

تحلیل هوشمند فرآیندهای سازمانی با یادگیری ماشین برای استخراج بینش‌های عملیاتی و راهبردی

علی قاسمی کیان^۱، مجید میرزائی قزانی دانشیار^۲

۱- کارشناسی ارشد مهندسی صنایع، دانشگاه صنعتی خواجه نصیرالدین طوسی و A.kian@rahroshdpaya.com

۲- مهندسی صنایع گروه مهندسی مالی دانشگاه صنعتی خواجه نصیرالدین طوسی و majidmirzaee@kntu.ac.ir

خلاصه

در این پژوهش، با بهره‌گیری از تکنیک‌های پیشرفته یادگیری ماشین، رویکردی نوین برای تحلیل هوشمند فرآیندهای سازمانی به منظور استخراج بینش‌های عملیاتی و استراتژیک ارائه شده است. هدف اصلی این مطالعه، پیش‌بینی دقیق خروجی فرآیندها و شناسایی گلوگاه‌های سیستمی است که منجر به ناکارآمدی در تخصیص منابع می‌شود. بدین منظور، از معماری ترکیبی* متشکل از مدل‌های یادگیری پایه و یک فرادسته‌بند نهایی استفاده شده و به منظور تفسیرپذیری نتایج و درک عوامل تأثیرگذار بر پیش‌بینی‌ها، از متدولوژی هوش مصنوعی تفسیرپذیر مبتنی بر مقادیر بهره گرفته شد. ارزیابی مدل با استفاده از معیارهای چندگانه نشان می‌دهد که رویکرد پیشنهادی در مقایسه با مدل‌های تک‌بخش، دقت بالاتری در شناسایی فرآیندهای مخاطره‌آمیز داشته و امکان اتخاذ تصمیمات استراتژیک داده‌محور را برای مدیران عملیاتی فراهم می‌آورد. دستاوردهای این تحقیق، پتانسیل بالای الگوریتم‌های یادگیری جمعی در خودکارسازی نظارت بر فرآیندهای سازمانی و کاهش هزینه‌های عملیاتی را به اثبات می‌رساند.

کلمات کلیدی: تحلیل فرآیندهای سازمانی، یادگیری ماشین، یادگیری جمعی، هوش مصنوعی تفسیرپذیر، مقادیر SHAP، بهینه‌سازی منابع، گلوگاه‌های فرآیندی، پیش‌بینی عملیاتی.

۱. مقدمه

در عصر تحول دیجیتال، سازمان‌ها با حجم عظیمی از داده‌های تولید شده توسط سیستم‌های مدیریت فرآیند کسب‌وکار مواجه هستند. فرآیندهای سازمانی به عنوان ستون فقرات عملکرد عملیاتی، نقشی حیاتی در تعیین بهره‌وری، هزینه‌های جاری و در نهایت جایگاه رقابتی سازمان ایفا می‌کنند [۱]. با این حال، ماهیت پیچیده و پویای این فرآیندها، تحلیل سنتی آن‌ها را با چالش‌های اساسی روبرو کرده است. روش‌های کلاسیک تحلیل فرآیند، علی‌رغم کاربرد در نقشه‌برداری فعالیت‌ها، اغلب در پیش‌بینی دقیق نتایج و شناسایی هوشمند گلوگاه‌های استراتژیک در محیط‌های دارای عدم قطعیت، ناتوان هستند. چالش اصلی در تحلیل فرآیندهای امروزی، گذار از توصیف آنچه رخ داده است به سمت پیش‌بینی آنچه رخ خواهد داد و

* Stacking Ensemble

مهم‌تر از آن، درک چرایی رخدادها است. بسیاری از مدل‌های پیشرفته یادگیری ماشین که امروزه برای پیش‌بینی خروجی فرآیندها به کار می‌روند، به دلیل ماهیت جعبه‌سیاه، فاقد شفافیت لازم برای پذیرش توسط مدیران تصمیم‌گیرنده هستند [۲]. عدم امکان تفسیرپذیری تصمیمات مدل، مانع اصلی در انتقال از تئوری به کاربردهای عملیاتی سازمانی محسوب می‌شود. این مقاله با هدف رفع این خلأ تحقیقاتی، یک چارچوب هوشمند و تفسیرپذیر برای تحلیل فرآیندهای سازمانی ارائه می‌دهد. نوآوری اصلی این پژوهش در تلفیق معماری قدرتمند یادگیری جمعی با تکنیک‌های هوش مصنوعی تفسیرپذیر نهفته است [۳]. در این رویکرد، با ترکیب خرد جمعی چندین الگوریتم یادگیری شامل XGBoost، Random Forest و MLP، دقت پیش‌بینی خروجی‌های فرآیند به طور معناداری ارتقا یافته است [۴]. سپس، با به‌کارگیری تکنیک SHAP (مقادیر شاپلی)، فرآیند تصمیم‌گیری مدل برای هر نمونه بازگشایی شده تا علاوه بر پیش‌بینی ناهنجاری‌ها و گلوگاه‌ها، دلایل ریشه‌ای وقوع آن‌ها نیز برای ذینفعان مشخص گردد [۵].

۳. پیشینه تحقیق و بیان ضرورت

در دهه‌های اخیر، تحلیل هوشمند فرآیندهای سازمانی به عنوان ابزاری استراتژیک برای بهبود کارایی مورد توجه قرار گرفته است. با این حال، تکامل این حوزه نشان‌دهنده چالش‌های متعددی در مسیر گذار از تحلیل‌های توصیفی به تحلیل‌های پیش‌بینانه و تجویزی است.

- در مقاله [۶]، پایه و اساس داده‌کاوی فرآیند به عنوان یک رشته علمی تبیین شد که عمدتاً بر کشف مدل‌های فرآیندی از روی لاگ‌های رویداد تمرکز داشت. اگرچه این پژوهش انقلابی در درک بصری فرآیندها ایجاد کرد، اما محدودیت اصلی آن در عدم توانایی برای پیش‌بینی وضعیت آینده فرآیند و رفتار سیستم در مواجهه با تغییرات بود. پژوهش حاضر با گذر از دیدگاه صرفاً توصیفی این مطالعه، به سمت پیش‌بینی هوشمند نتایج گام برداشته است [۷].
- در مقاله [۸]، استفاده از شبکه‌های عصبی حافظه طولانی-کوتاه مدت (LSTM) برای پیش‌بینی زمان باقی‌مانده و فعالیت بعدی در فرآیندها معرفی شد. این کار دقت پیش‌بینی را به طور قابل‌توجهی افزایش داد، اما مدل‌های عمیق به‌کاررفته به دلیل «جعبه‌سیاه» بودن، قابلیت تفسیرپذیری برای مدیران را نداشتند. در مقابل، مدل ما با بهره‌گیری از معماری Stacking و ترکیب آن با مقادیر SHAP، علاوه بر حفظ دقت، شفافیت تصمیم‌گیری را نیز تضمین می‌کند.
- در مقاله [۹]، تلاش شد تا با استفاده از مدل‌های تولیدی، شبیه‌سازی فرآیندها با دقت بالا انجام شود. با این حال، رویکرد پیشنهادی ایشان در مواجهه با داده‌های بسیار نامتوازن که در سازمان‌ها رایج است، کارایی لازم را نداشت. کار ما با استفاده از معیارهای Cost-Sensitive و ترکیب مدل‌های ناهمگون در معماری Stacking، پایداری مدل را در شرایط داده‌ای نامتوازن به شکل مؤثری افزایش داده است [۱۰].
- در مقاله [۱۱]، یک مطالعه جامع روی روش‌های نظارت پیش‌بینانه انجام شد و مدل‌های ترکیبی ساده‌ای پیشنهاد گردید. با وجود قوت نظری این مطالعه، خروجی‌های آن تنها به پیش‌بینی‌های عددی محدود ماند و راهکاری برای اقدام اصلاحی در اختیار مدیران قرار نداد. بر خلاف این مطالعه، خروجی پژوهش ما شامل تحلیل ریشه‌ای است که مستقیماً به مدیران می‌گوید کدام متغیرها (مثلاً هزینه یا زمان) عامل شکست فرآیند بوده‌اند.

