 0.065x_D + 0.049x_E \geq z, \\ & 0.047x_A + 0.034x_B + 0.036x_C - 0.043x_D + 0.037x_E \geq z, \\ & 0.055x_A + 0.063x_B + 0.017x_C + 0.062x_D + 0.055x_E \geq z, \\ & 0.036x_A + 0.048x_B + 0.047x_C + 0.034x_D + 0.025x_E \geq z, \\ & -0.039x_A + 0.025x_B - 0.059x_C + 0.035x_D + 0.052x_E \geq z, \\ & -0.043x_A + 0.040x_B + 0.047x_C + 0.056x_D + 0.020x_E \geq z, \\ & 0.046x_A + 0.036x_B + 0.040x_C + 0.057x_D + 0.045x_E \geq z, \\ & 0.052x_A - 0.017x_B + 0.032x_C + 0.025x_D + 0.040x_E \geq z, \\ & 0.0207x_A + 0.0316x_B + 0.0323x_C + 0.0337x_D + 0.0376x_E \geq 0.03, \\ & x_A + x_B + x_C + x_D + x_E = 1, \\ & 0 \leq x_A, x_B, x_C, x_D, x_E \leq 0.75, \\ & z \geq 0. \end{array} \right.$$

جواب بهینه این مسأله $x^* = (0, 0, 0.4596, 0, 0.5404)^T$ است. تأثیر پارامتر κ شبکه فوق روی

$\|x(t) - x^*\|^2$ در شکل (۶.۳) نشان داده شده است.

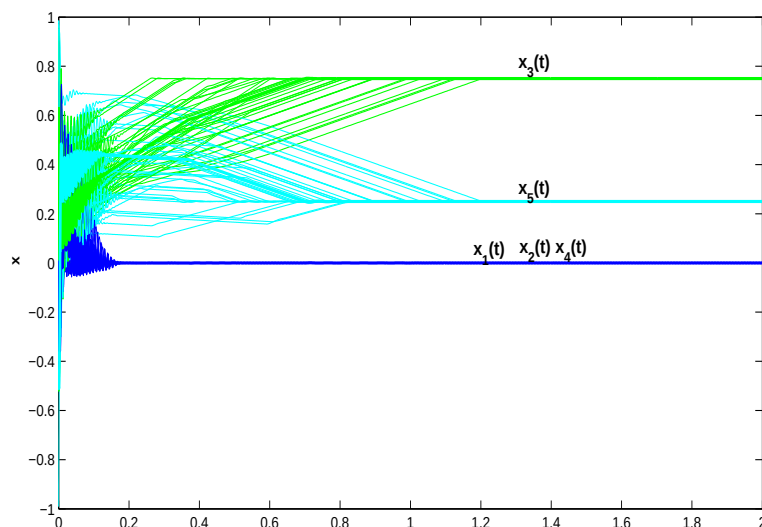
مثال ۴.۶.۳. مسأله برنامه‌ریزی خطی به فرم (۲۱.۳)-(۲۲.۳) را در نظر بگیرید:

$$\left\{ \begin{array}{l} \text{Minimize} \quad \frac{1}{12} \sum_{t=1}^{12} y_t \\ \text{s.t.} \\ 0.333x_A + 0.004x_B + 0.083x_C + 0.043x_D + 0.114x_E \leq y_1, \\ 0.243x_A + 0.234x_B + 0.203x_C + 0.283x_D + 0.294x_E \leq y_2, \\ 0.507x_A + 0.676x_B + 0.123x_C + 0.707x_D + 0.766x_E \leq y_3, \\ 0.387x_A + 0.204x_B + 0.087x_C + 0.163x_D + 0.134x_E \leq y_4, \\ 0.223x_A + 0.154x_B + 0.173x_C + 0.313x_D + 0.114x_E \leq y_5, \\ 0.263x_A + 0.002x_B + 0.003x_C + 0.767x_D + 0.006x_E \leq y_6, \\ 0.334x_A + 0.314x_B + 0.013x_C + 0.283x_D + 0.174x_E \leq y_7, \\ 0.153x_A + 0.164x_B + 0.113x_C + 0.003x_D + 0.126x_E \leq y_8, \\ 0.597x_A + 0.006x_B + 0.747x_C + 0.013x_D + 0.144x_E \leq y_9, \\ 0.637x_A + 0.008x_B + 0.113x_C + 0.223x_D + 0.176x_E \leq y_{10}, \\ 0.253x_A + 0.004x_B + 0.043x_C + 0.233x_D + 0.074x_E \leq y_{11}, \\ 0.313x_A + 0.486x_B + 0.037x_C + 0.087x_D + 0.024x_E \leq y_{12}, \\ 0.207x_A + 0.316x_B + 0.323x_C + 0.337x_D + 0.376x_E \leq 0.3, \\ x_A + x_B + x_C + x_D + x_E = 1, \\ 0 \leq x_A, x_B, x_C, x_D, x_E \leq 0.75, \\ y_t \geq 0, \quad t = 1, \dots, 12. \end{array} \right.$$

جواب بهینه این مسأله $x^* = (0, 0, 0.75, 0, 0.25)^T$. ما برای حل این مسأله از مدل شبکه عصبی

پیشنهاد شده در (۲۸.۳)-(۲۹.۳) استفاده می‌کنیم. شکل (۷.۳) نشان می‌دهد که مسیرهای شبکه عصبی

پیشنهادی با ۴۰ نقطه آغازین متفاوت، همگرا به جواب بهینه مسأله است.



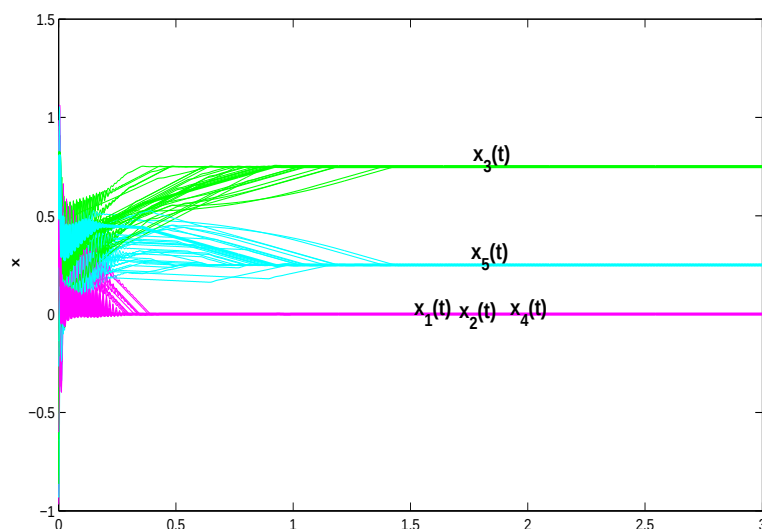
شکل ۷.۳: رفتار گذر شبکه عصبی (۲۸.۳)–(۲۹.۳) با ۴۰ نقطه آغازین دلخواه در مثال ۴.۶.۳.

مثال ۵.۶.۳. مسأله برنامه‌ریزی درجه دوم به فرم (۹.۳)–(۱۰.۳) را در نظر بگیرید:

$$\begin{cases}
 \text{Minimize} & \sum_{t=1}^{12} y_t^2 \\
 \text{s.t.} & \\
 & y_1 = \frac{1}{\sqrt{2}}(0.333x_A + 0.0004x_B + 0.0083x_C + 0.0043x_D + 0.0114x_E), \\
 & y_2 = \frac{1}{\sqrt{2}}(+0.243x_A + 0.234x_B + 0.203x_C + 0.283x_D + 0.294x_E), \\
 & y_3 = \frac{1}{\sqrt{2}}(-0.507x_A - 0.676x_B + 0.123x_C - 0.707x_D - 0.766x_E), \\
 & y_4 = \frac{1}{\sqrt{2}}(-0.387x_A + 0.204x_B - 0.087x_C + 0.163x_D + 0.134x_E), \\
 & y_5 = \frac{1}{\sqrt{2}}(+0.223x_A + 0.154x_B + 0.173x_C + 0.313x_D + 0.114x_E), \\
 & y_6 = \frac{1}{\sqrt{2}}(+0.263x_A + 0.0024x_B + 0.0003x_C - 0.767x_D - 0.0006x_E), \\
 & y_7 = \frac{1}{\sqrt{2}}(+0.334x_A + 0.314x_B + 0.0013x_C + 0.283x_D + 0.174x_E), \\
 & y_8 = \frac{1}{\sqrt{2}}(+0.153x_A + 0.164x_B + 0.113x_C + 0.0003x_D - 0.126x_E), \\
 & y_9 = \frac{1}{\sqrt{2}}(-0.597x_A - 0.0066x_B - 0.747x_C + 0.0013x_D + 0.144x_E), \\
 & y_{10} = \frac{1}{\sqrt{2}}(-0.637x_A + 0.0084x_B + 0.113x_C + 0.223x_D - 0.176x_E), \\
 & y_{11} = \frac{1}{\sqrt{2}}(+0.253x_A + 0.0044x_B + 0.0043x_C + 0.233x_D + 0.0074x_E), \\
 & y_{12} = \frac{1}{\sqrt{2}}(+0.313x_A - 0.486x_B - 0.0037x_C - 0.0087x_D + 0.0024x_E), \\
 & 0.207x_A + 0.316x_B + 0.323x_C + 0.337x_D + 0.376x_E \geq 0.3, \\
 & x_A + x_B + x_C + x_D + x_E = 1, \\
 & 0 \leq x_A, x_B, x_C, x_D, x_E \leq 0.75, \\
 & y_t \geq 0, \quad t = 1, \dots, 12.
 \end{cases}$$

جواب بهینه این مسأله $x^* = (0, 0, 0.5155, 0, 0.4845)^T$ است. ما برای حل این مسأله شبکه عصبی

پیشنهاد شده در (۲۸.۳)–(۲۹.۳) را به کار می‌بریم. شکل (۸.۳) نشان می‌دهد که مسیرهای شبکه عصبی



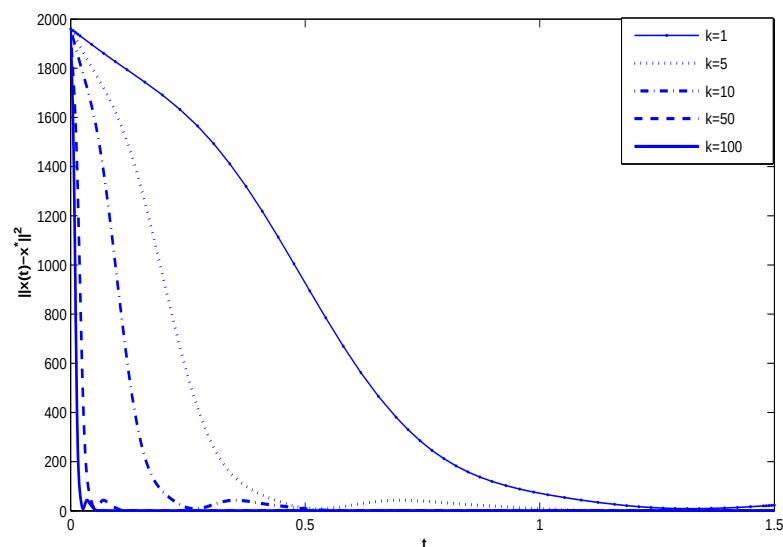
شکل ۸.۳: رفتار گذرا شبکه عصبی (۲۸.۳)–(۲۹.۳) با ۳۰ نقطه آغازین دلخواه در مثال ۵.۶.۳.

پیشنهادی با ۳۰ نقطه آغازین متفاوت، همگرا به جواب بهینه مسأله است. تأثیر پارامتر κ شبکه فوق روی $\|x(t) - x^*\|^2$ هم در شکل (۹.۳) نشان داده شده است.

مسأله II

مسأله فرضی دیگر به این صورت بیان می‌شود، فرض کنید سبد شامل سهام‌های US ، اوراق قرضه و وجوه نقد باشد. ما می‌خواهیم با استفاده از داده‌های بازدهی تاریخی برای این ۳ کلاس از دارایی‌ها، بازده مورد انتظار آتی آنها را تخمین بزنیم. از شاخص $S\&P 500$ برای بازدهی سهام‌ها، شاخص اوراق خزانه ۱۰ ساله^{۴۴} برای بازدهی اوراق استفاده کردیم و فرض می‌کنیم وجه ما در بازار پول سرمایه‌گذاری شده که بازدهی آن با شاخص $(1 - \text{day federal fund rate})$ داده شده است. سری زمانی سالانه برای بازدهی کل، میانگین و ماتریس کواریانس هر کلاس دارایی بین سال‌های ۱۹۶۰ و ۲۰۰۳ در [۵۷] داده شده

^{۴۴}10-year Treasury Bond Index



شکل ۹.۳: رفتار همگرایی از $\|x(t) - x^*\|^2$ در مثال ۵.۶.۳

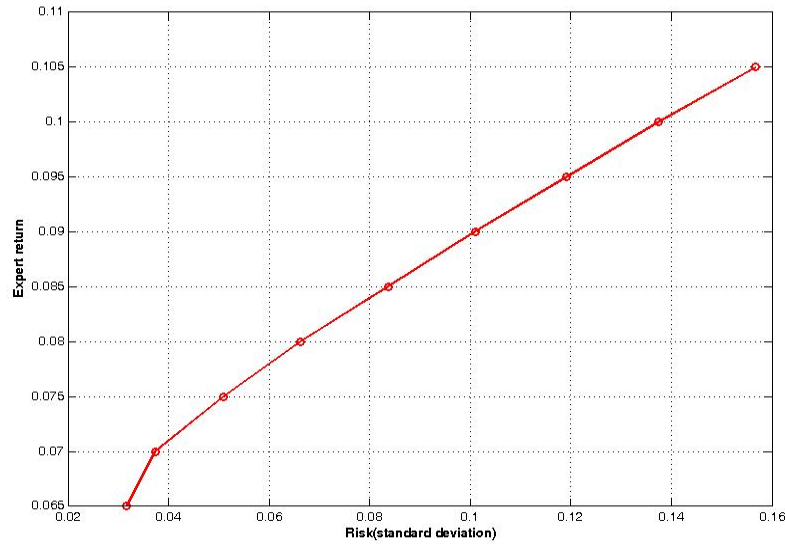
است.

مثال ۶.۶.۳. مسأله برنامه‌ریزی درجه دوم به فرم (۱.۳)–(۲.۳) را در نظر بگیرید:

$$\left\{ \begin{array}{l} \text{Minimize} \\ 0.02778x_S^2 + 0.01112x_B^2 + 0.00115x_M^2 + 2(0.00387x_Sx_B + 0.00021x_Sx_M - 0.00020x_Bx_M) \\ \text{s.t.} \\ 0.1073x_S + 0.0737x_B + 0.0627x_M \geq R, \\ x_S + x_B + x_M = 1, \\ x_S, x_B, x_M \geq 0. \end{array} \right.$$

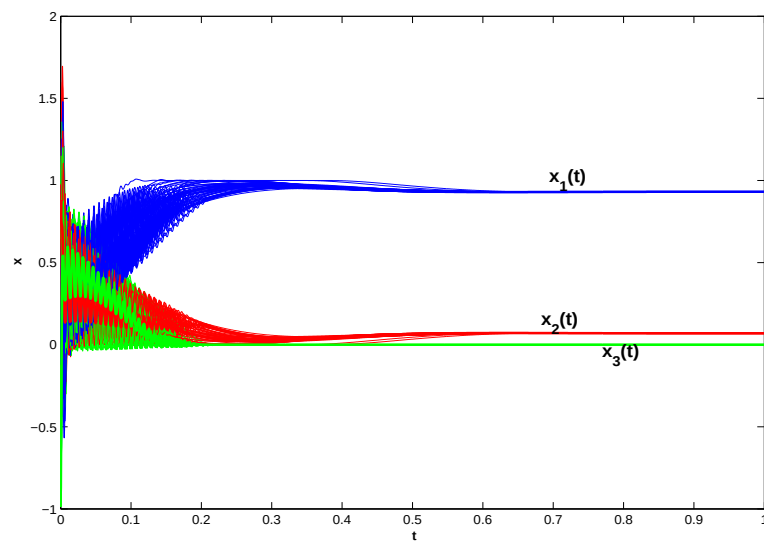
جدول ۱.۳: سبد سرمایه‌های کارا

MM	$Bonds$	$Stocks$	واریانس	نرخ بازدهی (R)
۰.۸۷	۰.۱۰	۰.۰۳	۰.۰۰۱۰	۰.۰۶۵
۰.۷۵	۰.۱۲	۰.۱۳	۰.۰۰۱۴	۰.۰۷۰
۰.۶۲	۰.۱۴	۰.۲۴	۰.۰۰۲۶	۰.۰۷۵
۰.۴۹	۰.۱۶	۰.۳۵	۰.۰۰۴۴	۰.۰۸۰
۰.۳۷	۰.۱۸	۰.۴۵	۰.۰۰۷۰	۰.۰۸۵
۰.۲۴	۰.۲۰	۰.۵۶	۰.۰۱۰۲	۰.۰۹۰
۰.۱۱	۰.۲۲	۰.۶۷	۰.۰۱۴۲	۰.۰۹۵
۰	۰.۲۲	۰.۷۸	۰.۰۱۸۹	۰.۱۰۰
۰	۰.۰۷	۰.۹۳	۰.۰۲۴۶	۰.۱۰۵

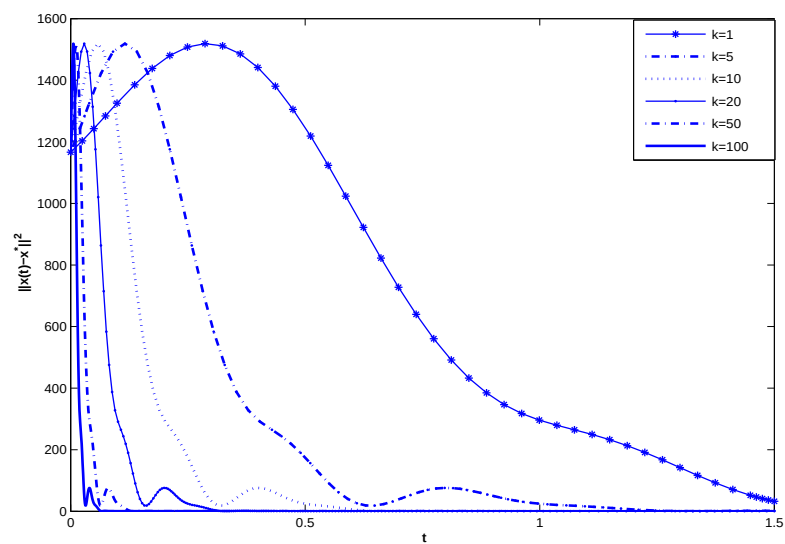


شکل ۱۰.۳: مرز کارایی سبد سرمایه‌های بدست آمده در مثال ۶.۶.۳.

به ازای $R = ۶/۵\%$ الی $R = ۱۱/۵\%$ با افزایش $R = ۵\%$ درصد و با حل این مدل به ازای مقادیر مختلف R ، سبد سرمایه‌های بهینه مطابق جدول (۱.۳) بدست می‌آید. شکل (۱۰.۳) مرز کارایی این سبد سرمایه‌ها را نشان می‌دهد. جواب بهینه این مسأله برای $R = ۱۰/۵\%$ ، $x^* = (۰, ۰/۰۶۸۵, ۰/۰۹۳۱۵)^T$ است. ما برای حل این مسأله شبکه عصبی پیشنهاد شده در (۲۸.۳)–(۲۹.۳) را به کار می‌بریم. شکل (۱۱.۳) نشان می‌دهد که مسیرهای شبکه عصبی پیشنهادی با ۵۰ نقطه آغازین متفاوت، همگرا به جواب بهینه مسأله است. تأثیر پارامتر κ شبکه فوق روی $\|x(t) - x^*\|^2$ با نقطه آغازین متفاوت، همگرا به جواب بهینه مسأله است. در شکل (۱۲.۳) نشان داده شده است.



شکل ۱۱.۳: رفتار گذرا شبکه عصبی (۲۸.۳)–(۲۹.۳) با ۳۰ نقطه آغازین دلخواه در مثال ۶.۶.۳.



شکل ۱۲.۳: رفتار همگرایی از $\|x(t) - x^*\|^2$ در مثال ۶.۶.۳.

مثال ۷.۶.۳. مسأله برنامه‌ریزی درجه دوم به فرم (۷.۳)-(۸.۳) را در نظر بگیرید:

$$\left\{ \begin{array}{l} \text{Minimize} \quad (y_1)^2 + (y_2)^2 + (y_3)^2 \\ \text{s.t.} \\ y_1 = 0.16667x_S, \\ y_2 = 0.02322x_S + 0.10286x_B, \\ y_3 = 0.00126x_S + 0.02468x_M - 0.00224x_B, \\ 0.1073x_S + 0.0737x_B + 0.0627x_M \geq 0.105, \\ x_S + x_B + x_M = 1, \\ x_S, x_B, x_M \geq 0, \\ y_1, y_2, y_3 \geq 0. \end{array} \right.$$

جواب بهینه این مسأله $x^* = (0.9315, 0.0685, 0)^T$ است. ما برای حل این مسأله شبکه عصبی پیشنهاد

شده در (۲۸.۳)-(۲۹.۳) را به کار می‌بریم.

مقایسه عملکرد مدل‌ها با یکدیگر

در مورد مقایسه عملکرد این ۵ مدل، با ۲ فرضیه در نظر می‌گیریم:

۱. سرمایه گذار، خواهان کمترین ریسک بدون توجه به میزان بازدهی دریافتی خود است. با توجه به

این شرط طبق جدول (۲.۳)، مدل MV بهترین انتخاب برای وی است.

۲. سرمایه گذار خواهان بیشترین بازدهی است، بنابراین مدل $Maximin$ را انتخاب می‌کند.

جدول ۲.۳: مقایسه مدل‌ها

$\sigma(\%)$	$E(r)(\%)$	وزن‌ها	مدل
۰.۰۳۸۳	۳.۴۰	$x^* = (0, 0.144, 0.417, 0.127, 0.312)$	MV
۰.۰۵۹۴	۳.۳۶	$x^* = (0, 0, 0.75, 0, 0.25)$	MAD
۰.۰۴۳۲	۳.۵۲	$x^* = (0, 0, 0.4596, 0, 0.5404)$	$Maximin$
۰.۰۵۹۴	۳.۳۶	$x^* = (0, 0, 0.75, 0, 0.25)$	$Minimax$
۰.۰۴۳۱	۳.۴۹	$x^* = (0, 0, 0.5155, 0, 0.4845)$	$Diago2$

فصل ۴

مسائل برنامه‌ریزی تصادفی

۱.۴ مقدمه

در این قطعه از پایان نامه ابتدا برنامه‌ریزی خطی تصادفی را معرفی کرده و مروری بر کارهای انجام شده داریم، سپس مسائل بهینه‌سازی با قیود تصادفی را مطرح می‌کنیم. یک مدل شبکه عصبی برای مسائل برنامه‌ریزی مخروط مرتبه دوم ارائه می‌دهیم و تحلیلی بر پایداری و همگرایی مدل پیشنهاد شده می‌آوریم. برای نشان دادن کارایی مدل شبکه عصبی پیشنهاد شده، ۴ مثال عددی از مسائل برنامه‌ریزی تصادفی که حالت خاصی از مسائل برنامه‌ریزی مخروط مرتبه دوم محسوب هستند، را حل می‌کنیم.

۱.۱.۴ برنامه‌ریزی خطی تصادفی

مسئله برنامه‌ریزی خطی تصادفی را می‌توان به صورت زیر فرمولبندی کرد:

$$\text{Minimize } f(x) = \sum_{j=1}^n c_j x_j, \quad (1.4)$$

subject to

$$\sum_{j=1}^n a_{ij} x_j \geq b_i, \quad i = 1, 2, \dots, m, \quad (2.4)$$

$$x_j \geq 0, \quad j = 1, 2, \dots, n. \quad (3.4)$$

وقتی که c_j و a_{ij} و b_i متغیرهای تصادفی با توزیعهای احتمال معلوم هستند. (برای سادگی، متغیرهای تصمیم x_j ، قطعی فرض می‌شوند)، مدل‌های برنامه‌ریزی تصادفی [۵۸، ۵۹، ۶۰] سبب ایجاد طرح‌هایی می‌شوند که در برابر زیان‌ها و عدم موفقیت محصول (نقص‌های فاجعه آمیز) بهتر می‌توانند مقاومت کنند. چنین مدل‌هایی برای کاربردهای مختلف شامل تولید انرژی الکتریکی^۱ [۶۱]، نقشه‌های مالی^۲ [۶۲، ۶۳، ۶۴]، ارتباط تلفنی^۳ [۶۵، ۶۶]، مدیریت زنجیره تأمین^۴ [۶۷]، صنایع نفت [۶۸]، تولید

^۱Electric power generation

^۲Financial planning

^۳Telecommunication network planning

^۴Supply chain management

کنندگان وسایل نقلیه [۶۹]، تأمین کنندگان الکتریسیته [۷۰، ۷۱]، محیط [۷۲]، حمل و نقل [۷۳، ۷۴]، ساختمان، انرژی، فرآورده‌های شیمیایی [۷۵]، هوا-فضا و سیستم نظامی [۷۶]^۵ به کار می‌روند.

مدل‌های برنامه‌ریزی تصادفی شامل متغیرهای تصمیم قابل پیش‌بینی^۶ و تطبیقی^۷ هستند. متغیرهای قابل پیش‌بینی مربوط به تصمیماتی می‌شوند که باید در زمان حال گرفته شوند و نمی‌توانند به پارامترهای تصادفی تشخیص جزئی و مشاهدات آینده وابسته باشند. متغیرهای تطابقی به تصمیمات آتی پس از مشاهده تمام یا برخی از پارامترهای تصادفی بستگی دارند. مدل‌های قابل پیش‌بینی می‌توانند در برنامه‌ریزی تصادفی با منابع^۸ و یا مدل‌های برنامه‌ریزی با قیود تصادفی^۹ فرمول بندی شوند.

در فرمول‌بندی برنامه‌ریزی با قیود تصادفی که توسط چارنس و کوپر^{۱۰} در [۵۹] ارائه شد، ما به یک متغیر خاص (تابع) درون دامنه هدف به همراه احتمال مشخص نیاز داریم. این مدل‌ها معمولاً به قیود غیرخطی و نامحدب منتهی می‌شوند. به کارگیری اولیه محاسباتی مسائل برنامه‌ریزی با قیود تصادفی محدود به متغیرهای تصادفی نرمال می‌شود. پریکوپا^{۱۱} [۷۷] یک الگوریتم دوگانه برای حل مسأله کلی قیود تصادفی پیشنهاد داد، او در [۷۸] به بحث درباره روشهای مختلف برای حل مدل‌های قیود تصادفی شامل روش‌های گرادیانی و استفاده از توابع جریمه می‌پردازد. هیلر^{۱۲} در [۷۹] فرآیندی برای تقریب قیود تصادفی با قیود خطی را پیشنهاد داد. وی قیود تصادفی را با قیود تفکیک‌پذیر محدب جایگزین کرد و سپس مسأله را با روش برنامه‌ریزی تفکیک‌پذیر محدب با استفاده از الگوریتم سیمپلکس حل کرد. سپالا در [۸۰] مجموعه‌ای از قیود خطی یکنواخت سخت^{۱۳} که جایگزین یک قید تصادفی منفرد می‌شود را معرفی کرد. براساس نظر سپالا این روش بسیار دقیق تر است اما نسبت به روش هیلر کارایی کمتری دارد.

^۵Military system

^۶Anticipative

^۷Adaptive

^۸Stochastic Programming with Recourse (SPR)

^۹Chance Constrained Programming (CCP)

^{۱۰}Charnes and Cooper

^{۱۱}Prekopa

^{۱۲}Hillier

^{۱۳}uniformly tighter linear constraints

این مدل محدود به مسائل دارای متغیر کران‌دار نمی‌شود، با این حال به محدودیت‌های بیشتری نسبت به روش هیلر نیازمند است. سپالا و اورپانا^{۱۴} در [۸۱] کارایی و دقت الگوریتم سپالا را امتحان کردند. اولسون و سونس^{۱۵} در [۸۲] یک فرمول تقریبی برای حل مسائل برنامه‌ریزی با قیود تصادفی پیشنهاد دادند که سبب ایجاد کران برای قیود تصادفی می‌گردد، که این (هم) ممکن است فرم غیر خطی داشته باشد، با این وجود این تکنیک خیلی دقیق نیست ([۸۳] را ببینید). وینتراوب و ورا^{۱۶} در [۸۴] روش صفحه برشی برای حل معادل قطعی مسائل برنامه‌ریزی با قیود تصادفی برای ضرایب محدود شده تصادفی با توزیع نرمال را ارائه کردند. دنچوا^{۱۷} و همکاران در [۸۵] ترتیب جزئی روی مجموعه‌ای از سناریوها تعریف کردند و نشان دادند که تعداد محدودی از سناریوها (که نقاط مؤثر $(1 - \alpha)$ نامیده می‌شوند)، نقش کلیدی دارند. آنها الگوریتمی ایجاد کردند که به‌طور مداوم مجموعه‌ای از نقاط مؤثر را به‌منظور تولید کران‌های بالا و پائین به‌روز رسانی می‌کند. این الگوریتم تا حد برنامه‌ریزی محدب کلی با قیود تصادفی در [۸۶] توسعه یافته است. برالدی و روژنسکی^{۱۸} در [۸۷] به بحث درباره روش‌هایی پرداختند که براساس شمارش جزئی یا کامل نقاط مؤثر برای برنامه‌ریزی صحیح تصادفی است. روژنسکی^{۱۹} در [۸۸] مسائل برنامه‌ریزی با قیود تصادفی چون مسائل برنامه‌ریزی صحیح آمیخته بزرگ مقیاس^{۲۰} را با در نظر گرفتن محدودیت‌های کوله پشتی^{۲۱} دوباره فرمول‌سازی کرد و نامعادله‌های صحیح برای مسأله برنامه‌ریزی صحیح آمیخته ایجاد کرد. آرینگری^{۲۲} در [۸۹] جستجوی ممنوع^{۲۳} و فرآیند شبیه‌سازی را برای دو مسأله بهینه‌سازی شبکه ارتباط تلفنی با فرمول‌بندی از برنامه‌ریزی با قیود تصادفی توسعه داد.

