

ارزیابی و مقایسه روش‌های یادگیری ماشین برای پیش‌بینی میزان شوری رودخانه کارون

هانی محبوبی^{۱*}، محمد پیرایش^۲، یحیی جمور^۳

۱- استادیار گروه مهندسی نقشه‌برداری، دانشکده مهندسی عمران، آب و محیط‌زیست، دانشگاه شهید بهشتی

۲- دانشجوی کارشناسی ارشد هیدروگرافی، دانشکده مهندسی عمران، آب و محیط‌زیست، دانشگاه شهید بهشتی

۳- دانشیار گروه مهندسی نقشه‌برداری، دانشکده مهندسی عمران، آب و محیط‌زیست، دانشگاه شهید بهشتی

تاریخ دریافت مقاله: ۱۴۰۴/۰۲/۲۱ تاریخ پذیرش مقاله: ۱۴۰۴/۰۵/۰۷

چکیده

بیشتر پهنه ایران دارای اقلیم گرم و خشک همراه با بارش اندک است. مهم‌ترین منابع آب مصرفی کشورمان آب‌های زیرزمینی و آب رودخانه می‌باشد. رودخانه‌ها در سال‌های کم بارش و خشک اخیر از حیث کنترل پارامترهای کیفی آب بیشتر موردتوجه قرار گرفته‌اند. در این راستا، شوری آب رودخانه یکی از مهم‌ترین پارامترهای کیفی است که بایست توجه ویژه‌ای به آن داشت. در این تحقیق تغییرات بلندمدت شوری رودخانه کارون مورد ارزیابی قرار گرفته و با روش‌های متنوع یادگیری ماشین ازجمله شبکه عصبی عمیق، جنگل تصادفی و تقویت گرادیان مضاعف به مدل‌سازی متوسط شوری رودخانه کارون در بازه زمانی ۱۳۸۰ تا ۱۳۹۵ پرداخته و در ادامه تا میانه سال ۱۳۹۷ به مدت ۱۸ ماه پیش‌بینی متوسط شوری ارائه خواهد شد. پیش‌بینی‌های ارائه‌شده با داده‌های برداشت شده در بازه زمانی ۱۸ ماه مقایسه شده و روش‌های مختلف به‌کاررفته در این تحقیق از لحاظ دقت و کارایی مقایسه شد. نتایج نشان داده که اولاً در بازه زمانی مطالعه میزان متوسط ماهیانه شوری در رود کارون با نرخ تقریبی 10 ppm/year اضافه شده که این روند افزایشی برای اکوسیستم منطقه مخاطره‌آمیز خواهد بود. روش شبکه عصبی عمیق و جنگل تصادفی از لحاظ دقت پیش‌بینی‌ها تاحدی شبیه هم عمل کرده و هر دو روش با تنظیم بهینه پارامترهای مؤثر خود توانستند متوسط شوری رودخانه را با دقت حدود $170-180 \text{ ppm}$ پیش‌بینی کنند. روش جنگل تصادفی نسبت به شبکه عصبی عمیق در پیش‌بینی بلندمدت بهتر عمل می‌کند. اما روش تقویت گرادیان مضاعف نسبت به سایر روش‌ها موفق‌تر عمل کرده و دقت پیش‌بینی‌های ارائه‌شده با این روش 150 ppm بود که در مقایسه با جنگل تصادفی ۱۳ درصد و در مقایسه با شبکه عصبی عمیق ۱۸ درصد کاهش خطای نسبی را نشان می‌دهد.

کلیدواژه‌ها: شوری رودخانه، شبکه عصبی عمیق، جنگل تصادفی، تقویت گرادیان مضاعف.

* نویسنده مکاتبه‌کننده: دانشگاه شهید بهشتی، پردیس فنی و مهندسی شهید عباسپور، دانشکده مهندسی عمران، آب و محیط‌زیست.