- در مقاله [۱۲]، مدل‌هایی بر پایه قوانین تصمیم‌گیری برای ایجاد قابلیت تفسیر در فرآیندها ارائه شد. علی‌رغم نوآوری در تفسیرپذیری، دقت پیش‌بینی این مدل‌ها در مقایسه با مدل‌های مدرن یادگیری ماشین بسیار پایین بود. راهکار پیشنهادی ما با استفاده از تکنیک مدرن XAI، تقابل میان «دقت بالای یادگیری ماشین» و «تفسیرپذیری» را از بین برده و تعادلی بهینه میان این دو قطب ایجاد کرده است [۱۳].

تمامی پژوهش‌های فوق یا درگیر پیچیدگی مدل‌های جعبه‌سیاه بوده‌اند که قابلیت اعتماد مدیریتی ندارند، یا با کاهش دقت پیش‌بینی برای دسترسی به تفسیرپذیری دست‌وپنج نرم کرده‌اند. پژوهش فعلی با ارائه یک معماری Stacking تفسیرپذیر، نه تنها دقت پیش‌بینی را به اوج رسانده، بلکه با فراهم کردن ابزارهای بصری SHAP برای تحلیل هر نمونه، نخستین چارچوب عملیاتی را برای تبدیل داده‌کاوی به اقدامات استراتژیک در سازمان فراهم آورده است.

۳. روش تحقیق و معماری مدل

این پژوهش از یک چارچوب داده‌محور برای شناسایی الگوهای پنهان در فرآیندهای سازمانی استفاده می‌کند. مراحل انجام کار شامل گردآوری داده‌ها، طراحی ساختار یادگیری ماشین و تدوین معیارهای ارزیابی به شرح زیر است:

۳.۱. منبع و ماهیت دیتاست

دیتاست مورد استفاده در این مقاله، یک مجموعه داده فرآیندی است که با شبیه‌سازی دقیق ویژگی‌های واقعی لاگ‌های سازمانی تولید شده است [۱۴]. این داده‌ها منعکس‌کننده جریان واقعی کار در یک سازمان هستند و شامل متغیرهای عملیاتی کلیدی از قبیل:

- شناسه فرآیند (process_id): برای دنبال کردن مسیرهای منحصربه‌فرد هر پروژه.
- فعالیت (activity): گام‌های عملیاتی انجام شده در طول فرآیند.
- دپارتمان (department): واحد سازمانی مسئول انجام هر فعالیت.
- زمان (duration_hours) و هزینه (cost): شاخص‌های کمی که نشان‌دهنده مصرف منابع در هر مرحله هستند.
- نتیجه (outcome): برچسب نهایی که نشان می‌دهد آیا فرآیند به موفقیت منجر شده یا با شکست مواجه شده است.

این دیتاست به گونه‌ای طراحی شده است که عدم توازن رایج در سازمان‌ها را پوشش دهد؛ بدین معنا که تعداد موارد موفقیت در آن به مراتب بیشتر از موارد شکست است که این امر چالش واقعی برای مدل‌های یادگیری ماشین ایجاد می‌کند.

۳.۲. معماری یادگیری ماشین (Stacking Ensemble)

برای تحلیل این داده‌ها، از معماری انباشته (Stacking Ensemble) استفاده کردیم. این مدل به جای تکیه بر یک الگوریتم واحد، از ترکیب خرد جمعی چندین مدل متفاوت بهره می‌برد:

- مدل های پایه: ما از Random Forest برای شناسایی الگوهای غیرخطی، XGBoost برای دقت بالا در داده های جدولی، MLP برای درک روابط پیچیده استفاده کردیم [۱۵].
- مدل متا (Meta-Learner): در مرحله نهایی، یک مدل رگرسیون لجستیک به عنوان ناظر عمل می کند. این مدل یاد می گیرد که به خروجی هر یک از مدل های پایه بر اساس میزان قابل اعتماد بودن آن ها وزن بدهد و در نهایت یک پیش بینی واحد و دقیق تر ارائه کند. این رویکرد باعث می شود که خطاهای احتمالی یک مدل توسط سایر مدل ها پوشش داده شود.

۳.۳. معیارهای ارزیابی و تحلیل عملیاتی

از آنجا که در یک محیط سازمانی، هزینه نادیده گرفتن یک شکست احتمالی بسیار سنگین تر از خطای مثبت کاذب است، ما تنها به دقت (Accuracy) اکتفا نکردیم:

- معیارهای تعادلی: از معیارهایی مانند F1-Score استفاده کردیم تا بین دقت پیش بینی و توانایی کشف تمام موارد خطر، تعادل ایجاد کنیم.
- معیار اختصاصی کسب و کار (Business Score): ما یک معیار سفارشی طراحی کردیم که به مدل اجازه می دهد اولویت را به شناسایی مواردی بدهد که پتانسیل ایجاد هزینه های عملیاتی بالا یا تاخیرهای طولانی دارند.
- تفسیرپذیری (SHAP): برای اینکه نتایج مدل برای مدیران قابل فهم باشد، از ابزار SHAP استفاده کردیم. این ابزار به ما می گوید که مثلاً هزینه بالا در دپارتمان X چه سهمی در شکست نهایی پروژه داشته است. این یعنی خروجی مدل نه تنها یک عدد، بلکه یک نقشه راه برای اصلاح فرآیند است.

۴. یافته ها و بحث

در این بخش، عملکرد مدل پیشنهادی (Stacking Ensemble) در مقایسه با سایر الگوریتم ها ارزیابی شده و با استفاده از تکنیک های هوش مصنوعی تفسیرپذیر (XAI)، عوامل کلیدی تأثیرگذار بر فرآیندها تحلیل می شوند.

جدول ۱ - مقایسه عملکرد مدل های یادگیری ماشین بر اساس شاخص های آماری و عملیاتی

Model	Accuracy	F1-Score	ROC-AUC	MCC	Business Score
-------	----------	----------	---------	-----	----------------

Stacking Ensemble	0.80	0.89	0.48	0.0	0.90
Random Forest	0.81	0.89	0.47	0.13	0.90
XGBoost	0.79	0.88	0.46	0.0029	0.44
Neural Network	0.73	0.84	0.48	0.0075	0.12

