

رویکرد یادگیری ماشین برای تقریب مارتینگل بازساز در مسائل مالی با ابعاد بالا

سعید سقایی طوقچی^{ib} بهاره اختری^{ib*}

چکیده. مسئله مدیریت ریسک و قیمت گذاری سبدهای مالی، از مسائل مهم در ریاضیات مالی به شمار می آید. روش های رایج به محاسبه مستقیم مقدار مورد انتظار شرطی در هر گام زمانی (به عنوان مثال با استفاده از روش های مبتنی بر مونت کارلو) می پردازند. چالش اصلی در این رویکرد، زمانی است که بعد مسئله افزایش می یابد و روش های شبیه سازی از نوع مونت کارلو با هزینه های محاسباتی سنگین مواجه می شوند. در این مقاله با استفاده از مفهوم مارتینگل بازساز که فرآیند ارزش سبد را از طریق اطلاعات نهایی بیان می کند، به جای محاسبه مستقیم مقدار مورد انتظار شرطی در هر گام زمانی، ابتدا تابع ارزش نهایی سبد تقریب زده می شود و سپس با تکیه بر ویژگی مارتینگل، مقادیر میانی محاسبه می شوند. تقریب مذکور با استفاده از روش یادگیری نظارت شده و بر پایه داده های شبیه سازی شده انجام می شود. نتایج روش پیشنهادی در مدل های گاوسی، با مقایسه نتایج آن با روش های مونت کارلوی تودرتو و مونت کارلوی کمترین مربعات از نوع رگرسیون-اکنون نشان می دهند که از نظر دقت و کارایی محاسباتی عملکرد بهتری دارد.

۱. مقدمه

در سال های اخیر، مدیریت ریسک به یکی از ارکان اصلی فعالیت مؤسسات مالی، بانک ها و شرکت های بیمه تبدیل شده است. ارزش سبدهای مالی و تعهدات بیمه ای تحت تأثیر عوامل ریسکی متعددی از جمله ریسک بازار، ریسک اعتباری، ریسک نرخ بهره و ریسک طول عمر قرار دارد. از این رو، ارزیابی دقیق توزیع زیان و محاسبه سبدهای ریسک مستلزم توسعه روش های عددی کارآمد است. با افزایش پیچیدگی محصولات مالی و رشد ابعاد مدل های مورد استفاده در عمل، نیاز به روش هایی که بتوانند این کمیت ها را با دقت و هزینه محاسباتی مناسب برآورد کنند، بیش از پیش احساس می شود.

یکی از چالش های اصلی در مدیریت ریسک و قیمت گذاری محصولات مالی، محاسبه مقدار مورد انتظار شرطی است که در بسیاری از مسائل مالی به عنوان مؤلفه اصلی ارزش گذاری و اندازه گیری ریسک ظاهر می شود. یکی از متداول ترین رویکردها برای این منظور، روش مونت کارلوی تودرتو است که در آن برای هر سناریوی بیرونی، مجموعه ای از شبیه سازی های درونی تولید می شود. با وجود دقت مناسب، هزینه محاسباتی بالای این روش، به ویژه در مسائل با ابعاد بالا، استفاده از آن را در بسیاری از کاربردهای عملی محدود می کند [۱۰، ۲۰].

رده بندی موضوعی ریاضیات ۲۰۲۰: 97M30, 91G60, 91G20

عبارات و کلمات کلیدی: یادگیری ماشین، مارتینگل بازساز، مونت کارلوی تودرتو، رگرسیون-اکنون، رگرسیون-آتی.

نوع مقاله: مروری

* نویسنده مسئول

دبیر تخصصی رابط: —

تاریخ دریافت: ۱۴۰X/XX/XX تاریخ پذیرش: ۱۴۰X/XX/XX تاریخ انتشار آنلاین: ۱۴۰X/XX/XX

ارجاع به مقاله: نام نویسندگان، عنوان مقاله، ریاضی و جامعه - (۱۴۰X)، شماره -، XX-XX. <https://dx.doi.org/> . /msci.xxxx

به منظور کاهش این هزینه، روش‌های مونت‌کارلوی کمترین مربعات توسعه یافته‌اند که مقدار مورد انتظار شرطی را از طریق رگرسیون تقریب می‌زنند. ایده‌ی اصلی این روش‌ها نخست در ارزش‌گذاری اختیار معامله‌های آمریکایی [۲۲] و در مسائل بیمه‌ی عمر به‌کار گرفته شد [۱، ۲]. در چارچوب روش‌های مونت‌کارلوی کمترین مربعات، دو رویکرد اصلی رگرسیون-اکنون و رگرسیون-آتی شکل گرفتند. در رگرسیون-اکنون، مقدار مورد انتظار شرطی مستقیماً بر حسب متغیرهای حالت در زمان فعلی تقریب زده می‌شود در حالی که در رگرسیون-آتی ابتدا پرداخت یا ارزش نهایی برآورد شده و سپس مقدار مورد انتظار شرطی به‌صورت تحلیلی محاسبه می‌شود. مطالعه‌ی [۹] نشان داد که رگرسیون-آتی، در شرایط یکسان، می‌تواند تقریب‌های دقیق‌تری نسبت به رگرسیون-اکنون تولید کند هرچند این مزیت معمولاً با افزایش پیچیدگی محاسباتی و نیاز به ارزیابی مقدار مورد انتظار شرطی توابع تقریب همراه است.

در سال‌های اخیر، پژوهشگران تلاش کرده‌اند با بهره‌گیری از ابزارهای یادگیری ماشین و روش‌های کاهش ابعاد، محدودیت‌های روش‌های مبتنی بر رگرسیون را برطرف کنند. در این راستا، مدل‌های مبتنی بر توابع پایه‌ای قطعه‌ای [۲۴]، توابع اسپلاین [۶]، شبکه‌های عصبی [۳]، انتخاب داده‌محور توابع پایه [۱۱]، روش‌های هسته‌ای [۱۷، ۴] و رگرسیون فرایند گاوسی [۱۴، ۲۶، ۲۱، ۱۳] پیشنهاد شده‌اند.

با وجود پیشرفت‌های حاصل‌شده، استفاده از روش‌های مبتنی بر رگرسیون در مسائل با ابعاد بالا همچنان با چالش‌های جدی روبه‌رو است. در بسیاری از کاربردهای عملی، افزایش تعداد عوامل ریسکی موجب رشد سریع تعداد توابع پایه مورد نیاز برای تقریب می‌شود. علاوه بر این، بسیاری از روش‌های موجود فاقد ساختاری هستند که امکان محاسبه فرم تحلیلی مقدار مورد انتظار شرطی را فراهم کند. در نتیجه، توسعه روش‌هایی که ضمن حفظ دقت، قابلیت استفاده در ابعاد بالا را داشته باشند، همچنان یک مسئله مهم محسوب می‌شود. در این پژوهش، برای مواجهه با این چالش‌ها، از چارچوب مارتینگل بازساز استفاده می‌کنیم. ایده اصلی این رویکرد، تقریب فرایند ارزش سید مالی با استفاده از ترکیب خطی توابع پایه آموزش‌دیده از داده‌ها می‌باشد. در این چارچوب، از ابزارهای یادگیری ماشین برای آموزش توابع تقریب استفاده شده و به‌منظور مقابله با مسئله ابعاد بالا، کاهش ابعاد خطی مبتنی بر داده‌ها به‌کار گرفته می‌شود. این ساختار امکان محاسبه مقدار مورد انتظار شرطی را به فرم تحلیلی برای کلاس مهمی از توابع پایه فراهم می‌کند و در نتیجه کارایی رویکرد رگرسیون-آتی را به‌طور قابل توجهی افزایش می‌دهد.

رویکرد مارتینگل بازساز می‌تواند پلی میان روش‌های کلاسیک مونت‌کارلوی کمترین مربعات و رویکردهای نوین یادگیری ماشین ایجاد کند. این چارچوب، ضمن حفظ تفسیرپذیری ریاضی مدل، امکان بهره‌گیری از تقریب‌های داده‌محور را برای محاسبه کمیتهای ریسکی در مسائل با ابعاد بالا فراهم می‌آورد. به‌ویژه نشان خواهیم داد که ترکیب کاهش ابعاد مبتنی بر داده‌ها با توابع پایه مناسب، نه تنها مشکل افزایش تعداد پارامترها را تا حد زیادی کنترل می‌کند، بلکه امکان محاسبه کارآمد مقدار مورد انتظار شرطی را نیز حفظ می‌نماید. برای ارزیابی رویکرد مزبور، عملکرد آن در دو مسئله‌ی شاخص شامل ارزش‌گذاری یک اختیار خرید اروپایی و محاسبه تعهدات یک محصول بیمه‌ی عمر وابسته به مسیر، بررسی شده و نتایج حاصل با روش‌های متداولی نظیر مونت‌کارلوی تودرتو و مونت‌کارلوی کمترین مربعات مبتنی بر رگرسیون-اکنون مقایسه می‌شود.

در ادامه، ابتدا صورت‌بندی ریاضی مسئله‌ی مارتینگل بازساز و چارچوب یادگیری ماشین ارائه می‌شود. سپس ساختار توابع تقریب، روش‌های عددی و معیارهای ارزیابی تشریح می‌گردند. در بخش مطالعات عددی، رفتار روش پیشنهادی در سناریوهای مالی و بیمه‌ای با ابعاد بالا مورد بررسی قرار گرفته و نقاط قوت و محدودیت‌های آن در مقایسه با رویکردهای مرسوم تحلیل می‌شود. این مسیر، از مبانی نظری تا نتایج محاسباتی، تصویری جامع از ظرفیت‌های رویکرد مارتینگل بازساز برای مدیریت ریسک و ارزش‌گذاری در مسائل پیچیده‌ی دنیای واقعی ارائه می‌کند. جزئیات نظری بیشتر و اثبات‌های مربوط به قضایا در [۷] ارائه شده است.

۲. مسئله‌ی مارتینگل بازساز

بردار رندم $X = (X_1, \dots, X_T)$ ، یک مولد سناریو اقتصادی با افق زمانی متناهی T است که در آن زمان در واحد سال اندازه‌گیری می‌شود. بردار X که در آن $X_t \in \mathbb{R}^d$ برای $t \in [0, T]$ ، یک فرآیند محرک تصادفی^۱ است و مولفه‌های X_t از یکدیگر مستقل هستند. پالایه‌ی $\mathcal{F}_t := \sigma(X_1, \dots, X_t)$ ، جریان اطلاعات را نشان می‌دهد. فضای مسیر را به صورت $\Omega = \mathbb{R}^{dT}$ در نظر می‌گیریم و توزیع این فرآیند را با نماد $\mathbb{Q}$ نمایش می‌دهیم که آن را اندازه‌ی ریسک-خنثی متناظر می‌نامیم. در این چارچوب تمامی مقادیر مالی و جریان‌های نقدی^۲ به صورت تنزیل شده نسبت به یک دارایی مبنا^۳ بیان می‌شوند مگر آنکه خلاف آن تصریح شود. سبده‌ی از دارایی‌ها^۴ و بدهی‌ها^۵ که ارزش فعلی آن مجموع جریان‌های نقدی است، در پایان افق زمانی دارای ارزش نهایی ذیل است:

$$(۱) \quad f(X) = \sum_{t=1}^T \zeta_t,$$

که در آن $\zeta_t = \zeta_t(X_1, \dots, X_t)$ جریان نقدی تا زمان t بوده که $\mathcal{F}_t$ -اندازه پذیر است. فرض می‌کنیم ζ_t و بنابراین تابع f ، به صورت برون‌زا^۶ و از پیش تعیین شده، عضو $L^2_{\mathbb{Q}}$ هستند. هدف، یافتن فرآیند ارزش سبد $V := (V_t)_{0 \leq t \leq T}$ است که به صورت ذیل تعریف می‌شود:

$$(۲) \quad V_t := \mathbb{E}^{\mathbb{Q}}[f(X)] = \underbrace{\sum_{s=1}^t \zeta_s}_{\text{جریان نقدی انباشته شده تا زمان } t} + \underbrace{\mathbb{E}^{\mathbb{Q}} \left[\sum_{s=t+1}^T \zeta_s \right]}_{\text{ارزش لحظه‌ای در زمان } t},$$

که در آن $\mathbb{E}^{\mathbb{Q}}[\cdot | \mathcal{F}_t] := \mathbb{E}^{\mathbb{Q}}_t[\cdot]$ مقدار مورد انتظار شرطی نسبت به اطلاعات تا زمان t است. با استفاده از ویژگی شرطی مکرر داریم:

$$\mathbb{E}^{\mathbb{Q}}[V_{t+1} | \mathcal{F}_t] = \mathbb{E}^{\mathbb{Q}}[\mathbb{E}^{\mathbb{Q}}[f(X) | \mathcal{F}_{t+1}] | \mathcal{F}_t] = \mathbb{E}^{\mathbb{Q}}[f(X) | \mathcal{F}_t] = V_t,$$

بنابراین V یک مارتینگل در فضای $L^2_{\mathbb{Q}}$ است که مقدار نهایی $f(X)$ را بازسازی می‌کند. در اغلب کاربردهای واقعی، مقدار V_t به صورت تحلیلی قابل محاسبه نیست و باید از طریق شبیه‌سازی تابع $f(X)$ تقریب زده شود. با توجه به اینکه برای هر مسیر شبیه‌سازی شده، نیاز به ارزیابی تابع داریم، این مرحله می‌تواند از نظر محاسباتی پرهزینه باشد.

یکی از کاربردهای مهم محاسبه‌ی فرآیند V در بحث مدیریت ریسک است. چارچوب‌های نظارتی توانگری مالی^۷ در صنعت بیمه و بانکداری، محاسبه‌ی کفایت سرمایه^۸ را بر مبنای اندازه‌گیری ریسک تغییرات ارزش سبد $\Delta V_t = V_t - V_{t-1}$ الزامی می‌دانند. در مقررات بیمه، ریسک معمولاً برای افق زمانی یک‌ساله سنجیده می‌شود و سرمایه اقتصادی مورد نیاز به صورت $\rho[-\Delta V_1]$ محاسبه می‌شود که در آن ρ یکی از معیارهای ریسک است. با توجه به اینکه هدف مدیریت ریسک، اندازه‌گیری زیان است، عبارت $-\Delta V_1$ به عنوان زیان بالقوه در نظر گرفته می‌شود. رایج‌ترین معیارهای ریسک شامل ارزش در معرض خطر^۹ و زیان مورد انتظار^{۱۰} است.

با وجود چارچوب نظری منسجم مارتینگل بازساز، در عمل محاسبه‌ی مقادیر مورد انتظار شرطی و ارزش‌گذاری‌های متناظر غالباً به صورت تحلیلی امکان‌پذیر نیست. از این رو، برای تقریب این مقادیر و برآورد معیارهای ریسک، نیاز به روش‌های عددی از جمله مونت‌کارلوی تودرتو^{۱۱} وجود دارد. در بخش بعد، به بررسی این روش و چالش‌های محاسباتی آن می‌پردازیم.

۳. مونت‌کارلو تودرتو و چالش‌های محاسباتی

در روش شبیه‌سازی مونت‌کارلوی تودرتو، هدف تخمین مقادیر مورد انتظار شرطی است که شامل دو مرحله شبیه‌سازی است:

¹Stochastic Driver Process ²Cash Account ³Nueraire ⁴Assets ⁵Liabilities ⁶Exogenously ⁷Solvency Regulatory Frameworks ⁸Capital Adequacy ⁹Value at Risk ¹⁰Expected Shortfall ¹¹Nested Monte Carlo

۱.۳. مرحله‌ی بیرونی. در مرحله‌ی بیرونی، ابتدا مجموعه‌ای از n_0 مسیر مستقل و هم‌توزیع از فرآیند تصادفی محرک X تا افق زمانی t تولید می‌شود. هر مسیر بیرونی را می‌توان به صورت ذیل نمایش داد:

$$X_{1:t}^{(i)} = (X_1^{(i)}, X_2^{(i)}, \dots, X_t^{(i)}), \quad i = 1, \dots, n_0,$$

۲.۳. مرحله‌ی درونی. برای هر مسیر بیرونی $X_{1:t}^{(i)}$ ارزش سید مالی در زمان t باید محاسبه شود. بنابراین برای هر مسیر بیرونی، یک مرحله‌ی شبیه‌سازی درونی اجرا می‌شود. این مرحله شامل تولید n_1 مسیر مستقل از ادامه‌ی فرآیند تصادفی پس از زمان t و محاسبه‌ی میانگین نمونه‌ای خروجی تابع ارزش سید مالی بر این مسیرهاست. به عبارت دیگر، اگر f تابع ارزش نهایی باشد آنگاه

$$V_t^{(i)} := \frac{1}{n_1} \sum_{j=1}^{n_1} f(X_{t+1:T}^{(i,j)} | X_{1:t}^{(i)}).$$

پس از تشکیل توزیع تجربی^{۱۲} بر اساس مقادیر $\{V_t^{(i)}\}_{i=1}^{n_0}$ ، زمانی که n_1 به اندازه‌ی کافی بزرگ باشد، افزایش n_0 موجب همگرایی توزیع تجربی به توزیع موردنظر می‌شود. در این حالت، بر اساس قضیه گلیونکو-کانتلی^{۱۳}، توزیع تجربی با افزایش تعداد نمونه‌های بیرونی به‌طور یکنواخت به تابع توزیع متناظر همگرا می‌شود. از این رو با انتخاب مناسب n_0 و n_1 ، می‌توان از توزیع تجربی به‌عنوان تقریب توزیع موردنظر استفاده کرد.

در این روش، دو منبع اصلی خطای شبیه‌سازی وجود دارد. نخست، خطای ناشی از شبیه‌سازی بیرونی است که به دلیل محدود بودن n_0 ایجاد می‌شود. با افزایش n_0 ، تعداد نمونه‌های مورد استفاده برای تشکیل توزیع تجربی افزایش یافته و طبق قانون اعداد بزرگ^{۱۴}، تقریب توزیع موردنظر دقیق‌تر می‌شود. دوم، برای هر مسیر بیرونی، مقدار $V_t^{(i)}$ از میانگین‌گیری روی n_1 شبیه‌سازی درونی به دست می‌آید. بنابراین این مقدار تنها یک تقریب از امید شرطی متناظر است. افزایش n_1 خطای شبیه‌سازی درونی را کاهش می‌دهد. در نتیجه، دقت نهایی روش مونت‌کارلوی تودرتو به هر دو تعداد n_0 و n_1 وابسته است.