۱- مقدمه

رودخانه‌ها و آب‌های زیرزمینی بزرگ‌ترین منابع آب مصرفی بشر هستند. اهمیت کیفیت آب رودخانه در منطقه ایران که دارای اقلیم خشک و غالباً بیابانی با بارش اندک است دوچندان می‌باشد. لذا، شوری آب رودخانه بایست به‌طور ویژه‌ای موردتوجه قرار گیرد. افزایش شوری رودخانه در درازمدت می‌تواند پیامدهای جدی برای محیط‌زیست، اقتصاد و سلامت انسان داشته باشد. ازجمله می‌توان به نابودی گونه‌های زیستی و کاهش تنوع آن‌ها [۱]، ایجاد هزینه مضاعف جهت نمک‌زدایی و بهره‌برداری از آب، کاهش حاصلخیزی خاک منطقه و تأثیرات منفی اقتصادی بر شیلات و صید ماهی اشاره نمود. رودخانه کارون^۱ یکی از مهم‌ترین رودهای ایران است که تأثیر بسزایی در زندگی مردم دارد. کارون از رشته‌کوه زاگرس و زرد کوه سرچشمه می‌گیرد و شاخه‌های فرعی گوناگونی نیز دارد. طول رودخانه حدود ۹۵۰ کیلومتر است و شاهد گونه‌های مختلف گیاهی و جانوری در این رودخانه و حوضه آبریز مربوطه می‌باشیم. لذا ارزش بالایی برای مطالعات کیفی آب داراست.

به‌جز مدل‌های تجربی^۲ ارائه‌شده برای شوری که با یافتن روابط ریاضی ساده بین عوامل محیطی و شوری ساخته می‌شوند و عموماً دقت کمی دارند، مدل‌سازی مبتنی بر فیزیک^۳ و مدل‌سازی مبتنی بر یادگیری ماشین^۴ نیز جهت پیش‌بینی شوری رودخانه کارایی فراوانی دارند. در مدل‌سازی مبتنی بر فیزیک از حل معادله دیفرانسیل انتقال-انتشار که بیانگر فرایندهای انتقال نمک در اثر جریان رودخانه و پخش نمک از ناحیه‌ای با تمرکز بیشتر به ناحیه‌ای با تمرکز کمتر است، استفاده می‌شود [۲،۳،۴،۵]. این روش چالش‌هایی نظیر پایداری طرح‌واره‌های عددی و میرایی

پیش‌بینی‌ها در بلندمدت را داراست.

امروزه تکنیک‌های مبتنی بر یادگیری ماشین جایگاه ویژه‌ای در علوم و مهندسی در زمینه پیش‌بینی بلندمدت پدیده‌ها یافته و از مزایایی نظیر دقت و سرعت بالا، مقیاس‌پذیری و توانایی مدیریت سیستم‌های پیچیده برخوردارند. در این راستا روش‌های شبکه‌های عصبی عمیق^۵، جنگل تصادفی^۶ و تقویت گرادیان مضاعف^۷ ازجمله روش‌های پرکاربرد در زمینه پیش‌بینی پدیده‌های فیزیکی هستند. چن و همکاران (۲۰۱۷)، هانتر و همکاران (۲۰۱۸)، اوبسه و همکاران (۲۰۲۱) و ژنگ و همکاران (۲۰۲۴) برای پیش‌بینی شوری در رودخانه‌های موری در استرالیا، نیوی در نیجریه و دانشوی در تایوان، از روش شبکه عصبی مصنوعی استفاده کردند [۶،۷،۸،۹]. همچنین هوانگ و فوو (۲۰۰۲)، کاربرد شبکه عصبی مصنوعی برای ارزیابی تغییرات شوری را در رودخانه آپالاچیکولا در فلوریدا بررسی کردند [۱۰]. هو و همکاران (۲۰۱۹)، ملسه و همکاران (۲۰۲۰)، از تکنیک جنگل تصادفی استفاده کردند [۱۱، ۱۲]. کولیس و همکاران (۲۰۲۴)، با استفاده از شبکه عصبی مصنوعی، هدایت الکتریکی را در رودخانه وارتا واقع در لهستان پیش‌بینی کردند و کمترین ضریب همبستگی مشاهده‌شده برای نمک‌های مختلف در این مطالعه ۰/۹۶ بوده است [۱۳]. همچنین ضریب همبستگی برای مطالعه‌ای مشابه توسط ربانی راشا (۲۰۲۲)، در رودخانه بهیراب در بنگلادش ۰/۹۸ بوده است [۱۴]. همچنین دوک و همکاران (۲۰۲۴)، به بررسی اثربخشی مدل‌های تقویت گرادیان مضاعف و حافظه کوتاه مدت-بلند مدت در پیش‌بینی شوری در دلتای مکونگ ویتنام پرداختند [۱۵]. کرباسی و همکاران (۲۰۲۴) با استفاده از ترکیب شبکه‌های