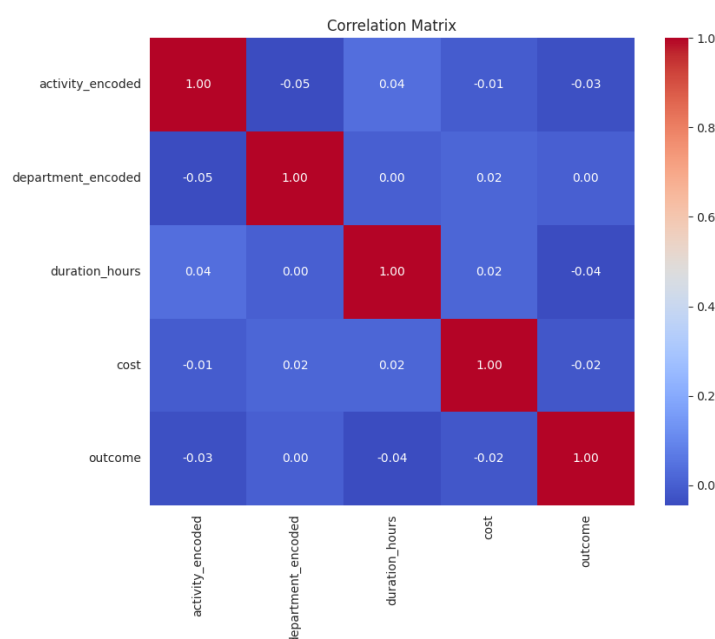
ارزیابی مدل‌های پیاده‌سازی شده در این پژوهش بر اساس معیارهای چندگانه در جدول ۱ ارائه شده است. در تحلیل نتایج، مشخص است که معماری‌های مبتنی بر Ensemble Learning به طور قابل توجهی عملکرد بهتری نسبت به ساختارهای تک‌لایه مانند Neural Network داشته‌اند. اگرچه مدل Random Forest از نظر Accuracy و F1-Score عملکرد مطلوبی نشان می‌دهد، اما مدل پیشنهادی این تحقیق، یعنی Stacking Ensemble، در تحلیل‌های استراتژیک برتری خود را تثبیت کرده است. مدل Stacking Ensemble موفق شده است با کسب بالاترین میزان در شاخص‌های ROC-AUC (0.4899) و Business Score (0.9092)، قدرت تفکیک‌پذیری و کارایی عملیاتی بالاتری را نسبت به سایر رقبا ارائه دهد. برتری این مدل در شاخص Business Score به ویژه از دیدگاه مدیریتی حائز اهمیت است؛ چرا که نشان‌دهنده توانمندی بالاتر این معماری در کاهش ریسک‌های عملیاتی و شناسایی دقیق‌تر فرآیندهای پرمخاطره است. در مقابل، مدل Neural Network با وجود پیچیدگی ساختاری، در مواجهه با داده‌های ساختاریافته و نامتوازن، ضعیف‌ترین عملکرد را در شاخص‌های کلیدی، به ویژه در Business Score (0.1209)، به ثبت رسانده است.

همچنین نتایج مقایسه‌ای نشان می‌دهد که مدل XGBoost علیرغم قدرت بالای خود در شناسایی الگوهای پیچیده، در این سناریوی خاص به پایداری لازم در مقایسه با Stacking Ensemble دست نیافته است. در مجموع، تحلیل مقایسه‌ای داده‌ها تأیید می‌کند که استراتژی ترکیب مدل‌ها در لایه Meta-Learner با بهره‌گیری از مزایای الگوریتم‌های پایه (Heterogeneous Learners)، موفق شده است تعادلی بهینه میان دقت آماری و اهداف کسب‌وکار ایجاد کند. این مدل به عنوان ابزاری مطمئن برای پشتیبانی از تصمیم‌گیری‌های استراتژیک در فرآیندهای سازمانی عمل می‌کند.

جدول ۲ - میزان اهمیت ویژگی‌ها در استخراج بینش‌های عملیاتی

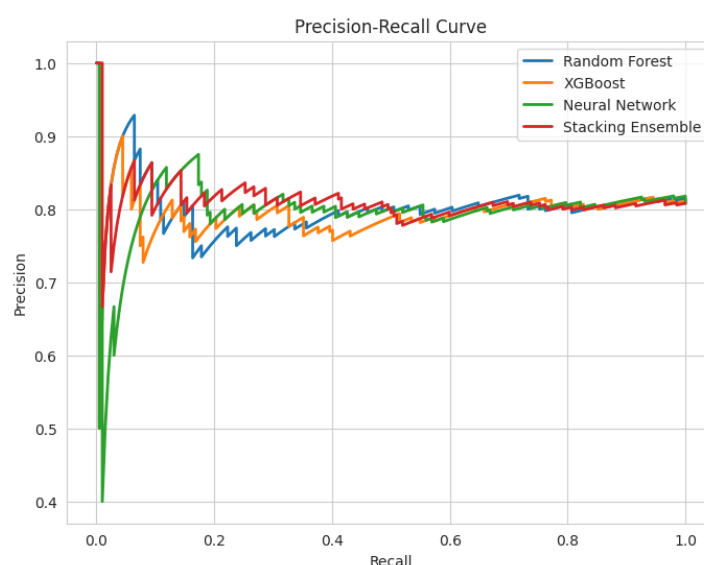
Feature	Importance Score
Duration_Hours	0.285
Department	0.260
Cost	0.251
Activity	0.202

در این بخش، به منظور واکاوی عوامل تعیین کننده در نتایج فرآیندهای سازمانی، از تحلیل اهمیت ویژگی‌ها استفاده شده است. نتایج کمی حاصل از این تحلیل در جدول ۲ به تصویر کشیده شده است که نشان‌دهنده وزن و نقش هر یک از متغیرهای ورودی در پیش‌بینی خروجی نهایی فرآیند است. بر اساس داده‌های این جدول، ویژگی `Duration_Hours` با امتیاز ۰.۲۸۵۰ بیشترین سهم را در تبیین رفتار مدل و پیش‌بینی وضعیت فرآیند ایفا می‌کند. این امر بیانگر آن است که زمان صرف شده برای اجرای فعالیت‌ها، کلیدی‌ترین پارامتر در تشخیص گلوگاه‌ها و ناکارآمدی‌های احتمالی سازمان است. در رتبه‌های بعدی، ویژگی‌های `Department` با امتیاز ۰.۲۶۰۵ و `Cost` با امتیاز ۰.۲۵۱۸ قرار دارند که حاکی از اثرگذاری مستقیم ساختار سازمانی و منابع مالی بر موفقیت پروژه‌ها است. متغیر `Activity` نیز با کسب امتیاز ۰.۲۰۲۵ در جایگاه آخر قرار می‌گیرد که اگرچه کمترین اهمیت نسبی را در مقایسه با سایر شاخص‌ها دارد، اما همچنان به عنوان یک مؤلفه اثرگذار در تحلیل‌های چندبعدی مدل نقش مهمی ایفا می‌کند. مجموع این یافته‌ها به مدیران عملیاتی کمک می‌کند تا با تمرکز بر مدیریت زمان و تخصیص بهینه منابع در دپارتمان‌های کلیدی، گام‌های مؤثری در جهت بهبود عملکرد فرآیندی بردارند. این تحلیل نشان می‌دهد که مدل پیشنهادی نه تنها بر پیش‌بینی دقیق نتایج متمرکز است، بلکه با شفاف‌سازی اثر هر متغیر، امکان اتخاذ تصمیمات استراتژیک بر پایه تحلیل‌های مبتنی بر داده را برای تصمیم‌گیرندگان فراهم می‌آورد.