در نهایت، معیار ریسک موردنظر با جایگزینی توزیع تجربی در تعریف نظری آن معیار تخمین زده می‌شود. اگرچه این روش از نظر پیاده‌سازی ساده است اما به‌کارگیری آن در چندین افق زمانی منجر به رشد نمایی تعداد کل شبیه‌سازی‌ها می‌شود. این ویژگی باعث می‌شود این روش تنها در موارد خاص یا برای اعتبارسنجی نتایج به‌کار گرفته شود.

۴. یادگیری ماشین

در سال‌های اخیر، پیشرفت‌های چشمگیر در توان محاسباتی و دسترسی به حجم عظیمی از داده‌ها، موجب شده است که روش‌های یادگیری ماشین به یکی از ابزارهای اصلی در حل مسائل پیچیده علمی، مهندسی و مالی تبدیل شوند. هدف اصلی یادگیری ماشین، استخراج الگوها و روابط موجود در داده‌ها و استفاده از این روابط برای پیش‌بینی، تقریب یا تصمیم‌گیری است. برخلاف روش‌های کلاسیک که در آن‌ها مدل ریاضی از پیش تعیین می‌شود، در یادگیری ماشین ساختار تابع تقریب‌زننده از داده‌ها آموخته می‌شود و مدل قادر است روابط پیچیده و غیرخطی میان متغیرها را شناسایی نماید.

از دیدگاه ریاضی، مسئله یادگیری ماشین را می‌توان به صورت تقریب یک نگاشت ناشناخته در نظر گرفت. فرض کنید $X \in \mathcal{X}$ یک متغیر ورودی و $Y \in \mathcal{Y}$ یک متغیر خروجی باشد. هدف، یافتن تابعی به صورت $f: \mathcal{X} \rightarrow \mathcal{Y}$ است که بتواند رابطه میان ورودی و خروجی را تا حد امکان بازسازی نماید. در عمل، این تابع ناشناخته بوده و تنها مجموعه‌ای از مشاهدات $\{(X^{(i)}, Y^{(i)})\}_{i=1}^N$ در اختیار قرار دارد. بر اساس این داده‌ها، یک خانواده از توابع انتخاب شده و پارامترهای آن به گونه‌ای تعیین می‌شوند که خروجی مدل با داده‌های مشاهده‌شده بیشترین سازگاری را داشته باشد.

¹²Empirical Distribution ¹³Glivenko–Cantelli Theorem ¹⁴Law of Large Numbers

یکی از رایج‌ترین شاخه‌های یادگیری ماشین، یادگیری نظارت‌شده^{۱۵} است. در این چارچوب، برای هر مشاهده ورودی، مقدار متناظر خروجی نیز در دسترس است و هدف یادگیری رابطه میان آن‌ها می‌باشد. در حوزه‌ی ریاضیات مالی، پیش‌بینی قیمت دارایی‌ها و ارزش‌گذاری ابزارهای مشتقه در این دسته قرار می‌گیرند.

۱.۴. صورت‌بندی رگرسیون. بسیاری از مسائل یادگیری ماشین را می‌توان به صورت یک مسئله رگرسیونی در فضای L^2 تفسیر کرد. فضای L^2 مجموعه‌ای از متغیرهای تصادفی است که دارای گشتاور مرتبه دوم متناهی باشند، یعنی $L^2 = \{X : \mathbb{E}[X^2] < \infty\}$. این فضا یک فضای هیلبرت^{۱۶} است و به کمک ضرب داخلی $\langle X, Y \rangle = \mathbb{E}[XY]$ ساختار هندسی مناسبی برای تعریف فاصله، تصویر عمودی و تقریب بهینه فراهم می‌کند. در این چارچوب، هدف یافتن تابعی است که بتواند متغیر هدف را بر اساس اطلاعات موجود در متغیرهای ورودی با کمترین خطای ممکن تقریب بزند. اگر X متغیر ورودی و Y متغیر هدف باشد، مسئله رگرسیون به صورت کمینه‌سازی میانگین مربعات خطا^{۱۷} تعریف می‌شود:

$$\hat{f} = \arg \min_{f \in L^2} \mathbb{E}[(Y - f(X))^2].$$

با استفاده از ویژگی‌های فضاهای هیلبرت می‌توان نشان داد که جواب این مسئله برابر است با $\hat{f}(X) = \mathbb{E}[Y | X]$. بنابراین بهترین تقریب در معنای L^2 ، همان امید شرطی متغیر هدف نسبت به اطلاعات موجود در متغیرهای ورودی است. از دیدگاه هندسی نیز امید شرطی را می‌توان تصویر عمودی متغیر تصادفی Y بر زیرفضای تولیدشده توسط اطلاعات X در فضای L^2 در نظر گرفت.

برای محاسبه کمیتهی به فرم $V_t = \mathbb{E}^Q[f(X_1, \dots, X_T) | \mathcal{F}_t]$ که در آن f یک پرداخت یا کمیت نهایی، $\mathcal{F}_t = \sigma(X_1, \dots, X_t)$ و X_t متغیرهای حالت در زمان t هستند، دو رویکرد اصلی وجود دارد:

در رگرسیون-اکنون، مقدار مورد انتظار شرطی مستقیماً بر حسب متغیرهای حالت در زمان فعلی تقریب زده می‌شود:

$$V_t(X_1, \dots, X_t) \approx \sum_{k=1}^K \beta_k \phi_k(X_1, \dots, X_t),$$

که در آن ϕ_k توابع پایه و β_k ضرایب مجهول هستند.

در مقابل، در رگرسیون-آتی ابتدا کمیت نهایی تقریب زده می‌شود:

$$f(X_1, \dots, X_T) \approx \sum_{k=1}^K \beta_k \phi_k(X_1, \dots, X_T),$$

و سپس با استفاده از خطی بودن امید شرطی خواهیم داشت:

$$V_t \approx \sum_{k=1}^K \beta_k \mathbb{E}^Q[\phi_k(X_1, \dots, X_T) | \mathcal{F}_t].$$

بنابراین در رگرسیون-اکنون، تقریب مستقیماً بر روی امید شرطی انجام می‌شود در حالی که در رگرسیون-آتی ابتدا کمیت نهایی تقریب زده شده و سپس ساختار شرطی مدل مورد استفاده قرار می‌گیرد. در این مقاله، محاسبه مارتینگل بازساز را می‌توان نمونه‌ای از رویکرد رگرسیون-آتی در نظر گرفت. زیرا ابتدا کمیت‌های نهایی تقریب زده شده و سپس ساختار شرطی مارتینگل برای بازسازی مقادیر زمانی مورد استفاده قرار می‌گیرد.

¹⁵Supervised Learning ¹⁶Hilbert Space ¹⁷Mean Squared Error

۲.۴. شبکه‌های عصبی مصنوعی^{۱۸}. شبکه‌های عصبی مصنوعی یکی از قدرتمندترین ابزارهای یادگیری نظارت‌شده برای تقریب توابع پیچیده و غیرخطی هستند. ایده اصلی این مدل‌ها از ساختار شبکه‌های عصبی زیستی الهام گرفته شده است، اگرچه مدل‌های ریاضی مورد استفاده در عمل بسیار ساده‌تر از سامانه عصبی واقعی هستند.

یک شبکه عصبی پیش‌خور^{۱۹} را می‌توان به عنوان یک تابع پارامتری $f_\theta: \mathbb{R}^d \rightarrow \mathbb{R}^m$ در نظر گرفت که در آن θ مجموعه پارامترهای مدل شامل وزن‌ها و اربیبی‌ها^{۲۰} است. در ساده‌ترین حالت، اگر $x = (x_1, \dots, x_d)^\top$ بردار ورودی باشد، خروجی اولین لایه پنهان^{۲۱} به صورت ذیل تعریف می‌شود:

$$h^{(1)} = \phi \left(W^{(1)}x + b^{(1)} \right)$$

که در آن $W^{(1)}$ ماتریس وزن‌ها، $b^{(1)}$ بردار اربیبی و ϕ تابع فعال‌سازی^{۲۲} است. توابع فعال‌سازی امکان مدل‌سازی روابط غیرخطی را فراهم می‌کنند و نقش اساسی در توانایی تقریب شبکه‌های عصبی دارند.

در شبکه‌های چندلایه، خروجی هر لایه به عنوان ورودی لایه بعدی مورد استفاده قرار می‌گیرد. اگر شبکه شامل L لایه باشد، خروجی نهایی به صورت $f_\theta(x) = h^{(L)}$ به دست می‌آید که در آن

$$h^{(\ell)} = \phi_\ell \left(W^{(\ell)}h^{(\ell-1)} + b^{(\ell)} \right), \quad \ell = 1, \dots, L.$$

وجود چندین لایه پنهان موجب می‌شود که شبکه بتواند ساختارهای پیچیده موجود در داده‌ها شناسایی کند. شبکه‌های دارای تعداد زیادی لایه پنهان را شبکه‌های عصبی عمیق^{۲۳} می‌نامند.

یکی از مهم‌ترین ویژگی‌های نظری شبکه‌های عصبی، توانایی آن‌ها در تقریب توابع غیرخطی است. بر اساس ویژگی تقریب‌پذیری سراسری، یک شبکه عصبی پیش‌خور با حداقل یک لایه پنهان و تعداد کافی نورون قادر است هر تابع پیوسته تعریف‌شده روی یک مجموعه فشرده را با دقت دلخواه تقریب زند. اهمیت این ویژگی در آن است که بسیاری از نگاشت‌های مورد نیاز در مسائل مالی، ساختار تحلیلی شناخته‌شده‌ای ندارند و معمولاً تنها از طریق شبیه‌سازی قابل محاسبه هستند. در چنین شرایطی، شبکه‌های عصبی می‌توانند به عنوان ابزارهای انعطاف‌پذیر برای تقریب این توابع مورد استفاده قرار گیرند.

۵. رویکرد یادگیری ماشین در مسئله‌ی مارتینگل بازساز

همان‌گونه که در بخش‌های پیشین مشاهده شد، چالش اصلی در محاسبه مارتینگل بازساز، ناشی از نیاز به برآورد امیدهای شرطی است که محاسبه آن‌ها غالباً به اجرای روش‌های مونت‌کارلوی تودرتو وابسته است. این موضوع به ویژه در مسائل با ابعاد بالا، هزینه محاسباتی قابل توجهی ایجاد می‌کند و استفاده از این روش‌ها را در کاربردهای عملی با محدودیت مواجه می‌سازد. از این رو، یافتن تقریب‌هایی کارآمد، گامی اساسی در جهت توسعه روش‌های محاسباتی به شمار می‌آید. برای دستیابی به این هدف، چارچوبی مبتنی بر یادگیری ماشین پیشنهاد می‌شود که ضمن حفظ دقت مناسب، می‌تواند هزینه محاسباتی مورد نیاز برای محاسبه مارتینگل بازساز را به طور قابل توجهی کاهش دهد.

در این چارچوب، مسئله به صورت یک مسئله‌ی رگرسیونی در فضای $L^2_{\mathbb{Q}}$ صورت‌بندی می‌شود. به طور مشخص، هدف یادگیری، تقریب نگاشت $X \mapsto \mathbb{E}_{\mathbb{Q}}[f(X_T) | X_t = X]$ است. ورودی مدل شامل نمونه‌هایی از متغیر حالت X در زمان‌های مختلف بوده و خروجی متناظر، مقادیر شبیه‌سازی‌شده‌ی تابع هدف است.

در این رویکرد، تابع ارزش نهایی f به صورت مستقیم با تصویر کردن آن بر روی یک زیرفضا با بعد متناهی در فضای $L^2_{\mathbb{Q}}$ تقریب زده می‌شود. این زیرفضا توسط مجموعه‌ای از توابع پایه تولید می‌گردد که همان نگاشت ویژگی^{۲۴} را تشکیل می‌دهند. عملکرد این روش به انتخاب مناسب توابع پایه وابسته است. فرض می‌کنیم که مقادیر مورد انتظار شرطی توابع پایه، به صورت تحلیلی وجود دارد. این تقریب

¹⁸ Artificial Neural Networks ¹⁹Feedforward Neural Network ²⁰Biases ²¹Hidden Layer ²²Activation Function ²³Deep Neural Networks ²⁴Feature Map

از طریق آموزش بر مبنای مجموعه‌ای متناهی از مشاهدات X حاصل می‌شود. همچنین مبتنی بر این مشاهدات، یک اندازه‌ی تجربی تعریف می‌شود که برآوردی از اندازه‌ی نظری مدل $\mathbb{Q}$ است که با افزایش تعداد مشاهدات، به اندازه‌ی واقعی مدل همگرا خواهد شد.

۱.۵. تقریب با بعد متناهی^{۲۵}. در این روش مجموعه‌ای از توابع پایه به صورت $\phi_{\theta,i}: \mathbb{R}^{dT} \rightarrow \mathbb{R}$ برای $i = 1, \dots, m$ و $\theta \in \Theta$ در فضای $L^2_{\mathbb{Q}}$ تعریف می‌شود. این توابع، نگاشت ویژگی $\phi_{\theta} = (\phi_{\theta,1}, \dots, \phi_{\theta,m})^T$ را تشکیل می‌دهند که ساختار آن وابسته به پارامتر θ است. پارامتر θ معمولاً شامل وزن‌ها و اریبی‌های قابل یادگیری است که در فرآیند آموزش، برای دستیابی به بهترین تقریب انتخاب می‌شوند.

برای هر $\theta \in \Theta$ ، تصویر تابع f در فضای $L^2_{\mathbb{Q}}$ در زیرفضای تولیدشده توسط $\{\phi_{\theta,1}, \dots, \phi_{\theta,m}\}$ به صورت $f_{\theta} = \sum_{i=1}^m \phi_{\theta,i} \beta_{\theta,i} = \phi_{\theta}^T \beta_{\theta}$ قرار می‌گیرد که در آن $\beta_{\theta} \in \mathbb{R}^m$ بردار ضرایب بهینه است و از طریق حل مسئله‌ی کمینه‌سازی ذیل به دست می‌آید:

$$(۳) \quad \min_{\beta \in \mathbb{R}^m} \|f - \phi_{\theta}^T \beta\|_{L^2_{\mathbb{Q}}}.$$

فرض می‌کنیم که مقدار مورد انتظار شرطی $\mathbb{E}_t^{\mathbb{Q}}[\phi_{\theta}(X)]$ به صورت تحلیلی قابل محاسبه است. اکنون تابع $G_{\theta,t}$ را به صورت ذیل تعریف می‌کنیم:

$$(۴) \quad \begin{aligned} G_{\theta,t}: \mathbb{R}^{d \times t} &\rightarrow \mathbb{R}^m, \\ G_{\theta,t}(x_1, \dots, x_t) &= \mathbb{E}^{\mathbb{Q}}[\phi_{\theta}(x_1, \dots, x_t, X_{t+1}, \dots, X_T)] \\ &= \mathbb{E}_t^{\mathbb{Q}}[\phi_{\theta}(X)]. \end{aligned}$$

در نتیجه، فرآیند ارزش سبد تقریب‌زده شده به صورت ذیل است:

$$(۵) \quad \begin{aligned} V_{\theta,t} &= \mathbb{E}_t^{\mathbb{Q}}[f_{\theta}(X)] \\ &= \mathbb{E}_t^{\mathbb{Q}}[\phi_{\theta}^T \beta_{\theta}] \\ &= G_{\theta,t}(X_1, \dots, X_t)^T \beta_{\theta} \approx V_t. \end{aligned}$$

در ادامه، فرض می‌کنیم که اندازه‌ی احتمال واقعی $\mathbb{P}$ نسبت به اندازه‌ی ریسک-خنثی $\mathbb{Q}$ مطلقاً پیوسته باشد یعنی $\mathbb{P} \ll \mathbb{Q}$ و مشتق رادون-نیکودیم^{۲۶} $\frac{d\mathbb{P}}{d\mathbb{Q}}$ در $L^2(\mathbb{Q})$ قرار گیرد. اندازه‌ی $\mathbb{Q}$ برای قیمت‌گذاری و تقریب تابع هدف به کار می‌رود در حالی که اندازه‌ی $\mathbb{P}$ بیانگر توزیع واقعی بازار بوده و برای ارزیابی ریسک و خطای مسیر مورد استفاده قرار می‌گیرد. بنابراین با تغییر اندازه از $\mathbb{Q}$ به $\mathbb{P}$ و نامساوی کوشی-شوارتز^{۲۷} می‌توان خطای مسیر را که در نرم $L^2_{\mathbb{P}}$ سنجیده می‌شود، به خطای تقریب در نرم $L^2_{\mathbb{Q}}$ مرتبط کرد. با توجه به اینکه فرآیندهای ارزش $(V_t)_{t=0}^T$ و $(V_{\theta,t})_{t=0}^T$ تحت اندازه‌ی $\mathbb{Q}$ مارتینگل هستند، فرآیند $M_t := V_t - V_{\theta,t}$ نیز تحت $\mathbb{Q}$ یک مارتینگل است. بنابراین، با استفاده از نابرابری دوب^{۲۸} می‌توان کران بالایی برای خطای پیشینه مسیر به دست آورد:

$$(۶) \quad \left\| \max_{t \leq T} |V_t - V_{\theta,t}| \right\|_{L^2_{\mathbb{P}}} \leq \left\| \frac{d\mathbb{P}}{d\mathbb{Q}} \right\|_{L^2_{\mathbb{Q}}} \left\| \max_{t \leq T} |V_t - V_{\theta,t}| \right\|_{L^2_{\mathbb{Q}}} \leq 2 \left\| \frac{d\mathbb{P}}{d\mathbb{Q}} \right\|_{L^2_{\mathbb{Q}}} \|f - f_{\theta}\|_{L^2_{\mathbb{Q}}}.$$

بنابراین هرچه اختلاف بین تابع هدف واقعی f و تقریب آن f_{θ} کمتر باشد، خطای کل مسیر نیز کاهش می‌یابد. توجه داشته باشید که کمیت $\left\| \frac{d\mathbb{P}}{d\mathbb{Q}} \right\|_{L^2_{\mathbb{Q}}}$ یک ضریب تغییر اندازه‌ی احتمال است که مقدار آن توسط مدل‌ساز و با توجه به ساختار تغییر اندازه‌ی ریسک قابل محاسبه

²⁵Finite-Dimensional Approximation ²⁶Radon-Nikodym ²⁷Cauchy-Schwarz ²⁸Doob's Inequality

است. همچنین عبارت $\|f - f_\theta\|_{L^2_{\mathbb{Q}}}$ در سمت راست نابرابری، همان تابع هدف مسأله یادگیری (۳) است. این کمیت را می توان با استفاده از شبیه سازی مونت کارلو و بر اساس داده های آموزشی که برای آموزش f_θ به کار رفته اند، برآورد کرد.