^۵ Deep neural networks

^۶ Random forest

^۷ Extreme gradient boosting (XGBoost)

^۱ Karun River

^۲ Empirical

^۳ Physics-based

^۴ Machine learning

گرفته‌شده و از سمت دیگر مقادیر معلوم هدایت الکتریکی^۱ که با کمک ضریب به شوری در واحد ppm ^۲ مبدل شده‌اند، جهت آموزش استفاده می‌شوند. برای تبدیل هدایت الکتریکی به مقدار نمک محلول (TDS)^۳ از ضریب $۰/۶۴$ استفاده شده است. برای توجیه ضریب فوق از داده‌های ثبت شده برای چندین نمونه از آب رودخانه استفاده شده است که هم هدایت الکتریکی آن مشخص بوده و هم مقدار نمک محلول آن به صورت آزمایشگاهی معلوم است. این نمونه‌های استفاده شده متعلق به مکان‌ها و زمان‌های مختلف است. درواقع، بهترین مقداری که می‌توانست داده‌های ثبت شده در مکان‌ها و زمان‌های مختلف را برهم انطباق دهد، با روش برآورد کمترین مربعات^۴ محاسبه شده و به عنوان ضریب تبدیل هدایت الکتریکی در نظر گرفته شد. شکل (۱) نشان‌دهنده ارتباط بین هدایت الکتریکی ثبت شده در ۵۰ نمونه آب مشاهداتی رودخانه کارون در مکان‌ها و زمان‌های مختلف می‌باشد. همانطور از شکل (۱-ب) مشهود است میزان هدایت الکتریکی (محور افقی) کمتر از $(\frac{\mu S}{cm}) 3000$ است و بوسیله ضریب $۰/۶۴$ بدست آمده از کمترین مربعات به میزان نمک محلول محاسبه شده که در محور عمودی نمایش داده شده، مبدل می‌گردد. شکل (۱-الف) میزان نمک محلول محاسبه شده با ضریب فوق را در مقابل نمک محلول مشاهداتی در نمونه‌های برداشت شده به تصویر می‌کشد. همانطور معلوم است، این دو کمیت انطباق خوبی با هم داشته به گونه ای که سنجه R^2 که معرف میزان انطباق این دو کمیت است برابر $۰/۹۹$ محاسبه شده است.

عصبی و الگوریتم تقویت گرادیان مضاعف به پیش‌بینی هدایت الکتریکی در رودخانه‌های آلبرت و باراتای استرالیا پرداختند [۱۶]. نیازکار و همکاران (۲۰۲۴)، به بررسی کاربردهای متنوع الگوریتم تقویت گرادیان مضاعف در مطالعات منابع آب پرداختند [۱۷]. همچنین در ایران نیز مطالعات مشابهی در رودخانه‌های آذربایجان صورت است. دهقانی و عباسپور (۲۰۱۴)، شوری آب رودخانه سیمینه رود واقع در استان آذربایجان غربی را با استفاده از شبکه عصبی مصنوعی تخمین زده و نتایج آن با روش‌های مرسوم آماری همچون رگرسیون خطی چند متغیره مقایسه شده است [۱۸]. همچنین، کنعانی و همکاران (۲۰۰۸)، شوری آب رودخانه آچه چای را با استفاده از شبکه عصبی مصنوعی پیش‌بینی نمودند [۱۹].