^{۱۴}Seppala and Orpana^{۱۵}Olson and Swenseth^{۱۶}Weintraub and Vera^{۱۷}Dentcheva^{۱۸}Beraldi and Ruszczyński^{۱۹}Ruszczyński^{۲۰}Large scale Mixed integer programming^{۲۱}Knapsack Constraints^{۲۲}Aringheri^{۲۳}Tabu Search

کالافیور و کامپی^{۲۴} در [۹۰]، دی‌فاریاس و ون‌روی^{۲۵} در [۹۱] روش‌های تقریبی بر مبنای نمونه‌گیری محدودیت و روش‌های یادگیری آماری را پیشنهاد دادند. کالافیور و کامپی از روش‌های یادگیری آماری برای ایجاد کردن کران‌هایی روی اندازه (احتمال یا حجم) از مجموعه قیود اصلی که احتمالاً از طریق حل تصادفی شده‌اند استفاده کردند. نتایج بدست آمده توسط آنها با استفاده از این اصل اساسی که جواب بهینه برای مسأله محدب از طریق تعداد متناهی از قیود پشتیبانی می‌شود. دی‌فاریاس و ون‌روی نشان دادند که احتمال وقوع می‌تواند برای یک مسأله برنامه‌ریزی خطی از طریق نمونه‌گیری تعداد قابل ردیابی^{۲۶} از قیود فراهم گردد. اندازه نمونه با زیاد شدن تعداد متغیرها به صورت چندجمله‌ای افزایش می‌یابد و مستقل از مجموعه قید اصلی است. اردوگان و لینگار^{۲۷} در [۹۲] نمونه نتایج پیچیدگی برای مسائل برنامه‌ریزی با قیود تصادفی با مسائل برنامه‌ریزی با قیود تصادفی مبهم گسترش دادند که در آن مجموعه غیر قطعی با گوی‌های تعریف شده در جملاتی از تابع متریک پروهورف^{۲۸} داده شده است. با استفاده از قضیه نمایش استراسن-دوادل، یک مسأله نمونه مقاوم^{۲۹} را ساختند که تقریب خوبی برای مسائل برنامه‌ریزی با قیود تصادفی مبهم با احتمال بالا است. ایوامورا و لیو^{۳۰} در [۹۳] براساس الگوریتم ژنتیک یک شبیه‌سازی تصادفی ایجاد کردند که در آن شبیه‌سازی، بررسی قیود تصادفی و محاسبه توابع هدف مورد استفاده قرار می‌گرفت. پوجاری و وارگس در [۹۴]^{۳۱} با الهام گرفتن از کار ایوامورا و لیو، چارچوبی شامل الگوریتم ژنتیک و شبیه‌سازی مونت کارلو^{۳۲} برای حل مسأله برنامه‌ریزی با قیود تصادفی پیشنهاد کردند.

در مهندسی و بسیاری از کاربردهای علمی، جوابهای زمان واقعی مسائل قیود تصادفی معمولاً مطلوب

^{۲۴}Calafiore and Campi

^{۲۵}De Farias and V. Roy

^{۲۶}tractable

^{۲۷}Erdogan and Iyengar

^{۲۸}Prohorov metric

^{۲۹}Robust

^{۳۰}Iwamura and Liu

^{۳۱}Poojari and Varghese

^{۳۲}Monte Carlo simulation

است. در نتیجه به روشهای مؤثر برای حل برنامه‌ریزی با قیود تصادفی زمان واقعی نیاز داریم. نتایج مطالعات قبلی [۹۵]–[۱۲۴] حاکی از آن است که از شبکه‌های عصبی می‌توان برای حل مسائل مختلف بهینه‌سازی استفاده کرد. ایده اصلی شبکه‌های عصبی برای بهینه‌سازی ایجاد یک تابع انرژی نامنفی و پایه‌گذاری یک سیستم دینامیکی است که یک شبکه عصبی مصنوعی را ایجاد می‌کند. سیستم دینامیکی معمولاً به فرم معادلات دیفرانسیل مرتبه اول است. به‌علاوه انتظار می‌رود که سیستم دینامیکی به یک وضعیت پایدار (یا نقطه تعادل) که متناظر با جواب مسأله بهینه‌سازی اصلی است که با یک نقطه اولیه آغاز می‌شود نزدیک گردد.

با توجه به مباحث فوق، در این فصل یک مدل از شبکه‌های عصبی برای حل مسائل برنامه‌ریزی تصادفی خطی ارائه می‌دهیم. ابتدا مسأله برنامه‌ریزی تصادفی به‌صورت یک مسأله برنامه ریزی مخروط مرتبه دوم محدب در می‌آید سپس این مسأله را تبدیل به برنامه‌ریزی غیرخطی محدب با قیود نامساوی می‌کنیم، زمانیکه با استفاده از تکنیک هموار سازی، قیود مخروط مرتبه دوم ناهموار به‌فرم هموار آن تبدیل شوند. براساس شرایط بهینگی ($K.K.T$) برنامه‌ریزی محدب، مدل شبکه عصبی برای حل مسائل برنامه‌ریزی غیرخطی محدب^{۳۳} پیشنهاد خواهد شد. همچنین ثابت می‌شود نقطه تعادل شبکه عصبی پیشنهاد شده برابر نقطه $K.K.T$ از مسأله برنامه ریزی مخروط مرتبه دوم محدب است. وجود و یکتایی نقطه تعادل شبکه عصبی نیز مورد آنالیز قرار خواهد گرفت. با ساختن یک تابع لیاپانوف مناسب، شرط کافی برای تضمین وجود و پایدار مجانبی سراسری برای نقطه تعادل یکتای شبکه پیشنهاد شده بدست خواهد آمد. مدل شبکه عصبی پیشنهاد شده برای مسائل $Minimax$ در [۱۰۸]، $maximum\ flow$ در [۱۰۹]، $shortest\ Path\ problem$ در [۱۱۰] و برای برنامه‌ریزی هندسی در [۱۱۱] نیز با موفقیت استفاده شده است.

^{۳۳}Convex Nonlinear Programming

۲.۴ مسائل بهینه‌سازی با قيود تصادفی

در روش برنامه‌ریزی با قيود تصادفی، مسأله برنامه‌ریزی تصادفی خطی (۱.۴)–(۳.۴) به صورت زیر بیان می‌شود:

$$\text{Minimize } f(x) = \sum_{j=1}^n c_j x_j \quad (۴.۴)$$

subject to

$$P\left[\sum_{j=1}^n a_{ij} x_j \leq b_i\right] \geq p_i, \quad i = 1, 2, \dots, m, \quad (۵.۴)$$

$$x_j \geq 0, \quad j = 1, 2, \dots, n. \quad (۶.۴)$$

وقتی که a_{ij} ، b_i ، c_j ، متغیرهای تصادفی و p_i ها احتمالات مشخص شده هستند. توجه کنید که (۲.۴) نشان می‌دهد که i امین قید

$$\sum_{j=1}^n a_{ij} x_j \leq b_i,$$

باید دست کم با احتمال p_i وقتی که $0 \leq p_i \leq 1$ است، برآورده شود. برای سادگی، فرض می‌کنیم که متغیرهای تصمیم x_j قطعی هستند. ابتدا حالات خاصی را که تنها c_j یا a_{ij} یا b_i ها متغیرهای تصادفی هستند را بررسی کرده، سپس به حالت کلی که a_{ij} ، c_j و b_i همگی متغیرهای تصادفی هستند، می‌پردازیم. علاوه بر این فرض می‌کنیم که همه متغیرهای تصادفی دارای توزیع نرمال با میانگین و انحراف معیار معلوم هستند.

هنگامی که تنها a_{ij} ها متغیرهای تصادفی هستند

فرض کنید $\bar{a}_{ij}$ ، میانگین و $\text{Var}(a_{ij}) = \sigma^2$ واریانس متغیر تصادفی a_{ij} با توزیع نرمال باشد. همچنین فرض کنید توزیع چند متغیره a_{ij} ها ($i = 1, \dots, m; j = 1, \dots, n$) و در نتیجه کواریانس بین متغیرهای

تصادفی a_{ij} و a_{kl} یعنی $\text{cov}(a_{ij}, a_{kl})$ معلوم باشد. کمیت d_i را به صورت زیر تعریف می‌کنیم:

$$d_i = \sum_{j=1}^n a_{ij}x_j, \quad i = 1, 2, \dots, m. \quad (۷.۴)$$

چون $a_{i1}, \dots, a_{in}$ دارای توزیع نرمال و $x_1, \dots, x_n$ مقادیر ثابت و نامعلوم هستند، d_i هم دارای توزیع

نرمال با میانگین و واریانس

$$\bar{d}_i = \sum_{j=1}^n \bar{a}_{ij}x_j, \quad i = 1, 2, \dots, m, \quad (۸.۴)$$

$$\text{Var}(d_i) = \sigma_{d_i}^2 = X^T V_i X, \quad (۹.۴)$$

است، که در آن V_i ، i امین ماتریس کواریانس، به صورت زیر تعریف می‌شود:

$$V_i = \begin{bmatrix} \text{Var}(a_{i1}) & \text{Cov}(a_{i1}, a_{i2}) & \dots & \text{Cov}(a_{i1}, a_{in}) \\ \text{Cov}(a_{i2}, a_{i1}) & \text{Var}(a_{i2}) & \dots & \text{Cov}(a_{i2}, a_{in}) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(a_{in}, a_{i1}) & \text{Cov}(a_{in}, a_{i2}) & \dots & \text{Var}(a_{in}) \end{bmatrix}. \quad (۱۰.۴)$$

قیود (۵.۴) را می‌توان به صورت زیر بیان کرد:

$$P[d_i \leq b_i] \geq p_i,$$

یعنی

$$P\left[\frac{d_i - \bar{d}_i}{\sqrt{\text{Var}(d_i)}} \leq \frac{b_i - \bar{d}_i}{\sqrt{\text{Var}(d_i)}}\right] \geq p_i \quad i = 1, 2, \dots, m, \quad (۱۱.۴)$$

وقتی $\frac{d_i - \bar{d}_i}{\sqrt{\text{Var}(d_i)}}$ یک متغیر تصادفی^{۳۴} نرمال استاندارد با میانگین صفر و واریانس یک است. بنابراین

احتمال اینکه d_i کوچکتر یا مساوی b_i باشد را می‌توان به صورت زیر نوشت:

$$P[d_i \leq b_i] = \phi\left(\frac{b_i - \bar{d}_i}{\sqrt{\text{Var}(d_i)}}\right), \quad i = 1, \dots, m, \quad (۱۲.۴)$$

^{۳۴}Variate

که در آن $\phi(x)$ بیانگر تابع توزیع تجمعی، توزیع نرمال استاندارد در نقطه x است. اگر e_i نشان دهنده مقدار توزیع نرمال استاندارد باشد که در آن

$$\phi(e_i) = p_i \geq 0.5 \quad (e_i \geq 0), \quad (13.4)$$

آنگاه قیود رابطه (۱۱.۴) را می‌توان به صورت زیر بیان کرد:

$$\phi\left(\frac{b_i - \bar{d}_i}{\sqrt{\text{Var}(d_i)}}\right) \geq \phi(e_i), \quad i = 1, 2, \dots, m, \quad (14.4)$$

این نامساوی تنها وقتی برقرار است که

$$\left(\frac{b_i - \bar{d}_i}{\sqrt{\text{Var}(d_i)}}\right) \geq e_i,$$

یا

$$\bar{d}_i + e_i \sqrt{\text{Var}(d_i)} - b_i \leq 0, \quad i = 1, 2, \dots, m. \quad (15.4)$$

با قرار دادن روابط (۸.۴) و (۹.۴) در (۱۵.۴)، داریم

$$\sum_{j=1}^n \bar{a}_{ij} x_j + e_i \sqrt{X^T V_i X} - b_i \leq 0, \quad i = 1, 2, \dots, m. \quad (16.4)$$

اینها قیود غیرخطی قطعی معادل با قیود خطی تصادفی هستند. بنابراین جواب مسأله برنامه‌ریزی احتمالی بیان شده در (۴.۴)–(۶.۴) را می‌توان با حل مسأله قطعی زیر به دست آورد.

$$\text{Minimize} \quad f(x) = \sum_{j=1}^n c_j x_j \quad (17.4)$$

subject to

$$\begin{cases} \sum_{j=1}^n \bar{a}_{ij} x_j + e_i \sqrt{X^T V_i X} - b_i \leq 0, & i = 1, 2, \dots, m, \\ x_j \geq 0, & j = 1, 2, \dots, n. \end{cases} \quad (18.4)$$

اگر متغیرهای تصادفی با توزیع نرمال a_{ij} مستقل باشند، رابطه (۱۰.۴) به یک ماتریس قطری به صورت زیر کاهش می‌یابد:

$$V_i = \begin{bmatrix} \text{Var}(a_{i1}) & \circ & \dots & \circ \\ \circ & \text{Var}(a_{i2}) & \dots & \circ \\ \vdots & \vdots & \ddots & \vdots \\ \circ & \circ & \dots & \text{Var}(a_{in}) \end{bmatrix}. \quad (۱۹.۴)$$

در این حالت قیود رابطه (۱۶.۴) به شکل زیر ساده می‌شوند:

$$\sum_{j=1}^n \bar{a}_{ij} x_j + e_i \sqrt{\sum_{j=1}^n [\text{Var}(a_{ij}) x_j]} - b_i \leq 0, \quad i = 1, 2, \dots, m. \quad (۲۰.۴)$$

هنگامی که تنها b_i ها متغیرهای تصادفی هستند

فرض کنید $\bar{b}_i$ و $\text{Var}(b_i)$ به ترتیب نشان دهنده میانگین و واریانس متغیر تصادفی b_i با توزیع نرمال باشند. ضرایب قیود (۵.۴) را می‌توان به صورت زیر بیان کرد:

$$\begin{aligned} P\left[\sum_{j=1}^n a_{ij} x_j \leq b_i\right] &= P\left[\frac{\sum_{j=1}^n a_{ij} x_j - \bar{b}_i}{\sqrt{\text{Var}(b_i)}} \leq \frac{b_i - \bar{b}_i}{\sqrt{\text{Var}(b_i)}}\right] \\ &= P\left[\frac{b_i - \bar{b}_i}{\sqrt{\text{Var}(b_i)}} \geq \frac{\sum_{j=1}^n a_{ij} x_j - \bar{b}_i}{\sqrt{\text{Var}(b_i)}}\right] \geq p_i \quad i = 1, 2, \dots, m, \end{aligned} \quad (۲۱.۴)$$

که در آن $\frac{b_i - \bar{b}_i}{\sqrt{\text{Var}(b_i)}}$ یک متغیر نرمال استاندارد با میانگین صفر و واریانس یک است. نامساوی (۲۱.۴) را به صورت زیر هم می‌توان نوشت:

$$P\left[\frac{\sum_{j=1}^n a_{ij} x_j - \bar{b}_i}{\sqrt{\text{Var}(b_i)}} \leq \frac{b_i - \bar{b}_i}{\sqrt{\text{Var}(b_i)}}\right] \leq 1 - p_i, \quad i = 1, 2, \dots, m. \quad (۲۲.۴)$$

اگر E_i بیانگر مقدار تغییر استاندارد باشد که در آن

$$\phi(E_i) = 1 - p_i < 0.5, \quad (E_i < 0)$$

قیود رابطه (۲۲.۴) را می‌توان به صورت زیر بیان کرد:

$$\phi\left(\frac{\sum_{j=1}^n a_{ij} x_j - \bar{b}_i}{\sqrt{\text{Var}(b_i)}}\right) \leq \phi(E_i), \quad i = 1, 2, \dots, m. \quad (۲۳.۴)$$

این نامساوی تنها وقتی برقرار است که:

$$\frac{\sum_{j=1}^n a_{ij}x_j - \bar{b}_i}{\sqrt{\text{Var}(b_i)}} \leq E_i < 0, \quad i = 1, 2, \dots, m,$$

یا

$$\sum_{j=1}^n a_{ij}x_j - \bar{b}_i - E_i\sqrt{\text{Var}(b_i)} \leq 0, \quad i = 1, 2, \dots, m. \quad (24.4)$$

بنابراین مسأله برنامه‌ریزی خطی بیان شده در (۴.۴)-(۶.۴) با مسأله قطعی زیر معادل است:

$$\text{Minimize} \quad f(x) = \sum_{j=1}^n c_j x_j \quad (25.4)$$

subject to

$$\begin{cases} \sum_{j=1}^n a_{ij}x_j - \bar{b}_i - E_i\sqrt{\text{Var}(b_i)} \leq 0, & i = 1, 2, \dots, m \\ x_j \geq 0, & j = 1, 2, \dots, n \end{cases} \quad (26.4)$$

هنگامی که تنها c_j ها متغیرهای تصادفی هستند

چون c_j ها متغیرهایی تصادفی با توزیع نرمال هستند، تابع هدف $f(x)$ هم یک متغیر تصادفی با توزیع

نرمال است. میانگین و واریانس f عبارتند از:

$$\bar{f} = \sum_{j=1}^n \bar{c}_j x_j, \quad (27.4)$$

$$\text{Var}(f) = X^T V X, \quad (28.4)$$

وقتی که $\bar{c}_j$ ، میانگین c_j و ماتریس کواریانس V به صورت زیر تعریف می‌شود:

$$V = \begin{bmatrix} \text{Var}(c_1) & \text{Cov}(c_1, c_2) & \dots & \text{Cov}(c_1, c_n) \\ \text{Cov}(c_2, c_1) & \text{Var}(c_2) & \dots & \text{Cov}(c_2, c_n) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(c_n, c_1) & \text{Cov}(c_n, c_2) & \dots & \text{Var}(c_n) \end{bmatrix}. \quad (29.4)$$

حال یک تابع هدف کمینه‌سازی قطعی جدید را می‌توان به صورت زیر فرمولبندی کرد:

$$F(x) = k_1 \bar{f} + k_2 \sqrt{\text{Var}(f)}, \quad (30.4)$$

وقتی که k_1 و k_2 مقادیر ثابت نامنفی هستند که مقادیر آنها نشان دهنده اهمیت نسبی $\bar{f}$ و انحراف معیار f برای کمینه سازی است. بنابراین $k_2 = 0$ نشان می‌دهد که مقدار مورد انتظار f باید بدون توجه به انحراف معیار f کمینه سازی شود. از طرف دیگر اگر $k_1 = 0$ باشد، نشان می‌دهد که باید کمینه سازی قابلیت تغییر f حول میانگینش را بدون توجه به مقدار میانگین f انجام دهیم. به همین ترتیب، اگر $k_1 = k_2 = 1$ باشد کمینه سازی میانگین و انحراف معیار f از اهمیت یکسانی برخوردار هستند. توجه کنید که تابع هدف جدید که در رابطه (۳۰.۴) بیان شده است، از دیدگاه واریانس f ، یک تابع غیر خطی از x است.

بنابراین جواب، مسأله برنامه‌ریزی خطی تصادفی (۴.۴)-(۶.۴) را می‌توان با حل مسأله برنامه‌ریزی

غیرخطی معادل زیر به دست آورد:

$$\begin{aligned} \text{Minimize} \quad & F(x) = k_1 \sum_{j=1}^n \bar{c}_j x_j + k_2 \sqrt{X^T V X} \\ \text{subject to} \quad & \end{aligned} \quad (31.4)$$

$$\begin{cases} \sum_{j=1}^n a_{ij} x_j - b_i \leq 0, & i = 1, 2, \dots, m, \\ x_j \geq 0, & j = 1, 2, \dots, n. \end{cases} \quad (32.4)$$

اگر همه متغیرهای تصادفی c_j مستقل خطی باشند، تابع هدف به صورت زیر در می‌آید:

$$F(x) = k_1 \sum_{j=1}^n \bar{c}_j x_j + k_2 \sqrt{\sum_{j=1}^n \text{Var}(c_j) x_j^2}. \quad (33.4)$$

هنگامی که c_j ، a_{ij} و b_i ها متغیرهای تصادفی هستند

چون متغیرهای تصادفی c_j تنها در تابع هدف ظاهر می‌شوند، می‌توانیم تابع هدف جدید $F(X)$ را به همان

صورت داده شده در رابطه (۳۰.۴) در نظر بگیریم. قیود (۵.۴) را به صورت زیر می‌توان بیان کرد:

$$P[h_i \leq 0] \geq p_i \quad i = 1, 2, \dots, m, \quad (34.4)$$

که در آن h_i یک متغیر تصادفی جدید است که به صورت زیر تعریف می‌شود

$$h_i = \sum_{j=1}^n a_{ij}x_j - b_i = \sum_{k=1}^{n+1} q_{ik}y_k, \quad (35.4)$$

که در آن

$$\begin{cases} q_{ik} = a_{ik}, & k = 1, 2, \dots, n, \\ q_{i,n+1} = b_i, \\ y_k = x_k, & k = 1, 2, \dots, n, \\ y_{n+1} = -1. \end{cases}$$

توجه کنید که ثابت y_{n+1} برای راحتی کار تعریف می‌شود. چون h_i به صورت ترکیب خطی متغیرهای

تصادفی q_{ik} با توزیع نرمال داده می‌شود، پس دارای توزیع نرمال است. میانگین و واریانس h_i عبارتند از:

$$\bar{h}_i = \sum_{k=1}^{n+1} \bar{q}_{ik}y_k = \sum_{j=1}^n \bar{a}_{ij}x_j - \bar{b}_i, \quad (36.4)$$

$$\text{Var}(h_i) = Y^T V_i Y, \quad (37.4)$$

که در آن

$$V_i = \begin{bmatrix} \text{Var}(q_{i1}) & \text{Cov}(q_{i1}, q_{i2}) & \dots & \text{Cov}(q_{i1}, q_{i,n+1}) \\ \text{Cov}(q_{i2}, q_{i1}) & \text{Var}(q_{i2}) & \dots & \text{Cov}(q_{i2}, q_{i,n+1}) \\ \vdots & & & \\ \text{Cov}(q_{i,n+1}, q_{i1}) & \text{Cov}(q_{i,n+1}, q_{i2}) & \dots & \text{Var}(q_{i,n+1}) \end{bmatrix}.$$

این را می‌توان به صورت صریح‌تر زیر نوشت:

$$\begin{aligned} \text{Var}(h_i) &= \sum_{k=1}^{n+1} [y_k^2 \text{Var}(q_{ik}) + 2 \sum_{l=k+1}^{n+1} y_k y_l \text{Cov}(q_{ik}, q_{il})] \\ &= \sum_{k=1}^n [y_k^2 \text{Var}(q_{ik}) + 2 \sum_{l=k+1}^n y_k y_l \text{Cov}(q_{ik}, q_{il})] \\ &\quad + y_{n+1}^2 \text{Var}(q_{i,n+1}) + \sum_{l=1}^n [2 y_l y_{n+1} \text{Cov}(q_{il}, q_{i,n+1})] \\ &= \sum_{k=1}^n [x_k^2 \text{Var}(a_{ik}) + 2 \sum_{l=k+1}^n x_k x_l \text{Cov}(a_{ik}, a_{il})] \\ &\quad + \text{Var}(b_i) - 2 \sum_{k=1}^n x_k \text{Cov}(a_{ik}, b_i). \end{aligned} \quad (38.4)$$

بنابراین قیود روابط (۳۴.۴) را می‌توان به‌صورت زیر بیان کرد:

$$P\left[\frac{h_i - \bar{h}_i}{\sqrt{\text{Var}(h_i)}} \leq \frac{-\bar{h}_i}{\sqrt{\text{Var}(h_i)}}\right] \geq p_i, \quad i = 1, 2, \dots, m, \quad (39.4)$$

که در آن $\left[\frac{h_i - \bar{h}_i}{\sqrt{\text{Var}(h_i)}}\right]$ بیان‌کننده یک متغیر نرمال استاندارد با میانگین صفر و واریانس یک است.

بنابراین اگر e_i نشان‌دهنده مقدار متغیر استاندارد باشد که در آن

$$\phi(e_i) = p_i \geq 0.5 \quad (e_i \geq 0), \quad (40.4)$$

قیود رابطه (۳۹.۴) قابل بیان به‌صورت زیر هستند:

$$\phi\left(\frac{-\bar{h}_i}{\sqrt{\text{Var}(h_i)}}\right) \geq \phi(e_i) \quad i = 1, 2, \dots, m. \quad (41.4)$$

این نامساویها تنها در صورتی برقرارند که، نامساویهای غیرخطی زیر برقرار باشند:

$$\frac{-\bar{h}_i}{\sqrt{\text{Var}(h_i)}} \geq e_i > 0, \quad i = 1, 2, \dots, m,$$

یا

$$\bar{h}_i + e_i \sqrt{\text{Var}(h_i)} \leq 0, \quad i = 1, 2, \dots, m. \quad (42.4)$$

بنابراین مسأله برنامه‌ریزی احتمالی (۴.۴)-(۶.۴) را می‌توان به‌عنوان یک مسأله برنامه‌ریزی غیرخطی

قطعی معادل زیر نوشت:

$$\text{Minimize} \quad F(x) = k_1 \sum_{j=1}^n \bar{c}_j x_j + k_2 \sqrt{X^T V X}, \quad k_1, k_2 \geq 0 \quad (43.4)$$

subject to

$$\begin{cases} \bar{h}_i + e_i \sqrt{\text{Var}(h_i)} \leq 0, & i = 1, 2, \dots, m, \\ x_j \geq 0, & j = 1, 2, \dots, n. \end{cases} \quad (44.4)$$

از تحلیل بالا و [۱۲۵]، واضح است که مسائل بهینه‌سازی با قیود تصادفی قابل تبدیل به مسائل برنامه‌ریزی مخروط مرتبه دو محدب^{۳۵} هستند. بنابراین، ما فرم کلی از مسائل برنامه‌ریزی مخروط مرتبه دو محدب را به صورت زیر در نظر می‌گیریم:

$$\text{Minimize } f(x) \quad (45.4)$$

subject to

$$Ex + d \succeq_{\mathcal{K}} 0, \quad (46.4)$$

$$x \geq 0. \quad (47.4)$$

وقتیکه $f: \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ ، تابع سره بسته محدب، E یک ماتریس $m \times n$ با $m \geq n$ ، d برداری در $\mathbb{R}^m$ و $x \succeq_{\mathcal{K}} 0$ یعنی $x \in \mathcal{K}$ و $\mathcal{K}$ حاصلضرب دکارتی از مخروط‌های مرتبه دوم که مخروط‌های لورنتس یا مخروط بستنی هم نامیده می‌شوند. به عبارت دیگر

$$\mathcal{K} = \mathcal{K}^{m_1} \times \mathcal{K}^{m_2} \times \dots \times \mathcal{K}^{m_r},$$

که در آن $m_1 + \dots + m_r = n$ با $r, m_1, \dots, m_r \geq 1$ و

$$\mathcal{K}^{m_i} := \{(x_1, x_2)^T \in \mathbb{R} \times \mathbb{R}^{m_i-1} \mid x_1 \geq \|x_2\|\}, \quad (48.4)$$

که در آن $\|\cdot\|$ نرم اقلیدسی و $\mathcal{K}^1$ مجموعه از $\mathbb{R}_+$ حقیقی نامنفی است. در قطعه بعد ما سعی داریم یک مدل شبکه عصبی با عملکرد بالا برای حل مسأله (۴۵.۴)–(۴۷.۴) پیشنهاد دهیم.

۳.۴ مدل شبکه عصبی

ما قیود غیرهموار مخروط مرتبه دوم $x \succeq_{\mathcal{K}^{m_i}} 0$ را برای $(i = 1, \dots, r)$ با یک هموار سازی از [۱۲۶] به صورت زیر جایگزین می‌کنیم:

$$\sqrt{\epsilon^2 + \|x_2\|^2} \leq x_1, \quad (49.4)$$

^{۳۵}CSOCP

که در آن ϵ ثابت کوچکی است که اغلب حدود 10^{-6} انتخاب می‌شود. اغتشاش به‌دست آمده در قید مخروط مرتبه دوم، باعث می‌شود قیود مسأله (۴۵.۴)–(۴۷.۴) محدب و هموار شود. بنابراین مسأله محدب و غیر هموار (۴۵.۴)–(۴۶.۴) می‌تواند به یک مسأله بهینه‌سازی محدب هموار معادل به‌صورت زیر تبدیل شود:

$$\text{Minimize} \quad f(x) \quad (50.4)$$

subject to

$$g(x) \leq 0. \quad (51.4)$$

که در آن $g(x) = \begin{bmatrix} h(x) \\ x \end{bmatrix}$ و $h(x) \in \mathbb{R}^m$ یک هموارسازی از $Ex + d \succeq_K 0$ با روش ϵ -اغتشاش است. واضح است که توابع برداری $h(x)$ محدب و دوبار مشتق‌پذیر هستند.

از [۱]، می‌بینیم که $x^* \in \mathbb{R}^n$ یک جواب بهینه از (۵۰.۴)–(۵۱.۴) اگر و تنها اگر بردار $u^* \in \mathbb{R}^m$ وجود

داشته باشد به‌طوری‌که $(x^{*T}, u^{*T})^T$ در سیستم $K.K.T$ زیر صدق کند:

$$\begin{cases} u^* \geq 0, & g(x^*) \leq 0, & u^{*T} g(x^*) = 0, \\ \nabla f(x^*) + \nabla g(x^*)^T u^* = 0. \end{cases} \quad (52.4)$$

x^* را نقطه $K.K.T$ از (۵۰.۴)–(۵۱.۴) و u^* را بردار ضریب لاگرانژ متناظر با x^* نامیم. واضح است که

x^* جواب بهینه از (۵۰.۴)–(۵۱.۴) است، اگر و تنها اگر x^* نقطه $K.K.T$ از (۵۰.۴)–(۵۱.۴) باشد.