۲.۵. یادگیری ویژگی^{۲۹}. در این بخش، علاوه بر ضرایب β ، پارامتر θ نیز مستقیماً بر اساس داده ها برآورد می شود. برای این منظور، پارامتر θ به گونه ای انتخاب می شود که خطای تقریب تابع هدف کمینه گردد. در نتیجه، مسئله بهینه سازی (۳) به صورت ذیل بر روزرسانی می شود:

$$(۷) \quad \min_{(\theta, \beta) \in \Theta \times \mathbb{R}^m} \|f - \phi_\theta^\top \beta\|_{L^2_{\mathbb{Q}}}.$$

معادله (۷) یک مسئله بهینه سازی غیرمحدب^{۳۰} است. این در حالی است که برای پارامتر β مسأله همواره محدب باقی می ماند. در ادامه، فرضیات لازم برای وجود جواب ارائه می شود.

لم ۱.۵. فرض کنید

(۱) Θ مجموعه ای فشرده است؛

(۲) $\phi_{\theta, i}: \Theta \rightarrow L^2_{\mathbb{Q}}$ برای همه i ها پیوسته است و

(۳) مجموعه $\{\phi_{\theta, 1}, \dots, \phi_{\theta, m}\}$ در $L^2_{\mathbb{Q}}$ برای همه $\theta \in \Theta$ مستقل خطی است. آنگاه معادله (۷) دارای جواب است.

اکنون نگاهی ویژگی ϕ_θ را به صورت ذیل در نظر می گیریم:

$$(۸) \quad \phi_{\theta, i}(x) := g_i(A^\top x + b), \quad i = 1, \dots, m.$$

که در آن $g_i: \mathbb{R}^p \rightarrow \mathbb{R}$ توابعی معلوم و $\theta = (A, b)$ شامل یک ماتریس وزن $A \in \mathbb{R}^{dT \times p}$ با رتبه ی کامل و یک بردار اربیبی $b \in \mathbb{R}^p$ است. به منظور انعکاس ساختار زمانی بردار محرک تصادفی $X \in \mathbb{R}^{dT}$ ، ماتریس A را به صورت بلوکی نمایش می دهیم:

$$A = \begin{pmatrix} A_1 \\ \vdots \\ A_T \end{pmatrix}, \quad A_t \in \mathbb{R}^{d \times p},$$

که در آن هر بلوک A_t متناظر با اطلاعات زمان t است. بنابراین

$$(۹) \quad \phi_{\theta, i}(X) = g_i\left(\sum_{t=1}^T A_t^\top X_t + b\right).$$

مقادیر مورد انتظار معرفی شده در رابطه (۴) به صورت ذیل نوشته می شوند:

$$G_{\theta, t, i}(x_1, \dots, x_t) = \mathbb{E}^{\mathbb{Q}} \left[\phi_{\theta, i}(x_1, \dots, x_t, X_{t+1}, \dots, X_T) \right]$$

$$(۱۰) \quad = \mathbb{E}^{\mathbb{Q}} \left[g_i \left(\sum_{s=1}^t A_s^\top x_s + \sum_{s=t+1}^T A_s^\top X_s + b \right) \right].$$

در ادامه، سه ساختار مختلف برای نداشت ویژگی در قالب رابطه ی (۸) مورد بررسی قرار خواهد گرفت: پایه چندجمله ای کامل^{۳۱}، نداشت چندجمله ای همراه با کاهش بعد خطی^{۳۲} و شبکه ی عصبی کم عمق^{۳۳}.

²⁹Feature Learning ³⁰Non-Convex Optimization ³¹Full Polynomial Basis ³²Polynomial Feature Map with Linear Dimensionality Reduction ³³Shallow Neural Network

۳.۵. پایه‌ی چندجمله‌ای کامل. در این بخش، حالت ساده‌ای را در نظر می‌گیریم که در آن هیچ‌گونه وزن و اریبی به‌عنوان پارامتر قابل یادگیری وجود ندارد. بنابراین مجموعه‌ی پارامترها به صورت $\Theta = \{(IdT, \circ)\}$ تعریف می‌شوند. بنابراین، مسأله‌ی بهینه‌سازی (۷) معادل با مسأله‌ی بهینه‌سازی (۳) است. نگاشت ویژگی ϕ از پایه‌ای تشکیل می‌شود که متعلق به فضای تمام چندجمله‌ای‌های با درجه‌ی حداکثر δ روی $\mathbb{R}^{dT}$ است و به‌صورت ذیل تعریف می‌شود:

$$(11) \quad \text{Pol}_\delta(\mathbb{R}^{dT}) = \text{span} \{x^\alpha \mid \alpha \in \mathbb{N}_0^{dT}, |\alpha| \leq \delta\},$$

که در آن $\mathbb{N}_0^{dT}$ مجموعه‌ی همه‌ی بردارهای $\alpha = (\alpha_1, \dots, \alpha_{dT})$ با طول dT و مؤلفه‌های نامنفی است. برای هر α ، تک‌عضوی پایه به‌شکل $x^\alpha = x_1^{\alpha_1} x_2^{\alpha_2} \dots x_{dT}^{\alpha_{dT}}$ تعریف می‌شود که در آن $\sum_{i=1}^{dT} \alpha_i \leq \delta$. برای اطمینان از این‌که هر تابع پایه ϕ_i عضو $L^2_{\mathbb{Q}}$ باشد، فرض می‌کنیم نرم بردار محرک تصادفی X در $L^2_{\mathbb{Q}}$ متناهی باشد:

$$(12) \quad \mathbb{E}^{\mathbb{Q}}[\|X\|^{2\delta}] < \infty.$$

در این مرحله کل مجموعه‌ی تک‌عضوی‌های چندجمله‌ای این فضا برای تقریب تابع هدف مورد استفاده قرار می‌گیرد. اگرچه این ساختار به استخراج فرم بسته‌ای برای کمیت $G_{t,i}$ در رابطه (۱۵) منجر می‌شود، اما تعداد جملات حاصل از بسط چندجمله‌ای با افزایش ابعاد مسئله به‌صورت نمایی رشد می‌کند؛ موضوعی که کارایی محاسباتی روش را در مسائل با ابعاد بالا به‌شدت محدود می‌سازد. به‌طور دقیق‌تر، تعداد توابع پایه از رابطه‌ی ذیل پیروی می‌کند:

$$(13) \quad m = \dim \text{Pol}_\delta(\mathbb{R}^{dT}) = \binom{dT + \delta}{dT}.$$

جدول ۱ نشان می‌دهد که بعد m از اندازه‌ی نمونه‌ی آموزشی در کاربردهای عملی فراتر می‌رود.

جدول ۱. بعد فضای $\text{Pol}_\delta(\mathbb{R}^{dT})$ برای $\delta = 3$.

$d = 5$	$d = 3$	T
۳۲۷۶	۸۱۶	۵
۱۳۷۳۷۰۱	۳۰۲۶۲۱	۴۰

فرم دقیق مقدار مورد انتظار شرطی در معادله (۴) وابسته به انتخاب نگاشت ویژگی ϕ و ساختار توزیع $\mathbb{Q} = \bigotimes_{t=1}^T \mathbb{Q}_t$ برای بردار تصادفی X است. در این ساختار، توزیع کل به‌صورت حاصل‌ضرب مستقل توزیع‌های حاشیه‌ای $\mathbb{Q}_t$ تعریف می‌شود و فضای هیلبرت متناظر نیز به‌صورت ضرب تانسوری فضاها‌ی تابعی تک‌دوره‌ای به‌صورت $L^2_{\mathbb{Q}} = \bigotimes_{t=1}^T L^2_{\mathbb{Q}_t}$ بیان می‌گردد. انتخاب مجموعه‌ای متعامد از چندجمله‌ای‌ها به‌عنوان پایه در این فضای تانسوری، ساختار محاسباتی مقدار مورد انتظار شرطی را به‌طور قابل ملاحظه‌ای ساده می‌کند. بر اساس قضیه‌ی ۸.۲۵ از [۲۷]، خانواده‌ای از چندجمله‌ای‌های متعامد $\{g_\alpha \mid \alpha \in \mathbb{N}_0^{dT}\}$ بر روی فضای $\mathbb{R}^{dT}$ نسبت به اندازه‌ی احتمال $\mathbb{Q}$ وجود دارد. بنابراین یک پایه‌ی متعامد $\{g_\alpha\}$ در $\mathbb{R}^{dT}$ نسبت به اندازه‌ی $\mathbb{Q} = \bigotimes_{t=1}^T \mathbb{Q}_t$ وجود دارد که ساختار آن به‌شکل ذیل تعریف می‌شود:

$$(14) \quad g_\alpha(X) = \prod_{t=1}^T h_{t,\alpha_t}(x_t), \quad X = (x_1, \dots, x_T), \quad \alpha = (\alpha_1, \dots, \alpha_T),$$

که در آن $\alpha_t \in \mathbb{N}_0^d$ و $x_t \in \mathbb{R}^d$ برای هر مؤلفه زمانی $t = 1, \dots, T$ یک خانواده‌ی متعامد از چندجمله‌ای‌های چندمتغیره روی $\mathbb{R}^d$ نسبت به اندازه‌ی احتمال $\mathbb{Q}_t$ است. در این ساختار، برای هر زمان t ، $\alpha_t = (\alpha_{t,1}, \dots, \alpha_{t,d})$ یک بردار شاخص با طول d است که نشان می‌دهد در چندجمله‌ای‌های متناظر با زمان t ، هر متغیر چه توانی دارد. به‌طور دقیق‌تر، h_{t,α_t} یک چندجمله‌ای

چندمتغیره روی $\mathbb{R}^d$ است که به شکل $h_{t,\alpha_t}(x_t) = \prod_{j=1}^d p_{t,j}^{(\alpha_{t,j})}(x_{t,j})$ تعریف می‌شود که در آن $p_{t,j}^{(\alpha_{t,j})}$ پایه‌ی متعامد تک‌متغیره درجه $\alpha_{t,j}$ در بعد j و زمان t است. در نتیجه:

$$\deg(g_\alpha) = |\alpha| = \sum_{t=1}^T |\alpha_t|, \quad |\alpha_t| = \sum_{j=1}^d \alpha_{t,j}, \quad \deg(h_{t,\alpha_t}) = |\alpha_t|.$$

در این ساختار هر پایه متعامد g_α با ترکیب دقیق پایه‌های تک‌متغیره متعامد ساخته می‌شود. بنابراین m تابع پایه‌ی متعامد g_{α_i} از فضای $\text{Pol}_\delta(\mathbb{R}^{dT})$ نسبت به اندازه‌ی $\mathbb{Q}$ در نظر گرفته می‌شود. مقادیر مورد انتظار شرطی در رابطه‌ی (۱۰) به واسطه‌ی متعامد بودن $\{h_{t,\alpha_t}\}_{\alpha_t}$ ها به صورت تحلیلی قابل محاسبه هستند. با توجه به اینکه $A = I_{dT}$ ، $b = 0$ و نظر به تعامد $\{h_{t,\alpha_t}\}_{\alpha_t}$ مقدار مورد انتظار شرطی هر تابع پایه ϕ_i نسبت به اطلاعات تا زمان t به صورت ذیل به دست می‌آید:

$$\begin{aligned} G_{t,i}(x_1, \dots, x_t) &= \mathbb{E}^{\mathbb{Q}} \left[g_{\alpha^i} \left(\sum_{s=1}^t A_s^\top x_s + \sum_{s=t+1}^T A_s^\top X_s + b \right) \right] \\ &= \prod_{s=1}^t h_{s,\alpha_s^i}(x_s) \times \prod_{s=t+1}^T \mathbb{E}^{\mathbb{Q}}[h_{s,\alpha_s^i}(X_s)] \\ &= \prod_{s=1}^t h_{s,\alpha_s^i}(x_s) \cdot \prod_{s=t+1}^T \mathbb{1}_{\alpha_s^i=0}, \end{aligned} \quad (15)$$

که در آن به دلیل تعامد توابع پایه نسبت به اندازه‌ی احتمال $\mathbb{Q}$ ، هر تابع پایه‌ی غیر ثابت نسبت به تابع ثابت متعامد است، بنابراین

$$\mathbb{E}^{\mathbb{Q}}[h_{s,\alpha_s^i}(X_s)] = 0, \quad \alpha_s^i \neq 0.$$

به طور مشابه، مقدار مورد انتظار هر تابع پایه نیز به صورت $\mathbb{E}^{\mathbb{Q}}[\phi_i(X)] = \phi_i(0) \cdot \mathbb{1}_{\alpha^i=0}$ محاسبه می‌شود.

مثال ۲.۵. فرض کنید $X \sim \mathcal{N}(0, I_{dT})$ یک متغیر نرمال چندمتغیره باشد. در این حالت، می‌توان پایه‌های کلی g_α را به صورت چندجمله‌ای‌های هرمیت چندمتغیره^{۳۴} تعریف کرد که نسبت به توزیع نرمال استاندارد متعامد هستند. به ویژه، پایه‌ی کلی با شاخص $\alpha \in \mathbb{N}_0^{dT}$ روی $\mathbb{R}^{dT}$ ساخته می‌شود در حالی که پایه‌ی هر مؤلفه $h_{t,\alpha_t} \equiv h_{\alpha_t}$ یک چندجمله‌ای هرمیت چندمتغیره در $\mathbb{R}^d$ همراه با بردار شاخص $\alpha_t \in \mathbb{N}_0^d$ است.

۴.۵. نگاشت ویژگی چندجمله‌ای همراه با کاهش بعد خطی. همان‌طور که در بخش قبل مشاهده کردیم، استفاده از پایه‌ی کامل چندجمله‌ای باعث افزایش تعداد توابع پایه و در نتیجه ابعاد مسئله می‌شود. در این بخش به راهکارهای لازم جهت کاهش اثر این موضوع می‌پردازیم. فرض می‌کنیم $p \leq dT$ و $\{g_1, \dots, g_m\}$ یک پایه از فضای $\text{Pol}_\delta(\mathbb{R}^p)$ است. نگاشت‌های ویژگی از نوع معادله (۸) به صورت ذیل تعریف می‌شوند:

$$\phi_{(A,b),i}(x) = g_i(A^\top x + b), \quad A \in \mathbb{R}^{dT \times p}, \quad b \in \mathbb{R}^p.$$

همچنین فرض می‌کنیم $\phi_{(A,b),i} \in L_2^{\mathbb{Q}}$. نکته‌ی مهم این است که بدون از دست دادن کلیت می‌توانیم ارببی b را صفر در نظر بگیریم. دلیل این کار آن است که فضای چندجمله‌ای‌ها شامل جمله ثابت است و در نتیجه هر انتقال خطی b را می‌توان در ترکیب پایه‌های چندجمله‌ای پوشش داد. علاوه بر این، ماتریس وزن A را طوری در نظر می‌گیریم که روی خمینه استیفل^{۳۵}

³⁴Multivariate Hermite Polynomials ³⁵Stiefel Manifold

$\mathbb{R}^{dT}$ است. $V_p(\mathbb{R}^{dT}) = \{A \in \mathbb{R}^{dT \times p} \mid A^\top A = I_p\}$ قرار گیرد. در این حالت، ستون‌های A چارچوبی متعامد از مرتبه p در فضای $\mathbb{R}^{dT}$ است.

قضیه ۳.۵. برای هر ماتریس $A, \tilde{A} \in \mathbb{R}^{dT \times p}$ با رتبه‌ی کامل و بردار $b \in \mathbb{R}^p$ ، شرایط ذیل معادل‌اند:

(۱) $\tilde{A} \in V_p(\mathbb{R}^{dT})$ و $\text{span}\{\phi_{(\tilde{A}, \circ), 1}, \dots, \phi_{(\tilde{A}, \circ), m}\} = \text{span}\{\phi_{(A, b), 1}, \dots, \phi_{(A, b), m}\}$

(۲) برای ماتریس متعامد $U_{p \times p}$ داریم: $\tilde{A} = A(A^\top A)^{-\frac{1}{2}}U$.

با توجه به قضیه ۳.۵، بدون کاستن از کلیت می‌توان مجموعه‌ی پارامترها را به‌صورت $\Theta = V_p(\mathbb{R}^{dT})$ در نظر گرفت. در واقع، دو ماتریس A و $\tilde{A}$ که تنها به‌وسیله‌ی یک تبدیل متعامد در فضای $\mathbb{R}^p$ به یکدیگر مرتبط‌اند، منجر به یک خانواده‌ی یکسان از توابع پایه‌ای می‌شوند. در این حالت، نگاشت ویژگی به فرم ساده‌ی $\phi_{(A, \circ), i}(x) := \phi_{A, i}(x) = g_i(A^\top x)$ بازنویسی می‌شود. قضیه بعد به بررسی وجود و یکتایی جواب برای مسأله‌ی بهینه‌سازی در چارچوب این ساختار پارامتری می‌پردازد.

قضیه ۴.۵. برای نگاشت ویژگی چندجمله‌ای، یک نقطه بهینه در مجموعه‌ی $\Theta = V_p(\mathbb{R}^{dT})$ برای مسئله بهینه‌سازی (۷) وجود دارد. در این حالت، یکتایی برقرار نیست. به عبارت دیگر زیرفضای بهینه‌ی $\text{span}\{\phi_{A, 1}, \dots, \phi_{A, m}\}$ به‌طور کلی یکتا نیست.