با استناد به نتایج پژوهش‌های پیشین می‌توان ادعا کرد که روش‌های یادگیری ماشین با جمع‌آوری داده‌های مناسب و بلندمدت می‌تواند به ارائه پیش‌بینی‌های دقیق و کارآمد منجر گردد. در این مطالعه هدف آن است که با استفاده از سری‌های زمانی ماهیانه متوسط شوری در چند ایستگاه از رودخانه کارون و به‌کارگیری روش‌های یادگیری ماشین از جمله تکنیک تقویت گرادیان مضاعف، جنگل تصادفی و شبکه عصبی عمیق، به پیش‌بینی بلندمدت و دقیق متوسط شوری در ایستگاه‌ها پرداخته شود.

۲- روش تحقیق

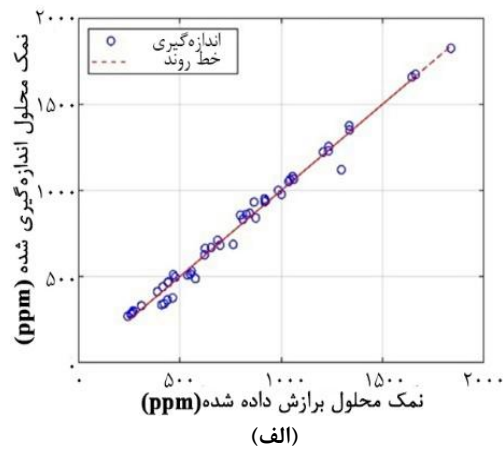
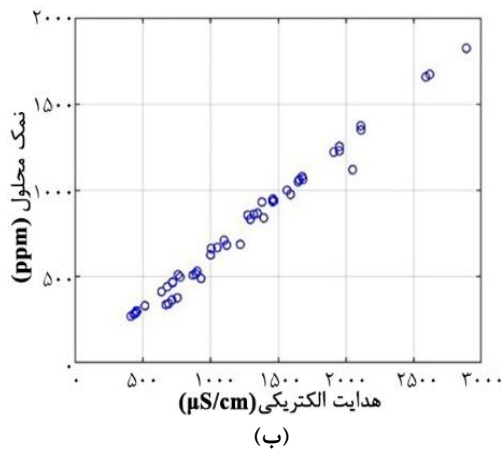
داده‌های مورد استفاده در این تحقیق شامل زمان، ویژگی‌های مکانی، کیفی و هیدرولوژیکی ایستگاه‌های هیدرومتری است که در یک بازه زمانی بلندمدت برداشت شده است. روش پیش‌بینی شوری استفاده از تکنیک‌های متنوع یادگیری ماشین است. از بخشی از داده‌های موجود برای آموزش استفاده شده و نهایتاً به پیش‌بینی شوری در بازه‌های زمانی مختلف پرداخته شده است. موقعیت و زمان برداشت نمونه به‌عنوان ورودی روش‌های مختلف یادگیری ماشین در نظر

¹ Electrical conductivity

² Part per million

³ Total dissolved solid

⁴ Least-squares estimation



شکل ۱: (الف) میزان نمک محلول محاسبه شده با ضریب تبدیل ۰/۶۴ در مقابل نمک محلول مشاهداتی در نمونه آب‌های برداشت شده، (ب) میزان نمک محلول در مقابل هدایت الکتریکی.

طول کارون از بالادست تا دلتای رودخانه که متوسط ماهیانه شوری را ثبت کرده‌اند، استفاده شده است. داده‌ها مربوط به شرکت مدیریت منابع آب بوده و به صورت سری زمانی ماهیانه بلندمدت (۱۷ ساله) از متوسط شوری در ایستگاه‌ها در اختیار کاربر قرار داده شده‌اند. در شکل (۲) محدوده جغرافیایی رودخانه و مکان ایستگاه‌ها نمایش داده شده‌اند.

برای حذف مقادیر پرت از داده‌ها، سری‌های زمانی ایستگاه‌ها به سه بخش روند کلی، نوسانات فصلی و باقیمانده^۴ تقسیم می‌شوند. با استفاده از روش میانگین متحرک^۳ به سادگی دو بخش اصلی روند کلی و نوسانات فصلی بازسازی شده و آنچه می‌ماند به عنوان باقیمانده لحاظ می‌گردد. در تحلیل سری‌های زمانی، میانگین متحرک یکی از روش‌های ساده و کاربردی برای هموارسازی داده‌ها و شناسایی ناهنجاری‌ها است.