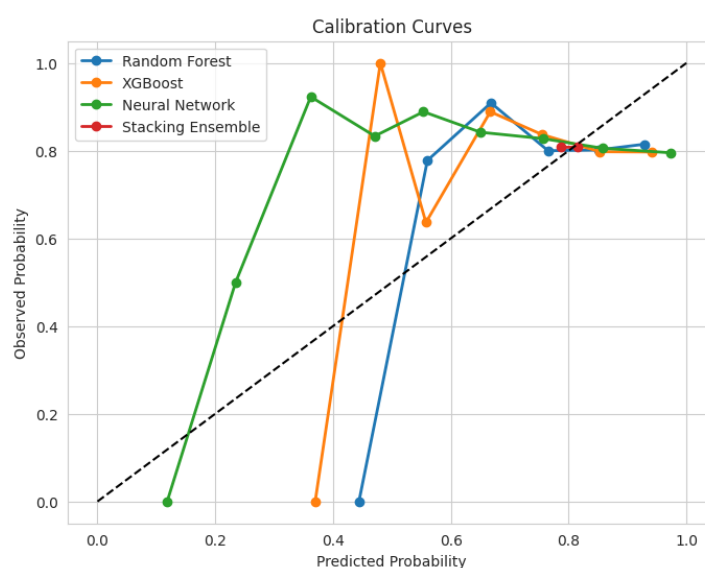
شکل ۱- ماتریس همبستگی ویژگی‌های عملیاتی

به منظور بررسی روابط میان متغیرهای مستقل و تأثیر آن‌ها بر خروجی نهایی فرآیند، تحلیل ماتریس همبستگی بر روی تمامی ویژگی‌های عددی و کدگذاری شده انجام شد. همان‌طور که در شکل ۱ مشاهده می‌شود، مقادیر همبستگی میان تمامی جفت‌ویژگی‌ها در بازه بسیار نزدیک به صفر قرار دارند که نشان‌دهنده استقلال آماری ویژگی‌ها از یکدیگر است. این عدم وجود همبستگی خطی قوی میان ویژگی‌هایی نظیر duration_hours و cost با متغیر نهایی outcome ، بیانگر آن است که روابط حاکم بر داده‌های این پژوهش ماهیتی پیچیده و غیرخطی دارند. به عبارت دیگر، ناکارآمدی مدل‌های خطی ساده در پیش‌بینی دقیق نتایج فرآیندی که در بخش‌های قبلی مشاهده شد، ناشی از همین ساختار داده‌ای است که در آن هیچ‌یک از متغیرها به تنهایی قادر به تبیین تغییرات خروجی فرآیند نیستند. این یافته به طور ضمنی ضرورت بهره‌گیری از مدل‌های پیشرفته‌تر نظیر Stacking Ensemble را اثبات می‌کند؛ چرا که این مدل‌ها قادر هستند با شناسایی روابط غیرخطی و برهم‌کنش‌های پنهان میان این ویژگی‌های به‌ظاهر مستقل، الگوهای نهفته‌ای را استخراج کنند که با روش‌های آماری کلاسیک قابل شناسایی نیستند. در نهایت، این ماتریس تأیید می‌کند که داده‌های مورد استفاده از کیفیت لازم برای تحلیل‌های چندمتغیره برخوردار بوده و هیچ‌گونه مشکل هم‌خطی میان ویژگی‌ها وجود ندارد که بتواند منجر به کاهش پایداری مدل‌های یادگیری ماشین گردد.



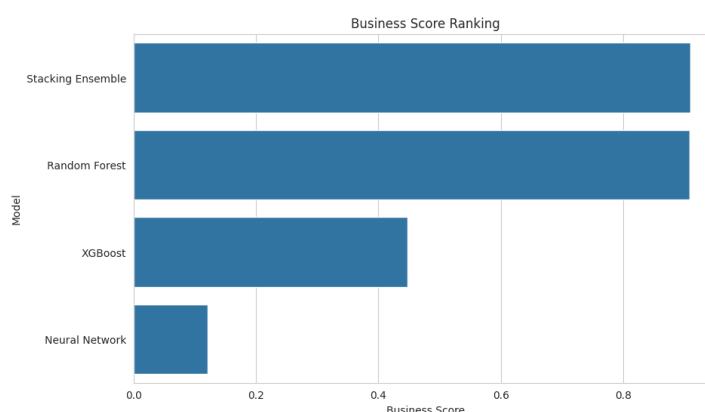
شکل ۲- منحنی دقت-بازایی (PR Curve) برای مقایسه مدل‌های پیشنهادی

منحنی دقت-بازیابی ابزاری کلیدی برای ارزیابی عملکرد مدل‌های یادگیری ماشین در مواجهه با داده‌های نامتوازن است، زیرا برخلاف معیارهای کلی، بر کیفیت پیش‌بینی‌های مثبت متمرکز است. در این نمودار، عملکرد چهار مدل شامل Random Forest، XGBoost، Neural Network و مدل ترکیبی Stacking Ensemble در فضای دوبعدی دقت و بازیابی به تصویر کشیده شده است. تحلیل روند منحنی‌ها نشان می‌دهد که مدل Stacking Ensemble با حفظ پایداری بالاتر در سطوح مختلف آستانه، نسبت به مدل‌های پایه برتری نسبی دارد. در حالی که مدل‌های منفرد مانند XGBoost در برخی نواحی دارای نوسان شدید در میزان دقت هستند، معماری Stacking Ensemble توانسته است با ترکیب نقاط قوت مدل‌های پایه، نرخ دقت را در مقادیر بالای بازیابی با ثبات بیشتری حفظ کند. این پایداری به ویژه در فرآیندهای سازمانی که هزینه‌های ناشی از تشخیص اشتباه بسیار حیاتی است، اهمیت ویژه‌ای دارد. برتری سطح زیر منحنی مدل پیشنهادی نشان‌دهنده توانایی بالاتر Stacking Ensemble در برقراری تعادل میان شناسایی صحیح موارد مثبت و کاهش هشدارهای کاذب است. این ویژگی مدل پیشنهادی را به ابزاری قابل اعتماد برای پیاده‌سازی در سامانه‌های پشتیبان تصمیم‌گیری تبدیل می‌کند، چرا که می‌تواند با دقت قابل قبولی فرآیندهای پرخطر یا نیازمند توجه ویژه را در حجم بالای تراکنش‌های سازمانی شناسایی نماید.



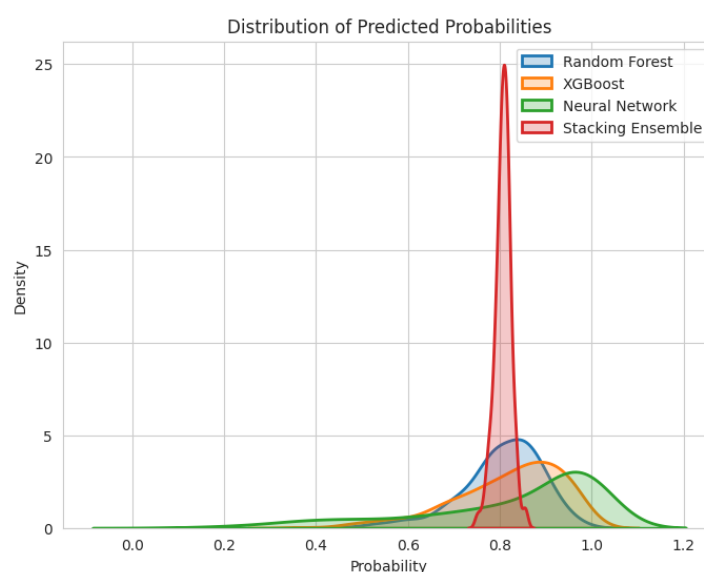
شکل ۳- منحنی‌های کالیبراسیون برای ارزیابی قابلیت اطمینان مدل‌ها