مسأله‌ی (۷) معادل یک مسأله‌ی کاهش بعد خطی روی خمینه‌ی $V_p(\mathbb{R}^{dT}) \times \mathbb{R}^m$ است. در این ساختار، نگاشت خطی $A: \mathbb{R}^{dT} \rightarrow \mathbb{R}^p$ صرفاً نقش کاهش‌دهنده‌ی بعد را ایفا می‌کند و هیچ محدودیتی بر فضای $\text{Pol}_\delta(\mathbb{R}^p)$ اعمال نمی‌شود. بنابراین کل مجموعه‌ی پایه‌های $\{g_i\}_{i=1}^m$ در فضای کاهش‌یافته به‌کار گرفته می‌شود. در واقع، بعد کل فضای بهینه‌سازی در این چارچوب برابر است با

$$\dim V_p(\mathbb{R}^{dT}) = dTp - \frac{1}{p}p(p+1) \quad \text{و} \quad m = \dim \text{Pol}_\delta(\mathbb{R}^p) = \binom{p+\delta}{p}.$$

جدول ۴.۵ نشان می‌دهد که بعد کل فضای بهینه‌سازی برای مقادیر مختلف d, T و p در مقایسه با جدول ۱ کمتر است.

جدول ۲. بعد کل: $\dim V_p(\mathbb{R}^{dT}) + \dim \text{Pol}_\delta(\mathbb{R}^p)$ برای $\delta = 3$.

$d = 5, p = 10$	$d = 3, p = 3$	T
$195 + 286 = 481$	$39 + 20 = 59$	۵
$1945 + 286 = 2231$	$354 + 20 = 374$	۴۰

محاسبه‌ی مقدار مورد انتظار شرطی در معادله (۴) به سادگی فرم بسته‌ی معادله (۱۵) نیست. بلکه باید مقادیر مورد انتظار مورد بحث (مطابق معادله (۱۰)) به‌صورت ذیل محاسبه شود:

$$(۱۶) \quad G_{A, t, i}(x_1, \dots, x_t) = \mathbb{E}^Q \left[g_i \left(\sum_{s=1}^t A_s^\top x_s + \sum_{s=t+1}^T A_s^\top X_s \right) \right],$$

که در آن $A \in V_p(\mathbb{R}^{dT})$ ارزیابی معادله (۱۶) در نهایت به محاسبه‌ی گشتاورهای چندمتغیره‌ی متغیر تصادفی $Y \in \mathbb{R}^p$ منتهی می‌شود که در آن $Y = \sum_{s=1}^t A_s^\top x_s + \sum_{s=t+1}^T A_s^\top X_s$. یکی از ابزارهای ترکیباتی مفید در محاسبه‌ی گشتاورهای چندمتغیره، رابطه‌ی بین تک‌جمله‌ای‌های چندمتغیره و ترکیبی از توان‌های جمع خطی متغیرها است که لم ذیل به آن می‌پردازد.

گزاره ۵.۵. [۱۸]. فرض کنید $x_1, \dots, x_n \in \mathbb{R}$ ، $s_1, \dots, s_n \in \mathbb{N}_0$ و $s = s_1 + s_2 + \dots + s_n$ در این صورت، داریم:

$$x_1^{s_1} x_2^{s_2} \dots x_n^{s_n} = \frac{1}{s!} \sum_{q_1=0}^{s_1} \dots \sum_{q_n=0}^{s_n} (-1)^{\sum_{i=1}^n q_i} \prod_{i=1}^n \binom{s_i}{q_i} \left(\sum_{i=1}^n h_i x_i \right)^s, \quad h_i = \frac{s_i}{s} - q_i.$$

که در آن عبارت $\left(\sum_{i=1}^n h_i x_i \right)^s$ نشان‌دهنده‌ی توان s -ام ترکیب خطی متغیرها است.

برای نمونه، اگر $n = 2$ و $s_1 = s_2 = 1$ باشد، داریم:

$$y_1 y_2 = \frac{1}{2!} \sum_{q_1=0}^1 \sum_{q_2=0}^1 (-1)^{q_1+q_2} \binom{1}{q_1} \binom{1}{q_2} (h_1 y_1 + h_2 y_2)^2, \quad h_i = \frac{1}{2} - q_i.$$

با توجه به اینکه $\binom{1}{q} = 1$ برای $q \in \{0, 1\}$ ، چهار حالت ممکن برای (q_1, q_2) را می‌نویسیم:

(q_1, q_2)	$(-1)^{q_1+q_2}$	(h_1, h_2)	$(h_1 y_1 + h_2 y_2)^2$
$(0, 0)$	$+1$	$(\frac{1}{2}, \frac{1}{2})$	$(\frac{1}{2} y_1 + \frac{1}{2} y_2)^2 = \frac{1}{4} (y_1 + y_2)^2$
$(1, 0)$	-1	$(-\frac{1}{2}, \frac{1}{2})$	$(-\frac{1}{2} y_1 + \frac{1}{2} y_2)^2 = \frac{1}{4} (y_2 - y_1)^2$
$(0, 1)$	-1	$(\frac{1}{2}, -\frac{1}{2})$	$(\frac{1}{2} y_1 - \frac{1}{2} y_2)^2 = \frac{1}{4} (y_1 - y_2)^2$
$(1, 1)$	$+1$	$(-\frac{1}{2}, -\frac{1}{2})$	$(-\frac{1}{2} y_1 - \frac{1}{2} y_2)^2 = \frac{1}{4} (y_1 + y_2)^2$

$$\begin{aligned} y_1 y_2 &= \frac{1}{2!} \left[\frac{1}{4} (y_1 + y_2)^2 - \frac{1}{4} (y_2 - y_1)^2 - \frac{1}{4} (y_1 - y_2)^2 + \frac{1}{4} (y_1 + y_2)^2 \right] \\ &= \frac{1}{2!} \left[\frac{1}{2} (y_1 + y_2)^2 - \frac{1}{2} (y_1 - y_2)^2 \right] = \frac{1}{2!} \left[(y_1 + y_2)^2 - (y_1 - y_2)^2 \right], \end{aligned}$$

در ادامه تعمیمی از رابطه‌ی $4 y_1 y_2 = (y_1 + y_2)^2 - (y_1 - y_2)^2$ ارائه می‌دهیم.

$$(17) \quad y^\alpha \equiv y_1^{\alpha_1} \dots y_p^{\alpha_p} = \frac{1}{|\alpha|!} \sum_{\nu=0}^{\alpha} (-1)^{|\nu|} \binom{\alpha_1}{\nu_1} \dots \binom{\alpha_p}{\nu_p} (h_{\alpha, \nu}^\top y)^{|\alpha|},$$

که در آن $h_{\alpha, \nu} = \left(\frac{\alpha_1}{\nu_1} - \nu_1, \dots, \frac{\alpha_p}{\nu_p} - \nu_p \right)^\top$ تعداد جملات در رابطه‌ی (۱۷) را می‌توان از طریق شمارش مقادیر ممکن برای $\nu = (\nu_1, \dots, \nu_p)$ به دست آورد. از این رو، تعداد جملات مؤثر برابر با $\frac{(\alpha_1+1)(\alpha_2+1)\dots(\alpha_p+1)}{p}$ است. بنابراین محاسبه‌ی $\mathbb{E}^\mathbb{Q}[Y^\alpha]$ در معادله (۱۶)، به محاسبه‌ی گشتاور مرتبه‌ی $|\alpha|$ متغیرهای تصادفی اسکالر $h_{\alpha, \nu}^\top Y$ در معادله (۱۷) کاهش می‌یابد. این گشتاورها برای توزیع‌های مختلف متغیر تصادفی X به صورت تحلیلی در دسترس هستند.

مثال ۶.۵. در حالت نرمال چندمتغیره، اگر $(X) \sim \mathcal{N}(0, I_{dT})$ ، آنگاه داریم:

$$h_{\alpha, \nu}^\top Y \sim \mathcal{N} \left(\sum_{s=1}^t h_{\alpha, \nu}^\top A_s^\top x_s, \sum_{s=t+1}^T \|A_s h_{\alpha, \nu}\|^2 \right).$$

این گشتاورهای تک‌متغیره همان‌گونه که در [۱۸] (گزاره ۲) بیان شده است؛ به صورت تحلیلی در دسترس هستند.

۵.۵. شبکه عصبی کم عمق. سومین ساختار نگاشت ویژگی، شبکه عصبی کم عمق (شامل یک لایه پنهان) با تابع فعال سازی رلو^{۳۶} است. برخلاف نگاشت های چند جمله ای که دارای ساختار تابعی از پیش تعیین شده هستند، شبکه های عصبی امکان تقریب خودکار توابع غیر خطی پیچیده را مستقیماً از داده ها فراهم می کنند. فرض می کنیم $p = m$ که در آن p بعد فضای بهینه سازی و m تعداد نورون های لایه پنهان است. همچنین تابع g_i به فرم $y_i^+ = \max\{0, y_i\} := g_i(y)$ برای $i = 1, \dots, m$ تعریف می شود. در نتیجه، نگاشت ویژگی از معادله (۸) به صورت ذیل بازنویسی خواهد شد:

$$g_i(y) := y_i^+ = \max\{0, y_i\}, \quad i = 1, \dots, m.$$

بنابراین نگاشت ویژگی در معادله (۸) به شکل ذیل بازنویسی می شود:

$$(۱۸) \quad \phi_{(A,b),i}(x) \equiv \phi_{(a_i,b_i)}(x) = (a_i^\top x + b_i)^+,$$

که در آن ماتریس وزن $A = (a_1, \dots, a_m) \in \mathbb{R}^{dT \times m}$ و بردار اریبی $b \in \mathbb{R}^m$. فرض می کنیم شرط (۱۲) برای $\delta = 1$ برقرار است لذا $\phi_{(a_i,b_i)} \in L_{\mathbb{Q}}^+$ با توجه به ویژگی همگنی مثبت^{۳۷} تابع فعال سازی رلو نسبت به پارامترها، یعنی $\phi(\lambda a_i, \lambda b_i) = \lambda \phi(a_i, b_i)$ برای هر $\lambda > 0$ بدون از دست دادن کلیت، فرض می کنیم که بردارهای پارامتر (a_i, b_i) روی کره واحد $\mathbb{S}_{dT} \subset \mathbb{R}^{dT+1}$ قرار دارند؛ یعنی $\|(a_i, b_i)\| = 1$. در نتیجه، فضای پارامترها را می توان به صورت خمینه $\Theta = (\mathbb{S}_{dT})^m$ در نظر گرفت که در آن حاصل ضرب $(\mathbb{S}_{dT})^m = \underbrace{\mathbb{S}_{dT} \times \dots \times \mathbb{S}_{dT}}_{m \text{ حاصل ضرب}}$ است.

در ارتباط با استقلال خطی توابع $\phi_{(a_i,b_i)}$ قضیه ای بنیادی در ادامه ارائه می شود که با وجود اهمیت نظری بالا، به نظر می رسد تاکنون در ادبیات علمی این حوزه کمتر مورد توجه قرار گرفته است.

قضیه ۷.۵. برای هر $(a_i, b_i), (\tilde{a}_i, \tilde{b}_i) \in \mathbb{S}_{dT}$ و $i = 1, \dots, m$ ، گزاره های ذیل برقرارند:

- (۱) اگر $\text{span}\{\phi_{(a_1,b_1)}, \dots, \phi_{(a_m,b_m)}\} = \text{span}\{\phi_{(\tilde{a}_1,\tilde{b}_1)}, \dots, \phi_{(\tilde{a}_m,\tilde{b}_m)}\}$ آنگاه $\{\pm(a_1, b_1), \dots, \pm(a_m, b_m)\} = \{\pm(\tilde{a}_1, \tilde{b}_1), \dots, \pm(\tilde{a}_m, \tilde{b}_m)\}$
- (۲) اگر برای هر $j \neq i$ $(a_i, b_i) \neq \pm(a_j, b_j)$ آنگاه $\{\phi_{(a_1,b_1)}, \dots, \phi_{(a_m,b_m)}\}$ مستقل خطی است.

توجه داشته باشید که حکم معکوس در قسمت ۱ از قضیه ۷.۵ برقرار نیست. علاوه بر این، فرض قسمت ۲ از قضیه ۷.۵ را نمی توان صرفاً به نابرابری جفتی $(a_i, b_i) \neq (a_j, b_j)$ برای هر $j \neq i$ تقلیل داد.

قضیه ۸.۵. برای شبکه عصبی کم عمق با تابع فعال سازی رلو، در حالت کلی کمیته کننده ای برای مسئله بهینه سازی (۷) وجود ندارد. علاوه بر این در صورت وجود جواب نیز، یکتایی برقرار نیست. به این معنا که زیر فضای بهینه $\text{span}\{\phi_{(a_1,b_1)}, \dots, \phi_{(a_m,b_m)}\}$ یکتا نیست.

مقدار مورد انتظار شرطی توابع رلو مطابق با رابطه (۱۰) برای $i = 1 \dots m$ به صورت ذیل محاسبه می شود:

$$(۱۹) \quad G_{(a_i,b_i),t}(x_1, \dots, x_t) = \mathbb{E}^{\mathbb{Q}} \left[\left(b_i + \sum_{s=1}^t a_{i,s}^\top x_s + \sum_{s=t+1}^T a_{i,s}^\top X_s \right)^+ \right],$$

که در آن $a_i^\top = (a_{i,1}^\top, a_{i,2}^\top, \dots, a_{i,T}^\top)$ و $a_{i,s}$ بردار وزن مربوط به مرحله زمانی s است.

³⁶ReLU ³⁷Positive Homogeneity

ارزیابی رابطه‌ی (۱۹) معادل با محاسبه‌ی عبارت $\mathbb{E}^{\mathbb{Q}}[Y^+]$ است که متغیر تصادفی اسکالر Y به شکل $Y = b_i + \sum_{s=1}^t a_{i,s}^\top x_s + \sum_{s=t+1}^T a_{i,s}^\top X_s$ تعریف می‌شود. در برخی موارد، با توجه به توزیع متغیر تصادفی X ، می‌توان این مقدار مورد انتظار را به صورت تحلیلی محاسبه کرد.

مثال ۹.۵. در حالت نرمال چندمتغیره، $X \sim \mathcal{N}(\circ, I_{dT})$ ، بخش تصادفی متغیر Y یعنی $\sum_{s=t+1}^T a_{i,s}^\top X_s$ ، یک ترکیب خطی از متغیرهای نرمال مستقل با میانگین صفر است و لذا خود نیز یک متغیر نرمال با میانگین صفر است. بنابراین متغیر Y دارای توزیع نرمال با پارامترهای ذیل خواهد بود:

$$(20) \quad Y \sim \mathcal{N}\left(b_i + \sum_{s=1}^t a_{i,s}^\top x_s, \sum_{s=t+1}^T \|a_{i,s}\|^2\right).$$

در این صورت، با ترکیب معادله (۲۰) و فرمول قیمت‌گذاری اختیار خرید بشیلیه^{۳۸} که برای متغیر تصادفی نرمال $Z \sim \mathcal{N}(\mu, \sigma^2)$ به صورت $\mathbb{E}^{\mathbb{Q}}[Z^+] = \mu \Phi\left(\frac{\mu}{\sigma}\right) + \sigma \Phi'\left(\frac{\mu}{\sigma}\right)$ است، می‌توان یک فرم تحلیلی برای $\mathbb{E}^{\mathbb{Q}}[Y^+]$ به صورت ذیل به دست آورد:

$$\begin{aligned} \mathbb{E}^{\mathbb{Q}}[Y^+] &= \left(b_i + \sum_{s=1}^t a_{i,s}^\top x_s\right) \Phi\left(\frac{b_i + \sum_{s=1}^t a_{i,s}^\top x_s}{\sqrt{\sum_{s=t+1}^T \|a_{i,s}\|^2}}\right) \\ &+ \sqrt{\sum_{s=t+1}^T \|a_{i,s}\|^2} \cdot \Phi'\left(\frac{b_i + \sum_{s=1}^t a_{i,s}^\top x_s}{\sqrt{\sum_{s=t+1}^T \|a_{i,s}\|^2}}\right), \end{aligned}$$

که در آن Φ تابع توزیع نرمال استاندارد و Φ' چگالی آن است [۵].

در مواردی که توزیع متغیر تصادفی پیچیده است اما تبدیل فوریه تعمیم یافته^{۳۹} توزیع‌های حاشیه‌ای $\hat{\mathbb{Q}}_t(u) = \mathbb{E}^{\mathbb{Q}}[e^{u^\top X_t}]$ به صورت تحلیلی برای دامنه‌ی مناسبی از بردارهای مختلط u در دسترس است، می‌توان برخی از توابع متغیر تصادفی را با استفاده از انتگرال‌گیری در حوزه‌ی مختلط تحلیل کرد. به‌ویژه تابع $Y^+ = \max(Y, \circ)$ را می‌توان به صورت نمایش انتگرالی بر حسب توابع نمایی مختلط بازنویسی کرد. برای هر $w > \circ$ ، تابع Y^+ دارای نمایش انتگرالی ذیل است:

$$(21) \quad Y^+ = \frac{1}{\sqrt{\pi}} \int_{\mathbb{R}} \frac{e^{(w+i\lambda)Y}}{(w+i\lambda)^2} d\lambda$$

بنابراین، ارزیابی معادله (۱۹) به محاسبه‌ی انتگرال خطی ذیل تقلیل می‌یابد:

$$\begin{aligned} G_{(a_i, b_i), t}(x_1, \dots, x_t) &= \mathbb{E}^{\mathbb{Q}}[Y^+] \\ &= \mathbb{E}^{\mathbb{Q}}\left[\frac{1}{\sqrt{\pi}} \int_{\mathbb{R}} e^{(w+i\lambda)Y} \cdot \frac{1}{(w+i\lambda)^2} d\lambda\right] \\ &= \frac{1}{\sqrt{\pi}} \int_{\mathbb{R}} \mathbb{E}^{\mathbb{Q}}\left[e^{(w+i\lambda)Y}\right] \cdot \frac{1}{(w+i\lambda)^2} d\lambda \\ &= \frac{1}{\sqrt{\pi}} \int_{\mathbb{R}} \hat{F}_Y(w+i\lambda) \cdot \frac{1}{(w+i\lambda)^2} d\lambda \end{aligned} \quad (22)$$

³⁸Bachelier ³⁹extended Fourier transform

که در آن $\hat{F}_Y(w + i\lambda)$ تبدیل فوریه تعمیم یافته متغیر Y است و فرم تحلیلی آن از ساختار خطی Y استخراج می شود:

$$\hat{F}_Y(w + i\lambda) = \mathbb{E}^{\mathbb{Q}} \left[e^{(w+i\lambda)Y} \right] = e^{(w+i\lambda)(b_i + \sum_{s=1}^t a_{i,s}^{\top} x_s)} \prod_{s=t+1}^T \hat{\mathbb{Q}}_s((w + i\lambda)a_{i,s}).$$

توجه داشته باشید که انتگرال هایی از نوع فوریه، در دنیای مالی به ویژه در مدل های مالی مدرن، از جمله مدل های با پرس ^{۴۰} یا مدل های افاین ^{۴۱}، برای قیمت گذاری اختیار معامله بسیار کاربرد دارد.