اکنون فرض کنید که $x(\cdot)$ و $u(\cdot)$ متغیرهای وابسته به زمان باشند. هدف ما بیان یک ساختار دینامیکی

زمان پیوسته است که به نقطه $K.K.T$ مسأله (۵۰.۴)–(۵۱.۴) و دوگان آن هدایت شود. ما مدل شبکه

عصبی بازگشتی برای حل (۵۰.۴)–(۵۱.۴) و دوگان آن را با معادله دینامیکی به‌صورت زیر بیان می‌کنیم:

$$\frac{dx}{dt} = -(\nabla f(x) + \nabla g(x)^T(u + g(x))^+), \quad (53.4)$$

$$\frac{du}{dt} = (u + g(x))^+ - u, \quad (54.4)$$

با نقطه آغازین $y_* = (x_*^T, u_*^T)^T$ که در آن

$$(u + g(x))^+ = ([u_1 + g_1(x)]^+, [u_2 + g_2(x)]^+, \dots, [u_m + g_m(x)]^+),$$

$$[u_k + g_k(x)]^+ = \max\{u_k + g_k(x), 0\}, \quad k = 1, 2, \dots, m.$$

برای سادگی تعریف می‌کنیم، $y = (x^T, u^T)^T \in \mathbb{R}^{n+m}$ ، D^* مجموعه نقاط بهینه از (۵۰.۴)–(۵۱.۴) و دوگان آن،

$$\Phi(y) = \begin{bmatrix} -(\nabla f(x) + \nabla g^T(u + g(x))^+) \\ (u + g(x))^+ - u \end{bmatrix}. \quad (۵۵.۴)$$

بنابراین شبکه عصبی (۵۳.۴)–(۵۴.۴) را می‌توان به صورت زیر نوشت:

$$\frac{dy}{dt} = \kappa \Phi(y), \quad (۵۶.۴)$$

$$y(t_*) = y_*. \quad (۵۷.۴)$$

که در آن κ پارامتر اسکالر^{۳۶} است و نرخ همگرایی شبکه عصبی (۵۶.۴)–(۵۷.۴) را نشان می‌دهد. برای سادگی، $\kappa = 1$ در نظر می‌گیریم. شکل (۱.۴) بیانگر این است که شبکه فوق روی سخت افزار چگونه اجرا می‌شود.

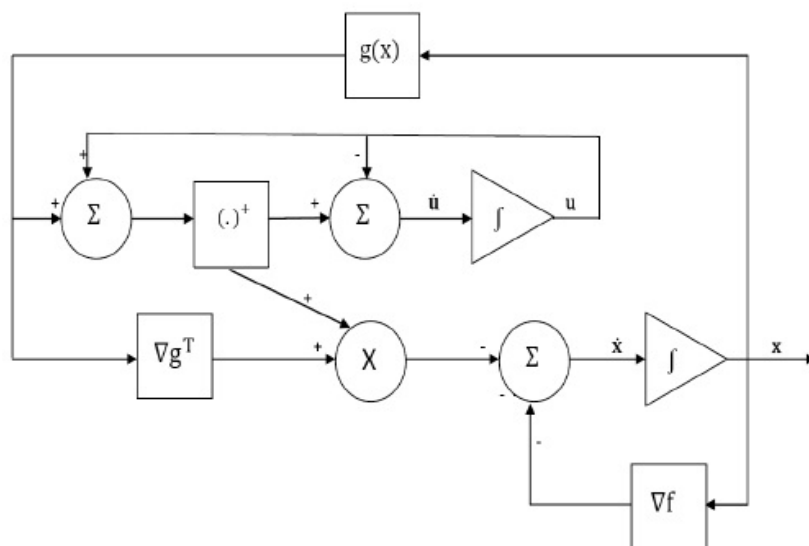
۴.۴ تحلیل پایداری و همگرایی

در این قطعه پایداری و همگرایی شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) بررسی می‌شود.

قضیه ۱.۴.۴. فرض کنید $y^* = (x^{*T}, u^{*T})^T$ ، نقطه تعادل از شبکه عصبی (۵۶.۴)–(۵۷.۴) باشد. آنگاه

x^* نقطه $K.K.T$ از مسأله (۵۱.۴)–(۵۰.۴)، برعکس اگر $x^* \in \mathbb{R}^n$ جواب بهینه مسأله (۵۱.۴)–(۵۰.۴)

^{۳۶}Scale



شکل ۱.۴: بلوک دیاگرام شبکه عصبی (۵۶.۴)–(۵۷.۴)

باشد، آنگاه بردار $u^* \in \mathbb{R}^m$ وجود دارد به‌طوری که $y^* = (x^{*T}, u^{*T})^T$ نقطه تعادل از شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) است.

اثبات. فرض کنید $y^* = (x^{*T}, u^{*T})^T$ نقطه تعادل شبکه عصبی (۵۶.۴)–(۵۷.۴) باشد، آنگاه $\frac{dx^*}{dt} = 0$ و $\frac{du^*}{dt} = 0$ در نتیجه

$$\nabla f(x^*) + \nabla g(x^*)^T (u^* + g(x^*))^+ = 0, \quad (58.4)$$

$$(u^* + g(x^*))^+ = u^*. \quad (59.4)$$

همچنین، داریم $u^* = (u^* + g(x^*))^+$ اگر و تنها اگر

$$u^* \geq 0, \quad g(x^*) \leq 0, \quad u^{*T} g(x^*) = 0. \quad (60.4)$$

بنابراین، با جایگذاری (۵۹.۴) در (۵۸.۴) داریم

$$\nabla f(x^*) + \nabla g(x^*)^T u^* = 0 \quad (61.4)$$

از (۶۰.۴) و (۶۱.۴)، نتیجه می‌شود که $y^* = (x^{*T}, u^{*T})^T$ در شرایط $K.K.T$ صدق می‌کند و این اثبات را کامل می‌کند. $\square$

لم ۲.۴.۴. نقطه تعادل مدل شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) یکتاست.

اثبات. چون مسأله (۵۰.۴)–(۵۱.۴) دارای جواب بهینه یکتای x^* است، شرایط لازم و کافی $K.K.T$ دارای جواب یکتای $y^* = (x^{*T}, u^{*T})^T$ هستند، بنابراین از قضیه (۱.۴.۴) می‌بینیم که نقطه تعادل مدل پیشنهاد شده (۵۶.۴)–(۵۷.۴) در سیستم $K.K.T$ ، صدق می‌کند. بنابراین نقطه تعادل شبکه عصبی (۵۶.۴)–(۵۷.۴) یکتاست. $\square$

لم ۳.۴.۴. برای هر نقطه آغازین $y(t_0) = (x(t_0)^T, u(t_0)^T)^T$ ، سیستم (۵۶.۴)–(۵۷.۴) دارای جواب یکتای پیوسته $y(t) = (x(t)^T, u(t)^T)^T$ است.

اثبات. چون $\nabla f(x)$ و $\nabla g_k(x)$ ، $(k = 1, \dots, m)$ ، توابع به‌طور پیوسته مشتق پذیر روی مجموعه باز $D \subseteq \mathbb{R}^{n+m}$ هستند، بنابراین $\nabla f(x) + \nabla g(x)^T(u + g(x))^+ - u$ و لیپ شیتز محلی پیوسته روی مجموعه باز $D \subseteq \mathbb{R}^{n+m+l}$ هستند، با توجه به قضیه (۸.۳.۱)، شبکه عصبی (۵۶.۴)–(۵۷.۴) دارای جواب یکتای پیوسته $y(t)$ برای $t \in [t_0, \tau)$ وقتی $\tau \rightarrow \infty$ است. $\square$

لم ۴.۴.۴. ماتریس ژاکوبی $\nabla \Phi(y)$ از نگاشت Φ تعریف شده در (۵۵.۴) نیمه معین منفی است.

اثبات. بدون کاستن از کلیت مسأله، فرض می‌کنیم $0 < p < m$ وجود دارد به‌طوری که

$$(u + g)^+ = (u_1 + g_1(x), u_2 + g_2(x), \dots, u_p + g_p(x), \underbrace{0, 0, \dots, 0}_{m-p}) \geq 0.$$

با محاسبات ساده می‌توان نشان داد که

$$\nabla \Phi(y) = \begin{bmatrix} -(\nabla^T f(x) + \sum_{k=1}^p ((u_k + g_k) \nabla^T g_k^p(x)) + \nabla g^p(x)^T \nabla g^p(x)) & -\nabla g^p(x)^T \\ \nabla g^p(x) & S_{m \times m} \end{bmatrix},$$

که در آن

$$\nabla g^p(x) = \begin{bmatrix} U_{p \times n} \\ O_{(m-p) \times n} \end{bmatrix} = \begin{bmatrix} \frac{\partial g_1}{\partial x_1} & \cdots & \frac{\partial g_1}{\partial x_{p-1}} & \frac{\partial g_1}{\partial x_p} & \frac{\partial g_1}{\partial x_{p+1}} & \cdots & \frac{\partial g_1}{\partial x_n} \\ \frac{\partial g_2}{\partial x_1} & \cdots & \frac{\partial g_2}{\partial x_{p-1}} & \frac{\partial g_2}{\partial x_p} & \frac{\partial g_2}{\partial x_{p+1}} & \cdots & \frac{\partial g_2}{\partial x_n} \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots \\ \frac{\partial g_p}{\partial x_1} & \cdots & \frac{\partial g_p}{\partial x_{p-1}} & \frac{\partial g_p}{\partial x_p} & \frac{\partial g_p}{\partial x_{p+1}} & \cdots & \frac{\partial g_p}{\partial x_n} \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots \end{bmatrix},$$

و $\nabla^2 g_k^p(x)$ برای $(k = 1, 2, \dots, p)$

$$\nabla^2 g_k^p(x) = \begin{bmatrix} \frac{\partial^2 g_k}{\partial x_1^2} & \cdots & \frac{\partial^2 g_k}{\partial x_1 \partial x_{p-1}} & \frac{\partial^2 g_k}{\partial x_1 \partial x_p} & \frac{\partial^2 g_k}{\partial x_1 \partial x_{p+1}} & \cdots & \frac{\partial^2 g_k}{\partial x_1 \partial x_n} \\ \frac{\partial^2 g_k}{\partial x_2 \partial x_1} & \cdots & \frac{\partial^2 g_k}{\partial x_2 \partial x_{p-1}} & \frac{\partial^2 g_k}{\partial x_2 \partial x_p} & \frac{\partial^2 g_k}{\partial x_2 \partial x_{p+1}} & \cdots & \frac{\partial^2 g_k}{\partial x_2 \partial x_n} \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots \\ \frac{\partial^2 g_k}{\partial x_p \partial x_1} & \cdots & \frac{\partial^2 g_k}{\partial x_p \partial x_{p-1}} & \frac{\partial^2 g_k}{\partial x_p \partial x_p} & \frac{\partial^2 g_k}{\partial x_p \partial x_{p+1}} & \cdots & \frac{\partial^2 g_k}{\partial x_p \partial x_n} \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots \end{bmatrix},$$

9

$$S_{m \times m} = \begin{bmatrix} O_{p \times p} & O_{p \times (m-p)} \\ O_{(m-p) \times p} & -I_{(m-p) \times (m-p)} \end{bmatrix},$$

که در آن O ماتریس صفر است. ماتریس $\nabla g^p(x)^T \nabla g^p(x)$ ، ماتریس نیمه معین مثبت است. چون توابع

$\nabla^2 g_k(x)$ و $\nabla^2 f(x)$ محذب و دوبار مشتق پذیر فرض شده، بنابراین ماتریس‌های هسین $\nabla^2 f(x)$ و $\nabla^2 g_k(x)$

برای $(k = 1, 2, \dots, p)$ ، ماتریس‌های نیمه معین مثبت هستند. به علاوه از معین مثبت بودن $\nabla^2 g_k(x)$

می‌توان گفت که ماتریس‌های $\nabla^2 g_k^p(x)$ نیمه معین مثبت هستند. همچنین واضح است که ماتریس

$S_{m \times m}$ ، نیمه معین منفی است.

طبق توضیحات گفته شده، ما می‌توانیم نتیجه بگیریم ماتریس ژاکوبی $\nabla\Phi(y)$ ، نیمه معین منفی است.

اگر $p = m$ یعنی $(u + g)^+ = (u_1 + g_1(x), u_2 + g_2(x), \dots, u_m + g_m(x))$ ، آنگاه داریم:

$$\nabla\Phi(y) = \begin{bmatrix} -(\nabla^* f(x) + \sum_{k=1}^m ((u_k + g_k) \nabla^* g_k(x)) + \nabla g(x)^T \nabla g(x)) & -\nabla g(x)^T \\ \nabla g(x) & O_{m \times m} \end{bmatrix}.$$

که در آن $\nabla g(x) = \nabla g^m(x)$ ماتریس ژاکوبی از توابع برداری $g(x)$ و $\nabla^* g_k(x) = \nabla^* g_k^m(x)$ ماتریس هسین از $g_k(x)$ برای $(k = 1, 2, \dots, m)$ هستند. مشابه حالت قبلی، می‌توان بسادگی نشان داد که $\nabla\Phi(y)$ ماتریس نیمه معین منفی است.

سرانجام، اگر $p = 0$ یعنی $(u + g)^+ = (\underbrace{0, 0, \dots, 0}_m)$ ، آنگاه داریم:

$$\nabla\Phi(y) = \begin{bmatrix} -\nabla^* f(x) & O_{n \times m} \\ O_{m \times n} & -I_{m \times m} \end{bmatrix}.$$

در این حالت هم، می‌توان نشان داد که $\nabla\Phi(y)$ ماتریس نیمه معین منفی است، این اثبات را کامل می‌کند. $\square$

قضیه ۵.۴.۴. مدل شبکه عصبی پیشنهاد شده در (۵۶.۴)–(۵۷.۴)، پایدار مجانبی به مفهوم لیاپانوف و همگرایی سراسری به $y^* = (x^{*T}, u^{*T})^T$ ، وقتی x^* جواب بهینه از (۵۰.۴)–(۵۱.۴) است.

اثبات. تابع لیاپانوف زیر را در نظر بگیرید:

$$E(y) = \|\Phi(y)\|^2 + \frac{1}{\gamma} \|y - y^*\|^2. \quad (۶۲.۴)$$

از (۵۵.۴)، داریم:

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

بنابراین

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

با بکار بردن لم (۴.۴.۴) داریم:

$$\Phi^T(y)(\nabla \Phi(y)^T + \nabla \Phi(y))\Phi(y) \leq 0, \quad \forall y \neq y^*. \quad (۴۳.۴)$$

با استفاده از تعریف (۶.۳.۱) و قضیه (۷.۳.۱) داریم:

$$(y - y^*)^T (\Phi(y) - \Phi(y^*)) = (y - y^*)^T \Phi(y) \leq 0, \quad \forall y \neq y^*.$$

بنابراین

$$\frac{dE(y(t))}{dt} \leq 0. \quad (۴۴.۴)$$

این بدان معنی است که شبکه عصبی (۵۷.۴)–(۵۶.۴)، پایدار مجانبی به مفهوم لیاپانوف است. در ادامه

چون

$$E(y) \geq \frac{1}{4} \|y - y^*\|^2, \quad (۴۵.۴)$$

زیردنباله‌های همگرای

$$\{(x(t_k)^T, u(t_k)^T)^T | t_0 < t_1 < \dots < t_k < t_{k+1}\},$$

وقتی $t_k \rightarrow \infty$ از $k \rightarrow \infty$ وجود دارد، به‌طوری‌که

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

وقتی $(\bar{x}^T, \bar{u}^T)^T$ در رابطه زیر صدق می‌کند:

$$\frac{dE(y(t))}{dt} = 0,$$

این نشان می‌دهد که $(\bar{x}^T, \bar{u}^T)^T$ ، یک نقطه $-\omega$ حدی از $\{(x(t)^T, u(t)^T)^T | t \geq t_0\}$ است. با استفاده از قضیه (۴.۴.۱) داریم $\{(x(t)^T, u(t)^T)^T \rightarrow M\}$ برای $t \rightarrow \infty$ ، که در آن M بزرگترین مجموعه ناورد

در $\{ \circ \} = \{(x(t)^T, u(t)^T)^T | \frac{dE(y(t))}{dt} = \circ\}$ است. از (۵۶.۴)–(۵۷.۴) و (۶۴.۴) داریم

$$\frac{dx}{dt} = \circ, \quad \frac{du}{dt} = \circ \Leftrightarrow \frac{dE(y(t))}{dt} = \circ.$$

بنابراین $(\bar{x}^T, \bar{u}^T)^T \in D^*$ با $M \subseteq K \subseteq D^*$.

با جایگذاری $x^* = \bar{x}$ و $u^* = \bar{u}$ در (۶۲.۴)، ما تابع لیاپانوف دیگری به صورت زیر تعریف می‌کنیم،

$$\bar{E}(y) = \|\Phi(y)\|^2 + \frac{1}{\gamma} \|y - \bar{y}\|^2. \quad (۶۶.۴)$$

آنگاه $\bar{E}(y)$ به طور پیوسته مشتق پذیر و $\bar{E}(\bar{y}) = \circ$ از طرفی

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

از اینرو داریم

$$\lim_{k \rightarrow \infty} \bar{E}(x(t_k)^T, u(t_k)^T)^T = \bar{E}(\bar{x}, \bar{u}).$$

بنابراین، به ازای هر $\epsilon > \circ$ ، $q > \circ$ وجود دارد به طوری که برای هر $t \geq t_q$ ، داریم $\bar{E}(y(t)) < \epsilon$.

مشابه می‌توانیم به دست بیاوریم $\frac{d\bar{E}(y(t))}{dt} \leq \circ$ برای $t \geq t_q$ داریم،

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

در نتیجه $\lim_{t \rightarrow \infty} \|y(t) - \bar{y}\| = \circ$ یا $\lim_{t \rightarrow \infty} y(t) = \bar{y}$. بنابراین، شبکه عصبی پیشنهاد شده در

(۵۶.۴)–(۵۷.۴) همگرایی سراسری به نقطه تعادل $(\bar{x}^T, \bar{u}^T)^T$ است که در آن $\bar{x}$ جواب بهینه از

□

(۵۰.۴)–(۵۱.۴) است.

ما از قضیه (۵.۴.۴)، گزاره زیر را نتیجه می‌گیریم.

گزاره ۶.۴.۴. اگر $D^* = \{(x^{*T}, u^{*T})^T\}$ ، آنگاه شبکه عصبی (۵۶.۴)–(۵۷.۴) برای حل کردن (۵۰.۴)–(۵۱.۴)، پایدار مجانبی سراسری به نقطه تعادل یکتای $y^* = (x^{*T}, u^{*T})^T$ است.

۵.۴ مثال‌های عددی

به منظور نشان دادن کارایی و تأثیر شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴)، در این قطعه، ما به ذکر چندین مثال عددی می‌پردازیم. برای برخی مثال‌ها، ما مقایسه عملکرد عددی از شبکه عصبی پیشنهادی را با مقادیر گوناگون از حالت اولیه $y(0)$ داریم. برای برخی دیگر از مسائل ما مقایسه عملکرد عددی از شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) با مقادیر گوناگون از نرخ همگرایی κ و مقادیر مختلف $\|x(t) - x^*\|^2$ با نقطه آغازین مشخص y انجام دادیم. نتایج شبیه‌سازی در پکیج *ode45*، نرم افزار *Matlab* انجام شده و جواب واقعی مسائل نیز توسط نرم افزار *Lingo ۱۱* محاسبه شده است.

مثال ۱.۵.۴. یک بنگاه تولیدی دو قطعه را با استفاده از دستگاههای تراش، فرز و مته تولید می‌کند. زمان موجود دستگاه‌های مختلف در هفته، سود هر قطعه در جدول زیر داده شده است. زمان ماشینکاری مورد نیاز هر قطعه بر روی دستگاه‌های مختلف به طور دقیق معلوم نیست (به کارگر مربوط بستگی دارد)، اما می‌دانیم که دارای توزیع نرمال با میانگین و انحراف معیار داده شده در جدول (۱.۴) است. تعداد قطعات I و II را که باید در هفته تولید شود به گونه‌ای تعیین کنید که سود بیشینه شود و در زمان صرف شده برای ماشینکاری بیش از ۱٪ از مقدار موجود تجاوز نکند.

حل: اگر تعداد قطعات I و II تولید شده در هفته را به ترتیب با x_1 و x_2 نشان دهیم، مسأله را می‌توان

به صورت زیر بیان کرد:

$$\text{Maximize} \quad f = 50x_1 + 100x_2 \quad (67.4)$$

جدول ۱.۴: مقادیر مثال ۱.۵.۴

نوع دستگاه		زمان ماشینکاری مورد نیاز هر قطعه (دقیقه)			
میانگین	انحراف معیار	قطعه I		قطعه II	
		میانگین	انحراف معیار	میانگین	انحراف معیار
تراش فرز مته	$\bar{a}_{11} = 10$	$\sigma_{a_{11}} = 6$	$\bar{a}_{12} = 5$	$\sigma_{a_{12}} = 4$	$b_1 = 2500$
	$\bar{a}_{21} = 4$	$\sigma_{a_{21}} = 4$	$\bar{a}_{22} = 10$	$\sigma_{a_{22}} = 7$	$b_2 = 2000$
	$\bar{a}_{31} = 1$	$\sigma_{a_{31}} = 2$	$\bar{a}_{32} = 1/5$	$\sigma_{a_{32}} = 3$	$b_3 = 450$
سود هر واحد (R_s)		$c_1 = 50$		$c_2 = 100$	

subject to

$$\begin{cases} P[a_{11}x_1 + a_{12}x_2 \leq 2500] \geq 0.99, \\ P[a_{21}x_1 + a_{22}x_2 \leq 2000] \geq 0.99, \\ P[a_{31}x_1 + a_{32}x_2 \leq 450] \geq 0.99, \\ x_1, x_2 \geq 0. \end{cases} \quad (68.4)$$

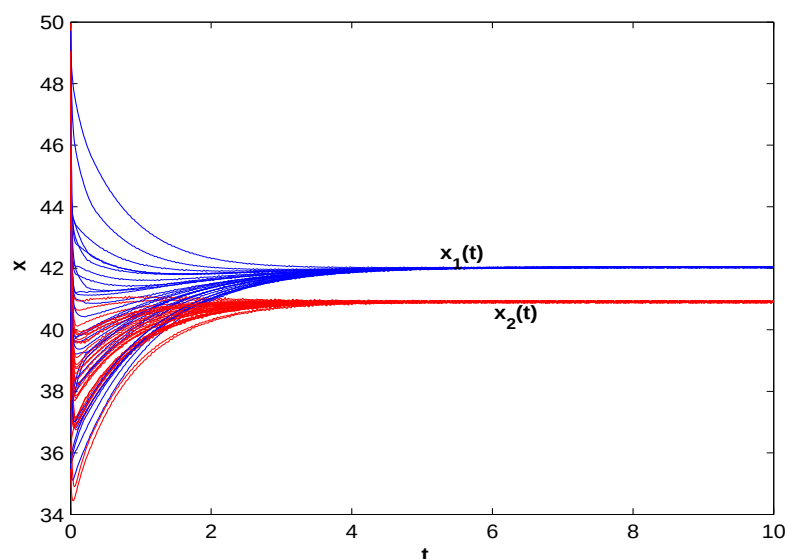
ار آنجا که اطلاعاتی درباره کواریانس داده نشده است، ضرایب a_{ij} را می‌توان به عنوان متغیرهای مستقل در نظر گرفت. مقدار متغیر تصادفی نرمال استاندارد e_i برای $(i = 1, \dots, 3)$ مربوط به احتمال ۹۹٪ را می‌توان از جدول توزیع نرمال استاندارد در [۱۲۷] به‌دست آورد. بنابراین معادل قطعی مسأله (۶۷.۴)–(۶۸.۴) با کمک (۱۷.۴)–(۱۸.۴) به‌صورت بیان می‌شود:

$$\text{Maximize } f = 50x_1 + 100x_2 \quad (69.4)$$

subject to

$$\begin{cases} 10x_1 + 5x_2 + 2/33\sqrt{36x_1^2 + 16x_2^2} - 2500 \leq 0, \\ 4x_1 + 10x_2 + 2/33\sqrt{16x_1^2 + 49x_2^2} - 2000 \leq 0, \\ x_1 + 1/5x_2 + 2/33\sqrt{4x_1^2 + 9x_2^2} - 450 \leq 0, \\ x_1, x_2 \geq 0. \end{cases} \quad (70.4)$$

جواب بهینه این مسأله $x^* = (42/0.3227, 40/0.893)^T$ است. ما برای حل این مسأله شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) را به‌کار می‌بریم. شکل (۲.۴) نشان می‌دهد که مسیرهای شبکه عصبی پیشنهادی با ۳۰ نقطه آغازین متفاوت و $\epsilon = 10^{-6}$ همگرا به جواب بهینه مسأله (۶۹.۴)–(۷۰.۴) است. منحنی فاز، شبکه پیشنهادی با ۱۴ حالت اولیه گوناگون در شکل (۳.۴) نشان داده شده است. ما همچنین



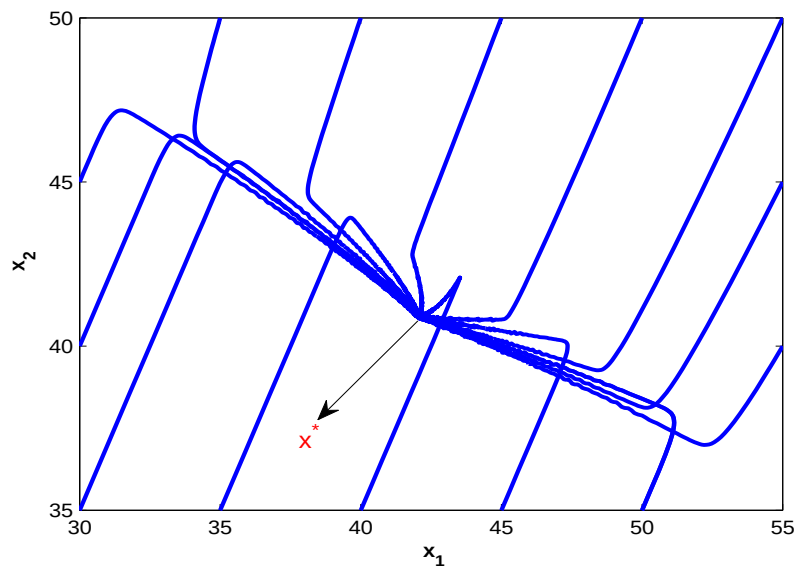
شکل ۲.۴: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۳۰ نقطه آغازین دلخواه در مثال ۱.۵.۴.

تأثیر پارامتر κ از شبکه عصبی (۵۶.۴)–(۵۷.۴) روی مقدار $\|x(t) - x^*\|^2$ بررسی می‌کنیم. در شکل (۴.۴)، مشاهده می‌کنیم κ بزرگتر، نرخ همگرایی بهتری برای $\|x(t) - x^*\|^2$ را نتیجه می‌دهد.

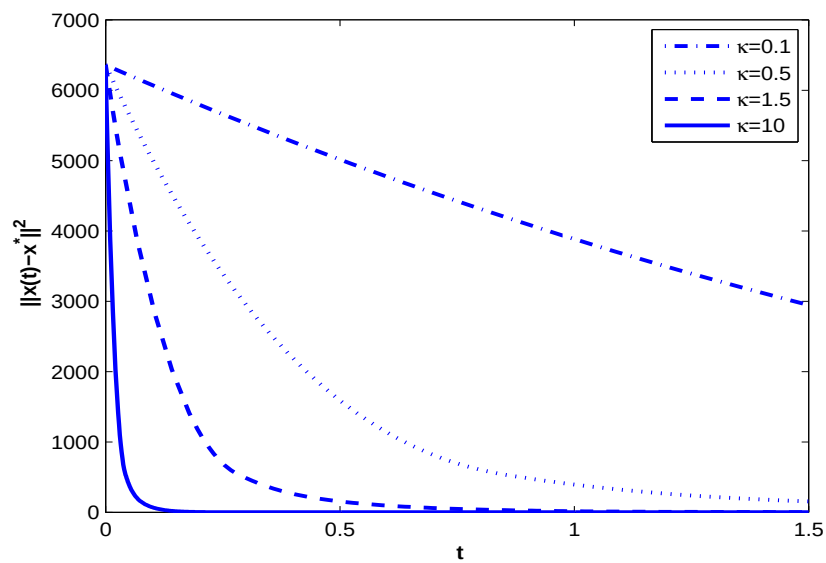
مثال ۲.۵.۴. اگر زمانهای ماشینکاری موجود بر روی دستگاه‌های مختلف، احتمالی (با توزیع نرمال) و پارامترهای مشخص شده در جدول (۲.۴) باشند، تعداد قطعات I و II را که باید در هفته ساخته شوند تا سود بیشینه شود، پیدا کنید.

جدول ۲.۴: مقادیر مثال ۲.۵.۴

نوع دستگاه		زمان ماشینکاری مورد نیاز هر قطعه (دقیقه)		حداکثر زمان موجود در هفته (دقیقه)	
		قطعه I	قطعه II	میانگین	انحراف معیار
تراش فرز مته سود هر واحد (Rs)	$a_{11} = 10$	$a_{12} = 5$	$\bar{b}_1 = 2500$	$\sigma_{b_1} = 500$	
	$a_{21} = 4$	$a_{22} = 10$	$\bar{b}_2 = 2000$	$\sigma_{b_2} = 400$	
	$a_{31} = 1$	$a_{32} = 1/5$	$\bar{b}_3 = 450$	$\sigma_{b_3} = 50$	
	$c_1 = 50$	$c_2 = 100$			



شکل ۳.۴: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۴ نقطه آغازین در مثال ۱.۵.۴.



شکل ۴.۴: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$ در مثال ۱.۵.۴.