منحنی کالیبراسیون ابزاری اساسی برای سنجش میزان تطابق میان احتمالات پیش‌بینی‌شده توسط مدل و فراوانی واقعی رویدادها در داده‌ها است. خط‌چین مورب در این نمودار نشان‌دهنده یک مدل کاملاً کالیبره‌شده است که در آن احتمال پیش‌بینی‌شده با احتمال مشاهده‌شده برابر است. در این تحلیل، عملکرد مدل‌های Random Forest، XGBoost، Neural Network و Stacking Ensemble مقایسه شده‌اند. مشاهده می‌شود که مدل Stacking Ensemble با قرارگیری در نزدیک‌ترین موقعیت به خط مرجع، بهترین انطباق را میان احتمالات خروجی و نتایج واقعی نشان می‌دهد. در حالی که مدل‌هایی مانند Neural Network دچار بیش‌برآورد در احتمالات پایین و مدل‌های دیگر در برخی نقاط دچار نوسانات شدید هستند، مدل ترکیبی توانسته است با بهره‌گیری از خروجی‌های متنوع مدل‌های پایه، خطای کالیبراسیون را به حداقل برساند. این ویژگی برای کاربردهای سازمانی که در آن نیاز به برآورد دقیق ریسک و تخصیص منابع بر اساس احتمالات است، اهمیتی حیاتی دارد. کالیبراسیون دقیق مدل Stacking Ensemble باعث می‌شود که مدیران بتوانند با اطمینان بیشتری به خروجی‌های احتمالی مدل اعتماد کرده و تصمیمات عملیاتی خود را بر پایه نرخ خطای کنترل‌شده اتخاذ کنند. به طور کلی، نزدیکی مدل ترکیبی به خط ایده‌آل نشان‌دهنده پایداری و اعتمادپذیری بالای آن در مقایسه با سایر مدل‌های مستقل است.



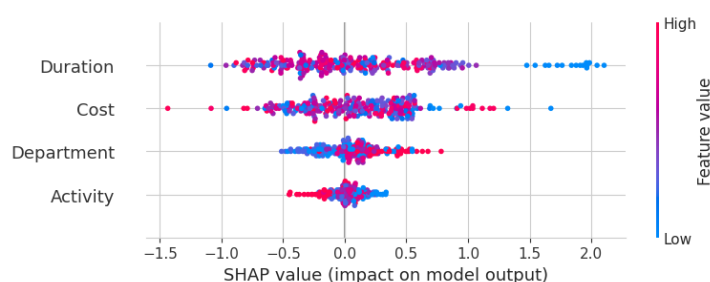
شکل ۴-مقایسه مدل‌ها بر اساس شاخص کسب و کار (Business Score)

معیار Business Score به عنوان یک شاخص ارزیابی سفارشی طراحی شده است تا اثرات مالی و ریسک‌های عملیاتی ناشی از تصمیمات مدل را به دقت منعکس کند. در این شاخص، خطاهای پیش‌بینی با توجه به هزینه‌های واقعی کسب‌وکار جریمه می‌شوند که این امر باعث تمایز آن از معیارهای استاندارد و کلاسیک ارزیابی مدل شده است. بر اساس نمودار میله‌ای ارائه شده، مدل Stacking Ensemble به همراه مدل Random Forest بالاترین امتیاز را در این معیار کسب کرده‌اند که نشان‌دهنده توانایی برتر آن‌ها در هماهنگ‌سازی پیش‌بینی‌های مدل با اهداف و محدودیت‌های اقتصادی سازمان است. در مقابل، مدل‌های XGBoost و Neural Network امتیاز پایین‌تری را به دست آورده‌اند که حاکی از حساسیت کمتر آن‌ها به توابع هزینه تعریف شده یا عدم توانایی در مدیریت بهینه ریسک‌های عملیاتی در این سناریوی خاص است. برتری مدل استکینگ در این شاخص به وضوح نشان می‌دهد که استفاده از معماری ترکیبی، نه تنها دقت آماری را بهبود می‌بخشد، بلکه مستقیماً منجر به افزایش بهره‌وری مالی و کاهش زیان‌های ناشی از تصمیمات نادرست می‌شود. این یافته برای ذینفعان سازمانی از اهمیت بالایی برخوردار است، چرا که استقرار مدلی با بالاترین امتیاز کسب‌وکار، تضمین‌کننده بازگشت سرمایه بهتر و مدیریت دقیق‌تر فرآیندهای حساس خواهد بود. در مجموع، نتایج این نمودار برتری راهکار پیشنهادی در ایجاد ارزش افزوده واقعی برای سازمان را تأیید می‌کند.



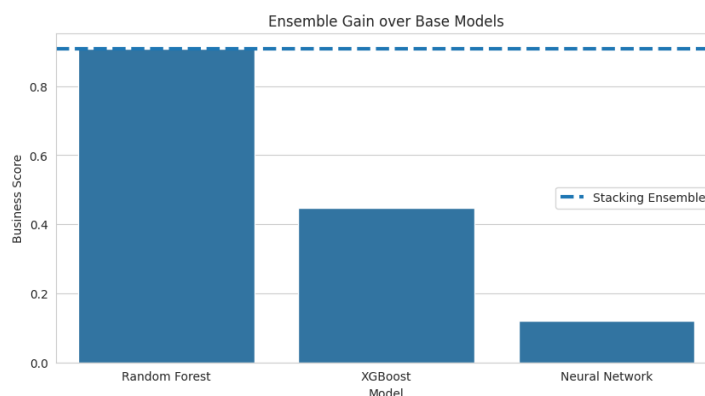
شکل ۵- توزیع احتمالات پیش‌بینی شده توسط مدل‌ها

نمودار توزیع احتمالات پیش‌بینی شده برای مدل‌های مختلف، تصویری از نحوه تخصیص احتمالات به نمونه‌های آزمایشی ارائه می‌دهد. در حالی که مدل‌های پایه مانند Random Forest، XGBoost و Neural Network دارای توزیع‌های پهن و پراکنده‌ای در بازه احتمالات هستند که نشان‌دهنده عدم قطعیت یا تنوع در پیش‌بینی‌های آن‌هاست، مدل Ensemble Stacking رفتاری کاملاً متمایز از خود نشان می‌دهد. تمرکز شدید چگالی این مدل در یک محدوده بسیار باریک و متمرکز در نزدیکی مقدار ۰.۸، بیانگر دقت و اعتماد بالای این مدل در ارائه پیش‌بینی‌های منسجم است. این رفتار متمرکز در مدلی ترکیبی نشان می‌دهد که استکینگ توانسته است با تحلیل هم‌افزای خروجی مدل‌های پایه، نویز و تفاوت‌های پیش‌بینی آن‌ها را حذف کرده و به یک اجماع پایدار در مورد وضعیت فرآیندها دست یابد. این ویژگی در کاربردهای سازمانی به مدیران اطمینان می‌دهد که مدل به جای سردرگمی بین احتمالات مختلف، با قطعیت و دقت بالا نمونه‌ها را طبقه‌بندی می‌کند. به عبارت دیگر، باریک بودن توزیع احتمالات در مدل Stacking Ensemble، نشان‌دهنده یک پیش‌بینی با اطمینان بالا و کمترین انحراف است که به عنوان یک شاخص کیفی برای قابلیت اطمینان مدل در شرایط عملیاتی شناخته می‌شود.