۶.۵. برآورد با نمونه های متناهی ^{۴۲}. علیرغم نتایج مربوط به عدم یکتایی و حتی عدم وجود جواب، در قضایای ۴.۵ و ۸.۵ برای نگاشت ویژگی چند جمله ای با کاهش بعد و شبکه های عصبی، در کاربردهای عملی می توان جواب های موضعی مسئله ی (۷) با استفاده از الگوریتم های شبه نیوتنی ^{۴۳} تقریب زد. در عمل به اندازه ی مدل $\mathbb{Q}$ که توزیع نظری زیربنای داده ها را توصیف می کند، دسترسی مستقیم نداریم. لذا آن را با توزیع تجربی $\hat{\mathbb{Q}}$ تقریب می زنیم. این توزیع، مبتنی بر یک نمونه ی آموزشی متشکل از مشاهدات مستقل و هم توزیع $x^{(1)}, \dots, x^{(n)} \sim \mathbb{Q}$ به صورت $\hat{\mathbb{Q}} = \frac{1}{n} \sum_{i=1}^n \delta_{x^{(i)}}$ تعریف می شود که در آن $\delta_{x^{(i)}}$ تابع دلتای دیراک ^{۴۴} در نقطه ی $x^{(i)}$ است. در روش تقریب با استفاده از پایه ی کامل چند جمله ای، مسئله ی (۷) به رابطه ی (۳) تقلیل می یابد و پارامتر بهینه ی تجربی به صورت بسته قابل محاسبه است:

$$\hat{\beta} = \left(\frac{1}{n} \Phi^{\top} \Phi \right)^{-1} \left(\frac{1}{n} \Phi^{\top} y \right), \quad (23)$$

که در آن $\Phi \in \mathbb{R}^{n \times m}$ ماتریس ویژگی با مؤلفه های $\Phi_{ij} = \phi_j(x^{(i)})$ است.

فرض می کنیم نمونه های آموزشی $\{x^{(i)}\}_{i=1}^n$ مستقل و هم توزیع بوده و ماتریس گرام ^{۴۵} مثبت معین و در نتیجه وارون پذیر است و همچنین تعداد پارامترها (m) ثابت است. تحت این فرضیات و در حالت کلاسیک $m < n$ برآوردگر کمترین مربعات طبق قانون اعداد بزرگ همگرای در احتمال است یعنی $\hat{\beta} \xrightarrow{P} \beta$. با فرض وجود گشتاورهای مرتبه ی دوم لازم، قضیه ی حد مرکزی ^{۴۶} بیان می دارد که $\sqrt{n}(\hat{\beta} - \beta) \xrightarrow{d} \mathcal{N}(0, \Sigma)$ که در آن Σ ماتریس کوواریانس حدی است. لازم به ذکر است که این نتایج کلاسیک در چارچوب بعد پارامتر ثابت و با فرض $m < n$ بیان می شوند و نمی توان آن ها را به طور مستقیم به حالت هایی که تعداد توابع پایه از تعداد نمونه ها بیشتر است (یعنی $m > n$) تعمیم داد. در این حالت، به ویژه در صورت عدم وارون پذیری ماتریس $\Phi^{\top} \Phi$ ، از منظم سازی رگرسیون ریب ^{۴۷} در پیاده سازی عددی استفاده می کنیم.

در حالی که در مبنای کامل چند جمله ای پارامترهای مدل تنها شامل ضرایب خطی هستند و برآورد آن ها در فرم بسته قابل محاسبه است، در دو ساختار دیگر مسئله به یک بهینه سازی غیر خطی تبدیل می شود.

برای نگاشت ویژگی چند جمله ای همراه با کاهش بعد خطی، مسئله ی (۷) روی فضای پارامترهای $V_p(\mathbb{R}^{dT}) \times \mathbb{R}^m$ تعریف می شود. مؤلفه ی اول شامل ماتریس هایی است که شرط تعامد ستون ها را برآورده می کنند. بنابراین به صورت یک خمینه ریمانی ^{۴۸} در نظر گرفته می شود. مؤلفه ی دوم فضای پارامترهای بردار وزن β را تشکیل می دهند. با توجه به اینکه فضای پارامتر غیر خطی و دارای ساختار خمینه است. برای بهینه سازی مؤثر در این فضا، از الگوریتم شبه نیوتنی RBFGS ^{۴۹} بهره می گیریم که نسخه ای تعمیم یافته از الگوریتم BFGS ^{۵۰} برای فضاهای ریمانی است. مقاله ی [۱۶] چارچوب نظری جامعی برای الگوریتم های شبه نیوتنی ریمانی ارائه داده است.

برای مدل شبکه ی عصبی کم عمق با تابع فعال سازی رلو، مسئله ی (۷) روی فضای پارامتر اقلیدسی $\mathbb{R}^{dT \times m} \times \mathbb{R}^m \times \mathbb{R}^m$ تعریف می شود. این فضا شامل وزن های بین ورودی و لایه ی پنهان، اربیبی های لایه ی پنهان و وزن های خروجی است. این مدل با وجود ساختار ساده اش، به دلیل ماهیت غیر محدب مسئله، بهینه سازی چالش برانگیزی دارد. از این رو، برای یافتن مینیمم موضعی، از الگوریتم شبه نیوتنی

⁴⁰Lévy-type ⁴¹Affine Models ⁴²Finite-Sample Estimation ⁴³Quasi-Newton ⁴⁴Dirac Delta Function

⁴⁵Gram Matrix ⁴⁶Central Limit Theorem ⁴⁷Ridge Regression ⁴⁸Riemannian Manifold ⁴⁹Riemannian

Broyden-Fletcher-Goldfarb-Shanno ⁵⁰Broyden-Fletcher-Goldfarb-Shanno

BFGS استفاده می‌کنیم که در کتابخانه‌ی Scikit-learn [۲۳] برای زبان برنامه‌نویسی پایتون پیاده‌سازی شده است. با توجه به نبود یکتایی و حتی وجود جواب برای نگاشت ویژگی چندجمله‌ای با کاهش بعد و شبکه‌ی عصبی رلو کم عمق، همچنان پرسشی باز در حوزه‌ی پژوهش باقی می‌ماند که آیا ویژگی همگرایی سازگار در این روش‌ها برقرار است و آیا می‌توان تضمین‌های نظری برای آن‌ها ارائه کرد یا خیر.

۶. نتایج عددی

در این بخش، معیارهای مورد استفاده برای ارزیابی عملکرد روش‌های مختلف معرفی می‌شوند. با توجه به اینکه در کاربردهای مدیریت ریسک، دقت مدل در نواحی بحرانی توزیع از اهمیت ویژه‌ای برخوردار است، کیفیت برازش در دم توزیع و در بدنه‌ی توزیع مورد بررسی قرار می‌گیرد.

برای سنجش عملکرد مدل در دم توزیع، از دو معیار متداول ارزش در معرض ریسک^{۵۱} و زیان مورد انتظار شرطی^{۵۲} استفاده می‌شود. برای یک متغیر تصادفی زیان L و سطح اطمینان $\alpha \in (0, 1)$ ، ارزش در معرض ریسک در سطح α به صورت ذیل تعریف می‌شود:

$$\text{VaR}_\alpha(L) := \inf\{x \in \mathbb{R} : \mathbb{P}(L \leq x) \geq \alpha\}.$$

این کمیت، صدک α -ام توزیع زیان را نشان می‌دهد.

همچنین زیان مورد انتظار شرطی در سطح α به صورت ذیل تعریف می‌شود:

$$\text{ES}_\alpha(L) := \frac{1}{1-\alpha} \int_\alpha^1 \text{VaR}_u(L) du.$$

در واقع ES_α میانگین زیان‌های فراتر از آستانه‌ی VaR_α را اندازه‌گیری می‌کند.

در این چارچوب، مدل‌ها به عنوان برآوردگرهای آماری در نظر گرفته می‌شوند. از این رو برای محاسبه‌ی VaR_α و ES_α ، با انجام R اجرای مستقل از کل زنجیره‌ی شبیه‌سازی، یک توزیع تجربی از نتایج به دست می‌آید. همچنین برای میانگین‌گیری این توزیع تجربی از معیار میانگین خطای مطلق درصدی^{۵۳} استفاده می‌شود.

۱.۶. میانگین خطای مطلق درصدی (MAPE).

فرض کنید مجموعه‌ای از n مشاهده‌ی مستقل از بردار تصادفی $X_{1:t}$ در اختیار داریم. برای تابع f ، با در نظر گرفتن R تکرار از برآوردگر آن، یعنی $\hat{f}$ ، مجموعه‌ای از توابع $\{\hat{f}_j\}_{j=1}^R$ به دست می‌آوریم. برای هر تابع در این مجموعه می‌توان یک توزیع تجربی از $\hat{V}_t = \mathbb{E}_t^\mathbb{Q}[\hat{f}(X)]$ با استفاده از $X_{1:t}$ تولید کرد. بنابراین مجموعه‌ای از مقادیر تجربی $\{\hat{V}_t^{(j)}\}_{j=1}^R$ به دست خواهد آمد. در نهایت میانگین خطای مطلق درصدی برای زیان مورد انتظار شرطی و ارزش فعلی به صورت ذیل تعریف می‌شود:

$$\text{MAPE ES}_\alpha = \frac{1}{R} \sum_{j=1}^R \frac{|ES_\alpha[-\Delta \hat{V}_t^{(j)}] - ES_\alpha[-\Delta V_t]|}{ES_\alpha[-\Delta V_t]}.$$

$$\text{MAPE PV} = \frac{1}{R} \sum_{j=1}^R \frac{|\hat{V}_t^{(j)} - V_t|}{V_t}.$$

⁵¹Value at Risk (VaR) ⁵²Expected Shortfall (ES) ⁵³Mean Absolute Percentage Error

۲.۶. میانگین خطای نسبی L^1 . برای ارزیابی عملکرد مدل در بدنه‌ی توزیع از میانگین خطای نسبی L^1 استفاده می‌کنیم که به صورت ذیل تعریف می‌شود:

$$(26) \quad \frac{1}{R} \sum_{j=1}^R \frac{\mathbb{E}[|\hat{V}_t^{(j)} - V_t|]}{\mathbb{E}[|V_t|]}.$$

۷. سبد مالی شامل اختیار خرید اروپایی^{۵۴}

در این بخش، به عنوان نخستین مطالعه‌ی عددی، سبدی شامل یک اختیار خرید اروپایی بررسی می‌شود که دارایی پایه‌ی آن یک شاخص سهام^{۵۵} است. در این تحلیل، فرض می‌شود دارنده‌ی اختیار در موقعیت فروش^{۵۶} قرار دارد. بنابراین، توجه ما معطوف به ناحیه‌ی زیان‌آور توزیع فرآیند ارزش سبد مالی است، یعنی زمانی که مقدار شاخص سهام افزایش می‌یابد. اگرچه سررسید قراردادها برابر $T = 5$ و $T = 40$ سال است، معیار ریسک برای افق یک‌ساله تعریف و گزارش می‌شود. یعنی هدف، سنجش ریسک ارزش سبد در پایان سال اول است. برای اندازه‌گیری ریسک دم، معیار $ES_{99\%}(-\Delta V_1)$ از توزیع تجربی سناریوهای شبیه‌سازی شده محاسبه می‌شود که در آن $-\Delta V_1 := V_0 - V_1$ زیان یک‌ساله است. به منظور بازتولیدپذیری نتایج، داده‌ها و سناریوهای شبیه‌سازی شده‌ی محرک تصادفی X در [A] و کدهای پایتون در [Profile GitHub](#) در دسترس هستند و می‌توان از آن‌ها برای اجرای مجدد محاسبات استفاده کرد. برای شبیه‌سازی مسیرها یک مولد سناریو اقتصادی در نظر می‌گیریم. این مولد وظیفه دارد مسیرهای متغیرهای مالی را برای مدل‌سازی پویایی وضعیت اقتصادی در طول افق زمانی مورد نظر شبیه‌سازی کند. مولد سناریو اقتصادی، بردار تصادفی X_t را به یک بردار از عوامل اقتصادی نگاشت می‌کند که شامل مؤلفه‌هایی نظیر شاخص سهام S_t ، نرخ بهره r_t و حساب نقدی C_t است.

برای محاسبه‌ی ارزش سبد مالی در سررسید T ، پرداختی اختیار خرید اروپایی به صورت $\max(S_T - K, 0)$ تعریف می‌شود. دارنده سبد مالی در این مثال فروشنده‌ی اختیار است، بنابراین این پرداخت به عنوان خروجی نقدی با علامت منفی نمایش داده شده است. از طرفی، برای تبدیل این پرداخت اسمی به مقدار تنزیل شده، از حساب نقدی C_T به عنوان دارایی مبنا استفاده می‌شود:

$$(27) \quad f(X) = - \frac{\max(S_T - K, 0)}{C_T}$$

در این مطالعه، قیمت اعمال برابر با مقدار اولیه‌ی شاخص سهام است ($K = S_0 = 100$) و مدل شامل سه محرک تصادفی است ($d = 3$).

شایان ذکر است که علامت منفی برای تاکید بر این نکته است که سودآوری نگهدارنده اختیار خرید از منظر فروشنده زیان خواهد بود. اما در عمل در همه محاسبات مربوط به V_t و V_0 مقدار $f(X)$ به صورت مثبت لحاظ می‌شود.

شبیه‌سازی مونت کارلو تودرتو برای سبد مالی شامل اختیار خرید اروپایی. همانطور که در بخش ۲ شرح دادیم، مقدار ارزش سبد مالی در زمان t برابر با مقدار مورد انتظار شرطی آن تحت اندازه‌ی ریسک خنثی $\mathbb{Q}$ است ($V_t = \mathbb{E}_t^{\mathbb{Q}}[f(X)]$). شبیه‌سازی تودرتو به گونه‌ای طراحی می‌شود که در هر تکرار از شبیه‌سازی بیرونی، یک شبیه‌سازی درونی نیز جهت برآورد مقدار مورد انتظار شرطی انجام گیرد. در این مقاله، برای به دست آوردن مقدار مبنا، تکنیک شبیه‌سازی مونت کارلو تودرتو در مقیاس بزرگ اجرا می‌شود تا برآورد دقیقی از مقدار مورد انتظار شرطی به دست آید. بدین منظور تعداد شبیه‌سازی‌های بیرونی را $n_0 = 1000000$ و شبیه‌سازی‌های درونی را $n_1 = 100000$ قرار می‌دهیم. حال معیار $ES_{99\%}(-\Delta V_1)$ را بر اساس توزیع تجربی داده‌های شبیه‌سازی شده محاسبه می‌کنیم. ابتدا داده‌های $-\Delta V_1$ به صورت نزولی مرتب می‌شوند و سپس چندک α -ام برای $\alpha = 0.99$ استخراج می‌گردد. اگر تعداد کل سناریوهای شبیه‌سازی شده n باشد، تعداد داده‌های مربوط به دنباله زیان‌آور برابر با $n_{\text{tail}} = \lfloor (1 - \alpha) \times n \rfloor$ است. در نتیجه، تخمین تجربی ES_{α}

⁵⁴European Call Option ⁵⁵Equity Index ⁵⁶Short Position

به صورت ذیل محاسبه می‌شود:

$$ES_{\alpha}(-\Delta V_1) = \frac{1}{n_{tail}} \sum_{i=1}^{n_{tail}} Largest_i(-\Delta V_1),$$

که در آن $Largest_i(-\Delta V_1)$ نشان‌دهنده i -مین مقدار بزرگ‌تر از نظر شدت زیان در بین داده‌های مرتب‌شده است. پیش از بررسی کیفیت تخمین‌گر مونت‌کارلو تودرتو در شرایط بودجه‌ی محدود شبیه‌سازی، لازم است تصمیم بگیریم این بودجه بین مرحله‌ی بیرونی و درونی چگونه تقسیم شود. ترکیب‌های مختلفی از n_0 و n_1 وجود دارند که هر کدام منجر به تخمین متفاوتی می‌شوند. جدول ۳ مقدار $MapE$ را برای ترکیب‌های مختلف از یک بودجه‌ی شبیه‌سازی ثابت $50000 = n_0 \times n_1$ نشان می‌دهد. همان‌طور که انتظار می‌رود، با افزایش تعداد شبیه‌سازی‌های درونی n_1 در شبیه‌سازی مونت‌کارلو تودرتو، خطای $MapE$ کاهش یافته و کنترل می‌شود.

به منظور دستیابی به نتایجی خوش‌بینانه‌تر برای روش شبیه‌سازی مونت‌کارلو تودرتو، بهترین ترکیب پارامترها بر اساس معیار کمینه‌سازی $MapE$ برای مقادیر برآوردشده‌ی زیان مورد انتظار انتخاب می‌شود. با این حال، این سوگیری احتمالی در جهت برآوردهای خوش‌بینانه، تأثیر معناداری بر تحلیل ندارد. زیرا روش پیشنهادی مبتنی بر مارتینگل بازساز در عمل دقت بالاتری نسبت به سایر رویکردها نشان می‌دهد.

جدول ۳. مقایسه‌ی $ES_{99\%}$ در روش مونت‌کارلو تودرتو برای افرازهای مختلف مجموع شبیه‌سازی‌های ثابت $n_0 \times n_1 = 50000$ که در آن تعداد شبیه‌سازی‌های داخلی n_1 در بازه ۱ تا ۵۰۰ تغییر می‌کند.