مقدار واریته نرمال استاندارد (E_i) برای $(i = 1, 2, 3)$ که در آن $\phi(E_i) = 1 - p_i = 0.01$ است را نمی‌توان مستقیماً از جدول توزیع نرمال استاندارد به‌دست آورد. لکن توجه داریم که چون $1 - p_i < 0.5$ در نتیجه $E_i < 0$ و بنابراین

$$\phi(-E_i) = 1 - \phi(E_i) = p_i = 0.99 \implies E_i = -2.33.$$

بنابراین از $(26.4)-(25.4)$ داریم:

$$\text{Maximize } f = 50x_1 + 100x_2 \quad (71.4)$$

subject to

$$\begin{cases} 10x_1 + 5x_2 - (-2.33)(500) - 2500 \leq 0, \\ 4x_1 + 10x_2 - (-2.33)(400) - 2000 \leq 0, \\ x_1 + 1.5x_2 - (-2.33)(50) - 450 \leq 0, \\ x_1, x_2 \geq 0. \end{cases} \quad (72.4)$$

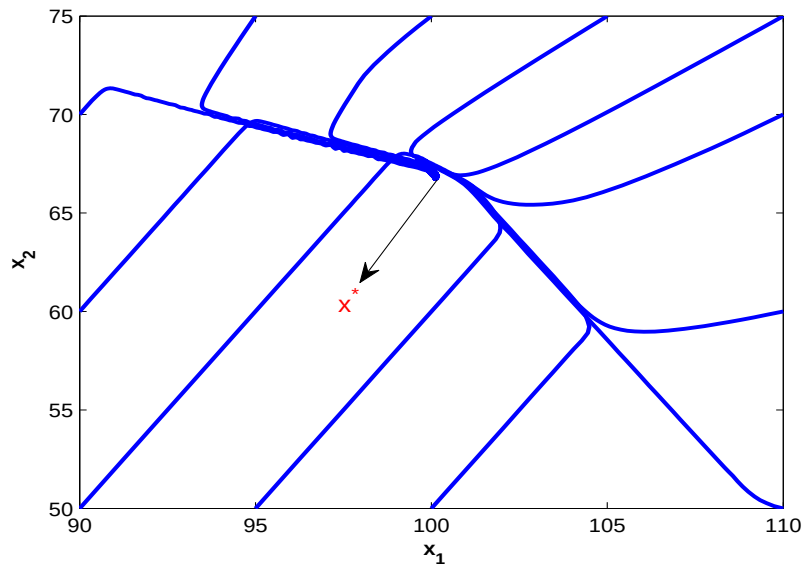
جواب بهینه این مسأله $x^* = (100/1250, 66/7500)^T$ است. ما شبکه عصبی پیشنهاد شده در $(56.4)-(57.4)$ را برای حل این مسأله بکار می‌بریم، شکل (5.4) منحنی فاز از $(56.4)-(57.4)$ با ۱۲ نقطه اولیه مختلف با $\epsilon = 10^{-6}$ نشان می‌دهد. واضح است حتی وقتی نقطه اولیه یک نقطه نشدنی باشد، مسیر از شبکه پیشنهادی هنوز به x^* همگراست. نرم خطای L_2 بین x و x^* با نقاط اولیه گوناگون هم در شکل (6.4) رسم شده است.

مثال ۳.۵.۴. فرض کنید که سود هر یک از قطعات I و II در مثال $(1.5.4)$ متغیری تصادفی با توزیع نرمال باشد. تعداد قطعاتی که باید در هر هفته تولید کرد تا سود بیشینه شود را بدست آورید. میانگین و انحراف سود به ترتیب برای هر قطعه I ، 20 ، 50 و برای هر قطعه II ، 50 ، 100 روپیه است.

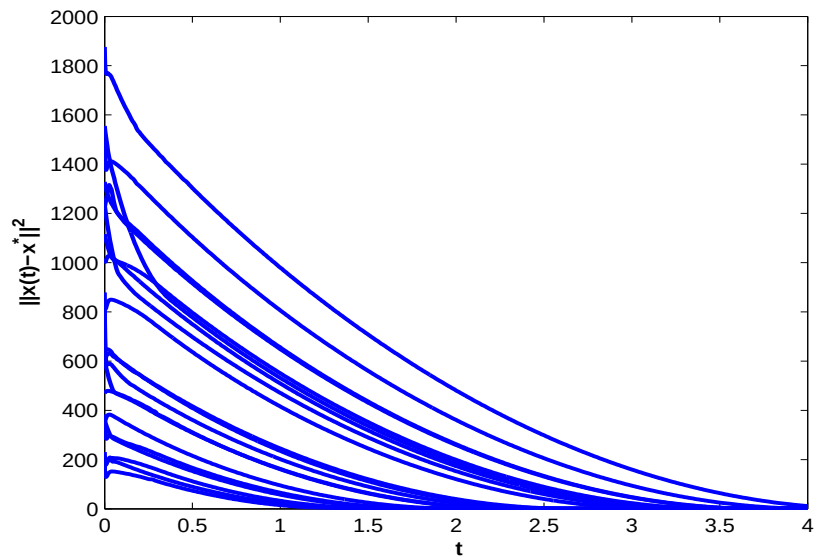
حل: با این فرض که سودها متغیرهای تصادفی مستقل هستند، مسأله قطعی معادل را با استفاده از

$(31.4)-(32.4)$ می‌توان به‌صورت زیر بیان کرد:

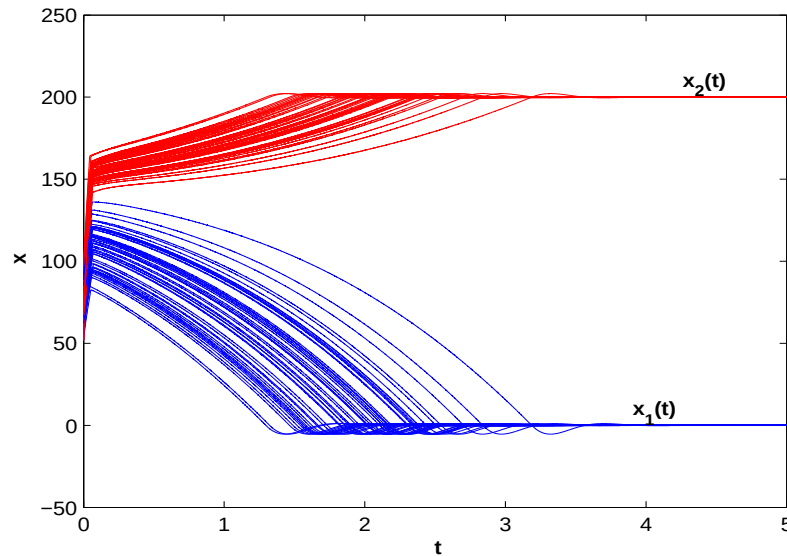
$$\text{Maximize } F = k_1(50x_1 + 100x_2) + k_2\sqrt{400x_1^2 + 2500x_2^2} \quad (73.4)$$



شکل ۴.۵: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۲ نقطه آغازین در مثال ۲.۵.۴.



شکل ۴.۶: رفتار همگرایی از $\|x(t) - x^*\|^2$ با ۲۰ نقطه اولیه دلخواه در مثال ۲.۵.۴.



شکل ۷.۴: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۵۰ نقطه آغازین دلخواه در مثال ۳.۵.۴.

subject to

$$\begin{cases} 10x_1 + 5x_2 - 2500 \leq 0, \\ 4x_1 + 10x_2 - 2000 \leq 0, \\ x_1 + 1/5x_2 - 450 \leq 0, \\ x_1, x_2 \geq 0. \end{cases} \quad (74.4)$$

وقتی که k_1, k_2 مقادیر ثابتی هستند که اهمیت نسبی میانگین و انحراف معیار سود را برای کمینه سازی نشان می‌دهند. جواب بهینه مسأله با $k_1 = k_2 = 1$ و $\epsilon = 10^{-6}$ و $x^* = (0, 200)^T$ است. قضیه (۵.۴.۴) و گزاره (۶.۴.۴) تضمین می‌کند که مدل بیان شده در (۵۶.۴)–(۵۷.۴)، همگرای سراسری به جواب بهینه یکتاست. شکل (۷.۴) نشان می‌دهد که مسیرهای شبکه پیشنهاد شده برای حل مسأله بالا با ۵۰ نقطه اولیه تصادفی همگرا به جواب بهینه مسأله است.

مثال ۴.۵.۴. اگر زمانهای ماشینکاری مورد نیاز، حداکثر زمان موجود و سودهای هر واحد همه متغیرهای تصادفی با توزیع نرمال در داده های جدول (۳.۴) فرض کنیم، تعداد قطعاتی را که باید ساخته شوند تا سود بیشینه شود را پیدا کنید. قیود باید با احتمال دست کم ۹۹٪ برقرار باشند.

جدول ۳.۴: مقادیر مثال ۴.۵.۴

نوع دستگاه	زمان ماشینکاری مورد نیاز هر قطعه(دقیقه)				حداکثر زمان موجود در هفته (دقیقه)
	قطعه I		قطعه II		
	میانگین	انحراف معیار	میانگین	انحراف معیار	
تراش فرز مته	$\bar{a}_{11} = 10$	$\sigma_{a_{11}} = 6$	$\bar{a}_{12} = 5$	$\sigma_{a_{12}} = 4$	$b_1 = 2500 \quad \sigma_{b_1} = 500$
	$\bar{a}_{21} = 4$	$\sigma_{a_{21}} = 4$	$\bar{a}_{22} = 10$	$\sigma_{a_{22}} = 7$	$b_2 = 2000 \quad \sigma_{b_2} = 400$
	$\bar{a}_{31} = 1$	$\sigma_{a_{31}} = 2$	$\bar{a}_{32} = 1/5$	$\sigma_{a_{32}} = 3$	$b_3 = 450 \quad \sigma_{b_3} = 50$
سود هر واحد(R_s)	$c_1 = 50 \quad \sigma_{c_1} = 20$		$c_2 = 100 \quad \sigma_{c_2} = 50$		

حل: با تعریف متغیرهای تصادفی جدید h_i به صورت زیر

$$h_i = \sum_{j=1}^n a_{ij}x_j - b_i,$$

در می‌یابیم که h_i هم دارای توزیع نرمال است. با فرض اینکه بین a_{ij} و b_i ها همبستگی وجود ندارد،

میانگین ها و واریانس های h_i را می‌توان از روابط (۳۶.۴) و (۳۷.۴) بدست آورد. تابع هدف را همان F

در مثال قبل می‌گیریم و قیود هم به صورت زیر بیان می‌شوند:

$$P[h_i \leq 0] \geq 0.99, \quad i = 1, 2, 3,$$

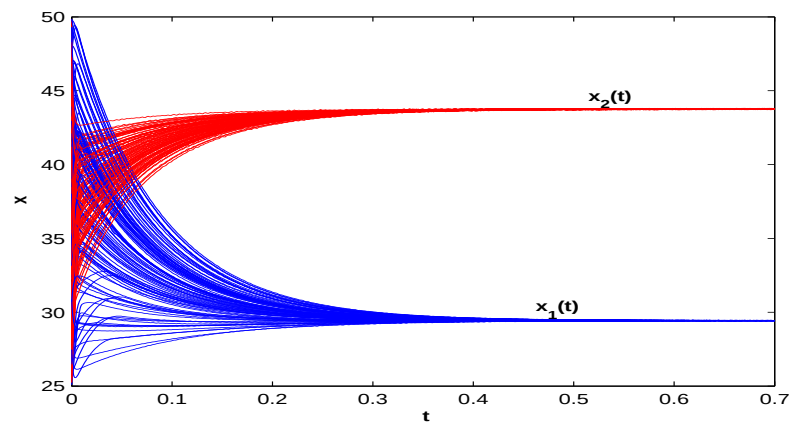
چون مقدار واریته نرمال استاندارد e_i مربوط به احتمال ۰/۹۹ برابر ۲/۳۳ است، می‌توانیم مسأله

بهینه‌سازی غیرخطی معادل را به صورت زیر بیان کنیم:

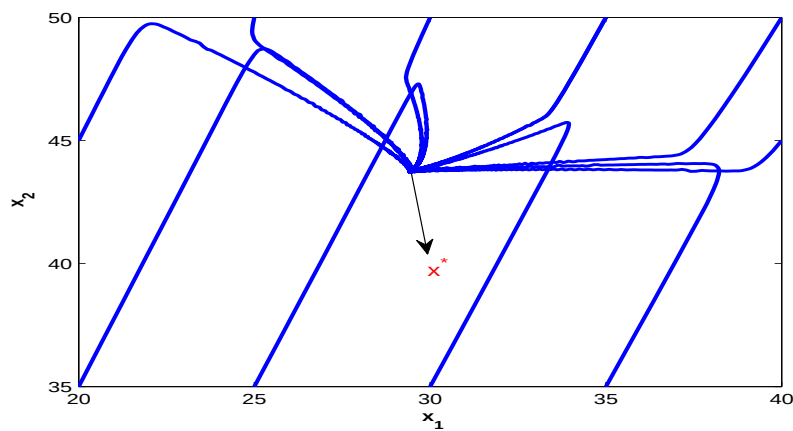
$$\begin{aligned} &\text{Maximize } F = k_1(50x_1 + 100x_2) + k_2\sqrt{400x_1^2 + 2500x_2^2} \\ &\text{subject to} \\ &\begin{cases} 10x_1 + 5x_2 + 2/33\sqrt{36x_1^2 + 16x_2^2} + 250000 - 2500 \leq 0, \\ 4x_1 + 10x_2 + 2/33\sqrt{16x_1^2 + 49x_2^2} + 160000 - 2000 \leq 0, \\ x_1 + 1/5x_2 + 2/33\sqrt{4x_1^2 + 9x_2^2} + 2500 - 450 \leq 0, \\ x_1, x_2 \geq 0. \end{cases} \end{aligned} \quad (75.4)$$

جواب بهینه این مسأله برای $k_1 = k_2 = 1$ ، $x^* = (29/43502, 43/76386)^T$ است. شکل (۸.۴)

نشان می‌دهد که مسیرهایی از شبکه پیشنهاد شده در (۵۶.۴)–(۵۷.۴) برای حل مسأله بالا با ۱۰۰ نقطه



شکل ۸.۴: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۰۰ نقطه آغازین دلخواه در مثال ۴.۵.۴.



شکل ۹.۴: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۰۰ نقطه آغازین در مثال ۴.۵.۴.

آغازین دلخواه همگرا به جواب بهینه مسأله است. منحنی فاز این شبکه عصبی نیز، در شکل (۹.۴) رسم شده است.

فصل ۵

مسائل برنامه‌ریزی کسری

۱.۵ مقدمه

در این فصل ابتدا برنامه‌ریزی کسری^۱ را معرفی می‌کنیم. مروری بر کاربردهای آنها داریم سپس به معرفی رده خاصی از برنامه‌ریزی کسری که می‌توان آنها را به برنامه‌ریزی مخروط مرتبه دوم تبدیل کرد، می‌پردازیم. در انتها چندین مثال عددی از مسائل کسری را با شبکه عصبی (۵۶.۴)–(۵۷.۴) که برای حل مسائل برنامه‌ریزی مخروط مرتبه دوم پیشنهاد کردیم، حل می‌کنیم.

۲.۵ برنامه‌ریزی کسری

برنامه‌ریزی کسری یک موضوع جالب در بسیاری از مسائل بهینه‌سازی از جمله انتخاب سبد سرمایه، تولید، تئوری اطلاعات^۲ و تصمیم‌گیری‌های بسیار پیچیده در علم مدیریت است، به‌طور خاص در مهندسی و اقتصاد برای حداقل‌سازی توابع اقتصادی یا فیزیکی یا هردو از جمله هزینه/زمان، هزینه/حجم، هزینه/سود و غیره به‌منظور اندازه‌کارایی و بهره‌وری سیستم کاربرد دارد. بسیاری کاربردهای دیگر، مسائل برنامه‌ریزی کسری توسط کراون^۳ در [۱۲۹]، ایباراکی^۴ و همکارش در [۱۳۰] ارائه شده است. برای بررسی گسترده‌تر برنامه‌ریزی کسری می‌تواند مراجع [۱۳۱]–[۱۳۵] را ببینید.

در مهندسی و بسیاری از کاربردهای علمی، جوابهای زمان واقعی مسائل برنامه‌ریزی کسری مطلوب است. در نتیجه به روشهای مؤثر برای حل برنامه‌ریزی کسری زمان واقعی نیاز داریم. نتایج مطالعات قبلی حاکی از آن است که از شبکه‌های عصبی می‌توان برای حل مسائل مختلف بهینه‌سازی استفاده کرد. ایده اصلی شبکه‌های عصبی برای بهینه‌سازی ایجاد یک تابع انرژی نامنفی و پایه‌گذاری یک سیستم دینامیکی است که یک شبکه عصبی مصنوعی را ایجاد می‌کند. سیستم دینامیکی معمولاً به فرم معادلات دیفرانسیل مرتبه اول است. به‌علاوه انتظار می‌رود که سیستم دینامیکی به یک وضعیت پایدار (یا نقطه تعادل) که

^۱Fractional Programming

^۲Information theory

^۳Craven

^۴Ibaraki

متناظر با جواب مسأله بهینه‌سازی اصلی است که با یک نقطه اولیه آغاز می‌شود نزدیک گردد. با توجه به مباحث فوق، در این فصل یک مدل از شبکه‌های عصبی برای حل مسائل برنامه‌ریزی کسری ارائه می‌دهیم.

ابتدا رده‌ای از مسائل برنامه‌ریزی کسری را تبدیل به مسائل برنامه‌ریزی مخروط مرتبه دوم محدب کرده، سپس این مسائل را تبدیل به برنامه‌ریزی غیرخطی محدب با قیود نامساوی می‌کنیم، زمانی که با استفاده از تکنیک هموار سازی، قیود مخروط مرتبه دوم ناهموار به‌فرم هموار آن تبدیل شوند. براساس شرایط بهینگی $(K.K.T)$ برنامه‌ریزی محدب، مدل شبکه عصبی برای حل مسائل برنامه‌ریزی غیرخطی محدب پیشنهاد خواهد شد. همچنین ثابت می‌شود نقطه تعادل شبکه عصبی پیشنهاد شده برابر نقطه $K.K.T$ مسأله برنامه‌ریزی مخروط مرتبه دوم محدب است. وجود و یکتایی نقطه تعادل شبکه عصبی نیز مورد آنالیز قرار خواهد گرفت. با ساختن یک تابع لیاپانوف مناسب، شرط کافی برای تضمین وجود و پایدار مجانبی سراسری برای نقطه تعادل یکتای شبکه پیشنهاد شده بدست خواهد آمد.

در این بخش رده‌های خاصی از مسائل برنامه‌ریزی کسری را مورد بررسی قرار می‌دهیم که بتوان آن‌ها را به مسائلی از نوع برنامه‌ریزی مخروط مرتبه دوم تبدیل کرد. برای تبدیل برنامه‌ریزی کسری به برنامه‌ریزی مخروط مرتبه دوم به نامساوی زیر نیاز داریم: [۱۲۵]

اگر $z, y \geq 0$ و w یک عدد حقیقی دلخواه باشد، آنگاه داریم:

$$w^2 \leq zy \iff \left\| \begin{bmatrix} 2w \\ z - y \end{bmatrix} \right\| \leq z + y, \quad (۱.۵)$$

و در صورت کلی، وقتی w یک بردار حقیقی است:

$$w^T w \leq zy \iff \left\| \begin{bmatrix} 2w \\ z - y \end{bmatrix} \right\| \leq z + y. \quad (۲.۵)$$

۱.۲.۵ برنامه‌ریزی کسری با صورت و مخرج خطی

اولین حالتی را که در برنامه‌ریزی کسری مورد بحث قرار می‌دهیم به صورت زیر بیان می‌شود ([۱۲۸])

$$\begin{aligned} &\text{Minimize} \quad \sum_{i=1}^p \frac{c}{a_i^T x + b_i} \\ &\text{subject to} \end{aligned} \quad (۳.۵)$$

$$\begin{cases} a_i^T x + b_i > 0, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q. \end{cases} \quad (۴.۵)$$

برای تبدیل مسأله بالا به برنامه‌ریزی مخروط مرتبه دوم ابتدا، متغیرهای جدید u_i را معرفی می‌کنیم

بنابراین مسأله زیر را داریم:

$$\begin{aligned} &\text{Minimize} \quad \sum_{i=1}^p u_i \\ &\text{subject to} \end{aligned} \quad (۵.۵)$$

$$\begin{cases} \frac{c}{a_i^T x + b_i} \leq u_i, & i = 1, \dots, p, \\ u_i > 0, & i = 1, \dots, p, \\ a_i^T x + b_i > 0, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q, \end{cases} \quad (۶.۵)$$

با فرض اینکه

$$z_i = a_i^T x + b_i, \quad y_i = u_i, \quad w_i = \sqrt{c}, \quad i = 1, \dots, p, \quad (۷.۵)$$

و با به کار بردن نامساوی (۱.۵) داریم:

$$c \leq (a_i^T x + b_i)u_i \iff \left\| \begin{bmatrix} \sqrt{c} \\ a_i^T x + b_i - u_i \end{bmatrix} \right\| \leq a_i^T x + b_i + u_i, \quad i = 1, \dots, p.$$

بنابراین جواب مسأله برنامه‌ریزی کسری (۳.۵)-(۴.۵) از حل مسأله معادل برنامه‌ریزی مخروط مرتبه دوم

زیر به دست می‌آید:

$$\begin{aligned} &\text{Minimize} \quad \sum_{i=1}^p u_i \\ &\text{subject to} \end{aligned} \quad (۸.۵)$$

$$\begin{cases} \left\| \begin{bmatrix} \sqrt{c} \\ a_i^T x + b_i - u_i \end{bmatrix} \right\| \leq a_i^T x + b_i + u_i, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q. \end{cases} \quad (۹.۵)$$

۲.۲.۵ برنامه‌ریزی کسری با تابع هدف درجه دوم

مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\begin{aligned} \text{Minimize} \quad & \sum_{i=1}^p \frac{c}{a_i^T x + b_i} + x^T p_* x + 2q_*^T x + r \\ \text{subject to} \quad & \end{aligned} \quad (10.5)$$

$$\begin{cases} a_i^T x + b_i > 0, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q, \end{cases} \quad (11.5)$$

که در آن $q_* \in \mathbb{R}^n$ و $p_* \in \mathbb{R}^{n \times n}$ یک ماتریس معین مثبت است. چون p_* معین مثبت است، بنابراین

$$x^T p_* x + 2q_*^T x + r \quad \text{قابل تبدیل به فرم زیر است} \quad ([12.5])$$

$$\| p_*^{\frac{1}{2}} x + p_*^{-\frac{1}{2}} q_* \|^2 + (r - q_*^T p_*^{-1} q_*). \quad (12.5)$$

اکنون با معرفی متغیرهای t_* و u_i برای $i = 1, \dots, p$ به‌طوری‌که

$$\frac{c}{a_i^T x + b_i} \leq u_i, \quad i = 1, \dots, p, \quad (13.5)$$

$$x^T p_* x + 2q_*^T x + r \leq t_*, \quad (14.5)$$

و با بکارگیری نامساوی (۲.۵) و (۱۲.۵)–(۱۴.۵)، مسأله برنامه‌ریزی کسری (۱۰.۵)–(۱۱.۵) به فرم معادل

برنامه‌ریزی مخروط مرتبه دوم زیر تبدیل می‌شود:

$$\begin{aligned} \text{Minimize} \quad & \sum_{i=1}^p u_i + t_* \\ \text{subject to} \quad & \end{aligned} \quad (15.5)$$

$$\begin{cases} \left\| \begin{bmatrix} \sqrt{c} \\ a_i^T x + b_i - u_i \end{bmatrix} \right\| \leq a_i^T x + b_i + u_i, & i = 1, \dots, p, \\ \left\| \begin{bmatrix} 2(p_*^{\frac{1}{2}} x + p_*^{-\frac{1}{2}} q_*) \\ 1 - t_* \end{bmatrix} \right\| \leq t_* + 1, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q. \end{cases} \quad (16.5)$$

۳.۲.۵ برنامه‌ریزی کسری محدب با صورت درجه دوم

مسئله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \sum_{i=1}^p \frac{\|F_i x + g_i\|^2}{a_i^T x + b_i} \quad (17.5)$$

subject to

$$\begin{cases} a_i^T x + b_i > 0, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q, \end{cases} \quad (18.5)$$

که در آن $F_i \in \mathbb{R}^{q_i \times n}$, $g_i \in \mathbb{R}^{q_i}$ ($i = 1, \dots, p$) هستند. اکنون با معرفی متغیرهای جدید u_i به‌طوری‌که

$$\frac{\|F_i x + g_i\|^2}{a_i^T x + b_i} \leq u_i, \quad i = 1, \dots, p. \quad (19.5)$$

و با بکار بردن نامساوی (۲.۵)، مسئله (۱۷.۵)–(۱۸.۵) به فرم برنامه‌ریزی مخروط مرتبه دوم زیر تبدیل

می‌شود:

$$\text{Minimize} \quad \sum_{i=1}^p u_i \quad (20.5)$$

subject to

$$\begin{cases} \left\| \begin{bmatrix} \|F_i x + g_i\|^2 \\ a_i^T x + b_i - u_i \end{bmatrix} \right\| \leq a_i^T x + b_i + u_i, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q. \end{cases} \quad (21.5)$$

اکنون حالت خاصی از مسئله (۱۷.۵)–(۱۸.۵) با قیود درجه دوم به‌صورت زیر در نظر می‌گیریم:

$$\text{Minimize} \quad \sum_{i=1}^p \frac{x^T P_i x + q_i^T x + r_i}{a_i^T x + b_i} \quad (22.5)$$

subject to

$$\begin{cases} a_i^T x + b_i > 0, & i = 1, \dots, p, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q, \\ x^T Q_k x + m_k^T x + l_k \leq 0, & k = 1, \dots, s, \end{cases} \quad (23.5)$$

که در آن $m_k \in \mathbb{R}^n$ و ماتریس‌های $P_i \in \mathbb{R}^{n \times n}$ ($i = 1, \dots, p$)، $Q_k \in \mathbb{R}^{n \times n}$ ($k = 1, \dots, s$)، ماتریس‌های معین مثبت هستند. فرم مخروط مرتبه دوم متناظر با توابع درجه دوم مسأله بنابراین

(۲۲.۵)–(۲۳.۵) به صورت زیر بیان می‌شوند: [۱۲۵]

$$\begin{aligned} & \left\| P_i^{\frac{1}{2}} x + P_i^{-\frac{1}{2}} q_i \right\|^2 + r_i - q_i^T P_i q_i = x^T P_i x + 2 q_i^T x + r_i, \\ & \left\| Q_k^{\frac{1}{2}} x + Q_k^{-\frac{1}{2}} m_k \right\|^2 \leq (m_k^T Q_k^{-1} m_k - l_k)^{\frac{1}{2}} \equiv x^T Q_k x + 2 m_k^T x + l_k \leq 0. \end{aligned}$$

بنابراین تابع هدف (۲۲.۵) را می‌توان به صورت زیر بیان کرد:

$$\frac{x^T P_i x + 2 q_i^T x + r_i}{a_i^T x + b_i} = \frac{\left\| P_i^{\frac{1}{2}} x + P_i^{-\frac{1}{2}} q_i \right\|^2}{a_i^T x + b_i} + \frac{r_i - q_i^T P_i q_i}{a_i^T x + b_i}, \quad i = 1, \dots, p.$$

با معرفی متغیرهای جدید u_i و v_i ، مشابه حالت قبلی، مسأله برنامه‌ریزی کسری (۲۲.۵)–(۲۳.۵) به فرم معادل زیر تبدیل می‌شود:

$$\text{Minimize} \quad \sum_{i=1}^p (u_i + v_i) \quad (24.5)$$

subject to

$$\begin{cases} \left\| \begin{bmatrix} 2 \left(P_i^{\frac{1}{2}} x + P_i^{-\frac{1}{2}} q_i \right) \\ a_i^T x + b_i - u_i \end{bmatrix} \right\| \leq a_i^T x + b_i + u_i, & i = 1, \dots, p, \\ \left\| \begin{bmatrix} 2 \left(r_i - q_i^T P_i^{-1} q_i \right)^{\frac{1}{2}} \\ a_i^T x + b_i - v_i \end{bmatrix} \right\| \leq a_i^T x + b_i + v_i, & i = 1, \dots, p, \\ \left\| Q_k^{\frac{1}{2}} x + Q_k^{-\frac{1}{2}} m_k \right\|^2 \leq (m_k^T Q_k^{-1} m_k - l_k)^{\frac{1}{2}}, & k = 1, \dots, s, \\ c_j^T x + d_j \geq 0, & j = 1, \dots, q. \end{cases} \quad (25.5)$$

۴.۲.۵ برنامه‌ریزی با تابع هدف مختلط از کسری و درجه دوم و همراه با قیود درجه دوم

مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \sum_{i=1}^p \frac{\|F_i x + g_i\|^2}{a_i^T x + b_i} + \sum_{j=1}^l \frac{c}{c_j^T x + d_j} + x^T P_* x + 2 q_*^T x + r \quad (26.5)$$

subject to

$$\begin{cases} a_i^T x + b_i > 0, & i = 1, \dots, p, \\ c_j^T x + d_j > 0, & j = 1, \dots, q, \\ x^T P_k x + 2q_k^T x + r_k \leq 0, & k = 1, \dots, s, \end{cases} \quad (27.5)$$

که در آن ماتریس‌های $P_k (k = 1, \dots, s)$ ، p معین مثبت هستند. با معرفی متغیرهای جدید $u_i (i = 1, \dots, p)$ و $v_j (j = 1, \dots, q)$ و با استفاده از روابط به کاررفته در حالت‌های قبل، مسأله (۲۶.۵)–(۲۷.۵)

قابل بیان به صورت معادل زیر است:

$$\text{Minimize} \quad \sum_{i=1}^p u_i + \sum_{j=1}^q v_j + t. \quad (28.5)$$

subject to

$$\begin{cases} \left\| \begin{bmatrix} 2(F_i x + g_i) \\ a_i^T x + b_i - u_i \end{bmatrix} \right\| \leq a_i^T x + b_i + u_i, & i = 1, \dots, p, \\ \left\| \begin{bmatrix} 2(p_*^{\frac{1}{2}} x + p_*^{-\frac{1}{2}} q_*) \\ 1 - t_* \end{bmatrix} \right\| \leq t_* + 1, \\ \left\| P_k^{\frac{1}{2}} x + P_k^{-\frac{1}{2}} q_k \right\| \leq (q_k^T P_k^{-1} q_k - r_k)^{\frac{1}{2}}, & k = 1, \dots, s, \\ \left\| \begin{bmatrix} 2\sqrt{c} \\ c_j^T x + d_j - v_j \end{bmatrix} \right\| \leq c_j^T x + d_j + v_j, & j = 1, \dots, l. \end{cases} \quad (29.5)$$

از تحلیل بالا و [۱۲۵، ۱۲۸]، واضح است که مسائل برنامه‌ریزی کسری قابل تبدیل به مسائل برنامه‌ریزی مخروط مرتبه دو محدب هستند. بنابراین، ما فرم کلی از مسائل برنامه‌ریزی مخروط مرتبه دو محدب را به صورت زیر در نظر می‌گیریم:

$$\text{Minimize} \quad f(x) \quad (30.5)$$

subject to

$$Ex + d \succeq_{\mathcal{K}} 0, \quad (31.5)$$

$$x \geq 0. \quad (32.5)$$

وقتیکه $f: \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ ، تابع سره بسته محدب، E یک ماتریس $m \times n$ با $m \geq n$ ، برداری در $\mathbb{R}^m$ و $x \succeq_K 0$ یعنی $x \in K$ و حاصلضرب دکارتی از مخروط‌های مرتبه دوم که مخروط‌های لورنتس یا مخروط بستنی هم نامیده می‌شوند. به عبارت دیگر،

$$K = K^{m_1} \times K^{m_2} \times \dots \times K^{m_r},$$

که در آن $m_1 + \dots + m_r = n$ با $r, m_1, \dots, m_r \geq 1$ و

$$K^{m_i} := \{(x_1, x_2)^T \in \mathbb{R} \times \mathbb{R}^{m_i-1} \mid x_1 \geq \|x_2\|\}, \quad (33.5)$$

که در آن $\|\cdot\|$ نرم اقلیدسی و K^1 مجموعه از $\mathbb{R}_+$ حقیقی نامنفی است. در فصل قبل مدل شبکه عصبی (۵۶.۴)–(۵۷.۴) با عملکرد بالا برای حل مسأله (۳۰.۵)–(۳۲.۵) پیشنهاد کردیم.