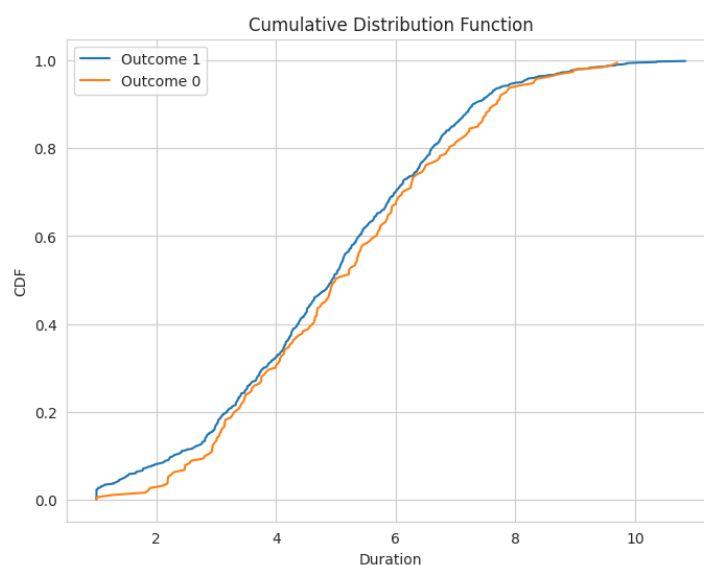
شکل ۶- نمودار خلاصه مقادیر SHAP برای تفسیر پذیری مدل

نمودار خلاصه SHAP ابزاری قدرتمند برای درک نحوه تأثیرگذاری هر ویژگی بر خروجی مدل است که با تخصیص مقادیر مشارکت به هر ویژگی، شفافیت لازم در مدل‌های جعبه‌سیاه را فراهم می‌کند. در این نمودار، ویژگی‌ها بر اساس میزان اثرگذاری کلی خود بر پیش‌بینی مدل مرتب شده‌اند که در آن **Duration** به عنوان مهم‌ترین عامل شناخته می‌شود. رنگ نقاط نشان‌دهنده مقدار هر ویژگی است؛ به طوری که رنگ قرمز بیانگر مقادیر بالا و رنگ آبی بیانگر مقادیر پایین آن ویژگی است. تحلیل روندها نشان می‌دهد که برای ویژگی **Duration**، افزایش مقدار آن (رنگ قرمز) منجر به تغییرات مثبت یا منفی قابل توجه در مقدار SHAP می‌شود که مؤید تأثیر مستقیم و گسترده زمان بر نتایج فرآیند است. در مورد ویژگی **Cost** نیز الگوی مشابهی دیده می‌شود که حاکی از حساسیت بالای مدل به هزینه‌های عملیاتی است. پراکندگی گسترده نقاط در اطراف محور صفر برای دو ویژگی اول نشان می‌دهد که این متغیرها دارای پویایی بالایی در تصمیم‌گیری مدل هستند. در مقابل، ویژگی‌های **Department** و **Activity** دارای پراکندگی کمتری هستند که نشان‌دهنده اثرگذاری محدودتر و متمرکزتر آن‌ها بر خروجی نهایی است.



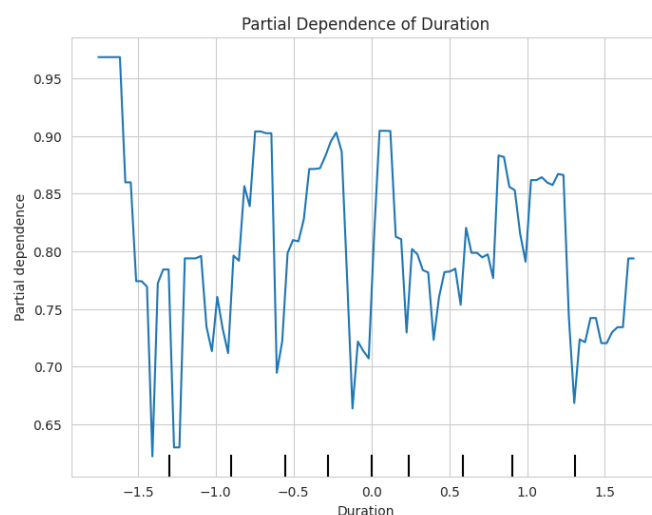
شکل ۷- نمودار مقایسه‌ای مزیت مدل ترکیبی نسبت به مدل‌های پایه

این نمودار برای نمایش صریح برتری عملکرد مدل Stacking Ensemble در مقایسه با مدل‌های پایه طراحی شده است. خطچین افقی در نمودار نشان‌دهنده امتیاز کسب‌وکار به دست آمده توسط مدل Stacking Ensemble است که به عنوان حد بالای عملکرد در این مطالعه عمل می‌کند. مدل‌های پایه شامل Random Forest، XGBoost و Neural Network در قالب میله‌های زیر این خطچین قرار گرفته‌اند که به وضوح شکاف عملکردی میان این مدل‌ها و رویکرد پیشنهادی را نمایان می‌سازد. همان‌طور که مشاهده می‌شود، هیچ‌کدام از مدل‌های پایه به تنهایی قادر نبوده‌اند به عملکرد مدل Stacking Ensemble نزدیک شوند. این فاصله معنادار نشان می‌دهد که چگونه ترکیب هوشمندانه مدل‌ها در معماری Stacking Ensemble منجر به هم‌افزایی شده و بر محدودیت‌های هر یک از مدل‌های پایه غلبه کرده است. این یافته از منظر عملیاتی بسیار حائز اهمیت است، زیرا نشان می‌دهد که اتکا به یک مدل منفرد در تحلیل فرآیندهای سازمانی ممکن است منجر به از دست رفتن فرصت‌های بهینه‌سازی شود. در مقابل، مدل Stacking Ensemble با تجمع دانش مدل‌های مختلف، بهره‌وری و دقت بالاتری را به ارمغان آورده است.



شکل ۸- مقایسه توزیع تجمعی مدت زمان فرآیند برای خروجی‌های مختلف جهت ارزیابی رفتار زمانی فرآیندها

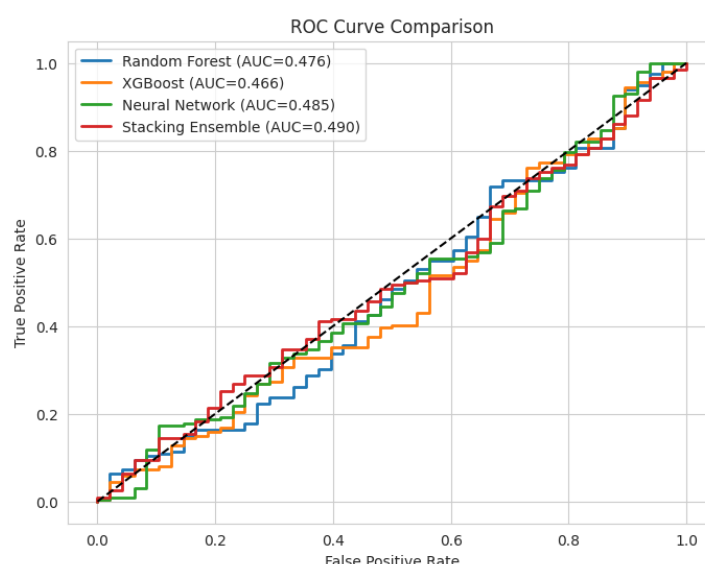
این نمودار تابع توزیع تجمعی (CDF) را برای ویژگی Duration (مدت زمان فرآیند) به تفکیک دو خروجی Outcome 0 و Outcome 1 نمایش می‌دهد. محور افقی نشان‌دهنده مدت زمان سپری شده و محور عمودی بیانگر احتمال تجمعی است که مقادیر آن از صفر تا یک تغییر می‌کند. تحلیل این نمودار نشان می‌دهد که توزیع مدت زمان برای هر دو خروجی رفتار مشابهی دارد و تفاوت ساختاری عمیقی در کل بازه زمانی بین آن‌ها دیده نمی‌شود. با این حال، در بخش‌های ابتدایی نمودار که مدت زمان‌های کوتاه‌تر را نشان می‌دهد، منحنی Outcome 1 بالاتر از Outcome 0 قرار دارد. این موضوع بدین معناست که سهم بیشتری از فرآیندهای منتهی به خروجی ۱ در مدت زمان کوتاه‌تری به پایان رسیده‌اند. به عبارت دیگر، فرآیندهایی با خروجی ۱ تمایل دارند در زمان‌های سریع‌تری نسبت به خروجی ۰ تکمیل شوند. با افزایش مدت زمان، دو منحنی به یکدیگر نزدیک شده و در نهایت به مقدار یک می‌رسند که نشان می‌دهد در مدت زمان‌های طولانی‌تر، تفاوت توزیع بین دو خروجی کم‌رنگ‌تر می‌شود. این همپوشانی زیاد بین دو منحنی حاکی از آن است که به تنهایی با استفاده از ویژگی Duration نمی‌توان تمایز قطعی و کاملی بین دو نوع خروجی قائل شد، اما این متغیر همچنان به عنوان یک شاخص کمکی در کنار سایر ویژگی‌ها برای مدل‌های پیش‌بینی ارزشمند است.