مونت کارلو تودرتو (خطای $MapE$ برای شبیه‌سازی‌های داخلی مختلف)								سررسید	مقدار مبنا
۵۰۰	۴۰۰	۲۵۰	۱۰۰	۵۰	۲۵	۱۰	۱		
۲۰/۴٪	۱۹/۰٪	۱۴/۴٪	۹/۳٪	۷/۹٪	۱۱/۸٪	۲۹/۵٪	۲۳۳/۴٪	۵۶/۸۵۸۸	۵
۱۹/۱٪	۲۱/۴٪	۲۴/۶٪	۶۶/۶٪	۱۲۹/۳٪	۲۳۲/۱٪	۴۶۳/۶٪	۲۰۲۷/۰٪	۶۲/۶۲۰۵	۴۰

در جدول ۴ مشاهده می‌کنیم که با افزایش بودجه‌ی شبیه‌سازی، خطای برآورد زیان مورد انتظار $ES_{99\%}(-\Delta V_1)$ با سرعت بیشتری کاهش می‌یابد؛ در حالی‌که بهبود دقت ارزش فعلی V_0 کندتر است. در شرایطی که هدف صرفاً برآورد مقدار اولیه V_0 باشد، اختصاص کل بودجه‌ی محاسباتی به مرحله‌ی شبیه‌سازی بیرونی منطقی‌تر خواهد بود. برای نمونه، در حالتی که اندازه‌ی نمونه برابر با ۵۰۰۰۰ باشد، $MapE$ برای V_0 در ساختار تک‌مرحله‌ای تنها حدود ۰/۶٪ تا ۱٪ گزارش شده است؛ در حالی‌که با توزیع همین بودجه بین دو مرحله‌ی شبیه‌سازی بیرونی و درونی (در قالب شبیه‌سازی تودرتو)، این مقدار به حدود ۱/۷٪ تا ۲/۴٪ افزایش می‌یابد.

جدول ۴. مقایسه‌ی ارزش فعلی و زیان مورد انتظار در روش مونت‌کارلو تودرتو با استفاده از خطای $MapE$ ؛ با تقسیم بهینه‌ی بودجه شبیه‌سازی درونی بیرونی ($n_0 \times n_1$) برای اندازه نمونه‌های مختلف. اعداد بر حسب درصد هستند.

تعداد نمونه‌ها ($n_0 \times n_1$)		ارزش فعلی (V_0)		زیان مورد انتظار (ES)	
		$T = 40$	$T = 5$	$T = 40$	$T = 5$
۱۰۰۰		۶/۵	۷/۵	۲۶/۷	۴۰۳/۶
۵۰۰۰		۴/۰	۳/۹	۱۵/۹	۱۱۰/۹
۱۰۰۰۰		۳/۸	۳/۱	۱۴/۵	۵۶/۶
۵۰۰۰۰		۱/۷	۲/۴	۷/۹	۱۹/۱

رویکرد یادگیری ماشین برای سبد مالی شامل اختیار خرید اروپایی. در این بخش، توابع پایه معرفی شده در بخش‌های ۳.۵ تا ۵.۵ را به کار می‌گیریم. برای هر یک از خانواده‌های توابع پایه، سنجه‌های ارزیابی معرفی شده در بخش ۷ به کار می‌روند و نتایج

روش‌های پیشنهادی با برآورد مبنا حاصل از شبیه‌سازی بزرگ مونت‌کارلوی تودرتو مقایسه می‌شوند. در ارزیابی روش‌های مبتنی بر رگرسیون، از میانگین خطای L_1 بر روی توزیع تجربی ΔV_1 نیز استفاده می‌شود.

در جداول ۵ و ۶ مشاهده می‌شود که روش رگرسیون-آتی (مارتینگل بازساز) در هر دو معیار ارزش فعلی و زیان مورد انتظار ۹۹٪ عملکرد بهتری دارد. برای مقایسه‌ی جامع‌تر، میانگین خطای L_1 در جدول ۷ نیز بررسی شده است. این موضوع، هم‌راستا با یافته‌های پیشین در حوزه روش‌های رگرسیون-آتی است. با این حال، جداول ۵ و ۶ نتایجی برای پایه‌ی کامل چندجمله‌ای تحت روش مارتینگل بازساز در افق زمانی $T = 40$ ارائه نمی‌کنند. دلیل این امر، افزایش غیرقابل کنترل تعداد توابع پایه با افزایش بعد مسئله است. این دشواری‌های محاسباتی در مسائل با بعد بالا، انگیزه‌ای برای استفاده از روش کاهش بعد خطی فراهم می‌سازد.

برای مثال اختیار خرید اروپایی، مقدار $p = 3$ انتخاب شده است. تحلیل حساسیت بر روی این پارامتر انجام گرفت و مشخص شد که مقادیر بزرگ‌تر گاهی و نه در همه موارد عملکرد بهتری دارند. در بخش ۹، تحلیلی از حساسیت نسبت به نقطه‌ی شروع بهینه‌سازی نیز ارائه می‌کنیم. با توجه به اینکه الگوریتم BFGS یک روش شبه‌نیوتنی است، در حالت کلی تضمینی برای یافتن کمینه‌ی سراسری وجود ندارد.

در ادامه، نتایج حاصل از مدل شبکه‌ی عصبی کم‌عمق معرفی‌شده را گزارش می‌کنیم. تعداد کل نورون‌های لایه‌ی پنهان برابر $m = 101$ در نظر گرفته شده است که یکی از آن‌ها صرفاً برای ارببی به‌کار می‌رود به‌طوری‌که $A_{101} = 0$. تعداد پارامترهای مدل شامل دو بخش است:

- وزن‌های بین ورودی و لایه‌ی پنهان: برای هر نورون فعال در لایه‌ی پنهان (یعنی $m - 1$ نورون) به تعداد ابعاد ورودی $(dT + 1)$ وزن وجود دارد. بنابراین تعداد وزن‌های بین ورودی و لایه‌ی پنهان برابر است با $(m - 1) \times (dT + 1)$.
- وزن‌های بین لایه‌ی پنهان و خروجی: با توجه به اینکه خروجی یک اسکالر است، برای هر یک از m نورون پنهان، یک وزن خروجی در نظر گرفته می‌شود. این بخش در مجموع m پارامتر خواهد داشت.

در نتیجه، تعداد کل پارامترها $m + (m - 1) \times (dT + 1)$ است. تعداد نورون‌ها و ساختار کلی مدل (شامل تعداد و نوع لایه‌ها، نحوه‌ی اتصال آن‌ها و توابع فعال‌سازی) از طریق فرآیند اعتبارسنجی متقابل انجام شده است. بهینه‌سازی شبکه‌ی عصبی با بهره‌گیری از الگوریتم BFGS و از طریق روش پس‌انتشار خطا^{۵۷} انجام شده است. نتایج جداول ۵ و ۶ نشان می‌دهد که مدل مارتینگل بازساز با شبکه‌ی عصبی در این مثال عملکرد بسیار خوبی دارد و در بیشتر موارد از سایر روش‌ها بهتر عمل می‌کند. اگرچه وقتی تعداد نمونه‌ها بسیار پایین است، همان‌گونه که در مدل چندجمله‌ای مشاهده شد، روش رگرسیون-آتی از رگرسیون-کنون عملکرد بهتری دارد. مقایسه‌ی زمان اجرای روش‌های مختلف در بخش ۱۰ آمده است.

جدول ۵. مقایسه‌ی خطای MAPE برای ارزش فعلی V_0 در مثال اختیار خرید اروپایی. اعداد بر حسب درصد هستند.

⁵⁷Error Backpropagation

نمونه‌ها	مونت کارلو تودرتو	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی	شبکه‌ی عصبی	
		رگرسیون-اکنون	رگرسیون-آتی		رگرسیون-اکنون	رگرسیون-آتی
T = ۵						
۱۰۰۰	۶/۵	۴/۲	۱/۲	۰/۴	۴/۳	۰/۲
۵۰۰۰	۴/۰	۱/۶	۰/۲	۰/۲	۲/۱	۰/۱
۱۰۰۰۰	۳/۸	۱/۲	۰/۱	۰/۱	۱/۳	۰/۱
۵۰۰۰۰	۱/۷	۰/۶	۰/۱	۰/۱	۰/۵	< ۰/۱
T = ۴۰						
۱۰۰۰	۷/۵	۶/۸		۴/۷	۸/۲	۵/۷
۵۰۰۰	۳/۹	۳/۴		۲/۳	۳/۸	۱/۷
۱۰۰۰۰	۳/۱	۲/۳		۱/۳	۲/۲	۰/۶
۵۰۰۰۰	۲/۴	۱/۰		۰/۵	۰/۹	۰/۲

جدول ۶. مقایسه‌ی خطای (MAPE) برای زیان مورد انتظار $ES_{99\%}$ در مثال اختیار خرید اروپایی. اعداد بر حسب درصد هستند.

نمونه‌ها	مونت کارلو تودرتو	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی	شبکه‌ی عصبی	
		رگرسیون-اکنون	رگرسیون-آتی		رگرسیون-اکنون	رگرسیون-آتی
T = ۵						
۱۰۰۰	۲۶/۷	۱۹/۲	۳/۱	۱/۶	۱۶۹/۱	۱/۰
۵۰۰۰	۱۵/۹	۸/۵	۱/۲	۱/۰	۴۱/۱	۰/۲
۱۰۰۰۰	۱۴/۵	۵/۹	۰/۸	۰/۸	۱۸/۴	۰/۱
۵۰۰۰۰	۷/۹	۲/۸	۰/۵	۰/۶	۵/۳	< ۰/۱
T = ۴۰						
۱۰۰۰	۴۰۳/۶	۱۴۶/۰		۱۴/۶	۷۶۸/۸	۱۹/۴
۵۰۰۰	۱۱۰/۹	۴۰/۲		۵/۹	۲۸۱/۲	۶/۲
۱۰۰۰۰	۵۶/۶	۲۴/۲		۴/۹	۱۷۳/۰	۳/۵
۵۰۰۰۰	۱۹/۱	۱۰/۳		۲/۴	۴۸/۸	۰/۷

جدول ۷. مقایسه‌ی میانگین خطای نسبی L^1 در مثال اختیار خرید اروپایی. اعداد بر حسب درصد هستند.

نمونه‌ها	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی		شبکه‌ی عصبی	
	رگرسیون-اکنون	رگرسیون-آتی	رگرسیون-آتی	رگرسیون-آتی	رگرسیون-اکنون	رگرسیون-آتی
$T = 5$						
۱۰۰۰	۱۵/۴	۴/۵	۱/۱	۸۳/۸	۰/۵	
۵۰۰۰	۶/۸	۰/۹	۰/۶	۳۴/۶	۰/۴	
۱۰۰۰۰	۴/۶	۰/۷	۰/۵	۲۳/۶	۰/۳	
۵۰۰۰۰	۲/۱	۰/۵	۰/۴	۱۱/۰	۰/۳	
$T = 40$						
۱۰۰۰	۲۶/۵		۵/۲	۱۱۴/۲	۷/۵	
۵۰۰۰	۱۱/۳		۲/۶	۵۴/۰	۲/۲	
۱۰۰۰۰	۸/۱		۱/۹	۳۷/۹	۱/۲	
۵۰۰۰۰	۳/۷		۱/۳	۱۷/۲	۰/۷	

۸. بیمه عمر

در این بخش، بعد از بررسی کارایی یادگیری مارتینگل بازساز در قیمت‌گذاری اختیار خرید اروپایی، قرارداد بیمه‌ی عمر متغیر تضمین‌شده^{۵۸} را بررسی می‌کنیم. برخلاف اختیار خرید اروپایی، این مسئله شامل جریان‌های نقدی وابسته به مسیر در چندین نقطه‌ی زمانی است و علاوه بر متغیرهای تصادفی بازار به یک مدل تصادفی مرگ‌ومیر^{۵۹} نیز وابسته است. در این مقاله، داده‌های مربوط به بیمه‌گذاران واقعی در دسترس نبوده و به‌جای آن از یک جمعیت فرضی برای انجام محاسبات استفاده شده است.

قرارداد بیمه‌ی عمر متغیر تضمین‌شده، حساب سرمایه‌گذاری بیمه‌گذار با تضمین بازده در صورت فوت است. در هر دوره زمانی، بیمه‌گذاران حق بیمه پرداخت می‌کنند که برای خرید دارایی‌ها استفاده می‌شود. این دارایی‌ها در صندوق سرمایه‌گذاری، سپرده‌گذاری شده و بر اساس یک تخصیص دارایی ثابت (مطابق با تخصیص اولیه) میان دارایی‌ها تقسیم می‌شود. زمانی که بیمه‌گذار فوت کند، ارزش حساب سرمایه‌گذاری به ذی‌نفعان آن پرداخت می‌شود. در تاریخ سررسید، بیشینه‌ی میان ارزش حساب سرمایه‌گذاری و مبلغ تضمین‌شده به تمامی بازماندگان پرداخت می‌شود. یعنی اگر ارزش حساب سرمایه‌گذاری کمتر از مبلغ تضمین‌شده باشد، شرکت بیمه مابه‌التفاوت را پرداخت می‌کند. مبلغ تضمین‌شده برابر با مجموع تمامی حق بیمه‌های پرداخت‌شده است.

در ادامه یک سبد سرمایه‌گذاری مشخص در نظر می‌گیریم که شامل چهار دارایی است: اوراق قرضه بدون کوپن^{۶۰} ۱۰ ساله، اوراق قرضه بدون کوپن ۲۰ ساله، شاخص سهام^{۶۱} و شاخص املاک و مستغلات^{۶۲}. در این مدل، اوراق قرضه به‌صورت سالانه بازتخصیص می‌یابند به‌گونه‌ای که سبد همواره شامل اوراقی با سررسید ثابت باقی بماند. تمامی متغیرهای مالی به‌صورت اسمی بیان شده‌اند مگر خلاف آن ذکر شود. جریان نقدی تنزیل‌شده که بدهی کل شرکت را نشان می‌دهد؛ به‌صورت ذیل تعریف می‌شود:

$$(28) \quad \zeta_t = \begin{cases} D_t \frac{\max(A_t, G_t)}{C_t}, & t < T, \\ L_{T-1} \frac{\max(A_T, G_T)}{C_T}, & t = T, \end{cases}$$

برای $t = 1, \dots, T$ ، در رابطه (۲۸)، D_t تعداد فوت‌شدگان در بازه $(t-1, t]$ ، L_t تعداد بیمه‌گذاران زنده در زمان t ، A_t ارزش دارایی‌ها در زمان t ، G_t ارزش تضمین‌شده در زمان t و C_t ارزش حساب نقدی در زمان t است.

⁵⁸Variable Annuity Guarantee ⁵⁹Mortality ⁶⁰Zero Coupon Bond ⁶¹Equity Index ⁶²Real Estate Index

ارزش سبد دارایی‌ها برابر است با

$$(29) \quad A_t = \sum_{i=1}^4 U_t^i V_t^i = \sum_{i=1}^4 U_{t-1}^i V_t^{i-} + P_t, \quad t = 0, \dots, T.$$

با

$$(30) \quad V_t^i = \begin{cases} B(t, t+10), & i=1, \\ B(t, t+20), & i=2, \\ S_t, & i=3, \\ RE_t, & i=4, \end{cases} \quad V_t^{i-} = \begin{cases} B(t, t+9), & i=1, \\ B(t, t+19), & i=2, \\ V_t^i, & i=3, 4, \end{cases}$$

که در آن V_t^i قیمت دارایی i -ام در زمان t است و V_t^{i-} بازتخصیص اوراق قرضه با سررسید ثابت را برای $i=1, 2$ نمایش می‌دهد. همچنین $U_t^i = \frac{U_{t-1}^i V_t^{i-} + M_i P_t}{V_t^i}$ تعداد واحدهای دارایی i -ام نگهداری‌شده در بازه $[t, t+1)$ است ($U_{-1}^i := 0$). همچنین $B(t, s)$ ارزش اوراق قرضه در زمان t با سررسید s ، S_t ارزش شاخص سهام در زمان t ، RE_t ارزش شاخص املاک و مستغلات در زمان t ، M_i نسبت تخصیص دارایی‌ها به صورت $M = (\frac{1}{3}, \frac{1}{3}, \frac{1}{5}, \frac{2}{5})$ و P_t حق بیمه پرداختی در زمان t برای دوره $[t, t+1)$ است. در ادامه قرار می‌دهیم $P_t = 100$.

تضمین بیمه‌نامه به صورت ذیل تعریف می‌شود:

$$(31) \quad G_t = \begin{cases} 0, & t=0, \\ G_{t-1} + P_{t-1}, & t \geq 1. \end{cases}$$

در مورد متغیرهای جمعیت‌شناختی^{۶۳}، تعداد کل فوت‌شدگان و تعداد کل افراد زنده به صورت ذیل تعریف می‌شوند:

$$(32) \quad D_t = \sum_x D_t^x, \quad L_t = \sum_x L_t^x,$$

که در آن $D_t^x = L_{t-1}^x q_x(t)$ تعداد فوت‌شدگان در سن x در بازه $[t-1, t)$ وقتی $x \in (30, 70)$ است. همچنین $q_x(t)$ نرخ مرگ‌ومیر برای سن x در بازه $[t-1, t)$ و $L_t^x = L_{t-1}^x - D_t^x$ تعداد افراد زنده در سن x در زمان t با $L_0^x = 1000$ است. محرک تصادفی X دارای $d=5$ مؤلفه است که دو مؤلفه مربوط به مدل نرخ بهره و دیگر مؤلفه‌ها در ارتباط با سهام، املاک و مستغلات و مرگ‌ومیر تصادفی است. مدل املاک و مستغلات از ساختاری مشابه مدل سهام پیروی می‌کند، با این تفاوت که از محرک تصادفی مستقل و نوسان‌پذیری کمتری استفاده می‌کند. همچنین، مدل مرگ‌ومیر تصادفی بر اساس مدل لی-کارتر^{۶۴} تدوین شده است تا روند و نوسانات تصادفی را در طول زمان بازنمایی کند. نرخ مرگ‌ومیر به صورت ذیل مدل‌سازی می‌شوند:

$$(33) \quad q_x(t) = 1 - e^{-m_x(t)}, \quad m_x(t) = e^{a_x + b_x k(t)},$$

$$(34) \quad k(t) = k(t-1) - 0.365 + \varepsilon_t, \quad \varepsilon_t = 0.621 X_t^{(lc)}, \quad k(0) = -11.41,$$

که در آن $m_x(t)$ شدت مرگ‌ومیر^{۶۵} در زمان t برای سن x ، همچنین $X_t^{(lc)}$ مؤلفه‌ی مربوط به مرگ‌ومیر از محرک تصادفی X و a_x و b_x پارامترهای مدل لی-کارتر هستند [۱۹].