۳.۵ مثال‌های عددی

در این بخش برای روشن شدن مطالب شرح داده شده و به منظور نشان دادن کارایی و تأثیر شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴)، ما به ذکر چندین مثال از مسائل برنامه‌ریزی کسری می‌پردازیم. برای برخی مثال‌ها، ما همچنین مقایسه عملکرد عددی از شبکه عصبی پیشنهادی را با مقادیر گوناگون از حالت اولیه $y(0)$ داریم. برای برخی دیگر از مسائل ما عملکرد عددی از شبکه عصبی پیشنهاد شده را با مقادیر گوناگون از نرخ همگرایی κ مقایسه می‌کنیم. در شبیه‌سازی همه مثال‌ها از دستور `ode45` در نرم افزار *Matlab* استفاده شده و جواب واقعی مسائل نیز توسط نرم افزار *Lingo* ۱۱ محاسبه شده است.

مثال ۱.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\begin{aligned} \text{Minimize} \quad & \frac{1}{x_1 + x_2 + 2} + \frac{1}{2x_1 - 3x_2 - 2} \\ \text{subject to} \quad & \end{aligned} \quad (34.5)$$

$$\begin{cases} x_1 + x_2 \leq 4, \\ x_1 + x_2 + 2 > 0, \\ 2x_1 - 3x_2 - 2 > 0, \\ x_1, x_2 \geq 0. \end{cases} \quad (35.5)$$

در این مثال با معرفی متغیرهای جدید u_1, u_2 و بکار بردن نامساوی (۱.۵) مسأله برنامه‌ریزی مخروط مرتبه دوم محدب زیر به دست می‌آید:

$$\text{Minimize} \quad u_1 + u_2 \quad (36.5)$$

subject to

$$\begin{cases} \left\| \begin{bmatrix} 2 \\ x_1 + x_2 + 2 - u_1 \end{bmatrix} \right\| \leq x_1 + x_2 + 2 + u_1, \\ \left\| \begin{bmatrix} 2 \\ 2x_1 - 3x_2 - 2 - u_2 \end{bmatrix} \right\| \leq 2x_1 - 3x_2 - 2 + u_2, \\ x_1 + x_2 \leq 4, \\ x_1, x_2 \geq 0. \end{cases} \quad (37.5)$$

جواب بهینه مسأله (۳۵.۵)–(۳۴.۵) $x^* = (4, 0)^T$ است. ما برای حل این مسأله شبکه عصبی پیشنهادی

(۵۶.۴)–(۵۷.۴) را به کار می‌بریم، شکل (۱.۵) نشان می‌دهد که مسیرهای شبکه عصبی پیشنهادی با ۵۰

نقطه آغازین متفاوت و $\epsilon = 10^{-6}$ ، همگرا به جواب بهینه مسأله است. همچنین در شکل (۲.۵) تأثیر پارامتر

κ از شبکه عصبی روی مقدار $\|x(t) - x^*\|^2$ با حالت اولیه $y_0 = (1, -1, 1, -1, 1, -1, 1, -1, 1)^T$

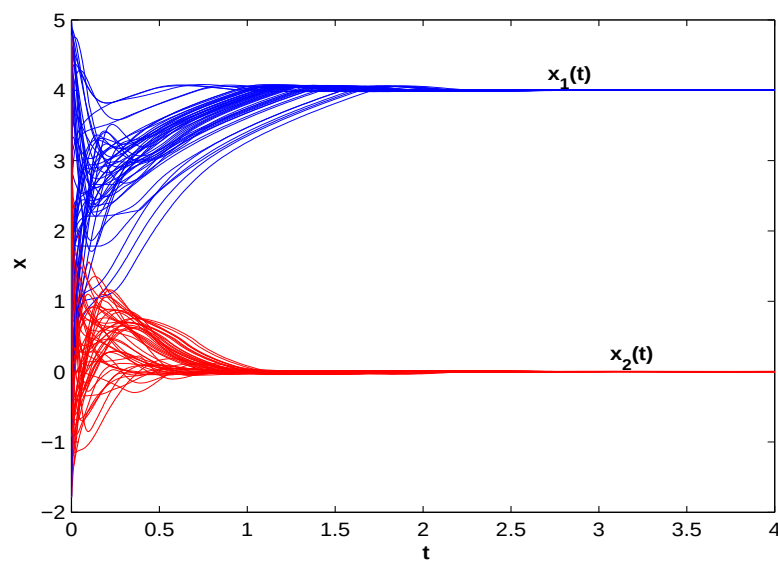
بررسی شده است، در شکل فوق مشاهده می‌کنیم κ بزرگتر، نرخ همگرایی بهتری از خطای $\|x(t) - x^*\|^2$

به دست می‌آورد.

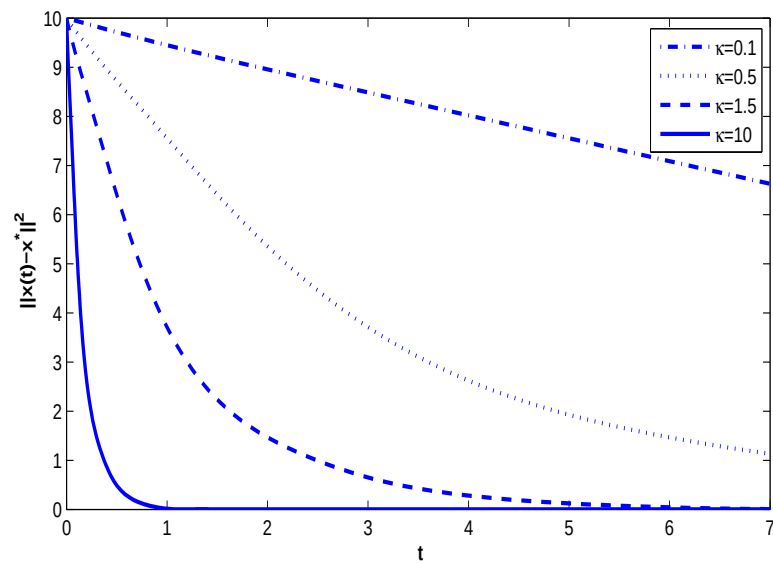
مثال ۲.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \frac{1}{3x_1 - 4x_2 + 5} + \frac{1}{2x_1 + 5x_2 + 7} \quad (38.5)$$

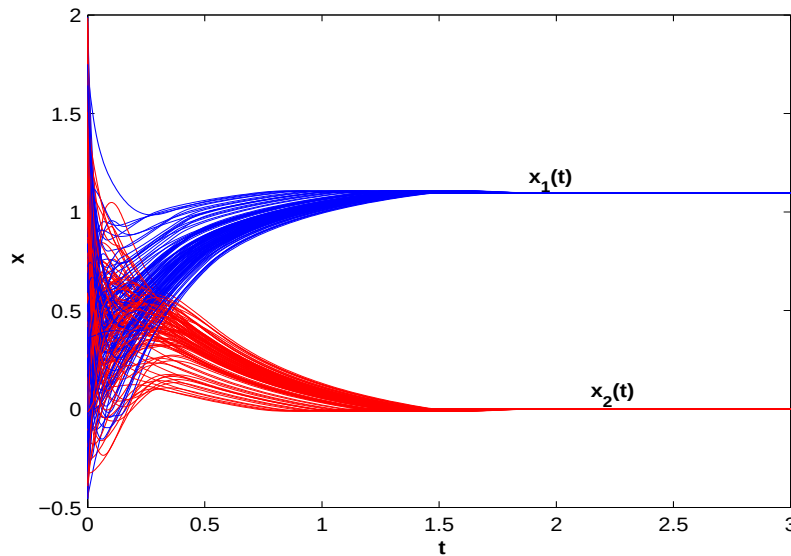
subject to



شکل ۱.۵: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۵۰ نقطه آغازین دلخواه در مثال ۱.۳.۵.



شکل ۲.۵: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$ در مثال ۱.۳.۵.



شکل ۳.۵: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۰۰ نقطه آغازین دلخواه در مثال ۲.۳.۵.

$$\begin{cases} 3x_1 - 4x_2 + 5 > 0, \\ 2x_1 + 5x_2 + 7 > 0, \\ 4x_1^2 + 2x_2^2 + 4x_1x_2 + 2x_1 + 4x_2 - 7 \leq 0. \end{cases} \quad (39.5)$$

با تبدیل مسأله به برنامه‌ریزی مخروط مرتبه دوم محدب و حل آن، جواب بهینه مسأله،

$$x^* = (39.5) - (38.5)^T \text{ به دست می‌آید، بنابراین جواب مسأله } (39.5) - (38.5)^T$$

است. نتایج شبیه سازی نشان می‌دهد که مسیرهایی از (۵۶.۴)–(۵۷.۴) با هر نقطه اولیه

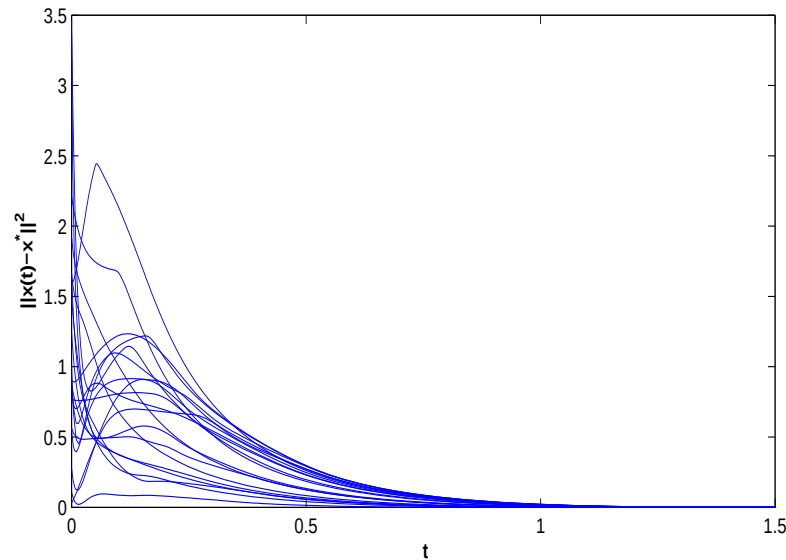
همگرا به $y^* = (x^{*T}, u^{*T})^T$ است. برای مثال شکل (۳.۵) نشان دهنده رفتار گذرا شبکه عصبی با ۱۰۰

نقطه آغازین متفاوت و $\epsilon = 10^{-6}$ است، نرم خطای L_2 بین x و x^* با ۲۰ نقطه اولیه متفاوت در شکل

(۴.۵) رسم شده است.

مثال ۳.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \frac{1}{5x_1 + x_2 + 1} + \frac{1}{7x_1 - x_2 + 4} + x_1^2 + 3x_2^2 + 3x_1x_2 + x_1 - 4x_2 - 8$$



شکل ۴.۵: رفتار همگرایی از $\|x(t) - x^*\|^2$ با ۲۰ نقطه اولیه متفاوت در مثال ۲.۳.۵.

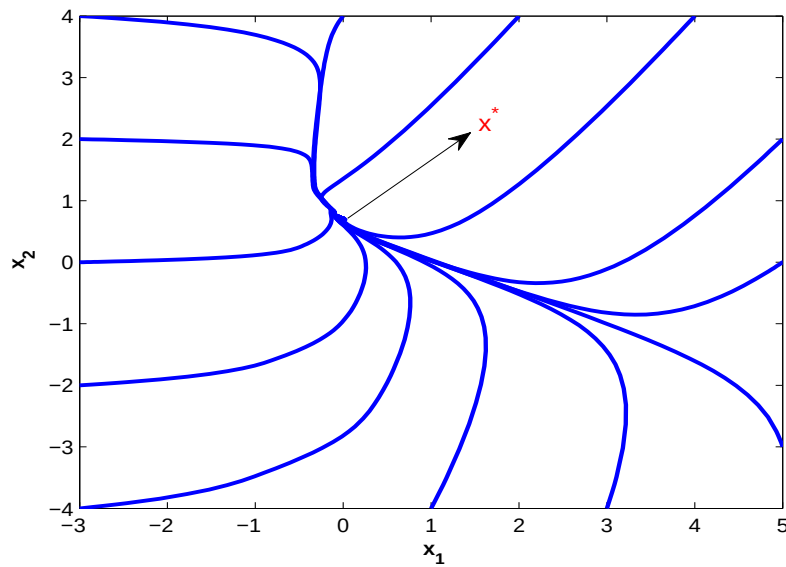
subject to

$$\begin{cases} 5x_1 + x_2 + 1 > 0, \\ 7x_1 - x_2 + 4 > 0, \\ x_1, x_2 \geq 0. \end{cases} \quad (40.5)$$

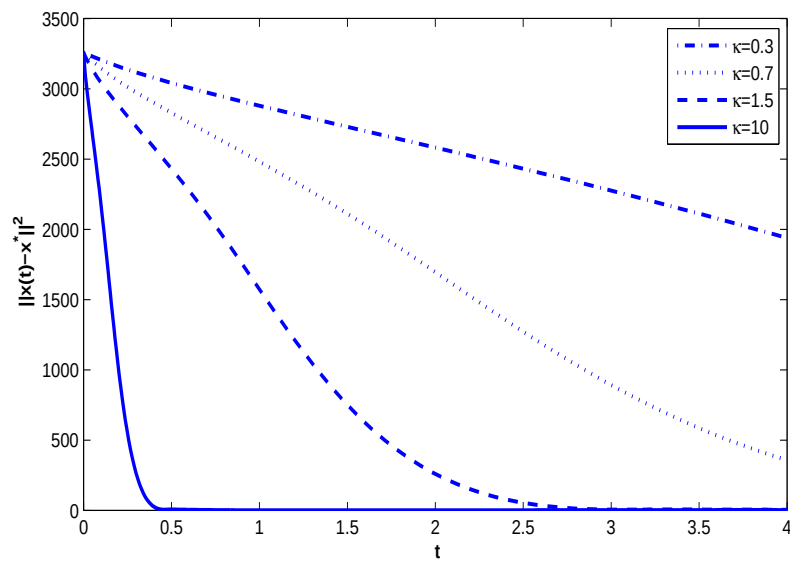
جواب بهینه این مسأله، $x^* = (0, 0.7084)^T$ است. قضیه (۵.۴.۴) و گزاره (۶.۴.۴) تضمین می‌کند که مدل بیان شده در (۵۶.۴)–(۵۷.۴)، همگرای سراسری به جواب بهینه یکتاست. شکل (۵.۵) منحنی فاز متغیرهای حالت $(x_1(t), x_2(t))^T$ ، با ۱۳ نقطه اولیه متفاوت و $\epsilon = 10^{-6}$ نشان می‌دهد. این بردارها همگرای سراسری به جواب دقیق x^* هستند. شکل (۶.۵) رفتار همگرایی، خطای $\|x(t) - x^*\|^2$ با مقادیر مختلف κ و حالت اولیه $y_0 = (40, -40, 40, -40, 40, -40, 40, -40, 40, -40)^T$ نشان می‌دهد. واضح است که κ بزرگتر منجر به نرخ همگرایی بهتری شده است.

مثال ۴.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \frac{\|3x_1 + x_2 + 1\|^2}{x_1 + x_2 + 2} + \frac{\|x_1 + x_2 + 1\|^2}{2x_1 - 3x_2 + 1} \quad (41.5)$$



شکل ۵.۵: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۳ نقطه آغازین در مثال ۳.۳.۵.



شکل ۶.۵: رفتار همگرایی از $\|x(t) - x^*\|^2$ با y اولیه در مثال ۳.۳.۵.

subject to

$$\begin{cases} x_1 + x_2 + 2 > 0, \\ 2x_1 - 3x_2 + 1 > 0, \\ x_1 + x_2 \leq 4, \\ x_1, x_2 \geq 0. \end{cases} \quad (42.5)$$

با معرفی متغیرهای جدید u_1 و u_2 با فرض اینکه

$$\frac{\|2x_1 + x_2 + 1\|^2}{x_1 + x_2 + 2} \leq u_1, \quad \frac{\|x_1 + x_2 + 1\|^2}{2x_1 - 3x_2 + 1} \leq u_2.$$

بنابراین مسأله بالا تبدیل به معادل برنامه‌ریزی مخروط مرتبه دوم محدب زیر می‌شود:

$$\text{Minimize} \quad u_1 + u_2 \quad (43.5)$$

subject to

$$\begin{cases} \left\| \begin{pmatrix} 2(2x_1 + x_2 + 1) \\ x_1 + x_2 + 2 - u_1 \end{pmatrix} \right\| \leq x_1 + x_2 + 2 + u_1, \\ \left\| \begin{pmatrix} 2(x_1 + x_2 + 1) \\ 2x_1 - 3x_2 + 1 - u_1 \end{pmatrix} \right\| \leq 2x_1 - 3x_2 + 1 + u_2, \\ x_1 + x_2 \leq 4, \\ x_1, x_2 \geq 0. \end{cases} \quad (44.5)$$

جواب بهینه این مسأله $x^* = (0, 0)^T$ است. شکل (۷.۵) بیان کننده منحنی فاز متغیرهای حالت

$(x_1(t), x_2(t))^T$ ، با ۱۷ حالت اولیه متفاوت و $\epsilon = 10^{-6}$ است. شکل (۸.۵) رفتار همگرا از خطای

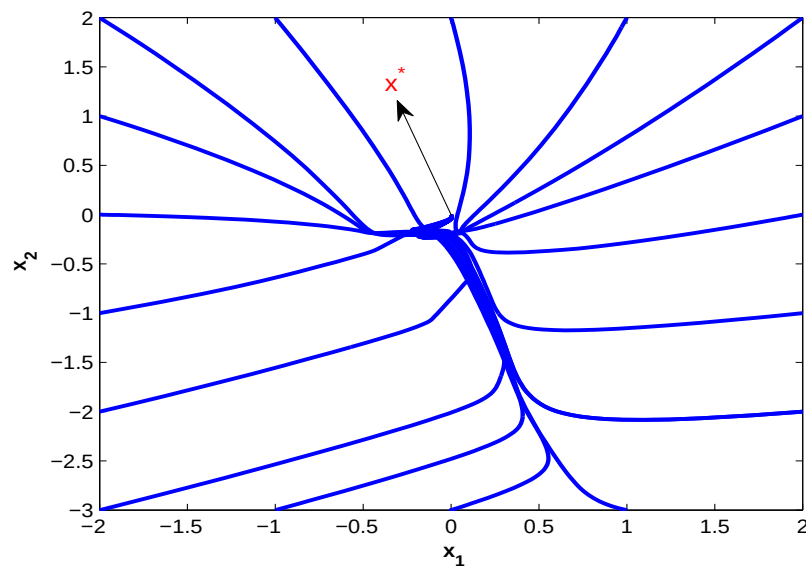
$\|x(t) - x^*\|^2$ با κ های متفاوت و نقطه آغازین $y_0 = (0, -1, -2, -3, -4, -5, -6, -7, -8)^T$ را

نشان می‌دهد.

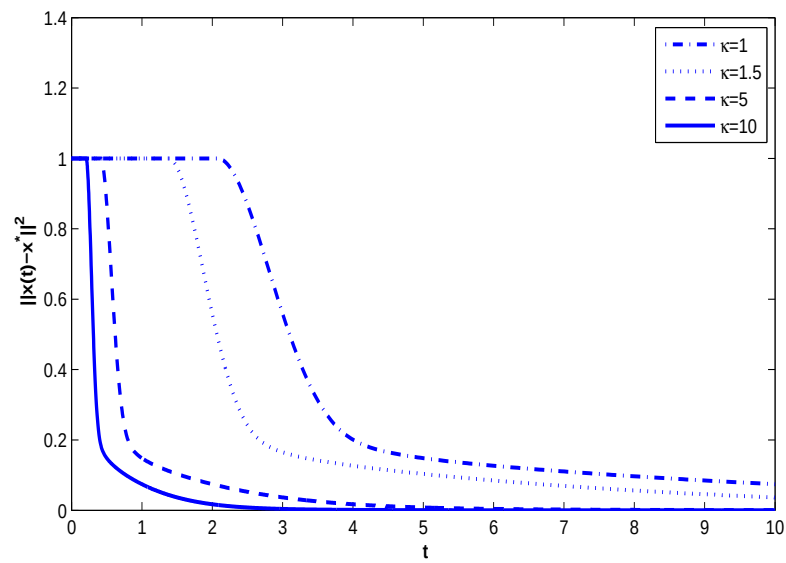
مثال ۵.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \frac{x_1^2 + 3x_2^2 + 3x_1x_2 + x_1 + x_2 + 2}{x_1 + x_2 + 2} + \frac{\|x_1 + x_2 + 1\|^2}{2x_1 - 3x_2 + 1} \quad (45.5)$$

subject to



شکل ۷.۵: منحنی فاز شبکه عصبی $(-۵۶.۴) - (۵۷.۴)$ با ۱۷ نقطه آغازین در مثال ۴.۳.۵.



شکل ۸.۵: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (0, -1, -2, -3, -4, -5, -6, -7, -8)^T$ در مثال ۴.۳.۵.

$$\begin{cases} x_1 + x_2 + 2 > 0, \\ 2x_1 - 3x_2 + 1 > 0, \\ x_1 + x_2 \leq 4, \\ x_1, x_2 \geq 0. \end{cases} \quad (46.5)$$

برنامه ریزی کسری در این مثال قابل بیان به صورت زیر است:

$$\text{Minimize} \quad \frac{\left\| \begin{bmatrix} 0.07794x_1 + 0.6265x_2 \\ 0.6265x_1 + 1.6148x_2 \end{bmatrix} + \begin{bmatrix} 0.5706 \\ 0.0883 \end{bmatrix} \right\|^2}{x_1 - x_2 + 2} + \frac{1.66675}{x_1 - x_2 + 2} + \frac{\|x_1 + x_2 + 1\|^2}{2x_1 - 3x_2 + 2}$$

subject to

$$\begin{cases} x_1 - x_2 + 2 > 0, \\ 2x_1 - 3x_2 + 1 > 0, \\ x_1 + x_2 \leq 4, \\ x_1, x_2 \geq 0. \end{cases} \quad (47.5)$$

جواب بهینه مسأله (۴۵.۵)-(۴۶.۵) ، $x^* = (0, 0)^T$ ، بیان کننده مسیرهای $x(t)$ با

۵۰ نقطه اولیه متفاوت و $\epsilon = 10^{-6}$ است. واضح است، وقتی نقاط اولیه از نقاط نشدنی انتخاب شوند،

مسیرهای شبکه همچنان همگرا به x^* هستند. منحنی فاز متغیرهای حالت $(x_1(t), x_2(t))^T$ ، با ۱۷

حالت اولیه متفاوت در شکل (۱۰.۵) نشان داده شده. نرم خطای L_2 بین x و x^* با ۳۰ نقطه اولیه متفاوت

در شکل (۱۱.۵) رسم شده است.

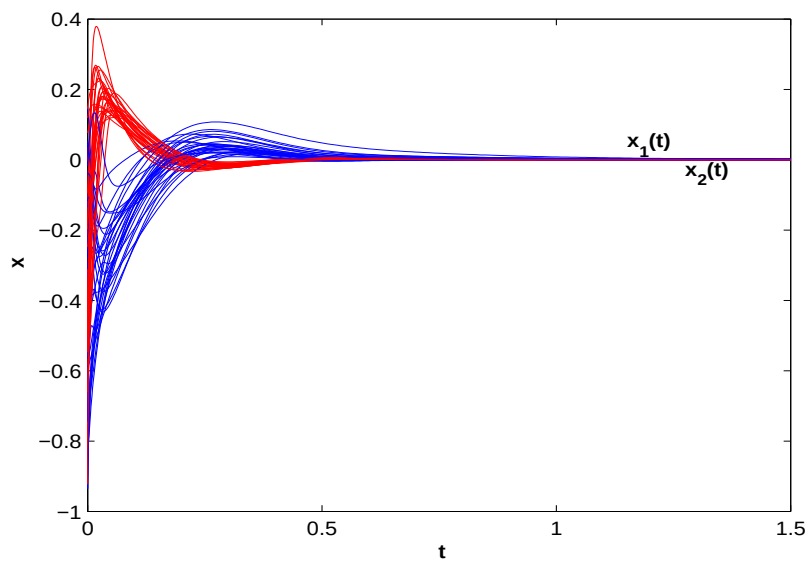
مثال ۶.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

$$\text{Minimize} \quad \frac{\| -3x_1 + 4x_2 + 5 \|^2}{-x_1 + 2x_2 + 4} + \frac{\| 4x_1 - 3x_2 + 4 \|^2}{x_1 + x_2 + 5} \quad (48.5)$$

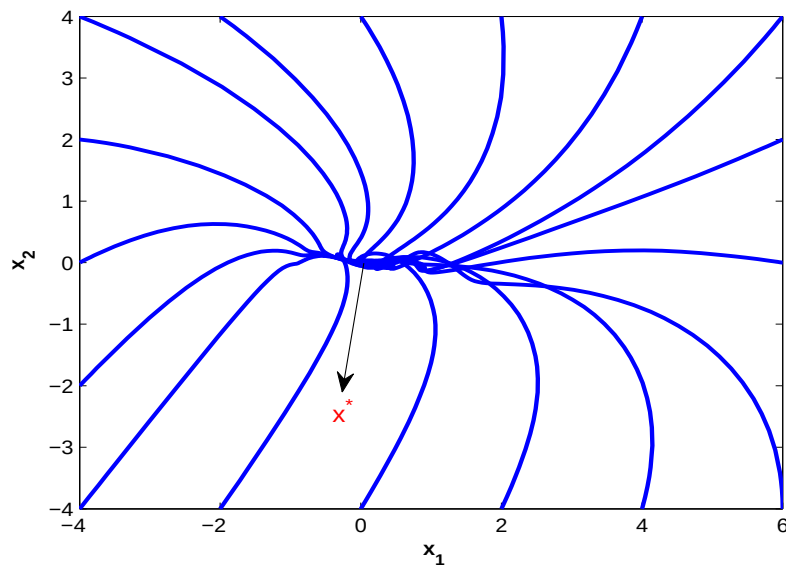
subject to

$$\begin{cases} -x_1 + 2x_2 + 4 > 0, \\ x_1 + x_2 + 5 > 0, \\ 2x_1^2 + x_2^2 - 2x_1x_2 + 2x_1 - x_2 - 6 \leq 0, \\ x_1, x_2 \geq 0. \end{cases} \quad (49.5)$$

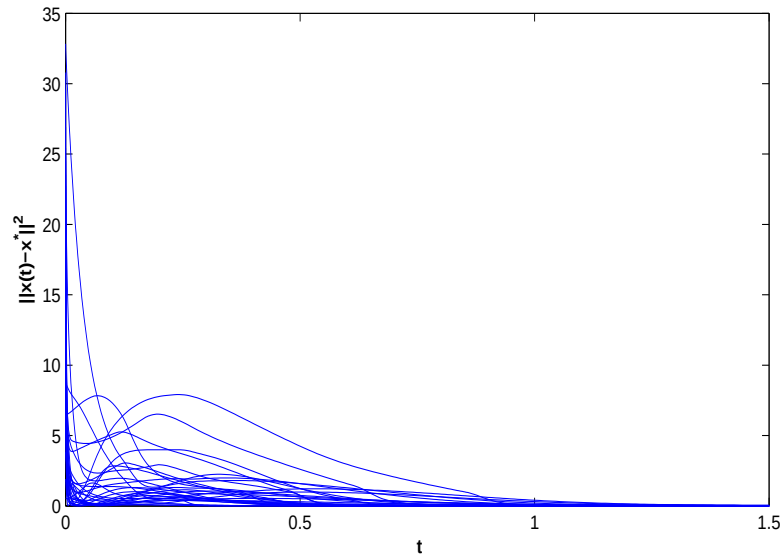
با تبدیل مسأله بالا به برنامه‌ریزی مخروط مرتبه دوم محدب جواب بهینه مسأله



شکل ۹.۵: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۵۰ نقطه آغازین دلخواه در مثال ۵.۳.۵.



شکل ۱۰.۵: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۷ نقطه آغازین در مثال ۵.۳.۵.



شکل ۵.۱۱: رفتار همگرایی از $\|x(t) - x^*\|^2$ با ۳۰ نقطه اولیه متفاوت در مثال ۵.۳.۵.