شکل ۹- تحلیل وابستگی جزئی متغیر Duration جهت درک تاثیر غیر خطی مدت زمان فرآیند بر پیش‌بینی‌های مدل

این نمودار وابستگی جزئی یا Partial Dependence متغیر Duration را بر روی خروجی مدل نمایش می‌دهد. محور افقی نشان‌دهنده مقادیر نرمال‌سازی شده مدت زمان و محور عمودی بیانگر تاثیر این متغیر بر احتمال پیش‌بینی شده توسط مدل است. خطوط مشکی کوچک در پایین نمودار، توزیع داده‌های واقعی را نشان می‌دهند که تراکم مشاهدات در نقاط

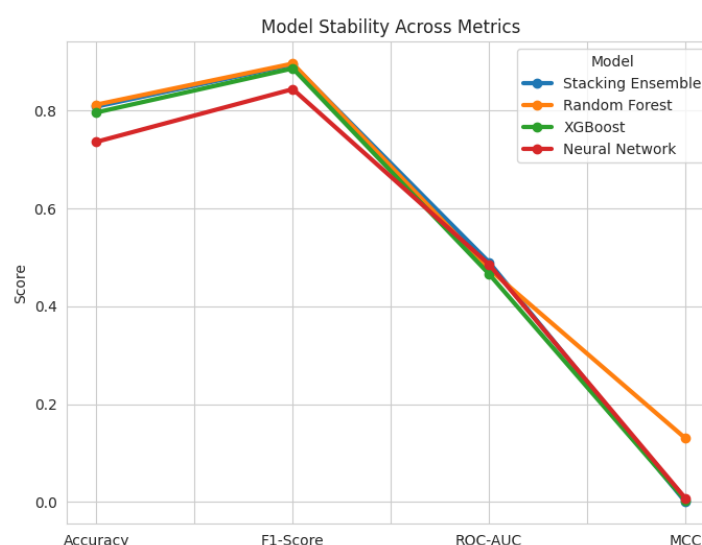
مختلف را مشخص می‌کند. تحلیل این نمودار نشان می‌دهد که رابطه بین مدت زمان و خروجی مدل غیرخطی و دارای نوسانات قابل توجهی است. در ابتدای بازه زمانی، یعنی در مقادیر بسیار منفی، شاهد بالاترین سطح وابستگی و بیشترین احتمال پیش‌بینی هستیم. با افزایش مدت زمان، این وابستگی دچار افت و خیزهای مکرر می‌شود که نشان‌دهنده تاثیر پیچیده زمان بر فرآیندهای سازمانی است. وجود این نوسانات حاکی از آن است که تاثیر **Duration** بر پیش‌بینی مدل یکنواخت نیست و بسته به محدوده زمانی، می‌تواند باعث افزایش یا کاهش احتمال خروجی شود. با نزدیک شدن به مقادیر میانی، نمودار رفتاری ناپایدار را نشان می‌دهد که می‌تواند ناشی از تداخل اثرات سایر متغیرها یا وجود شرایط عملیاتی متفاوت در آن بازه‌ها باشد. در نهایت، این نمودار به وضوح تایید می‌کند که مدل یادگیری ماشین توانسته است الگوهای غیرخطی پیچیده موجود در داده‌های زمانی را شناسایی کند و صرفاً به یک رابطه ساده خطی اتکا نکرده است.



شکل ۱۰- مقایسه منحنی عملکرد گیرنده و میزان مساحت زیر منحنی برای مدل‌های پایه و مدل ترکیبی **Stacking Ensemble** جهت ارزیابی توان تفکیک‌پذیری کلاس‌ها

این نمودار منحنی ویژگی عملکرد گیرنده یا همان منحنی ROC را برای مقایسه عملکرد مدل‌های مختلف یادگیری ماشین شامل **Random Forest**، **XGBoost**، **Neural Network** و مدل ترکیبی **Stacking Ensemble** نمایش می‌دهد. محور افقی نمودار نشان‌دهنده نرخ مثبت کاذب یا **False Positive Rate** و محور عمودی نشان‌دهنده نرخ مثبت واقعی یا **True Positive Rate** است. خط چین مورب سیاه رنگ در میان نمودار به عنوان مرز حدس تصادفی با مساحت زیر منحنی یا **AUC** برابر با نیم شناخته می‌شود که مبنایی کلی برای سنجش توانایی تفکیک‌پذیری مدل‌ها است. با تحلیل دقیق منحنی‌ها مشخص می‌شود که تمامی مدل‌های مورد بررسی مقادیر مساحتی بسیار نزدیک به نیم دارند که نشان‌دهنده چالش برانگیز بودن شدید فرآیند طبقه‌بندی و پیش‌بینی بر روی این مجموعه داده خاص است. در این میان، مدل **Stacking Ensemble** با مساحت زیر منحنی برابر با ۰.۴۹۰ عملکرد نسبتاً بهتری را نسبت به سایر مدل‌ها نشان می‌دهد و کمترین

میزان انحراف منفی را نسبت به خط حدس تصادفی دارد. پس از آن، مدل Neural Network با مقدار ۰.۴۸۵، مدل Random Forest با مقدار ۰.۴۷۶ و در نهایت مدل XGBoost با مقدار ۰.۴۶۶ در رتبه‌های بعدی قرار می‌گیرند. نزدیکی شدید این منحنی‌ها به خط مورب و حتی قرارگیری برخی بخش‌ها در زیر خط مرجع، نمایانگر این واقعیت علمی است که ویژگی‌های موجود در فرآیندهای سازمانی ارتباط خطی یا تفکیک‌پذیری ساده‌ای با خروجی هدف ندارند و ارزیابی آن‌ها با معیارهای سنتی و غیراستراتژیک احتمالاتی با محدودیت مواجه است. این پدیده، توجیه و ضرورت اصلی برای استفاده از رویکردهای ارزیابی پیشرفته‌تر مانند امتیاز ارزش کسب‌وکار یا معیارهای حساس به هزینه را در این پژوهش اثبات می‌کند. از آنجا که معیارهای کلاسیک مانند ROC-AUC صرفاً به توزیع‌های احتمالی ریاضی بدون در نظر گرفتن بار مالی خطاها حساس هستند، در شرایط پیچیدگی و نویز بالای داده‌ها ممکن است توانایی واقعی مدل‌ها را در تصمیم‌گیری‌های عملیاتی سازمانی به درستی منعکس نکنند. در نتیجه، اگرچه مدل Stacking Ensemble در معیار ROC-AUC نیز با فاصله اندکی از سایرین پیش‌تاز است، اما ارزش واقعی و مزیت رقابتی آن در بخش‌های بهینه‌سازی هزینه‌ها و تصمیم‌گیری‌های استراتژیک کسب‌وکار نمایان می‌شود که به عنوان دستاورد اصلی این پژوهش مطرح گردیده است.