نتایج حاصل از قیمت‌گذاری بیمه‌ی عمر، نتایج مربوط به اختیار خرید اروپایی را تأیید می‌کنند. روش مارتینگل بازسازی عملکرد بسیار خوبی دارد. به‌ویژه مدل شبکه‌ی عصبی که در اکثر موارد بهترین نتایج را ارائه می‌دهد. اگرچه این قرارداد، برخی محدودیت‌های روش‌ها را بهتر نشان می‌دهد.

⁶³Demographic Variables ⁶⁴Lee-Carter Model ⁶⁵Force Of Mortality

در برآورد ارزش فعلی، جدول ۸ نشان می‌دهد که روش مونت کارلوی تودرتو همچنان بسیار کارآمد است. اما روش‌های مبتنی بر رگرسیون، دقت بهتری دارند. مدل شبکه‌ی عصبی در حالتی که تعداد نمونه‌ها بسیار پایین (۱۰۰۰) و بعد بالا ($T = 40$) باشد، عملکرد ضعیفی دارد و لذا بدترین نتایج به دست می‌آید. علت این امر بیش‌برازش است. مشابه مثال اختیار خرید اروپایی، از طریق اعتبارسنجی متقابل، $m = 101$ نوروں را انتخاب کرده‌ایم که یکی از آن‌ها صرفاً برای اربیی است به طوری که $A_{101} = 0$. بنابراین تعداد کل پارامترها برای $T = 5$ برابر با $2701 + 101 = 2600$ و برای $T = 40$ برابر است با $20201 + 101 = 20100$. همچنین مشاهده می‌کنیم که روش چندجمله‌ای کاهش بعد خطی (که با $p = 10$ محاسبه شده است) برتری خود را نسبت به مبنای کامل چندجمله‌ای نشان می‌دهد. مبنای چندجمله‌ای کامل دارای MAPE برابر ۶۱٪ است. زیرا این مبنای شامل ۳۲۷۶ پارامتر می‌شود (جدول ۱)، در حالی که تنها ۱۰۰۰ نمونه در دسترس بوده است. در مقابل، روش کاهش بعد خطی چندجمله‌ای دارای MAPE کمتر از ۱٪ است. زیرا تنها شامل ۲۸۶ پارامتر است.

برخلاف مثال اختیار خرید اروپایی که در آن شبکه‌های عصبی بهترین نتایج را داشتند؛ در این مثال روش کاهش بعد خطی چندجمله‌ای در برخی موارد نتایج بهتری ارائه می‌دهد. در مجموع، شبکه‌های عصبی ظرفیت بیشتری برای بهبود از طریق انتخاب ابرپارامترهای جایگزین دارند و می‌توان انتظار داشت که در مثال بیمه‌ی عمر نتایج بهتری به دست آید. در برخی موارد، ناکافی بودن تعداد نمونه‌های آموزشی منجر به بیش‌برازش و نتایج ضعیف خارج از نمونه می‌شود. برای مثال، در مبنای کامل چندجمله‌ای با $T = 5$ ، بهبود قابل توجهی در نتایج هنگام افزایش داده‌ی آموزشی از ۱۰۰۰ به ۵۰۰۰ نمونه مشاهده می‌شود. این مبنای دارای ۳۲۷۶ پارامتر است (جدول ۱)، به این معنا که در حالت ۱۰۰۰ نمونه، تعداد پارامترها بیشتر از تعداد نمونه‌هاست. همین اثر در مدل شبکه‌ی عصبی مارتینگل بازساز با $T = 40$ نیز دیده می‌شود زمانی که اندازه نمونه از ۱۰۰۰۰ به ۵۰۰۰۰ افزایش می‌یابد. این امر با توجه به تعداد بالای پارامترها (۲۰۲۰۱) قابل توضیح است.

در برخی حالت‌ها، به طور مثال در برآوردگر رگرسیون-آتی شبکه‌ی عصبی با $T = 5$ ، مقدار MAPE برای معیار زیان مورد انتظار، هنگام افزایش اندازه نمونه از ۵۰۰۰ به ۱۰۰۰۰ و ۵۰۰۰۰ نه تنها کاهش پیدا نمی‌کند، بلکه افزایش نیز می‌یابد (جدول ۹). این نتیجه در نگاه اول متناقض به نظر می‌رسد، زیرا انتظار می‌رود که با افزایش تعداد نمونه‌ها، دقت روش بهبود یابد. اما باید توجه داشت که معیاری که در فرآیند آموزش مدل کمینه می‌شود، با معیاری که در نهایت گزارش و مقایسه می‌شود یکسان نیست. در آموزش شبکه‌ی عصبی، تابع هدف به گونه‌ای طراحی شده است که خطا را بر روی کل توزیع جریان‌های نقدی کاهش دهد. یعنی مدل تلاش می‌کند تا در تمامی مقاطع توزیع، تقریب مناسبی ارائه دهد. در مقابل، خطای مورد استفاده در ارزیابی زیان مورد انتظار تنها بر بخش دم توزیع تمرکز دارد. به بیان دیگر، بهبود عملکرد مدل در کل توزیع لزوماً به معنای بهبود دقت در نواحی دم توزیع نیست. در واقع ممکن است مدل با افزایش نمونه‌ها در تقریب کلی بهتر عمل کند اما دقت آن در بخش‌های انتهایی توزیع کاهش یابد. زیرا این نواحی وزن کمتری در فرآیند بهینه‌سازی داشته‌اند. به همین دلیل است که در جدول ۹ با معیار MAPE افزایش خطا مشاهده می‌شود، در حالی که در خطای میانگین نسبی L^1 چنین رفتاری دیده نمی‌شود (جدول ۱۰). این پدیده نشان‌دهنده اهمیت انتخاب دقیق معیار خطای بهینه‌سازی در مسائل مدیریت ریسک است. زیرا با توجه به معیار انتخاب‌شده، عملکرد مدل در نواحی حساس توزیع می‌تواند بسیار متفاوت باشد.

جدول ۸. مقایسه‌ی خطای MAPE برای ارزش فعلی در مثال بدهی بیمه. اعداد بر حسب درصد هستند.

نمونه‌ها	مونت کارلو تودرتو	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی	شبکه‌ی عصبی	
		رگرسیون-اکنون	رگرسیون-آتی		رگرسیون-اکنون	رگرسیون-آتی
$T = 5$						
۱۰۰۰	۰/۳	۰/۲	۶۱/۰	< ۰/۱	۰/۳	۰/۱
۵۰۰۰	۰/۳	۰/۱	< ۰/۱	< ۰/۱	۰/۱	< ۰/۱
۱۰۰۰۰	۰/۲	۰/۱	< ۰/۱	< ۰/۱	۰/۱	< ۰/۱
۵۰۰۰۰	۰/۱	< ۰/۱	< ۰/۱	< ۰/۱	< ۰/۱	< ۰/۱
$T = ۴۰$						
۱۰۰۰	۰/۵	۰/۵		۰/۲	۰/۶	۶/۱
۵۰۰۰	۰/۳	۰/۲		۰/۱	۰/۲	۰/۳
۱۰۰۰۰	۰/۳	۰/۲		۰/۱	۰/۲	۰/۴
۵۰۰۰۰	۰/۲	۰/۱		< ۰/۱	۰/۱	۰/۱

جدول ۹. مقایسه‌ی خطای MAPE برای زیان مورد انتظار در مثال بدهی بیمه. اعداد بر حسب درصد هستند.

نمونه‌ها	مونت کارلو تودرتو	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی	شبکه‌ی عصبی	
		رگرسیون-اکنون	رگرسیون-آتی		رگرسیون-اکنون	رگرسیون-آتی
T = ۵						
۱۰۰۰	۳۰/۷	۸۰/۶	۶۱۳/۹	۷/۹	۱۹۸/۱	۲/۹
۵۰۰۰	۱۱/۱	۱۸/۶	۰/۵	۵/۳	۴۸/۷	۰/۵
۱۰۰۰۰	۹/۲	۱۰/۳	۰/۳	۵/۱	۲۵/۸	۰/۶
۵۰۰۰۰	۵/۷	۳/۳	۰/۲	۳/۷	۶/۶	۰/۸
T = ۴۰						
۱۰۰۰	۱۰۵/۴	۲۲۶/۴		۲۲/۹	۵۲۰/۷	۱۴/۴
۵۰۰۰	۳۲/۷	۶۴/۹		۵/۷	۱۴۷/۳	۱۱/۱
۱۰۰۰۰	۱۸/۷	۳۵/۹		۵/۹	۸۴/۴	۱۰/۵
۵۰۰۰۰	۱۰/۵	۱۰/۰		۷/۲	۱۳/۲	۰/۵

جدول ۱۰. مقایسه‌ی میانگین خطای L^1 در مثال بدهی بیمه. اعداد بر حسب درصد هستند.

نمونه‌ها	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی		شبکه‌ی عصبی	
	رگرسیون-آکون	رگرسیون-آتی	رگرسیون-آتی	رگرسیون-آتی	رگرسیون-آکون	رگرسیون-آتی
$T = 5$						
۱۰۰۰	۱/۶	۶۱/۰	۰/۲	۴/۷	۰/۱	۰/۱
۵۰۰۰	۰/۷	< ۰/۱	۰/۲	۱/۶	۰/۱	۰/۱
۱۰,۰۰۰	۰/۵	< ۰/۱	۰/۲	۱/۰	< ۰/۱	< ۰/۱
۵۰,۰۰۰	۰/۲	< ۰/۱	۰/۲	۰/۳	< ۰/۱	< ۰/۱
$T = 40$						
۱۰۰۰	۳/۳		۰/۸	۱۰/۰	۶/۱	۰/۱
۵۰۰۰	۱/۴		۰/۳	۳/۴	۰/۶	۰/۱
۱۰,۰۰۰	۱/۰		۰/۴	۲/۲	۰/۵	۰/۱
۵۰,۰۰۰	۰/۴		۰/۴	۰/۶	۰/۱	۰/۱

۹. تحلیل حساسیت نسبت به ابرپارامترها

تحلیل حساسیت روش کاهش بعد خطی چندجمله‌ای. روش کاهش بعد خطی چندجمله‌ای معرفی شده در بخش ۴.۵ دارای دو ابرپارامتر اصلی است: بعد هدف (p) و درجه چندجمله‌ای (δ). علاوه بر این، الگوریتم RBFGS که برای حل مسئله بهینه‌سازی (۷) در این روش به کار می‌رود، خود دارای پارامترهای متعددی است که مهم‌ترین آن‌ها نقطه‌ی شروع برای ماتریس A است. پارامتر درجه‌ی چندجمله‌ای δ همانند تمام تقریب‌های چندجمله‌ای، اثری قابل انتظار بر نتایج دارد. افزایش δ باعث غنی‌تر شدن فضای توابع تقریب می‌شود اما در مقابل تعداد پارامترهای مدل و هزینه‌ی محاسباتی افزایش می‌یابد. در ادامه، تمرکز بر دو پارامتر حساس‌تر روش کاهش بعد خطی چندجمله‌ای یعنی p و A قرار گرفته است.

انتخاب p . افزایش بعد هدف (p) نگاشت ویژگی ϕ_θ را غنی‌تر می‌کند و خطای تقریب را می‌کاهد. اما در عوض پارامترهای بیشتری را برای آموزش معرفی می‌کند و به داده‌ی آموزشی بیشتری نیاز دارد. افزایش p با تعداد داده‌ی آموزشی کم می‌تواند به افزایش خطا منجر شود. برای بررسی اثر این پارامتر در مثال بیمه عمر، یکی از روش‌های نقطه شروع که در ادامه با آن آشنا می‌شویم (روش قطری) و یک سررسید ثابت ($T = 5$) انتخاب شده است. نتایج در جداول ۱۱ و ۱۲ ارائه شده‌اند.

جدول ۱۱. مقایسه‌ی خطای MAPE برای زیان مورد انتظار در مثال بدهی بیمه برای مقادیر مختلف p در مبنای کاهش بعد خطی چندجمله‌ای. اعداد بر حسب درصد هستند.

نمونه‌ها	$p = 5$	$p = 10$	$p = 15$	$p = 20$
$T = 5$				
۱۰۰۰	۲۲/۳	۱۴/۹	۲۲۸/۰	۷۶۷/۹
۵۰۰۰	۲۹/۳	۱۷/۵	۸/۲	۹/۷
۱۰۰۰۰	۳۳/۳	۲۳/۶	۱۳/۲	۹/۰
۵۰۰۰۰	۳۷/۵	۳۱/۰	۲۱/۰	۱۰/۷

جدول ۱۲. مقایسه‌ی میانگین خطای L_1 در مثال بدهی بیمه برای مقادیر مختلف p در مبنای کاهش بعد خطی چندجمله‌ای. اعداد بر حسب درصد هستند.

نمونه‌ها	$p = 5$	$p = 10$	$p = 15$	$p = 20$
$T = 5$				
۱۰۰۰	۱/۰	۱/۰	۳/۱	۳۷/۳
۵۰۰۰	۱/۰	۰/۸	۰/۶	۰/۵
۱۰۰۰۰	۱/۰	۰/۹	۰/۶	۰/۶
۵۰۰۰۰	۱/۰	۰/۹	۰/۷	۰/۵

انتخاب نقطه‌ی شروع A_0 . در ادامه سه روش ذیل برای ساخت A_0 بررسی می‌شوند.

(۱) نقطه‌ی شروع تصادفی

در ساده‌ترین حالت، نقطه‌ی شروع A_0 به صورت تصادفی ساخته می‌شود. برای این منظور ابتدا dT نمونه‌ی مستقل از توزیع نرمال استاندارد $\mathcal{N}(0, 1)$ تولید کرده و آن‌ها را در یک ماتریس $B \in \mathbb{R}^{dT \times p}$ قرار می‌دهیم. سپس با نرمال‌سازی ستون‌ها، ماتریس شروع به صورت $A_0 = B(B^T B)^{-1/2}$ تعریف می‌شود. این ساختار تضمین می‌کند که A_0 روی خمینه استیفل $V_p(\mathbb{R}^{dT})$ قرار گیرد. هدف از این انتخاب آن است که الگوریتم RFBGS از نقطه‌ای خنثی و بدون سوگیری اولیه آغاز شود. با این حال، در عمل چنین نقطه‌ی شروعی اغلب عملکرد ضعیف‌تری دارد. زیرا در مسائل مالی، اهمیت زمانی جریان‌های نقدی یکسان نیست.

(۲) ماتریس قطری مستطیلی

در این روش، نقطه‌ی شروع A_0 به گونه‌ای ساخته می‌شود که تضمین کند هر یک از dT بعد ورودی، دست‌کم در یکی از p بعد خروجی، وزنی غیر صفر داشته باشد. به طور دقیق، ماتریس $B := A_0$ تنظیم می‌شود که در آن

$$B_{ij} = \begin{cases} 1 & i = j \wedge i \neq p, \\ \frac{1}{\sqrt{dT - p + 1}} & i \geq p, \\ 0 & \text{در غیر این صورت} \end{cases}$$

این نقطه‌ی شروع بازتاب‌دهنده‌ی این ایده است که ابعاد نزدیک‌تر به زمان حال همچنان به صورت جداگانه نمایان می‌شوند، در حالی که ابعاد دورتر به طور جمعی نمایش داده می‌شوند.

(۳) فشرده‌سازی زمانی^{۶۶}

ایده همان تجمیع ابعاد دورتر در یک بعد خروجی است. با این تفاوت که تنها در امتداد زمان (بعد T) تجمیع انجام می‌شود، نه در ابعاد d . برای مثال اختیار خرید اروپایی $T = 5$ ، $d = 3$ و $p = 3$ یک ساختار نرمال‌شده به صورت ذیل است:

$$A_0 = \begin{bmatrix} \frac{1}{\sqrt{T}} & 0 & 0 \\ 0 & \frac{1}{\sqrt{T}} & 0 \\ 0 & 0 & \frac{1}{\sqrt{T}} \\ \vdots & \vdots & \vdots \\ \frac{1}{\sqrt{T}} & 0 & 0 \\ 0 & \frac{1}{\sqrt{T}} & 0 \\ 0 & 0 & \frac{1}{\sqrt{T}} \end{bmatrix} \quad (35)$$

⁶⁶Folding

در مثال بیمه عمر، در سناریوی $T = 5$ ، $d = 5$ و $p = 10$ ، یک نقطه شروع مناسب در نمایش بلوکی به صورت ذیل است:

$$(26) \quad A_0 = \begin{bmatrix} \mathbf{I}_d & [\circ]_{d \times (p-d)} \\ [\circ]_{d \times d} & \frac{1}{\sqrt{T-1}} \mathbf{I}_{(p-d)} \\ \vdots & \vdots \end{bmatrix} \quad (\text{در } T-1 \text{ سطر زمانی تکرار می شود}).$$

که در آن $\mathbf{I}_d$ ماتریس همانی و $[\circ]$ ماتریس صفر با ابعاد مشخص است. این ساختار دو ویژگی مهم دارد: (۱) هر بعد ورودی دستکم یک ستون خروجی با وزن غیر صفر دارد که این یعنی پوشش کامل ابعاد. (۲) فشردسازی فقط در بعد زمان انجام می شود و نرمال سازی باعث می شود که نرم این ستون های تجمعی برابر ۱ باقی بماند. در جداول ۱۳ و ۱۴ عملکرد سه نوع نقطه شروع را مقایسه کردیم. نتایج نشان می دهند که روش فشردسازی زمانی بهترین عملکرد را در میان تمامی نقاط شروع دارد.