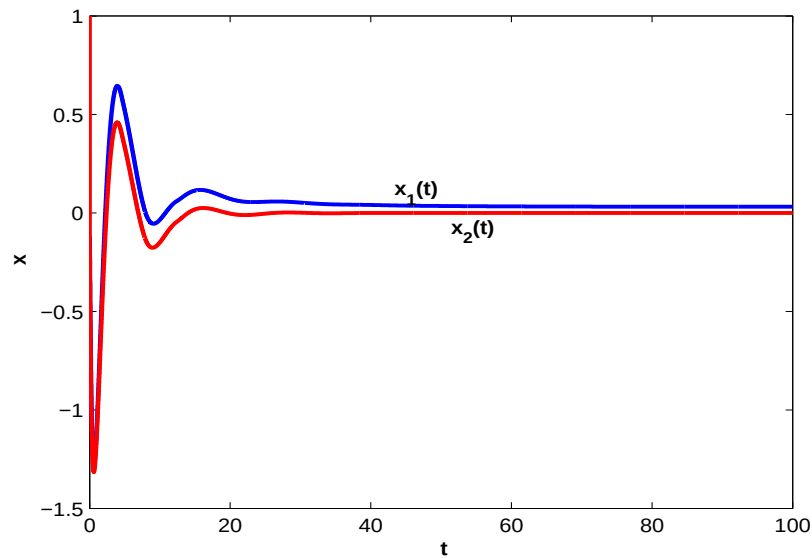
شکل ۵.۱۳: رسم شده است. شکل (۱۴.۵)، تأثیر پارامتر κ از شبکه عصبی روی مقدار $\|x(t) - x^*\|^2$ با $x_0 = (-1, 1)^T$ و $\epsilon = 10^{-6}$ در شکل (۱۲.۵) و منحنی فاز این مدل با ۱۸ نقطه اولیه متفاوت در شکل (۱۳.۵) را گذرا از شبکه عصبی (۵۶.۴)–(۵۷.۴) $x^* = (0.318, 0.628, 3.3844)^T$ به دست می‌آید. رفتار گذرا از شبکه عصبی (۵۶.۴)–(۵۷.۴) نشان می‌دهد.

مثال ۷.۳.۵. مسأله برنامه‌ریزی کسری زیر را در نظر بگیرید:

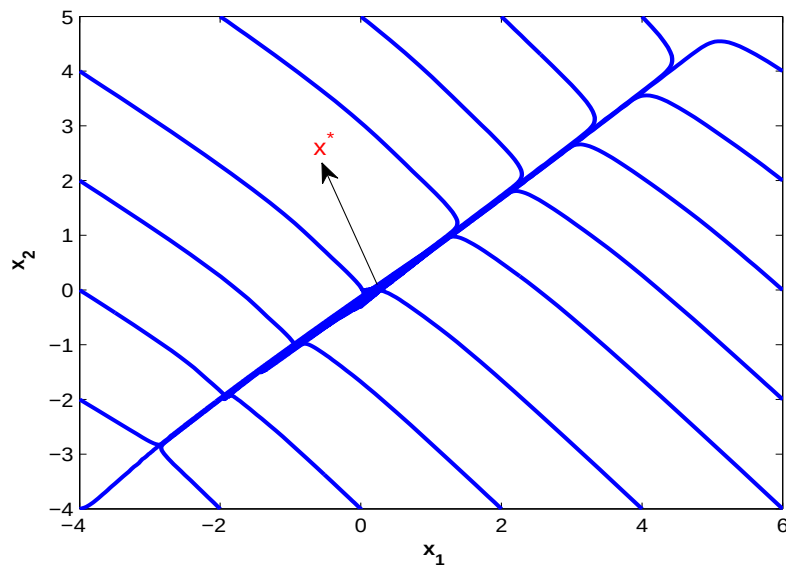
$$\begin{aligned} \text{Minimize} \quad & \frac{4x_1^2 + 4x_2^2 - 4x_1x_2 - x_1 - 4x_2 + 6}{3x_1 - 7x_2 + 4} + 3x_1^2 + 2x_2^2 + 2x_1x_2 - 4x_1 - 2x_2 + 4 \\ \text{subject to} \quad & \begin{cases} 3x_1 - 7x_2 + 4 > 0, \\ x_1, x_2 \geq 0. \end{cases} \end{aligned} \quad (50.5)$$

ابتدا مسأله را به برنامه‌ریزی مخروط مرتبه دوم محدب تبدیل کرده و آن را با شبکه عصبی (۵۶.۴)–

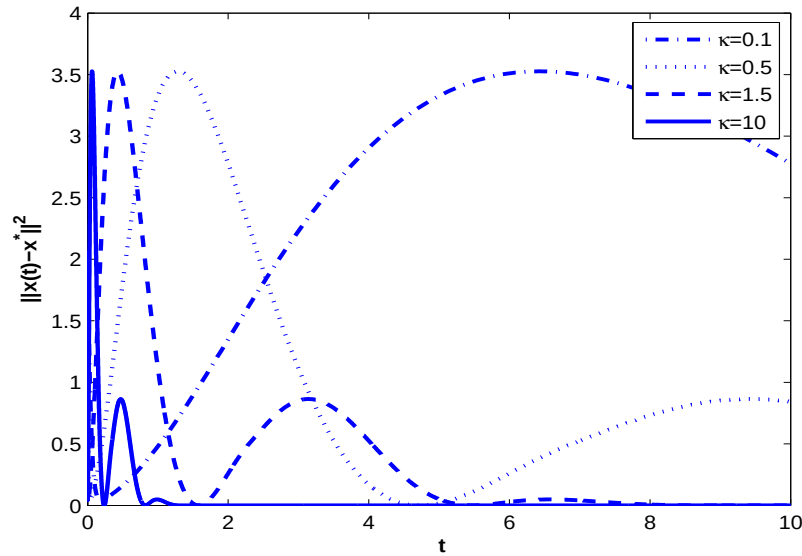
(۵۷.۴) حل کردیم. شکل (۱۵.۵) نشان می‌دهد که مسیرهای شبکه پیشنهادی در (۵۶.۴)–(۵۷.۴) با



شکل ۱۲.۵: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با $x_0 = (-1, 1)^T$ در مثال ۶.۳.۵.

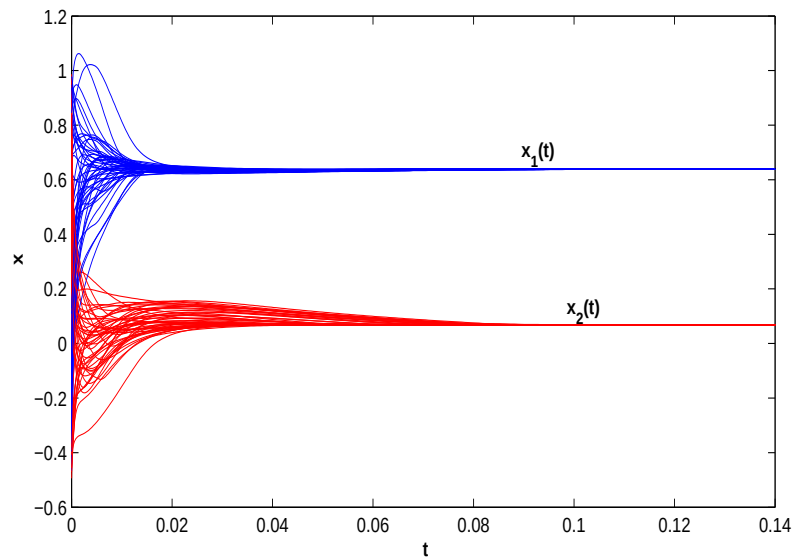


شکل ۱۳.۵: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۸ نقطه آغازین در مثال ۶.۳.۵.

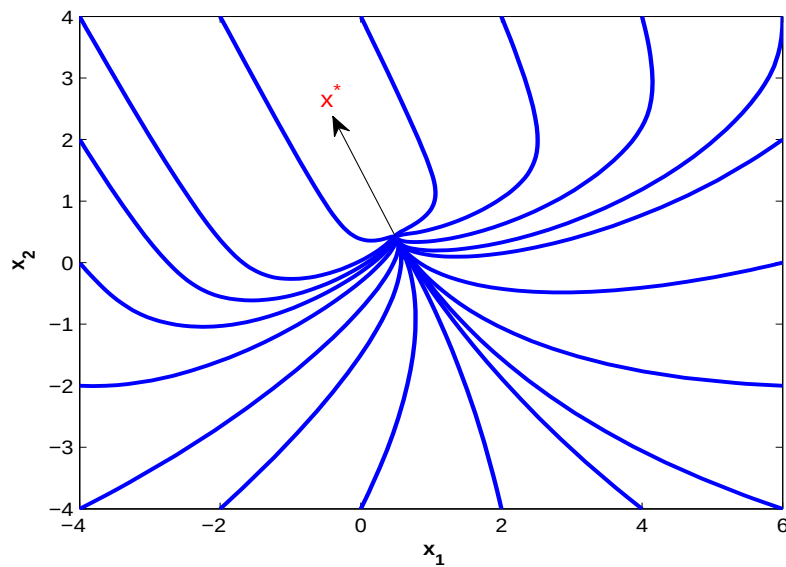


شکل ۱۴.۵: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (-1, 1, -1, 1, -1, 1, -1, 1, -1)^T$ در مثال ۶.۳.۵.

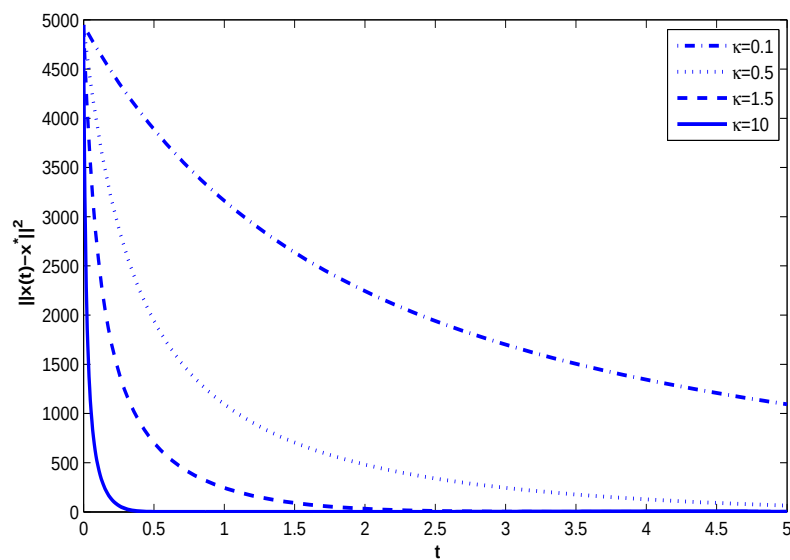
$\epsilon = 10^{-6}$ همگرا به $x^* = (0.6376, 0.687)^T$ است. منحنی فاز متغیرهای حالت $(x_1(t), x_2(t))^T$ با ۱۸ حالت اولیه متفاوت و $\epsilon = 10^{-6}$ نیز در شکل (۱۶.۵) رسم شده است. واضح است که نقاط اولیه داخل و خارج ناحیه شدنی، مدل شبکه عصبی پیشنهادی همواره همگرا به جواب بهینه x^* است. رفتار همگرایی از $\|x(t) - x^*\|^2$ با نقطه اولیه $y_0 = (50, -50, 50, -50, 50, -50, 50, -50, 50, -50)^T$ هم در شکل (۱۷.۵) نشان داده شده است.



شکل ۱۵.۵: رفتار گذرا شبکه عصبی با ۵۰ نقطه آغازین دلخواه در مثال ۷.۳.۵.



شکل ۱۶.۵: منحنی فاز شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۱۸ نقطه آغازین در مثال ۷.۳.۵.



شکل ۱۷.۵: رفتار همگرایی از $\|x(t) - x^*\|^2$ با y در مثال ۷.۳.۵.

فصل ۶

بهینه‌سازی مقاوم

۱.۶ مقدمه

در این فصل از پایان نامه، ابتدا بهینه‌سازی داده‌های غیرقطعی و رویکردهای مقابله با عدم قطعیت در مدل‌ها را بیان می‌کنیم سپس چند نوع عمده از مجموعه‌های غیرقطعی را عنوان کرده و در ادامه با معرفی اشکال مختلف مقاوم بودن پرداخته و مروری بر کارهای انجام شده داریم. در ادامه با مسأله برنامه‌ریزی خطی مقاوم آشنا می‌شویم و در انتها به معرفی مسأله بهینه‌سازی مقاوم سبد سرمایه و ذکر یک مثال از آن می‌پردازیم.

۲.۶ عدم قطعیت در پارامترهای مدل

در هنگام ارائه یک مدل غالباً پارامترهای ورودی در آینده مشخص می‌شوند یا مقدار دقیق آن در زمان مدل‌سازی مسأله مشخص نیستند. عدم قطعیت^۱ در داده‌ها و پارامترهای مسأله، مشکلاتی در بسیاری از موارد عملی بخصوص مقادیر مالی مانند بازده سرمایه‌گذاری، ریسک و غیره به وجود می‌آورد. دو رویکرد کلی برای برخورد با این عدم قطعیت وجود دارد:

۱. برخورد واکنشی یا غیر فعال:

در این رویکرد که روش شناخته شده آن تحلیل حساسیت^۲ است، ابتدا مسأله را براساس تخمین نقطه‌ای پارامترهای غیرقطعی تحت عنوان مقادیر اسمی^۳ حل می‌کنند و بعد از حل به اندازه‌گیری حساسیت جواب مسأله نسبت به پارامترهای غیرقطعی می‌پردازند و محدوده تغییرات پارامترهای مذکور را طوری به دست می‌آورند که جواب مسأله همچنان شدنی باقی بماند. برای برآورد پارامترهای غیرقطعی کارهای زیادی صورت گرفته است. در این زمینه چاپرا^۴ و همکاران استفاده

^۱Uncertainty

^۲Sensitivity Analysis

^۳Nominal Values

^۴Chapra

از برآوردگر جیمز-استاین را برای میانگین‌ها پیشنهاد کردند، در حالیکه کلین و باوا^۵، فراست و ساوارینو^۶، بلک و لیترمن^۷ استفاده از برآوردگر بیزی^۸ را برای میانگین‌ها و کواریانس‌ها ارائه کردند اگر چه این روش‌ها حساسیت سبد انتخابی را به مقدار برآورد شده پارامترها کاهش می‌دهد، ولی تضمینی برای کارایی میانگین-واریانس سبد سرمایه فراهم نمی‌کند. میچاد^۹ روشی را که مبتنی بر چند بار نمونه‌گیری از میانگین بازده‌ها (بوت استرپ^{۱۰}) و ماتریس کواریانس دارایی‌ها از فاصله اطمینان حول مقادیر اسمی پارامترها و سپس یکی کردن سبدهای به‌دست آمده از حل مسأله مارکویتز برای هر نمونه است پیشنهاد کرده است.

روش دیگر در برخورد واکنشی، حل مسأله با پارامتری در نظر گرفتن پارامترهای غیرقطعی است که به آن برنامه‌ریزی پارامتری می‌گویند.

در این رویکرد در فرایند حل مدل اقدامات لازم برای غیرقطعی بودن پارامترها در نظر گرفته نمی‌شود و کیفیت جواب از لحاظ برخورد با اختلالات و آشفتگی‌های موجود در پارامترهای قطعی بالا نمی‌رود. بنابراین به‌طور کلی جواب بهینه به‌دست آمده از رویکرد قطعی هم از نظر شدنی بودن و هم از نظر بهینگی تابع هدف مورد سوال است.

۲. برخورد سازنده یا فعال:

در این رویکرد به‌دنبال روشهایی هستیم که برخورد با عدم قطعیت در داده‌ها را قبل از حل و در متن مدل پیش‌بینی کند و چاره‌ای برای پایداری آنها ارائه کند. این روشها سعی در حل مسائل برنامه‌ریزی در شرایط عدم قطعیت دارد که هر یک مشکلات خاص خود را دارا هستند. لذا اعتقاد بر این است که به رویکردی نیاز است که در مقابل داده‌های غیرقطعی علاوه بر برخورد فعال و

^۵Klein and Bava

^۶Frost and Savarynv

^۷Black and Lytrmn

^۸Bayesian estimator

^۹Mychad

^{۱۰}Bootstrap

سازنده دارای خطا و هزینه کمی باشد. تکنیک مدل سازی اساسا متفاوت، برنامه‌ریزی تصادفی و بهینه سازی مقاوم^{۱۱} در این رویکرد وجود دارد که در پایین شرح داده شده‌اند.

پس از رشد و گسترش مدل‌های برنامه‌ریزی خطی و روشهای حل آن، اولین برخورد فعال با داده‌های غیرقطعی رویکردی بود که اطلاعات احتمالی را درباره داده‌های مسأله به‌صورت برنامه‌ریزی تصادفی مطرح کرد که سابقه آن دنتزیگ^{۱۲} در سال ۱۹۵۵ باز می‌گردد[۶۰]. در این روش مدل‌سازی، عدم قطعیت داده‌های مدل از همان ابتدا در مدل‌سازی در نظر گرفته می‌شوند. این روش برای حالتی استفاده می‌شود که عدم قطعیت آن طبیعت احتمالی داشته باشد و از توانایی تعیین توزیع احتمالی آن نیز برخوردار باشیم. لازم به ذکر است که در این نوع مدل‌سازی به دلیل عدم قطعیت موجود، امکان نقض محدودیت‌های مدل توسط جواب حاصله وجود دارد. حالت مهم آن برنامه‌ریزی خطی احتمالی است که در فصل قبل به بحث در مورد آن پرداختیم. همچنین برنامه‌ریزی پویای احتمالی نیز روش دیگری است که به‌نوعی به برنامه‌ریزی دارای روابط بازگشتی اطلاق می‌شود که برای نخستین بار توسط ریچارد بلمن^{۱۳} در [۱۳۷] ارائه و توسعه داده شده است. این روش از روی مسائل برنامه‌ریزی سرچشمه گرفته است که تغییرات در طول زمان در آن‌ها دارای اهمیت است و برنامه‌ریزی پویا^{۱۴} نامیده شده است. ایده این تکنیک تقسیم مسأله به چند مرحله بهینه‌سازی بازگشتی است که با ترکیب با عناصر احتمالی به برنامه‌ریزی پویای احتمالی تبدیل می‌شود.

در روش‌های بیان شده یکی از اصلی‌ترین مسائلی که در عمل به آن برخورد می‌کنیم دشواری جمع آوری اطلاعات جزئی درباره توزیع احتمالی عدم قطعیت‌های موجود در مدل است. در روشهای برنامه‌ریزی پویا و احتمالی با توجه به امکان افزایش تعداد سناریوهای (طرح‌های)^{۱۵} درگیر در

^{۱۱}Robust Optimization

^{۱۲}Dantzing

^{۱۳}Bellman

^{۱۴}Dynamic programming

^{۱۵}Scenarios

مسأله، معمولاً جمع آوری این اطلاعات به‌شدت پرهزینه می‌شود به‌نحوی که استفاده از این روش مدل‌سازی را دچار مشکل می‌کند. با توجه به این مطلب این روشها در عمل به‌طور گسترده مورد استفاده قرار نمی‌گیرند. بهینه‌سازی مقاوم یک تکنیک جدید توسعه داده شده برای مسائل مشابه مسائل برنامه‌ریزی احتمالی است. در این تکنیک یک فرض عمومی برای توزیع احتمالی پارامترهای دارای عدم قطعیت در نظر گرفته می‌شود تا بدین وسیله کار با مدل مسأله از نظر محاسباتی قابل کنترل شود.

در مرور ادبیات، معمولاً بهینه‌سازی مقاوم به حوزه‌ای از بهینه‌سازی گفته می‌شود که ریشه در ادبیات کنترل مقاوم مهندسی^{۱۶} دارد. در این حوزه با توجه خاص به عدم نقض محدودیت‌های مدل (شرط موجه بودن مدل) به حل هم‌تای مقاوم^{۱۷} مدل اصلی با در نظر گرفتن مجموعه‌های عدم قطعیت پارامترها می‌پردازد. در حقیقت هم‌تای مقاوم این مسائل، بدترین مدل‌سازی مسأله اصلی را در حوزه انحرافات پارامترها از مقدار اسمی آن‌ها بیان می‌کند. به هر حال به‌نوعی بدترین سناریوها به روشی هوشمندانه ایجاد می‌شود به‌نحوی که منجر به مدلی به‌شدت محافظه‌کارانه نشود. در بهینه‌سازی مقاوم، مدل‌سازی مسائل بادر نظر گرفتن مجموعه عدم قطعیت، تضمین کننده مناسب بودن جواب حاصل از بهینه‌سازی برای کلیه مقادیر ممکن پارامترهای مدل است. در این حالت، روش مذکور دارای فرضیات غیر تصادفی موجود در بهینه‌سازی احتمالی برای پارامترهای مدل نیست و اهمیت یکسانی برای تمام مقادیر ممکن از پارامترهای مدل در نظر می‌گیرد و به این ترتیب جواب به‌دست آمده در بهینه‌سازی به ازای تمام مقادیر پارامترهای مدل در فضای موجه^{۱۸} قرار می‌گیرد. برنامه‌ریزی مقاوم برای هنگامی به کار می‌رود که تحلیلگر به دنبال رفتار مطلوب جواب به ازای تمامی رخدادهای ممکن برای داده‌هایی که دچار عدم اطمینان (قطعیت) هستند باشد

^{۱۶}Robust Contorol Engineering

^{۱۷}Robust Counterpart

^{۱۸}Feasible Space

و براین اساس اهمیت یکسانی را برای تمام نموده‌های مسأله قائل می‌شود. به عبارتی این تکنیک را برای زمانی که برخی از پارامترها با ریسک برآورد مواجه هستند، جواب بهینه به تغییرات کوچک حساس باشد و یا تصمیم‌گیرنده ریسک‌پذیری پایینی داشته باشد بسیار مناسب است.

۳.۶ انواع مجموعه‌های عدم قطعیت در بهینه‌سازی مقاوم

عدم اطمینان در پارامترها توسط مجموعه‌های عدم قطعیت منظور می‌شود، که شامل همه یا بیشتر مقادیر ممکن است که می‌توان توسط پارامترهای نامعین رخ دهد. این مجموعه‌ها می‌توان براساس نظرات و تجربیات کارشناسان و یا برآوردهای گوناگون تکنیکهای آماری و یا بیزی بر مبنای داده‌های تاریخی تعیین کرد. دانستن چگونگی توزیع احتمالی داده‌ها برای تعیین مجموعه عدم قطعیت می‌تواند مفید باشد. یک هدف مهم در رویکرد بهینه‌سازی مقاوم این است که معادل قطعی این مسائل به‌فرم صریحی در آیند که قابل حل با روش‌های حل موجود باشد. امکان وجود چنین فرمی برای هر مسأله مورد بررسی، علاوه بر ساختار مسأله به نوع مجموعه عدم قطعیت $\mathcal{U}$ نیز بستگی دارد. شکل و هندسه مجموعه عدم قطعیت، علاوه بر آن که بر خصوصیات معادل مقاوم مسأله تأثیر فراوانی می‌گذارد، می‌تواند در کنترل ماهیت محافظه کارانه رویکرد بهینه‌سازی مقاوم نیز نقش اساسی ایفا کند. تلاش‌های زیادی در جهت کنترل درجه محافظه کارانه بودن صورت مقاوم مسأله برنامه‌ریزی خطی به کمک تعریف مجموعه‌های عدم قطعیت جدید انجام شده است که در ادامه به چند مورد از آن‌ها اشاره خواهد شد.

۱. فضای نمونه گسسته: در این شکل بیشتر روشهای بهینه‌سازی با طرح‌های مجزا برای مقادیر ممکن

پارامتر قابل طرح است که به‌صورت زیر بیان می‌شوند:

$$\mathcal{U} = \{p_1, p_2, \dots, p_k\},$$

که در آن p_k بردارهای پارامترهای غیرقطعی برای سناریو k ام است.

۲. فضای نمونه پیوسته: عدم قطعیت بازه‌ای به واسطه سادگی و کاربرد فراوانش در شرایط عملی مورد توجه بسیاری از محققین قرار گرفته است. در این حالت، برای هر داده نامعلوم یک بازه تغییرات در نظر می‌گیریم که این بازه‌ها می‌توانند از روی فواصل اطمینان حاصل از نتایج آماری به دست آمده باشند:

$$\mathcal{U} = \{p : l \leq p \leq m\},$$

که در آن p بردار پارامترهای غیرقطعی است.

۳. مجموعه‌های عدم قطعیت براساس فضای نرمی: می‌توان فضای نرمی را به صورت زیر بیان کرد:

$$B_r(\theta) = \{x \in \mathbb{R} \mid \|x\|_r \leq \theta\}.$$

این فضا خصوصیات متقارن، محدب، بسته و کراندار بودن را داراست. در صورتی که θ را برابر یک در نظر بگیریم به آن دیسک واحد می‌گویند و اگر r را برابر دو فرض کنیم آنگاه این فضا یک بیضی‌گون را نشان می‌دهد. بنابراین مجموعه‌های عدم قطعیت را با استفاده از فضای نرمی به صورت زیر تعریف می‌کنیم:

$$\mathcal{U} = \{p : p = p_* + Mu \mid u \in B_r(\theta)\}.$$

در حقیقت این فضا میزان انحراف پارامترهای غیرقطعی از مقادیر اسمی خود یعنی بردار p_* را نشان می‌دهد. این انحراف تحت تأثیر شیب تغییرات، ماتریس M و به اندازه اختلالاتی که در فضای نرمی $B_r(\theta)$ می‌تواند رخ دهد، تعریف می‌شود.

۴. مجموعه‌های عدم قطعیت چندوجهی‌گونه (چند سقفی)^{۱۹}: این مجموعه‌ها پوسته محدبی^{۲۰} از تعداد محدودی سناریوی تولید شده برای مقادیر ممکن پارامترها هستند.

$$\mathcal{U} = \text{conv}\{p_1, p_2, \dots, p_k\},$$

^{۱۹}Polytopic

^{۲۰}Convex Hull

۵. مجموعه‌های عدم قطعیت بیضی‌گون: مهم ترین کارها در زمینه عدم قطعیت بیضی‌گون توسط بن تال^{۲۱} و نمیروفسکی^{۲۲} و به‌طور مستقل از آن‌ها ال‌گوی^{۲۳} و همکارانش انجام شده است. آن‌ها معادل مقاوم رده‌های مهمی از مسائل حوزه بهینه‌سازی محدب را تحت مجموعه عدم قطعیت توسعه دادند. در واقع می‌توان کارهای بن-تال و نمیروفسکی را هسته اولیه و اساسی تئوری بهینه‌سازی مقاوم قلمداد کرد. بن-تال و نمیروفسکی در [۱۳۸] اثبات کردند که تحت چندین فرض معقول، معادل مقاوم یک مسأله برنامه‌ریزی محدب خود یک برنامه‌ریزی محدب است. اما عمده کار آن‌ها در این مرجع و در [۱۳۹] به یافتن فرمی صریح برای معادل مقاوم چند رده عمده مسائل بهینه‌سازی محدب بر می‌گردد، آن‌ها نشان دادند که:

۱. معادل مقاوم یک مسأله برنامه‌ریزی خطی تحت یک مجموعه عدم قطعیت بیضی‌گون یا اشتراک تعداد محدودی بیضی‌گون دارای فرم صریح برنامه‌ریزی مخروط مرتبه دوم است.
۲. معادل مقاوم یک مسأله برنامه‌ریزی محدب درجه دوم با قیود درجه دوم تحت یک مجموعه عدم قطعیت بیضی‌گون دارای فرم صریح برنامه‌ریزی نیمه معین^{۲۴} است.
۳. معادل مقاوم یک مسأله برنامه‌ریزی مخروط مرتبه دوم تحت یک مجموعه عدم قطعیت بیضی‌گون و چند فرض نه چندان محدود کننده، یک برنامه‌ریزی نیمه معین است.

جزئیات قضایا و روش‌های حل ارائه شده توسط بن-تال و نمیروفسکی را می‌توانید در مراجع

[۱۳۸]-[۱۴۶] ببینید.

^{۲۱}Ben-Tal

^{۲۲}Nemirovski

^{۲۳}El Ghaoui

^{۲۴}Semi-Definite programming (SDP)

۴.۶ برنامه‌ریزی خطی مقاوم

مسأله زیر را در نظر بگیرید:

$$\begin{aligned} &\text{Minimize} && c^T x \\ &\text{subject to} && \\ &&& a_i^T x \leq b_i, \quad i = 1, \dots, m, \end{aligned} \quad (۱.۶)$$

که در آن برخی از پارامترهای c ، a_i و b_i غیرقطعی هستند. برای سادگی فرض می‌کنیم c و b_i ها ثابت و a_i ها هم با مجموعه غیرقطعی بیضی‌گونهای به فرم زیر تعریف می‌شوند:

$$a_i \in \varepsilon_i = \{\bar{a}_i + P_i u \mid \|u\| \leq 1\},$$

وقتی $P_i \succeq 0$ ، یعنی ماتریس‌های متقارن و نیمه معین مثبت هستند. در بدترین حالت، ما نیاز داریم که قيود برای همه مقادیر ممکن از پارامترهای a_i صدق کند که منجر به برنامه‌ریزی خطی مقاوم زیر می‌شود:

$$\begin{aligned} &\text{Minimize} && c^T x \\ &\text{subject to} && \\ &&& a_i^T x \leq b_i \quad \forall a_i \in \varepsilon_i, \quad i = 1, \dots, m. \end{aligned} \quad (۲.۶)$$

قيود مقاوم خطی بالا را می‌توان به صورت زیر بیان کرد:

$$\max\{a_i^T x \mid a_i \in \varepsilon_i\} = \bar{a}_i^T x + \|P_i x\| \leq b_i, \quad (۳.۶)$$

واضح است که قيود مخروط مرتبه دوم می‌باشد. بنابراین (۸.۶) را می‌توان به صورت برنامه‌ریزی مخروط مرتبه دوم زیر بیان کرد:

$$\text{Minimize} \quad c^T x$$

subject to

$$\bar{a}_i^T x + \|P_i x\| \leq b_i \quad i = 1, \dots, m. \quad (۴.۶)$$

همچنین برنامه ریزی خطی مقاوم را می‌توان در چارچوب آماری در نظر گرفت. در این حالت ما فرض می‌کنیم پارامترهای a_i متغیرهای تصادفی نرمال مستقل، با میانگین $\bar{a}_i$ و واریانس Σ_i باشند. ما نیاز داریم که هر قید با احتمال حداقل η وقتی $\eta \geq 0$ باشد، برقرار باشد یعنی:

$$P(a_i^T x \leq b_i) \geq \eta. \quad (۵.۶)$$

می‌خواهیم نشان دهیم که قید احتمالی فوق را می‌توان به صورت یک قید مخروط مرتبه دوم بیان کرد.