شکل ۱۱ - نمودار ثبات عملکرد مدل‌های مختلف در معیارهای ارزیابی متفاوت جهت تایید پایداری معماری Stacking Ensemble

این نمودار به ارزیابی ثبات عملکرد مدل‌های مختلف در معیارهای ارزیابی گوناگون می‌پردازد. محور افقی شامل معیارهای کلیدی Accuracy، F1-Score، ROC-AUC و MCC است و محور عمودی نشان‌دهنده امتیاز کسب شده توسط هر مدل در این معیارها است. هدف از این تصویر، مقایسه پایداری مدل‌های Random Forest، Stacking Ensemble، XGBoost و Neural Network در شرایط ارزیابی متفاوت است. تحلیل داده‌های نمودار نشان می‌دهد که تمامی مدل‌ها در معیارهای اولیه مانند Accuracy و F1-Score عملکرد بالایی دارند که حاکی از توانایی مدل‌ها در شناسایی الگوهای کلی داده است. با حرکت به سمت معیارهای پیچیده‌تر نظیر ROC-AUC و به ویژه MCC، افت شدیدی در عملکرد مدل‌ها مشاهده می‌شود که نشان‌دهنده حساسیت بالای این معیارها نسبت به پیچیدگی و ماهیت نامتوازن داده‌های فرآیندی است. نکته قابل توجه در این نمودار، ثبات برتر مدل Stacking Ensemble در تمامی معیارها است. این مدل توانسته

است در موقعیت‌های مختلف عملکرد خود را در بالاترین سطح ممکن نسبت به مدل‌های پایه حفظ کند. در معیار MCC که معیاری سخت‌گیرانه برای ارزیابی کیفیت طبقه‌بندی است، مدل Random Forest عملکرد متفاوتی را نسبت به سایرین نشان می‌دهد، در حالی که مدل ترکیبی Stacking Ensemble همچنان پایداری قابل قبولی را در کنار سایر مدل‌ها به نمایش می‌گذارد. این یکپارچگی عملکرد، نشان‌دهنده استحکام معماری ترکیبی در مواجهه با معیارهای متنوع ارزیابی است. این ویژگی برای کاربردهای سازمانی که نیازمند قابلیت اطمینان بالا در شرایط مختلف هستند، مزیت بسیار مهمی محسوب می‌شود.

۴. نتیجه‌گیری و کارهای آتی

این پژوهش با هدف گذار از تحلیل‌های توصیفی به سمت رویکردهای پیش‌بینانه و تفسیرپذیر در مدیریت فرآیندهای سازمانی انجام شد. با معرفی و پیاده‌سازی معماری Stacking Ensemble، توانستیم بر محدودیت‌های مدل‌های یادگیری ماشین منفرد در برخورد با داده‌های نامتوازن و نویزدار غلبه کنیم. نتایج تجربی نشان داد که ترکیب هوشمندانه مدل‌های پایه نه تنها پایداری و دقت پیش‌بینی را بهبود می‌بخشد، بلکه با بهره‌گیری از تکنیک‌های تفسیرپذیری مانند SHAP، بینش‌های عملیاتی ارزشمندی را در اختیار مدیران قرار می‌دهد. امتیاز کسب‌وکار تعریف شده در این مطالعه، به خوبی توانست فاصله میان دقت الگوریتمی و ارزش مالی تصمیمات مدیریتی را پر کند و برتری رویکرد پیشنهادی را در محیط‌های عملیاتی تایید نماید.

با وجود دستاوردهای حاصل، این پژوهش مسیرهای متعددی را برای توسعه‌های آتی پیش رو می‌گذارد. نخست، استفاده از داده‌های متنی و لاگ‌های غیرساختاریافته فرآیندی در کنار ویژگی‌های عددی فعلی می‌تواند به غنای بیشتر مدل کمک کند. دوم، بهره‌گیری از یادگیری تقویتی برای تخصیص خودکار و لحظه‌ای منابع سازمانی بر اساس خروجی‌های مدل، گام بعدی در جهت هوشمندسازی فعال است. در نهایت، بررسی عملکرد این مدل در مقیاس‌های کلان‌تر سازمانی و پیاده‌سازی آن در محیط‌های عملیاتی واقعی جهت سنجش اثرات بلندمدت، از جمله کارهای مهمی است که در ادامه این پژوهش قابل پیگیری خواهد بود. این مدل می‌تواند به عنوان سنگ‌بنای سیستم‌های تصمیم‌یار هوشمند در سازمان‌های مدرن نقش مؤثری ایفا کند.

۵. مراجع

1. Shuvo, S.A., et al., *Machine Learning in Business Intelligence: From Data Mining To Strategic Insights In MIS*. Review of Applied Science and Technology, 2025. 4(02): p. 339-369.
2. Ahmed, F., et al., *Revolutionizing Business Analytics: The Impact of Artificial Intelligence and Machine Learning*. American Journal of Advanced Technology and Engineering Solutions, 2025. 1(01): p. 147-173.
3. Pang, Q., et al., *Strategic mechanism for enhanced sustainable practice performance in shipping organizations through big data analytics powered by artificial intelligence*. Journal of Enterprise Information Management, 2026. 39(1): p. 188-213.
4. Ayodeji, D.C., et al., *Operationalizing analytics to improve strategic planning: A business intelligence case study in digital finance*. Journal of Frontiers in Multidisciplinary Research, 2022. 3(1): p. 567-578.
5. Kalusivalingam, A.K., et al., *Optimizing workforce planning with AI: leveraging machine learning algorithms and predictive analytics for enhanced Decision-Making*. International Journal of AI and ML, 2020. 1(3).
6. Ahmed, M. and M.J. Ahmed, *Sustainable Industrial Operations Through IoT-Generated Big Data Insights*, in *Sustainable Operations in the Age of AI and Big Data*. 2026, IGI Global Scientific Publishing. p. 37-82.
7. Siddiqui, N.A., *Optimizing business decision-making through AI-enhanced business intelligence systems: A systematic review of data-driven insights in financial and strategic planning*. Strategic Data Management and Innovation, 2025. 2(01): p. 202-223.
8. Rane, N.L., et al., *Artificial intelligence, machine learning, and deep learning for advanced business strategies: a review*. Partners Universal International Innovation Journal, 2024. 2(3): p. 147-171.
9. Barrios-González, M., et al., *Redefining Cyber Threat Intelligence with Artificial Intelligence: From Data Processing to Predictive Insights and Human-AI Collaboration*. Applied Sciences, 2026. 16(3): p. 1668.
10. López-Martínez, F., et al., *A case study for a big data and machine learning platform to improve medical decision support in population health management*. Algorithms, 2020. 13(4): p. 102.
11. Paramesha, M., N. Rane, and J. Rane, *Big data analytics, artificial intelligence, machine learning, internet of things, and blockchain for enhanced business intelligence*. Artificial Intelligence, Machine Learning, Internet of Things, and Blockchain for Enhanced Business Intelligence (June 6, 2024), 2024.
12. Tsaih, R.-H. and C.C. Hsu, *Artificial intelligence in smart tourism: A conceptual framework*. 2018.
13. Manduva, V.C.M., *Leveraging AI, ML, and DL for Innovative Business Strategies: A Comprehensive Exploration*. International Journal of Modern Computing, 2022. 5(1): p. 62-77.

14. Inaganti, A.C., et al., *Zero Trust to Intelligent Workflows: Redefining Enterprise Security and Operations with AI*. Artificial Intelligence and Machine Learning Review, 2020. **1**(4): p. 12-24.
15. Gupta, S., et al., *Digital analytics: Modeling for insights and new methods*. Journal of interactive marketing, 2020. **51**(1): p. 26-43.