جدول ۱۳. مقایسه ی خطای MAPE برای زیان مورد انتظار در مثال بدهی بیمه برای نقاط شروع مختلف در مبنای چند جمله ای با کاهش بعد خطی. اعداد بر حسب درصد هستند.

نمونه ها	فشردسازی زمانی	قطری	تصادفی
$T = 5$			
۱۰۰۰	۷/۹	۱۴/۹	۱۶/۶
۵۰۰۰	۵/۳	۱۷/۵	۳۱/۵
۱۰۰۰۰	۵/۱	۲۲/۶	۳۱/۸
۵۰۰۰۰	۳/۷	۳۱/۰	۳۲/۴
$T = 40$			
۱۰۰۰	۲۲/۹	۳۸/۰	۶۳/۹
۵۰۰۰	۵/۷	۵۴/۷	۶۹/۴
۱۰۰۰۰	۵/۹	۵۳/۰	۶۸/۲
۵۰۰۰۰	۷/۲	۴۹/۲	۶۷/۰

جدول ۱۴. مقایسه ی میانگین خطای L_1 برای نقاط شروع مختلف در مبنای چند جمله ای با کاهش بعد خطی در مثال بیمه عمر. اعداد بر حسب درصد هستند.

نمونه ها	فشردسازی زمانی	قطری	تصادفی
$T = 5$			
۱۰۰۰	۰/۲	۱/۰	۱/۲
۵۰۰۰	۰/۲	۰/۸	۱/۰
۱۰۰۰۰	۰/۲	۰/۹	۱/۰
۵۰۰۰۰	۰/۲	۰/۹	۱/۰
$T = 40$			
۱۰۰۰	۰/۸	۱/۷	۱/۷
۵۰۰۰	۰/۳	۱/۳	۱/۵
۱۰۰۰۰	۰/۴	۱/۳	۱/۵
۵۰۰۰۰	۰/۴	۱/۲	۱/۵

تحلیل حساسیت روش شبکه عصبی. روش شبکه عصبی دارای یک ابرپارامتر اصلی یعنی عرض لایه پنهان (h) است. در نتایج اصلی مثال بیمه عمر، تمامی حالت‌ها در لایه‌ی پنهان از پهنای ثابت $h = 100$ استفاده شده است. این مقدار بر اساس یک آزمون حساسیت روی مقادیر مختلف (۱۰، ۵۰، ۱۰۰ و ۲۰۰) انتخاب شده است. انتخاب پهنای لایه می‌تواند به روش‌های متفاوتی صورت گیرد و در این بخش برخی از این روش‌ها در مثال بیمه عمر بررسی می‌شوند. خلاصه‌ی این نتایج در جداول ۱۵ و ۱۶ آمده است. نتایج نشان می‌دهند که برخی از این روش‌ها حتی عملکرد بهتری نسبت به انتخاب ما در نتایج اصلی داشته‌اند؛ این امر بیانگر آن است که مبنای شبکه‌ی عصبی، که در مقایسه‌های ما بهترین عملکرد را داشته، هنوز ظرفیت بهبود و ارتقای بیشتر را داراست.

اولین روش استفاده از پهنای حداقل نظری^{۶۷} برای عرض لایه پنهان است که این روش در [۱۲] توصیف شده است. این روش عملکرد مناسبی نشان نمی‌دهد که این موضوع، تناقضی با حداقل نظری ندارد. زیرا تعریف حداقل نظری مبتنی بر شبکه‌هایی با عمق دلخواه است. روش دوم آن است که پهنای لایه‌ی شبکه را به‌گونه‌ای انتخاب کنیم که ابعاد فضای پارامترها با ابعاد متناظر در روش کاهش بعد خطی چندجمله‌ای یکسان شود. این روش در بعد بالا یعنی $T = 40$ نتایج بسیار خوبی دارد. روش سوم آن است که پهنای لایه‌ی شبکه‌ی عصبی به‌گونه‌ای انتخاب شود که با تعداد توابع پایه m مطابقت داشته باشد. نتایج حاصل از این انتخاب مشابه نتایج دو مورد قبل است و بهبود چشمگیری نسبت به آن‌ها مشاهده نمی‌شود.

جدول ۱۵. مقایسه‌ی خطای MAPE برای زیان مورد انتظار شرطی در مثال بدهی بیمه برای عرض‌های مختلف لایه پنهان شبکه‌ی عصبی. اعداد بر حسب درصد هستند.

نمونه‌ها	$h = 100$ (ثابت)	حداقل عرض	برابر با ابعاد روش کاهش بعد خطی	ابعاد برابر m
$T = 5$				
۱۰۰۰	۲٫۹	۲٫۲	۳٫۱	۳٫۹
۵۰۰۰	۰٫۵	۳٫۲	۴٫۱	۰٫۴
۱۰۰۰۰	۰٫۶	۳٫۳	۴٫۲	۰٫۲
۵۰۰۰۰	۰٫۸	۳٫۳	۴٫۲	۰٫۳
$T = 40$				
۱۰۰۰	۱۴٫۴	۱۴٫۵	۱۷٫۵	۱۴٫۱
۵۰۰۰	۱۱٫۱	۱۲٫۵	۲٫۷	۱۲٫۸
۱۰۰۰۰	۱۰٫۵	۱۲٫۶	۲٫۰	۱۳٫۲
۵۰۰۰۰	۰٫۵	۱٫۴	۲٫۲	۲٫۴

جدول ۱۶. مقایسه‌ی میانگین خطای L_1 در مثال بدهی بیمه برای پهنای لایه‌های مختلف در شبکه‌ی عصبی. اعداد بر حسب درصد هستند.

⁶⁷Minimal Width

نمونه‌ها	$h = 100$ (ثابت)	حداقل پهنا	برابر با ابعاد روش کاهش بعد خطی	ابعاد برابر m
$T = 5$				
۱۰۰۰	۰/۱	۰/۱	۰/۱	۰/۱
۵۰۰۰	۰/۱	۰/۱	۰/۱	۶/۱
۱۰۰۰۰	۰/۱	۰/۱	۰/۱	۶/۴
۵۰۰۰۰	۰/۱	۰/۱	۰/۱	۶/۵
$T = 40$				
۱۰۰۰	۰/۱	$< 0/1$	$< 0/1$	۰/۱
۵۰۰۰	۰/۱	۰/۱	۰/۱	۰/۶
۱۰۰۰۰	۰/۱	۰/۱	۰/۱	۰/۸
۵۰۰۰۰	۰/۱	۰/۱	۰/۱	۰/۵

۱۰. مقایسه زمان اجرا

در این بخش بررسی می‌کنیم که اجرای مراحل آموزش و پیش‌بینی برای هر مدل چه مدت زمان نیاز دارد. نتایج برای محاسبه $\hat{V}_1(x_i)$ برای یک میلیون نمونه اعتبارسنجی است. این نتایج به‌وضوح اثر کاهش بعد را در هزینه محاسباتی روش مارتینگل بازساز آشکار می‌کند. همان‌طور که انتظار می‌رود، روش‌های رگرسیون-اکنون سریع‌تر از روش‌های رگرسیون-آتی هستند. کندترین روش، کاهش بعد خطی چندجمله‌ای است که در مسائل با بعد بالا می‌تواند تا ۲۵ دقیقه زمان ببرد. زمان اجرای روش کاهش بعد خطی چندجمله‌ای به شدت به پارامتر p حساس است. به‌عنوان مثال، برای $T = 40$ و اندازه نمونه ۱۰۰۰، حل مسئله با $p = 5$ تنها ۳۳ ثانیه طول می‌کشد در حالی‌که با $p = 10$ این زمان به ۲۷۹/۱ ثانیه می‌رسد. این در حالیتیست که مدل شبکه عصبی را می‌توان نسبتاً سریع حل کرد.

جدول ۱۷. مقایسه‌ی زمان اجرا در مثال اختیار خرید اروپایی (بر حسب ثانیه).

نمونه‌ها	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی		شبکه‌ی عصبی
	رگرسیون-اکنون	رگرسیون-آتی	رگرسیون-آتی	رگرسیون-اکنون	رگرسیون-آتی
$T = 5$					
۱۰۰۰	۰/۹	۱/۵	۱/۴	۹/۵	۱۸/۴
۵۰۰۰	۰/۹	۲/۳	۱/۶	۱۶/۱	۲۱/۱
۱۰۰۰۰	۰/۹	۳/۰	۲/۱	۲۴/۰	۲۴/۴
۵۰۰۰۰	۱/۰	۱۳/۴	۷/۱	۹۲/۶	۴۹/۴
$T = 40$					
۱۰۰۰	۰/۹		۷/۲	۱۰/۰	۱۸/۵
۵۰۰۰	۰/۹		۲۲/۳	۱۵/۵	۲۴/۴
۱۰۰۰۰	۱/۰		۵۱/۷	۲۴/۰	۳۰/۰
۵۰۰۰۰	۱/۴		۲۷۱/۷	۱۹/۱	۷۶/۱

جدول ۱۸. مقایسه‌ی زمان اجرا در مثال بدهی بیمه (بر حسب ثانیه).

نمونه‌ها	پایه‌ی چندجمله‌ای کامل		کاهش بعد خطی		شبکه‌ی عصبی	
	رگرسیون-آکنون	رگرسیون-آتی	رگرسیون-آتی	رگرسیون-آکنون	رگرسیون-آتی	رگرسیون-آکنون
$T = 5$						
۱۰۰۰	۰/۸	۴/۳	۳۰/۴	۴/۰	۵/۵	
۵۰۰۰	۰/۹	۵۴/۴	۱۱۳/۲	۱۱/۱	۸/۲	
۱۰۰۰۰	۱/۱	۴۵/۳	۴۹/۶	۱۹/۸	۱۱/۹	
۵۰۰۰۰	۳/۰	۱۰۸/۱	۲۳۸/۶	۹۵/۱	۴۲/۷	
$T = 40$						
۱۰۰۰	۱/۳		۲۷۹/۱	۴/۶	۶/۳	
۵۰۰۰	۲/۶		۶۹۸/۹	۱۲/۹	۱۵/۷	
۱۰۰۰۰	۴/۳		۱۸۱۵/۱	۲۲/۹	۲۶/۵	
۵۰۰۰۰	۱۷/۹		۱۴۷۲/۸	۳۱/۷	۱۱۵/۵	

۱۱. نتیجه‌گیری

در چارچوب نیاز روزافزون به محاسبات دقیق و سریع برای قیمت‌گذاری سبدها و مدیریت ریسک، در این مقاله روشی داده‌محور برای ساخت مارتینگل‌های بازساز معرفی گردید. مدل پیشنهادی در چارچوب یادگیری نظارت‌شده، در مسائل با ابعاد بالا عملکرد مؤثری دارد. استفاده از مارتینگل بازساز مبتنی بر یادگیری ماشین، ضمن کاهش قابل توجه هزینه محاسباتی، دقتی قابل توجه و در بسیاری از موارد برتر از روش شبیه‌سازی مونت کارلوی تودرتوی استاندارد فراهم می‌کند.

نتایج نظری برای دو مثال کاربردی شامل یک اختیار خرید اروپایی و یک محصول بیمه‌ای با جریان‌های نقدی وابسته به مسیر مورد بررسی قرار گرفت. مارتینگل‌های بازساز در این دو مثال، نسبت به سایر روش‌های رایج در صنعت مالی، عملکرد بهتری از خود نشان دادند. این برتری با مقایسه‌های گسترده و تحلیل‌های حساسیت به‌خوبی نشان داده شد.

در نهایت، چالش نظری بنیادین در مدل پیشنهادی، تحلیل سازگاری آماری و رفتار همگرایی برآوردگرهای حاصل از نگاشت‌های کاهش‌بعد و ساختارهای مبتنی بر شبکه‌های عصبی است. این مسئله از آن جهت اهمیت دارد که در چنین ساختارهایی، نه تنها یکنایی جواب بهینه لزوماً برقرار نیست، بلکه حتی وجود جواب نیز قابل تضمین نخواهد بود. پرداختن به این پرسش‌ها می‌تواند زمینه‌ساز توسعه‌ی تضمین‌های نظری مستحکم‌تر و گسترش افق‌های پژوهشی جدید در این حوزه باشد.

مراجع

- [1] G. Andreatta and S. Corradin, Valuing mortality risk in life insurance contracts, Working paper, 2003.
- [2] J. F. Carriere, Valuation of the early-exercise price for options using simulations and nonparametric regression, *Insurance: Mathematics and Economics*, **19** (1996), 19–30.
- [3] C. Castellani, P. De Angelis, M. Peracchi and M. Vellekoop, Solvency capital requirement: Neural network versus least-squares Monte Carlo, *Insurance: Mathematics and Economics*, **80** (2018), 147–161.
- [4] J. De Spiegeleer, D. B. Madan, S. Reyners and W. Schoutens, Machine learning for quantitative finance: fast derivative pricing, hedging and fitting, *Quantitative Finance*, **18** (2018), 1635–1643.
- [5] F. Delbaen and W. Schachermayer, *The Mathematics of Arbitrage*, Springer Finance, Springer-Verlag, Berlin, 2006.
- [6] Q. D. Duong, Application of Bayesian penalized spline regression for internal modeling in life insurance, *European Actuarial Journal*, **9** (2019), 67–107.
- [7] L. Fernandez-Arjona and D. Filipovi'c, A machine learning approach to portfolio pricing and risk management for high-dimensional problems, *Mathematical Finance*, **32** (2022), 982–1019.

- [8] L. Fernandez-Arjona and D. Filipovi'c, Reproducibility data for tables in "A machine learning approach to portfolio pricing and risk management for high-dimensional problems", Dataset, Zenodo, 2022. <https://zenodo.org/records/3837381>
- [9] P. Glasserman and B. Yu, Simulation for American options: Regression now or regression later, in *Monte Carlo and Quasi-Monte Carlo Methods 2002*, Springer, 2004, 213–226.
- [10] M. B. Gordy and S. Juneja, Nested simulation in portfolio risk measurement, *Management Science*, **56** (2010), 1833–1848.
- [11] H. Ha and D. Bauer, Least-squares Monte Carlo for insurance liability valuation with data-driven basis selection, *European Actuarial Journal*, **10** (2020), 521–548.
- [12] B. Hanin and M. Sellke, Approximating continuous functions by ReLU nets of minimal width, arXiv preprint arXiv:1710.11278, 2017.
- [13] D. Hainaut, American option pricing with model constrained Gaussian process regressions, *Applied Mathematics and Computation*, **512** (2026), 129747.
- [14] L. J. Hong, Z. Hu and G. Liu, Monte Carlo methods for value-at-risk and conditional value-at-risk: A review, *ACM Transactions on Modeling and Computer Simulation*, **27** (2017), 1–37.
- [15] K. Hornik, M. Stinchcombe and H. White, Multilayer feedforward networks are universal approximators, *Neural Networks*, **2** (1989), 359–366.
- [16] W. Huang, K. A. Gallivan and P.-A. Absil, A Broyden class of quasi-Newton methods for Riemannian optimization, *SIAM Journal on Optimization*, **25** (2015), 1660–1685.
- [17] L. Jiang and C. Xu, A new options pricing method: semi-stochastic kernel regression method with constraints, *International Journal of Computer Mathematics*, **100** (2023), 1809–1820.
- [18] R. Kan, From moments of sum to moments of product, *Journal of Multivariate Analysis*, **99** (2008), 542–554.
- [19] R. D. Lee and L. R. Carter, Modeling and forecasting U.S. mortality, *Journal of the American Statistical Association*, **87** (1992), 659–671.
- [20] S. H. Lee and P. W. Glynn, Computing the distribution function of a conditional expectation via Monte Carlo: Discrete and continuous cases, in *Proceedings of the Winter Simulation Conference 2003*, (2003), 356–364.
- [21] N. Lehdili, P. Oswald and H. D. Nguyen, Performance-Enhancing Market Risk Calculation Through Gaussian Process Regression and Multi-Fidelity Modeling, *Computation*, **13** (2025), 134.
- [22] F. A. Longstaff and E. S. Schwartz, Valuing American options by simulation: A simple least-squares approach, *The Review of Financial Studies*, **14** (2001), 113–147.
- [23] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot and E. Duchesnay, Scikit-learn: Machine learning in Python, *Journal of Machine Learning Research*, **12** (2011), 2825–2830.
- [24] A. Pelsser and N. Schweizer, The difference between LSMC and replicating portfolio in insurance liability modelling, *European Actuarial Journal*, **6** (2016), 441–494.
- [25] P. Petersen, M. Raslan and F. Voigtlaender, Topological properties of the set of functions generated by neural networks of fixed size, *Foundations of Computational Mathematics*, **21** (2021), 375–444.
- [26] J. Risk and M. Ludkovski, Gaussian process regression for simulation metamodeling, *Operations Research*, **66** (2018), 802–823.
- [27] T. J. Sullivan, *Introduction to Uncertainty Quantification*, Texts in Applied Mathematics **63**, Springer, Cham, 2015.
- [28] Z. Zhang, S. Zohren and S. Roberts, Deep learning for portfolio optimization, *Journal of Financial Data Science*, **3** (2021), 8–20.

سعید سقایی طوقچی

اصفهان، خیابان هزار جریب، دانشگاه اصفهان، دانشکده ریاضی و آمار، گروه ریاضی کاربردی و علوم کامپیوتر
saeidsaghaei۱۰@gmail.com



سعید سقایی در سال ۱۳۹۷ وارد مقطع کارشناسی رشته ریاضیات و کاربردها در دانشگاه صنعتی اصفهان شد و در سال ۱۴۰۴ مدرک کارشناسی ارشد رشته ریاضیات مالی را از دانشگاه اصفهان اخذ کرده است.

بهاره اختری

اصفهان، خیابان هزار جریب، دانشگاه اصفهان، دانشکده ریاضی و آمار، گروه ریاضی کاربردی و علوم کامپیوتر
b.akhtari@mcs.ui.ac.ir
b.akhtari@ipm.ir



بهاره اختری عضو هیئت علمی گروه ریاضی مالی دانشکده ریاضی و آمار دانشگاه اصفهان است. وی همچنین پژوهشگر مقیم پژوهشگاه دانش‌های بنیادی است.