فرض کنید $u = a_i^T x$ ، با σ^2 که نشان دهنده واریانس آنها است، بنابراین داریم:

$$P\left(\frac{u - \bar{u}}{\sqrt{\sigma^2}} \leq \frac{b_i - \bar{u}}{\sqrt{\sigma^2}}\right) \geq \eta.$$

چون $z = \frac{u - \bar{u}}{\sqrt{\sigma^2}} \sim N(0, 1)$ و همچنین،

$$\Phi(z) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^z e^{-\frac{t^2}{2}} dt,$$

بنابراین قید (۵.۶) را می‌توان به صورت زیر بیان کرد:

$$\frac{b_i - \bar{u}}{\sqrt{\sigma^2}} \geq \Phi^{-1}(\eta),$$

یا

$$\bar{u} + \Phi^{-1}(\eta)\sqrt{\sigma^2} \leq b_i.$$

از $\bar{u} = \bar{a}_i^T x$ و $\sigma^2 = x^T \Sigma_i x$ داریم:

$$\bar{a}_i^T x + \Phi^{-1}(\eta)\|\Sigma_i^{\frac{1}{2}} x\| \leq b_i.$$

اکنون در صورتیکه $\eta \geq 0.5$ یعنی $\Phi^{-1}(\eta) \geq 0$ باشد، این قید یک قید مخروط درجه دوم است.

بنابراین، مسأله

$$\text{Minimize} \quad c^T x$$

subject to

$$P(\bar{a}_i^T x \leq b_i) \geq \eta \quad i = 1, \dots, m. \quad (۶.۶)$$

می‌تواند با برنامه‌ریزی مخروط مرتبه دوم زیر بیان شود:

$$\text{Minimize} \quad c^T x$$

subject to

$$\bar{a}_i^T x + \Phi^{-1}(\eta) \|\Sigma_i^{\frac{1}{2}} x\| \leq b_i \quad i = 1, \dots, m. \quad (۷.۶)$$

که این مسائل با روشهای نقطه درونی قابل حل بوده، ما نیز برای حل این مسائل از شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) کمک می‌گیریم.

۵.۶ اشکال گوناگون مقاوم بودن

۱.۵.۶ مقاوم بودن محدودیت‌ها

یک مفهوم مهم در این تکنیک مقاوم بودن محدودیت‌ها^{۲۵} است که اغلب به آن مقاوم بودن مدل^{۲۶} می‌گویند. این مقاوم بودن مربوط به به‌دست آوردن جوابهایی است که برای تمامی مقادیر ممکن ورودی‌های غیرقطعی شدنی باقی بماند. مدل زیر را در نظر بگیرید:

$$\text{Minimize}_x \quad f(x)$$

subject to

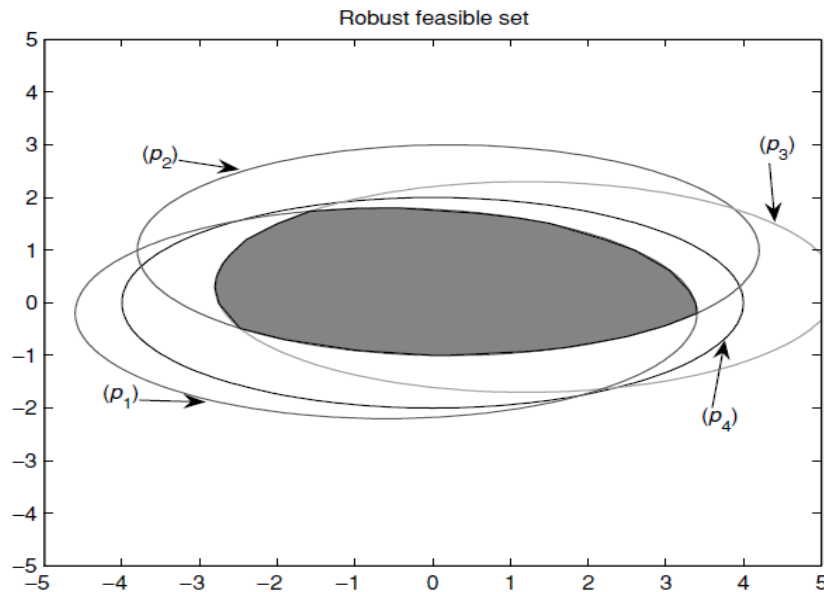
$$G(x, p) \in K. \quad (۸.۶)$$

در صورتیکه p را پارامترهای نامطمئن مسأله و $\mathcal{U}$ را مجموعه عدم قطعیت در نظر بگیریم، مدل مقاوم مسأله به‌صورت زیر خواهد بود:

$$\text{Minimize}_x \quad f(x)$$

^{۲۵}Constraint Robustness

^{۲۶}Model Robustness



شکل ۱.۶: مقاوم بودن محدودیتها

subject to

$$G(x, p) \in K, \quad \forall p \in \mathcal{U}. \quad (9.6)$$

مسأله (۸.۶) نشان می‌دهد، مجموعه شدنی مقاوم اشتراک مجموعه‌های شدنی

$S(p) = \{x : G(x, p) \in K\}$ ، با مجموعه عدم قطعیت $\mathcal{U}$ است. ما آن را در شکل (۱.۶) برای مجموعه

بیضی‌گون $\mathcal{U} = \{p_1, p_2, p_3, p_4\}$ ، و قتیکه p_i متناظر با مرکز عدم قطعیت بیضی است نشان داده‌ایم.

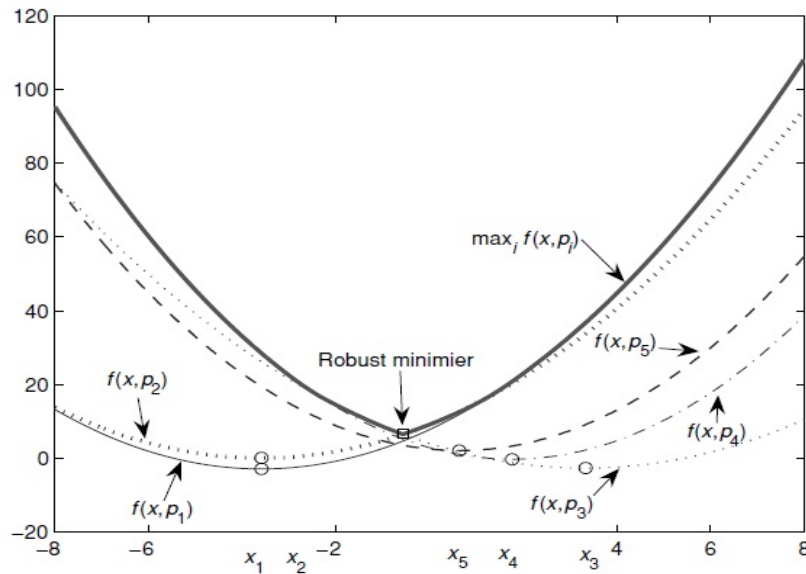
۲.۵.۶ مقاوم بودن تابع هدف

مفهوم دیگری که در مقاوم بودن مطرح است، مقاوم بودن تابع هدف است. در این حالت به دنبال جواب‌هایی

هستیم که برای تمام حالت‌های ممکن پارامترهای غیرقطعی، نزدیک به بهینگی باقی بماند. مدل زیر را

برای بردار p در تابع هدف در نظر بگیرید:

$$\text{Minimize}_x \quad f(x, p)$$



شکل ۲.۶: مقاوم بودن تابع هدف

subject to

$$x \in S, \quad (10.6)$$

که در آن S مجموعه شدنی و تابع هدف که وابسته به پارامترهای غیرقطعی p است. بنابراین با حل مسأله زیر جواب مقاوم تابع هدف به دست می‌آید:

$$\text{Minimize}_{x \in S} \text{Maximize}_{p \in \mathcal{U}} f(x, p) \quad (11.6)$$

مسأله مقاوم بودن تابع هدف (۱۱.۶) در شکل (۲.۶) نشان داده شده است. در این مثال، S مجموعه شدنی خط حقیقی، $\mathcal{U} = \{p_1, p_2, p_3, p_4, p_5\}$ و $f(x, p_i)$ تابع درجه دوم محدب است. برای بررسی اشکال دیگر مرجع [۱۴۷] را ببینید.

۶.۶ مدل های بهینه سازی مقاوم در مسائل مالی

بسیاری از مسائل مالی شامل مقادیر آتی از قیمت دارایی‌ها، نرخ بهره، نرخ ارز و غیره است که در زمان حال معلوم نیستند اما می‌توان آنها را پیش‌بینی یا برآورد کرد. از جمله فرمول‌های بهینه‌سازی مقاوم برای مسائل مالی می‌توان به مسائل انتخاب سبد سرمایه^{۲۷}، مدیریت ریسک^{۲۸}، پوشش و قیمت‌گذاری مشتقات مالی^{۲۹} اشاره کرد.

۱.۶.۶ انتخاب سبد سرمایه مقاوم

بن-تال و نمیروفسکی در [۱۴۸] یک مسأله انتخاب سبد سرمایه چند دوره‌ای مقاوم را با استفاده از رویکرد برنامه‌ریزی خطی فرمول بندی کردند که با روش سیمپلکس یا روشهای نقطه درونی قابل حل است. گلدفارب و آینگار^{۳۰} [۱۴۹] و نیز لوبو و بوید^{۳۱} در [۱۲۵] مسأله سبد مالی مقاوم را در چارچوب میانگین-واریانس مورد مطالعه قرار دادند و برای بازده دارایی‌های ریسکی، برخی از انواع عدم قطعیت مانند بیضی‌گون و بازه‌ای را معرفی کردند. آن‌ها مسأله را به برنامه‌ریزی مخروط مرتبه دوم و برنامه‌ریزی نیمه معین تبدیل کردند، که توسط الگوریتم‌های نقطه درونی به صورت کارایی قابل حل هستند. در این بخش با یادآوری مدل میانگین-واریانس مارکویتز می‌توان ترکیبی از بازدهی و ریسک را در تابع هدف به صورت زیر داشته باشیم:

$$\text{Max}_{x \in \mathcal{X}} \mu(x) - \lambda x^T \Sigma x \quad (۱۲.۶)$$

وقتی که μ_i برآورد از بازدهی موردانتظار از دارایی i ام، σ_{ij} کواریانس بین بازدهی دارایی i و j ام و λ هم نرخ ریسک‌گریزی^{۳۲} که تعادلی بین ریسک و بازده برقرار می‌سازد، مجموعه $\mathcal{X}$ ، مجموعه همه سبد

^{۲۷}Portfolio Selection

^{۲۸}Risk Managment

^{۲۹}Dreivatives Pricing/Hedging

^{۳۰}Goldfarb and Iyengar

^{۳۱}Lobo and Boyd

^{۳۲}Risk-Aversion

سرمایه‌های شدنی که شامل اطلاعاتی چون، محدود کردن موقعیت کوتاه، استفاده از کل بودجه و غیره است. چون قیود ما مشخص شده هستند، ما می‌توانیم فرض کنیم مجموعه $\mathcal{X}$ بدون هیچ عدم قطعیتی در زمانی که مسأله حل می‌شود. یکی از محدودیت‌های این مدل، برآورد دقیق و صحیح از بازدهی‌های مورد انتظار و واریانس‌هاست. در [۱۵۰] باوا^{۳۳} و همکارانش درباره استفاده برآوردهایی از بازدهی‌های مورد انتظار نامعلوم و کواریانس بحث کردند که منجر به برآورد ریسک در انتخاب سبد سرمایه شد. بنابراین، جواب بهینه به آشفتگی در پارامترهای ورودی-تغییرات کوچک در برآورد بازدهی و واریانس حساس است، که شاید منجر به تغییرات بزرگ در جواب متناظر باشد. برای مثال [۱۵۱، ۱۵۲] را ببینید. برای بهینه‌سازی مقاوم سبد سرمایه، ما مدلی که اطلاعات بازدهی و ماتریس کواریانس دارای فرم بازه‌ای هستند را به صورت زیر نظر می‌گیریم:

$$\mathcal{U} = \{(\mu, \Sigma) : \mu^L \leq \mu \leq \mu^U, \Sigma^L \leq \Sigma \leq \Sigma^U, \Sigma \succeq 0\}. \quad (۱۳.۶)$$

با استفاده از مدل مقاوم بودن تابع هدف در (۱۱.۶)، ما می‌توانیم سبد سرمایه را پیدا کنیم که تابع هدف در (۱۲.۶) را در بدترین حالت تشخیص پارامترهای ورودی از مجموعه $\mathcal{U}$ در (۱۳.۶) بیشینه کند. با در نظر گرفتن این شرایط مسأله بهینه‌سازی مقاوم به صورت زیر بیان می‌شود:

$$\text{Max}_{x \in \mathcal{X}} \{ \text{Min}_{(\mu, \Sigma) \in \mathcal{U}} -\mu(x) + \lambda x^T \Sigma x \}. \quad (۱۴.۶)$$

چون $\mathcal{U}$ کراندار است، با استفاده از نتایج کلاسیک از آنالیز محدب می‌توان نشان داد که (۱۴.۶) معادل دوگان آن است وقتی که جای Max و Min با هم عوض می‌شود:

$$\text{Min}_{(\mu, \Sigma) \in \mathcal{U}} \{ \text{Max}_{x \in \mathcal{X}} \mu(x) - \lambda x^T \Sigma x \}.$$

بنابراین، جواب مسأله (۱۴.۶) نقطه زینی تابع $f(x, \mu, \Sigma) = \mu(x) - \lambda x^T \Sigma x$ است و با استفاده از تکنیک‌های نقطه درونی در [۱۵۳] تعیین می‌شود.

^{۳۳}Bawa

۲.۶.۶ مثال عددی

در این بخش به ذکر یک مثال عددی که از [۱۵۴] گرفته شده می‌پردازیم، مسأله زیر را در نظر بگیرید:

$$\text{Maximize } \mu_1 x_1 + \mu_2 x_2 + \mu_3 x_3 \quad (15.6)$$

subject to

$$\begin{cases} TE(x_1, x_2, x_3) \leq 0.10, \\ x_1 + x_2 + x_3 = 1, \\ x_1, x_2, x_3 \geq 0, \end{cases} \quad (16.6)$$

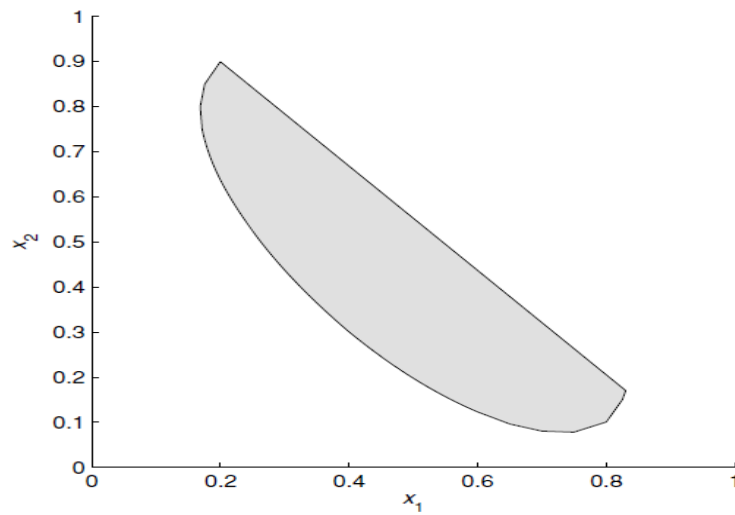
که در آن

$$TE(x_1, x_2, x_3) = \sqrt{\begin{bmatrix} x_1 - 0.5 \\ x_2 - 0.5 \\ x_3 \end{bmatrix}^T \begin{bmatrix} 0.1764 & 0.09702 & 0 \\ 0.09702 & 0.1089 & 0 \\ 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} x_1 - 0.5 \\ x_2 - 0.5 \\ x_3 \end{bmatrix}}.$$

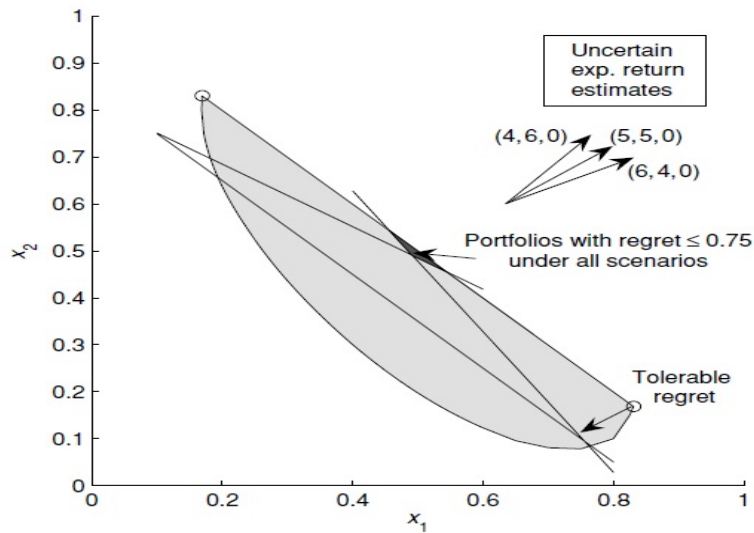
در این مسأله سبد سرمایه شامل ۳ دارایی است که دارایی سوم متناظر با x_3 ، نسبتی از بودجه است که سرمایه‌گذاری نشده است. دو دارایی اول دارای انحراف معیارهای ۴۲٪، ۳۳٪ و ضریب همبستگی ۰/۷ هستند. محک^{۳۴} در سبد سرمایه، سرمایه‌گذاری نصف-نصف میان دو دارایی اول است. تابع $TE(x)$ ، خطای ردیابی (پیگردی)^{۳۵} است که باید کمتر از ۱۰٪ باشد، این تابع که در آن x_{BM} مقدار محک است به صورت زیر تعریف می‌شود:

$$TE(x) = \sqrt{(x - x_{BM})^T \Sigma (x - x_{BM})}.$$

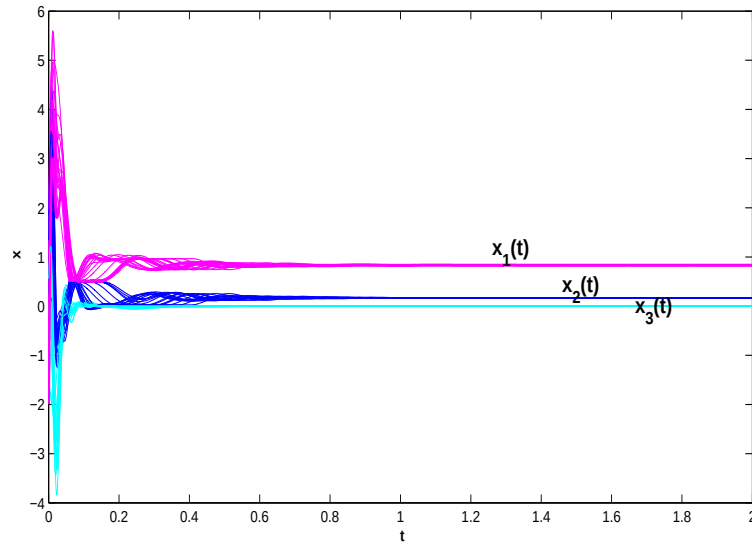
ما در شکل (۳.۶) ناحیه شدنی این مسأله را که توسط فضای متغیرهای x_1 و x_2 پوشانده شده را می‌بینیم. اکنون یک مدل مقاوم بودن نسبی برای این مسأله سبد سرمایه می‌سازیم، فرض کنید برآورد (تخمین) ماتریس کواریانس، قطعی است و مجموعه عدم قطعیت ساده‌ای شامل ۳ سناریو که در شکل (۴.۶) با پیکان‌هایی نشان داده شده است برای تخمین بازدهی موردانتظار در نظر می‌گیریم. این ۳ سناریو متناظر



شکل ۳.۶: ناحیه شدنی مسأله (۱۵.۶)–(۱۶.۶)



شکل ۴.۶: مجموعه جوابهایی با پشیمانی کمتر از ۷۵٪ مسأله (۱۵.۶)–(۱۶.۶)

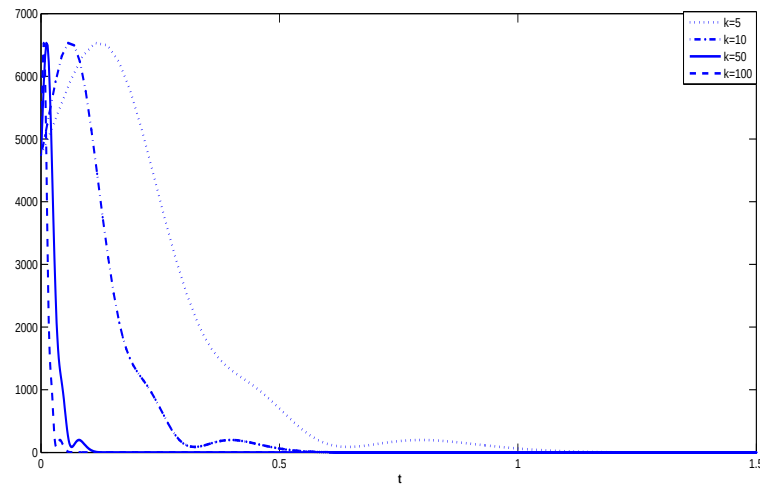


شکل ۵.۶: رفتار گذرا شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۲۰ نقطه آغازین دلخواه

با مقادیر $(\mu_1, \mu_2, \mu_3) = (6, 4, 0), (5, 5, 0), (4, 6, 0)$ هستند. به نظر می‌رسد مسأله فوق یک حالت خاص از برنامه‌ریزی مخروط مرتبه دوم است، بنابراین ما برای حل این مسأله از مدل شبکه عصبی پیشنهاد شده (۵۶.۴)–(۵۷.۴) استفاده می‌کنیم. جواب بهینه این مسأله وقتی $(\mu_1, \mu_2, \mu_3) = (6, 4, 0)$ باشد، $X^* = (0.831, 0.169, 0)$ است، با مقدار تابع هدف ۵/۶۶۲. شکل (۵.۶) نشان می‌دهد که مسیرهایی از شبکه عصبی (۵۶.۴)–(۵۷.۴) با ۲۰ نقطه اولیه متفاوت و $\epsilon = 10^{-6}$ همگرا به جواب بهینه است. همچنین تأثیر پارامتر κ از شبکه روی $\|x(t) - x^*\|^2$ در شکل (۶.۶) نشان داده شده است. مشابه، وقتی $(\mu_1, \mu_2, \mu_3) = (4, 6, 0)$ باشد، $X^* = (0.169, 0.831, 0)$ است، با مقدار تابع هدف ۵/۶۶۲، همچنین وقتی $(\mu_1, \mu_2, \mu_3) = (5, 5, 0)$ است تمام نقاط بین دو جواب بهینه قبلی، بهینه بوده با مقدار

^{۳۴}Benchmark

^{۳۵}Tracking Error



شکل ۶.۶: رفتار همگرایی از $\|x(t) - x^*\|^2$ با $y_0 = (40, -40, 40, -40, 40, -40, 40)^T$

هدف مشترک ۵. بنابراین فرم مقاوم بودن نسبی برای این مسأله می‌تواند به‌فرم زیر نوشته شود: ([۱۴۷])

$$\text{Minimize}_{x,t} \quad t \quad (17.6)$$

subject to

$$\begin{cases} 5/662 - (6x_1 + 4x_2) \leq t, \\ 5/662 - (4x_1 + 6x_2) \leq t, \\ 5 - (5x_1 + 5x_2) \leq t, \\ TE(x_1, x_2, x_3) \leq 0/10, \\ x_1 + x_2 + x_3 = 1, \\ x_1 \geq 0, x_2 \geq 0, x_3 \geq 0. \end{cases} \quad (18.6)$$

به‌جای حل کردن مسأله وقتی که سطح پیشیمانی^{۳۶} بهینه ما متغیر است (t), یک استراتژی ساده انتخاب،

سطحی از تأسف که می‌تواند پورتفوهای را بیابد که از سطح پیشیمانی در هر سناریو تجاوز نکند. برای

مثال حداکثر سطح پیشیمانی قابل تحمل^{۳۷} را ۰/۷۵ می‌گیریم بنابراین مسأله شدنی زیر را داریم:

$$\text{Find } x \quad (19.6)$$

subject to

^{۳۶}Regret Level

^{۳۷}Tolerable

$$\left\{ \begin{array}{l} ۵/۶۶۲ - (۶x_۱ + ۴x_۲) \leq ۰/۷۵, \\ ۵/۶۶۲ - (۴x_۱ + ۶x_۲) \leq ۰/۷۵, \\ ۵ - (۵x_۱ + ۵x_۲) \leq ۰/۷۵, \\ TE(x_۱, x_۲, x_۳) \leq ۰/۱۰, \\ x_۱ + x_۲ + x_۳ = ۱, \\ x_۱ \geq ۰, x_۲ \geq ۰, x_۳ \geq ۰ \end{array} \right. \quad (۲۰.۶)$$

این مسأله و مجموعه شدنی از جواب‌های آن در شکل (۴.۶) نشان داده شده است. گوشه مثلث که تیره‌تر نشان داده شده است بیانگر پورتفوهای است که دارای سطح پشیمانی کمتر و مساوی ۰/۷۵ تحت هر سه طرح هستند.

نتایج و پیشنهاداتی برای کارهای آتی

نتایج

در این پایان نامه سعی کرده‌ایم مروری بر شبکه‌های عصبی داشته و با استفاده از این نظریه به حل رده‌ای از مسائل بهینه‌سازی بپردازیم. یکی از انواع مسائل بهینه‌سازی، مسائل مالی هستند که به چهار دسته کلی مدیریت ریسک، انتخاب سبد سرمایه، قیمت‌گذاری اختیار معامله‌ها، مدیریت بدهی و دارایی تقسیم می‌شوند. این مسائل شامل برنامه‌ریزی خطی، برنامه‌ریزی درجه دوم، برنامه‌ریزی مخروط، بهینه‌سازی مقاوم و برنامه‌ریزی تصادفی هستند. برای مسائل انتخاب سبد سرمایه، مدل‌هایی با فرم برنامه‌ریزی خطی و برنامه‌ریزی درجه دوم معرفی کردیم و برای حل آنها به ارائه یک مدل از شبکه‌های عصبی پرداختیم. همچنین برای مسائل برنامه‌ریزی مخروط مرتبه دوم نیز مدل دیگری از شبکه‌های عصبی ارائه کردیم. این مدل پیشنهادی را برای حل مسائل قابل تبدیل به برنامه‌ریزی مخروط مرتبه دوم محدب از جمله برنامه‌ریزی تصادفی، رده‌ای از مسائل برنامه‌ریزی کسری و بهینه‌سازی مقاوم بکار بردیم. همچنین پایداری و همگرایی سراسری مدل‌ها مورد بررسی قرار گرفت و کارایی مدل‌های پیشنهاد شده را با حل چندین مثال عددی نشان دادیم.

از جمله مزیت‌های مدل‌های ارائه شده می‌توان به موارد زیر اشاره کرد:

* همگرایی سراسری مدل و وابسته نبودن به نقطه اولیه لزوماً شدنی برخلاف بیشتر روشهای بهینه سازی از جمله روشهای نقطه درونی و سیمپلکس که ممکن است مسأله جواب یکتا نداشته باشد

و در دور بیافتد.

* نداشتن پارامتر جریمه بر خلاف مدل‌های گرادینانی و سادگی مدل پیشنهاد شده.

* محدود نبودن مدل به تعداد متغیرها و قیود بر خلاف روش‌های دیگر و نرم افزارها.

پیشنهادهای برای کارهای آتی

از جمله تحقیقات دیگری که در ادامه می‌توان انجام داد عبارتند از:

۱. کلاس‌های مختلفی از مسائل بهینه‌سازی سبد و دیگر مسائل مالی از جمله مدیریت ریسک و غیره

که قابل بیان به فرم برنامه‌ریزی محدب درجه دوم باشند را با این مدل حل نماییم.

۲. مسائل مدیریت ریسک شامل تابع ارزش در معرض خطر و مسائل مدیریت بدهی و دارایی‌ها که

حالت خاصی از مسائل برنامه‌ریزی تصادفی هستند را با کمک مدل‌های شبکه عصبی حل نماییم.

۳. مسأله برنامه‌ریزی چند هدفه برای انتخاب سبد سرمایه که با روش بهینه‌سازی با قیود تصادفی

قابل تبدیل به فرم برنامه‌ریزی مخروط مرتبه دوم است را با این مدل پیشنهادی حل نماییم.

۴. امید می‌رود که بتوان مسائل برنامه‌ریزی کسری با قیود کسری، مسائل برنامه‌ریزی کسری با ضرایب

تصادفی، مسائل چندهدفه و مسائل مینیماکس کسری را با روش مذکور و یا با روشهای دیگر حل

کرد.

۵. مسائلی دیگر از بهینه‌سازی مقاوم و مسائل مالی مثل مدیریت ریسک، قیمت‌گذاری اختیار معامله‌ها،

آربیتراژ و غیره را که قابل بیان شدن به فرم برنامه‌ریزی مخروط مرتبه دوم محدب را حل نمود.

امید می‌رود بتوان مسأله بهینه‌سازی مقاوم سبد مالی یک (چند دوره‌ای) با استفاده از ارزش در

معرض خطر که همتای مقاوم آنها منجر به برنامه‌ریزی مخروط درجه دوم شوند را با مدل شبکه عصبی پیشنهاد شده حل کرد.

۶. همچنین با کمک مدل‌های ارائه شده می‌توان مسائل حرکت روبات‌ها و تابع رگرسیون را نیز حل نمود.

مراجع و مآخذ

- [1] M. S. Bazaraa, H. D. Sherali, C. M. Shetty, (1993), **"Nonlinear Programming-Theory and Algorithms"**, 2nd ed. New York: Wiley.
- [2] B. N. Datta,(2003), **"Numerical Methods for Linear Control Systems"**, Elsevier Academic Press, San Diego, CA.
- [3] G. Xing-Bao, (2003), "Exponential Stability of Globally Projected Dynamic Systems", **IEEE Trans.On Neural Networks**,14,2, 426-431.
- [4] S. Effati, A. R. Nazemi, (2006), "Neural network models and its application for solving linear and quadratic programming problems", **Appl. Math. Comput**, 172, 305-331.
- [5] G. Simmons, (2006), **"Differential Equations and Applications"**, McGraw-Hill Education (India).
- [6] J. C. Hull, (2006), **"Options,Futures,And Other Derivatives"**, 6th Edition, Perntice Hall,New Jersey.
- [7] C. P. Jones, (2001), **"Investments: Analysis and Management"**, New York: Wiley.
- [8] M. J. A. Berry, G. Linoff, (1997), **"Data Mining Techniques"**, John wiley and sons, 450-460.

- [9] B. Shelly, S. Stephenson, (2000), "The use of artificial neural network in completion simulation and design", **Computer and Geosciences**, 26, 941-951.
- [10] S. Haykin, (1999), "**Neural Network: A Comprehensive Foundation**", prentice hall, pp 842.
- [11] D. L. Hudson, M. E. Cohen, (1999), "Neural Network and Artificial Intelligence for Biomedical Engineering", **IEEE Press Series on Biomedical Engineering**.
- [12] J. J. Hopfield, D. W. Tank, (1986), "Simple neural optimization network: An A/D converter, signal decision circuit, and a linear programming circuit", **IEEE Trans on Circuits and Systems**, 33, 533-541.
- [13] M. P. Kennedy, L. O. Chua, (1988), "Neural networks for nonlinear programming", **IEEE Trans. on Neural Networks**, 35, 554-562.
- [14] A. Bouzerdoum, T. R. Pattison, (1993), "Neural network for quadratic optimization with bound constraints", **IEEE Trans. on Neural Networks**, 4, 293-304.
- [15] Y. Xia, (1996), "A neural network for solving linear and quadratic programming problems", **IEEE Trans. on Neural Networks**, 7, 1544-1547.
- [16] Y. Xia, G. Feng, (2005), "An improved network for convex quadratic optimization with application to real-time beam forming", **Neurocomputing**, 64, 359-374.
- [17] Y. Xia, G. Feng, J. Wang, (2004), "A recurrent neural network with exponential convergence for solving convex quadratic program and related linear piecewise equation", **Neural Networks**, 17, 1003-1015.
- [18] Y. Xia, J. Wang, (2004), "A general projection neural network for solving monotone variational inequality and related optimization problems", **IEEE Trans. On Neural Networks**, 15, 318-328.

- [19] Y. Xia, H. Leung, J. Wang, (2002), "A projection neural network and its application to constrained optimization problems", **IEEE Transactions On Circuits And Systems**, 44, 4, 447-458.
- [20] Y. Xia, J. Wang, (2000), "A recurrent neural network for solving linear projection equation", **Neural Networks**, 13, 337-350.
- [21] X. Wu, Y. Xia, J. Li and W. Chen, (1996), "A High Performance Neural Network for solving Linear and Quadratic Programming Problems", **IEEE Trans. On Neural Networks**, 7, 3, 643-651.
- [22] A. Rodriguez-Vazquez, R. Dommnguer-Castro, A. Rueda, J. L. Huertas and E. Sahnchez-Sinencio, (1990), "Nonlinear switched-capacitor neural network for optimization problems", **IEEE Trans on Circuits and Systems**, 384-390.
- [23] L. Qi, J. Sun, (1993), "A nonsmooth version of Newton method", **Math Program**, 58, 353-368.
- [24] S. Azarbek, A. R. Nazemi, D. SHahsavani, (2010), "An Application of Neural networks for solving linear and quadratic programming problems", **3th International Conference of Iranian Operations Research**.
- [25] S. Effati, A. Ghomashi, A. R. Nazemi, (2007), "Application of projection neural network in solving convex programming problems", **Appl. Math. Comput**, 188, 1103-1114.
- [26] W. F. Sharpe, (1964), "Capital Asset Price", **Journal of Finance**, 19, 3, 16-18.
- [27] J. lintner, (1965), "The Valuation of Risk Asset and the selection of risky Investments in stock: Portfolio and Capital Budgets", **Reviwe of Economics and Statistics**, 47, 22-26.

- [28] S.A. Ross, (1976), "The arbitrage theory of capital asset pricing", **Journal of Economic Theory**, 13, 341-360.
- [29] J. Mossin, (1966), "Equilibrium in a capital asset market", **Econometrica**, 34(4), 14-16.
- [30] F. Black, R. Litterman, (1990), "Asset allocation: combing investor views with market equilibrium", **Goldman Sachs Fixed Income Research**.
- [31] H. Markowitz, (1952), "Portfolio Selection", **Journal of Finance**, 7, 77-91.
- [32] H. Markowitz, (1959), "Portfolio selection: Efficient Diversification of investment", John Wiley and Sons, New York.
- [33] S. M. Lee, A. J. Lerro, (1973), "Optimization the portfolio selection for Mutual Funds", **The Journal of finance**, 28, 1087-1099.
- [34] S. M. Lee, D. L. Chesser, (1980), "Goal programming for portfolio", **The Journal of portfolio manegmant**, 22-26.
- [35] R. Mansini, M.G. Speranza, (1997), " Heuristic Algorithms for the portfolio selection problem with Minimum Transaction Lots", **European Journal of Opeartional Research**, 114, 219-233.
- [36] M.R. Young, (1998), "A minimax portfolio selection rule with linear programming solution", **Management Science**, 44, 673-683.
- [37] Y. Xia, B. Liu, W. Shoayang, K. Lai, (1999), "A model for portfoilo selection with order of Expected Returns", **Computers and Opeartionals Research**, 27, 409-424.

- [38] T. Chang, T. Meade, J. E. Beasley, Y. M. Sharaihan, (2000), " Heuristics for Cardinay Constrained portfoilo optimization", **Computers and Opeartionals Research**, 27, 1271-1302.
- [39] J. Palmquist, S. Uryaser, P. Krokmal, (1999), "**Portfolio optimization with conditional Value-at-Risk objective and constraints**".
- [40] S. M. Lobo, M. Fazel, S. Boyd, (2002), "**Portfolio optimization with Linear and Fixed Transaction Coasts**", California Institue of Technology.
- [41] C. Papahristodoulou, E. Dotzauer, (2004), "Optimal portfolios using linear programming models", **Journal of Operations Research Society**, 55, 1169-1177.
- [42] R. Mansini, W. Oyryczak, M. G. Speranza, (2003), "LP solvable models for portfolio optimization: Aclassification and Computational comparison", **IMA Journal of management mathematics**, 14, 187-220.
- [43] A. Fernandez, S. Gomez, (2007), "Portfolio selection using neural networks", **Computers and Operations Research**, 34, 1177–1191.
- [44] D. R. Kincard, (1942), "**Numerical Analysis Mathimatics of Scientific Computing**".
- [45] H. Konno, (1988), "**Portfolio Optimization using L_1 Risk Function**", Institute of Human and Social Sciences, IHSS Report Tokyo Institute of Technology, 88-89.
- [46] H. Konno, H. Yamazaki, (1991), "Mean-absolute deviation portfolio optimization model and its application to Tokyo Stock Market", **Management Science**, 37, 519-531.

- [47] Y. Simaan, (1997), "Estimation Risk in Portfolio Selection:The Mean Variance Model Versus the Mean Absolute Deviation Model", **Management Science**, 43, 1437-1446.
- [48] X. Q. Cai, K. L. Teo, X. Q. Yang, X. Y. Zhou, (2000), "Portfolio optimization under a minimax rule", **Management Science**, 46, 957-972.
- [49] M. Yu, H. Inouez, J. Shi, "**Portfolio optimization problems with linear programming models**", 14-16.
- [50] K. L. Teo, X. Q. Yang, (2001), "Portfolio selection problem with minimax type risk function", **Annals of Operations Research**, 101, 333-349.
- [51] O. Ogryczak, A. Ruszczyński, (1999), "From stochastic dominance mean-risk model: semideviation as risk measure", **European Journal of Operational Research**, 116, 33-50.
- [52] P. C. Fishburn, (1977), "Mean-risk analysis with risk associated with below-target returns", **American Economics Review**, 67, 116-126.
- [53] R. K. Miller, A. N. Michel, (1982), "**Ordinary Differential Equations**", Academic Press, New York.
- [54] H. Amann, (1990), "**Ordinary Differential Equations: An Introduction to Nonlinear Analysis**", Walter de Gruyter, Berlin.
- [55] J. M. Vegas, P. J. Zufiria, (2004), "Generalized neural networks for spectral analysis: dynamics and Liapunov functions", **Neural Networks**, 17, 233-245.
- [56] C. Papahristodoulou, E. Dotzauer, (2004), "Optimal portfolios using linear programming models", **Journal of Operations Research Society**, 55, 1169-1177.

- [57] G. Cornuejols, R. Tutuncu, (2007), **Optimization Methods in finance**, New York, 140-144.
- [58] E. Beale, (1955), "On minimizing a convex function subject to linear inequalities", **Journal of Royal Statistical Society**, Series B, 173–184.
- [59] A. Charnes, W. Cooper, (1959), "Chance constrained programming", **Management Science**, 5, 73–79.
- [60] G. Dantzig, (1955), "Linear programming under uncertainty", **Management Science**, 1, 197– 206.
- [61] F. Murphy, S. Sen, A. L. Soyster, (1982), "Electric utility capacity expansion planning with uncertain load forecasts", **AIIE Transaction**, 14, 52–59.
- [62] M. Dempster, G. Consigli, (1998), "The CALM stochastic programming model for dynamic assetliability management", **World Wide Asset and Liability Modelling**, Cambridge University Press, 464–500.
- [63] C. Dert, (1995), **Asset liability management for pension funds: A multistage chance constrained programming approach**, Ph.D. thesis, Erasmus University, Rotterdam, The Netherlands.
- [64] M. vander Vlerk, (2003), "Integrated chance constraints in an ALM model for pension funds", **Technical Report 03A21, University of Groningen, Research Institute SOM (Systems, Organisations and Management)**.
- [65] S. Sen, R. Doverspike, S. Cosares, (1994), "Network planning with random demand", **Telecommunication Systems**, 3, 11–30.
- [66] A. Tomasgard, S. Dye, S. Wallace, J. Audestad, L. Stougie, (1998), "Modelling aspects of distributed processing in telecommunication networks", **Annals of Operations Research**, 82, 161–184.

- [67] S. MirHassani, C. Lucas, G. Mitra, E. Messina, C. Poojari, Computational solution of capacity planning models under uncertainty, *Parallel Computing* 26 (2000) 511–538.
- [68] M. Dempster, N. H. Pedron, E. Medova, J. Scott, A. Sembos, (2000), "Planning logistic operations in the oil industry", **Journal of Operational Research Society**, 51, (11), 1271–1288.
- [69] G. Eppen, R. Martin, L. Schrage, (1989), "A scenario approach to capacity planning", **Operational Research**, 37, (4), 517–525.
- [70] J. R. Robinson, (1988), "Loaded questions: New approaches to utility forecasting", **Energy Policy** , 16, (1), 58–68.
- [71] S. Takriti, J. Birge, E. Long, (1996), "A stochastic model for the unit commitment problem", **IEEE Transactions on Power Systems**, 11, 1497–1508.
- [72] A. Kampas, B. White, (2003), "Probabilistic programming for nitrate pollution control: Comparing different probabilistic constraint approximations", **European Journal of Operational Research**, 147, (1), 217–228.
- [73] G. Laporte, F. Louveaux, L. V. Hamme, (1994), "Exact solution to a location problem with stochastic demands", **Transportation Science**, 28, 95–103.
- [74] S. E. Elmaghraby, H. Soewandi, M. J. Yao, (2001), "Chance-constrained programming in activity networks: A critical evaluation", **European Journal of Operational Research**, 131,(1), 440–458.
- [75] M. Ierapetritou, J. Acevedo, E. Pistikopoulos, (1996), "An optimization approach for process engineering problems under uncertainty", **Computers and Chemical Engineering**, 20, 703–709.

- [76] J. Smith, (1999), "Optimizing platform survivability using a chance constrained linear program", in: US Army Ground Vehicle Survivability Symposium. Survivability/lethality Analysis Directorate, Munitions and Platform Branch, White Sands Missile Range, NM.
- [77] A. Prekopa, (1990), "Dual method for a one-stage stochastic programming problem with random RHS obeying a discrete probability distribution", **Zeitschrift für Operations Research**, 34, 441–461.
- [78] A. Prekopa, (1993), "Programming under probabilistic constraint and maximizing a probability under constraints". center for Operations Research, Rutgers university, P.O.Box 5062, New Brunswick, New Jersey, 35-93.
- [79] F. Hillier, (1967), "Chance-constrained programming with 0–1 or bounded continuous decision variables", **Management Science**, 14, 34–57.
- [80] Y. Seppala, (1971), "Constructing sets of uniformly tighter linear approximations for a chanceconstraint", **Management Science**, 17, 736–749.
- [81] Y. Seppala, T. Orpana, (1984), "Experimental study on the efficiency and accuracy of a chanceconstrained programming algorithm", **European Journal of Operational Research**, 16, 345–357.
- [82] D. Olson, S. Swenseth, (1987), "A linear approximation for chance constrained programming", **Journal of Operational Research Society**, 38, 261–267.
- [83] M. Yahia, A. Daneshmand, (1995), "A linear approximation method for solving a special class of the chance constrained programming problem", **European Journal of Operational Research**, 80, 213–225.
- [84] A. Weintraub, J. Vera, (1991), "A cutting plane approach for chance constrained linear program", **Operations Research**, 39, (5), 770–776.

- [85] D. Dentcheva, A. Prekopa, A. Ruszczyński, (2000), "Concavity and efficient points of discrete distributions in probabilistic programming", **Mathematical Programming**, 89, 235–263.
- [86] D. Dentcheva, A. Prekopa, A. Ruszczyński, (2001), "On convex probabilistic programming with discrete distributions", **Nonlinear Analysis**, 47, 1997–2009.
- [87] P. Beraldi, A. Ruszczyński, (2002), "A branch and bound method for stochastic integer problems under probabilistic constraints", **Optimization Methods and Software**, 359–382.
- [88] A. Ruszczyński, (2002), "Probabilistic programs with discrete distributions and precedence constrained knapsack polyhedra", **Mathematical Programming**, 93, 195–215.
- [89] R. Aringheri, (2004), "Lecture Notes in Computer Science, Solving Chance-Constrained Program Combining Tabu Search and Simulation", Springer-Verlag, Berlin, Heidelberg, 30–41.
- [90] G. Calafiore, M. Campi, (2005), "Uncertain convex programs: randomized solutions and confidence levels", **Mathematical Programming**, 102, 1, 25–46.
- [91] D. de Farias, B. V. Roy, (2004), "On constraint sampling for the linear programming approach to approximate dynamic programming", **Mathematics of Operations Research**, 29, (3), 462–478.
- [92] E. Erdogan, G. Iyengar, (2004), "Ambiguous chance constrained problems and robust optimization", CORC Technical Report TR-2004-10, IEOR Department, Columbia University.

- [93] K. Iwamura, B. Liu, (1996), "A genetic algorithm for chance constrained programming", **Journal of Information and Optimization Sciences**, 17, (2), 409-422.
- [94] C. A. Poojari, B. Varghese, (2008), "Genetic Algorithm based technique for solving Chance Constrained Problems", **European Journal of Operational Research**, 185, 1128-1154.
- [95] A. Rodriguez-Vazquez, R. Dominguez-Castro, A. Rueda, J. L. Huertas and E. Sanchez- Sinencio, (1990), "Nonlinear switched-capacitor neural networks for optimization problems", **IEEE Transactions on Circuits and Systems**, 37, 384-397.
- [96] K. Z. Chen, Y. Leung, K. S. Leung, X. B. Gao, (2002), "A neural network for nonlinear programming problems", **Neural Computing and Application**, 11, 103-111.
- [97] T. L. Friesz, D. H. Bernstein, N. J. Mehta, R. L. Tobin, and S. Ganjizadeh, (1994), "Day-to-day dynamic network disequilibria and idealized traveler information systems", **Operation Research**, 42, 1120-1136.
- [98] M. Forti, P. Nistri, M. Quincampoix, (2004), "Generalized neural network for nonsmooth nonlinear programming problems", **IEEE Transactions on Circuits and System-I**, 51, 1741-1754.
- [99] X. Gao, (2004), "A novel neural network for nonlinear convex programming", **IEEE Transactions on Neural Networks**, 15, 613-621.
- [100] C. H. Ko, J. S. Chen, C. Y. Yang, (2011), "Recurrent neural networks for solving second-order cone programs", **Neurocomputing**, 74, 3646-3653.

- [101] Y. Leung, K. Chen, X. Gao, (2003), "A high-performance feedback neural network for solving convex nonlinear programming problems", **IEEE Transactions on Neural Networks**, 14, 1469–1477.
- [102] X. B. Liang, J. Wang, (2000), "A recurrent neural network for nonlinear optimization with a continuously differentiable objective function and bounded constraints", **IEEE Transactions on Neural Networks**, 11, 1251–1262.
- [103] L. Z. Liao, H. D. Qi, (1999), "A neural network for the linear complementarity problem", **Mathematical and Computer Modeling**, 29, (3), 9-18.
- [104] W. E. Lillo, M. H. Loh, S. Hui and S. H. Zak, (1993), "On solving constrained optimization problems with neural networks: A penalty method approach", **IEEE Trans. Neural Networks**, 4, 931-939.
- [105] A. Malek, N. Hosseini-pour-Mahani, S. Ezazipour, (2010), "Efficient recurrent neural network model for the solution of general nonlinear optimization problems", **Optimization Methods and Software**, 25, 1–18.
- [106] X. Mu, S. Liu, Y. Zhang, (2005), "A neural network algorithm for second-order conic programming", in: Proceedings of the **Second International Symposium on Neural Networks**, Chongqing, China, Part II, 718724.
- [107] A. R. Nazemi, (2012), "A dynamic system model for solving convex nonlinear optimization problems", **Communications in Nonlinear Science and Numerical Simulation**, 17, 1696–1705.
- [108] A. R. Nazemi, (2011), "A dynamical model for solving degenerate quadratic minimax problems with constraints", **Journal of Computational and Applied Mathematics**, 236, 1282-1295.

- [109] A. R. Nazemi, F. Omid, (2012), "A capable neural network model for solving the maximum flow problem", **Journal of Computational and Applied Mathematics**, 236, 14, 3498-3513.
- [110] A. R. Nazemi, F. Omid, (2012), "An efficient dynamic model for solving the shortest path problem", **Transportation Research Part C: Emerging Technologies**, (In press).
- [111] A. R. Nazemi, E. Sharifi, (2012), "Solving a class of geometric programming problems by an efficient dynamic model", **Communications in Nonlinear Science and Numerical Simulation**, (In press).
- [112] J. Sun, J. S. Chen, C. H. Ko, "Neural networks for solving second-order cone constrained variational inequality problem", **Computational Optimization and Applications**, (First online).
- [113] J. Sun, L. Zhang, (2009), "A globally convergent method based on Fischer–Burmeister operators for solving second-order cone constrained variational inequality problems", **Computers and Mathematics with Applications**, 58, 1936–1946.
- [114] Q. Tao, J. Cao, M. Xue, H. Qiao, (2001), "A high performance neural network for solving nonlinear programming problems with hybrid constraints", **Physics Letters A**, 288, 88–94.
- [115] J. Wang, Q. Hu, D. Jiang, (1999), "A Lagrangian neural network for kinematic control of redundant robot manipulators", **IEEE Transactions on Neural Networks**, 10, (5), 1123-1132.

- [116] Y. Xia, J. Wang, (2005), "A recurrent neural network for solving nonlinear convex programs subject to linear constraints", **IEEE Transactions on Neural Networks**, 16, 379-386.
- [117] Y. Xia, J. Wang, (2004), "A recurrent neural network for nonlinear convex optimization subject to nonlinear inequality constraints", **IEEE Transactions on Circuits and Systems I:Regular Papers** 51, (7), 1385-1394.
- [118] Y. Xia, G. Feng, (2007), "A new neural network for solving nonlinear projection equations", **Neural Networks**, 20, 577-589.
- [119] Y. Xia, J. Wang, L. M. Fok, (2004), "Grasping-force optimization formultifingered robotic hands using a recurrent neural network", **IEEE Transactions on Robotics and Automation**, 20, (3), 549-554.
- [120] X. Xue, W. Bian, (2007), "A project neural network for solving degenerate convex quadratic program", **Neural Networks**, 70, 2449-2459.
- [121] X. Xue, W. Bian, (2008), "Subgradient-based neural network for nonsmooth convex optimization problems", **IEEE Transactions on Circuits and System-I**, 55, 2378–2391.
- [122] Y. Yanga, J. Cao, (2010), "The optimization technique for solving a class of non-differentiable programming based on neural network method", **Nonlinear Analysis**,11, 1108–1114.
- [123] Y. Yang, X. Xu, (2007), "The projection neural network for solving convex nonlinear programming, D. S. Huang, L. Heutte, and M. Loog (Eds.): ICIC 2007, LNAI 4682, Springer-Verlag Berlin Heidelberg, 174-181.

- [124] Y. Yang, J. Cao, (2008), "A feedback neural network for solving convex constraint optimization problems", **Applied Mathematics and Computation**, 201, 340-350.
- [125] M. S. Lobo, L. Vandenberghe, S. Boyd, H. Lebre, (1998), "Applications of second-Order cone programming", **Linear Algebra and its Applications**, 284, 193-228.
- [126] H. Y. Benson, R. J. Vanderbei (2003), "Solving Problems with Semidefinite and Related Constraints Using Interior-Point Methods for Nonlinear Programming", **Mathematical Programming Series B**, 95, 279-302.
- [127] E. Chinlar,(1975), "**Interoduction to stochastic processes**", Northwestern university.
- [128] S. Effati, M. M. Moghadam, (2006), "Conversion of some classes of fractional programming to second-order cone programming and solving it by potential reduction interior point method", **Applied Mathematics and Computation**, 181, 563-578.
- [129] B.D. Craven, (1988), "Fractional Programming", **Sigma Series in Applied Mathematics**, 4, Helderman Verlag, Berlin, Germany, 144-146.
- [130] S. Schaible, T. Ibaraki, (1983), "Fractional programming", **European Journal of Operational Research**, 12, 325-338.
- [131] S. Schaible, (1995), "**Fractional programming**", Handbook of Global Optimization, Kluwer Academic Publishers, Dordrecht, Netherlands, 495-608.
- [132] I.M. Stancu-Minasian, (1997), "**Fractional Programming: Theory Methods and Applications**", Kluwer Academic Publishers, Dordrecht, Netherlands.

- [133] I.M. Stancu-Minasian,(1980), "Applications of the fractional programming, theory, methods and applications", **Economic Comput. Economic Cyber. Studies Res**,14, 69–86.
- [134] I.M. Stancu-Minasian,(1999), " A fifth bibliography of fractional programming", **Optimization**, 45, 343–367.
- [135] I.M. Stancu-Minasian,(2006), " A sixth bibliography of fractional programming", **Optimization**, 55, 405–428.
- [136] G. B. Dantzing, (1955), "**Linear Programming under Uncertainty**", *Magmt. Sci*, 197-206.
- [137] R. E. Bellman, (1957), "**Dynamic Programming**", Princeton University Press.
- [138] A. Ben-Tal, A. Nemirovski, (1998), "Robust convex optimization", **Mathematics of Operations Research**, 23, 769-805.
- [139] A. Ben-Tal, A. Nemirovski, (1995), "Robust solution to uncertain linear problems", Report No 6/95, Optimization Laboratory, Faculty of Industrial Engineering and Management, Technion- the Israel Institute of Technology, Technion City, Haifa 32000, Israel.
- [140] A. Ben-Tal, L. EL Ghaoui, A. Nemirovski, (2009), "**Robust optimization**", Princeton University, Princeton Series in Applied Mathematics.
- [141] A. Ben-Tal, A. Nemirovski, (1999), "Robust solutions to uncertain programs", **Operations Research Letters**, 25, 1-13.
- [142] A. Ben-Tal, A. Nemirovski, (2000), "Robust solutions of linear programming problems contaminated with uncertain data", **Mathematical Programming**, 88, 411-424.

- [143] A. Ben-Tal, L. EL Ghaoui, A. Nemirovski, (2000), "**Robust semidefinite programming**", Handbook on Semidefinite Programming and application, Kluwer Academic Publishers, 139-162.
- [144] A. Ben-Tal, A. Nemirovski, C. Roos, (2002), "Robust solutions to uncertain quadratic and conic-quadratic problems", **SIAM Journal on Optimization**, 13, 535-560.
- [145] A. Ben-Tal, A. Goryashko, E. Guslitzer, A. Nemirovski, (2004), "Adjustable robust solutions of uncertain linear programs", **Mathimatical Programming**, 99, 351-376.
- [146] L. El-Ghaoui, H. Lebret, (1997), "Robust solutions to least-squares problems with uncertain data matrices", **SIAM Journal on Matrix Analysis Application**, 18, 1035-1064.
- [147] G. Cornuejols, R. H. Tutuncu, (2007), "**Optimizatiom method in finance**", 295-301.
- [148] A. Ben-Tal, T. Margalit, A. N. Nemirovski, (2002), "Robust modeling of multi-stage portfolio problems", In H. Frenk, K. Roos, T. Terlaky, and S. Zhang, editors, High Performance OptimizatioNn, 303-328.
- [149] D. Goldfarb, G. Iyengar, (2003), "Robust portfolio selection problems", **Mathematics of Operations Research**, 28, 1-38.
- [150] V. S. Bawa, S. J. Brown, and R.W. Klein,(1976), "**Estimation Risk and Optimal Portfolio Choice**", (Amsterdam: North-Holland).
- [151] R. O. Michaud, (1989), "he Markowitz optimization enigma: is optimized optimal", **Financial Analysts Journal**, 45, 31-42.

- [152] R. O. Michaud,(1998), "**Efficient Asset Management**", (Boston, MA: Harvard Business School Press).
- [153] B. Halldorsson, R. H. Tutuncu, (2003), "An interior-point method for a class of saddle point problems", **Journal of Optimization Theory and Applications**, 116, 3, 559-590.
- [154] S. Ceria, R. Stubbs, (2006), "Incorporating estimation errors into portfolio selection. Robust portfolio selection", **Journal of Asset Management**.

واژه‌نامه انگلیسی به فارسی

Absolute Divation	انحراف مطلق
Activation Function	تابع تحریک
Adaptive	تطبیق پذیر
Anticipative	قابل پیش بینی
Asset	دارایی
Bias	اریبی
Benchmark	محک
Capital Gain	بازده ناشی از تغییر قیمت
Chance Constrained	قیود شانس
Combination Function	تابع ترکیب
Complementarity	مکمل
Convex	محدب
Convex Hull	پوسته محدب
Corrolation	همبستگی
Downside Risk	ریسک نامطلوب
Dreivatives	مشتقات
Dynamic	پویا
Efficent Frontier	مرز کارایی
Equilibrium Point	نقطه تعادل
Expected Return	بازده مورد انتظار
Fractional Programming	برنامه‌ریزی کسری
Globally Asymptotically Stable	پایداری مجانبی سراسری
Global Convergence	همگرای سراسری
Global Minimum	کمینه سراسری
Goal Programming	برنامه‌ریزی آرمانی
Invariant Set	مجموعه پایدار
Local Minimum	کمینه موضعی
Neuron	سلول عصبی
Nominal Values	مقادیر اسمی
Polytopic	چند سقفی

Portfolio	سبد سرمایه
Return	بازده
Regret Level	سطح پشیمانی
Realized Return	بازده تحقق یافته
Risk-Averse	ریسک گریزی
Robust	مقاوم (استوار)
Robust Conterpart	همتای مقاوم
Robustness	مقاوم بودن
Second Order Cone	مخروط مرتبه دوم
Semi Definite	نیمه معین
Semi Variance	نیم واریانس
Sensitivity Analysis	تحلیل حساسیت
Scenario	سناریو (طرح)
Short-Selling	فروش استقراضی
Smooth	هموار
Stochastic	تصادفی (احتمالی)
Strictly Monotonic	اکیدا یکنوا
Tolerable	قابل تحمل
Tracking Error	خطای پیگردی (ردیابی)
Transfor Function	تابع انتقال
Value-at-Risk	ارزش در معرض ریسک
Uncertainty	عدم قطعیت
Utility	مطلوبیت
yeild	بازده ناشی از دریافت سود

Abstract

This involves enlarging the size of the optimization problems that exist in practice. The necessary conditions of efficiency in the use of techniques that enable high-speed, very large problems solved with acceptable quality can be felt more than.

Recently methods of optimization based on artificial intelligence approaches have been developed remarkable success in solving optimization problems efficiently acquired. Methods such as Genetic Algorithms, Tabu Search, refrigeration simulation and neural networks, their ability to solve large problems have good action. Special rates available on the possible application of neural networks in a wide range of research has provided. It points to the possibility of learning and performance improvement based on the input data points. It also allows parallel computations in a neural network is another advantage of the parallel hardware, enabling very large problems by this approach possible.

In this thesis, we tried two different models of recursive neural network is presented to solve optimization problems in the traces. Analysis of uniqueness, stability and convergence of global solutions are examined and the performance of the proposed method using several examples of stochastic programming problems, Fractional programming, optimization and robust portfolio optimization is shown.

Finally, we provide conclusions and recommendations for future work.

Keywords: *Neural networks, stability, the mean-variance, portfolio selection problem, second order cone programming, stochastic programming, Fractional programming, Robust optimization*



Shahrood University of Technology

Faculty of Mathematical

Department of Mathematics

M.S Thesis

**An application of neural networks models for
solving portfolio optimization problem**

By:

Narges Tahmasbi

Supervisor:

P.hD. Alireza Nazemi

Date

Sep 2012