به منظور راست آزمایی و ارزیابی دقت نیز از مقایسه پیش‌بینی‌ها با مشاهدات در بازه زمانی ۱۸ ماه پس از آموزش که به عنوان داده تست کنار گذاشته شده‌اند، استفاده می‌گردد. علاوه بر این، تکنیک‌های مختلف تقویت گرادیان مضاعف، جنگل تصادفی و شبکه عصبی عمیق از حیث دقت و عملکردشان در پیش‌بینی‌های بلندمدت مقایسه می‌گردند.

۲-۱- داده‌ها و منطقه مورد مطالعه

منطقه مطالعه رودخانه کارون است. کارون طولانی‌ترین و پرآب‌ترین رودخانه کشور بوده که از زردکوه در استان چهارمحال و بختیاری سرچشمه می‌گیرد. این رودخانه از چندین استان عبور کرده که بخش عمده آن واقع در استان خوزستان است.

حوضه آبریز کارون مساحتی حدود ۶۷۰۰۰ کیلومتر مربع را دربر گرفته است. کارون یکی از منابع حیاتی آب آشامیدنی، کشاورزی و صنعتی در اقلیم خشک جنوب غربی کشور است. در سالیان اخیر افزایش شوری این رودخانه یک مسئله بحرانی قلمداد شده که منجر به تهدید اکوسیستم و حیات جانداران شده است.

در این مطالعه از داده‌های ۱۵ ایستگاه هیدرومتری در

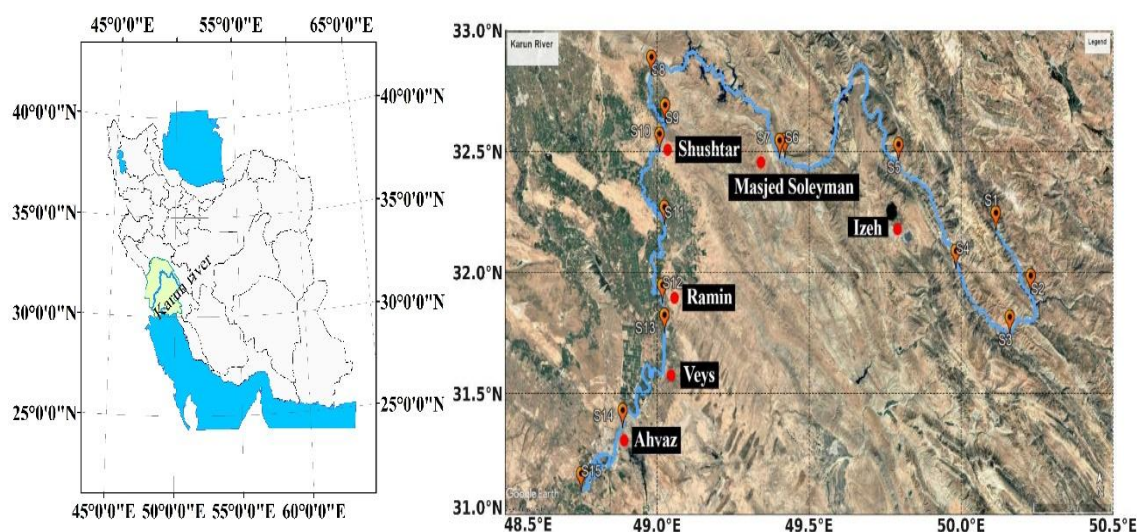
¹ Outlier

² Trend

³ Seasonality

⁴ Residual

⁵ Moving average



شکل ۲: منطقه مطالعه رودخانه کارون و ایستگاه‌های هیدرومتری سنجش شوری ماهیانه در طول رودخانه

رودخانه معادله انتقال-پخش^۲ است که یک معادله دیفرانسیل با مشتقات جزئی مکانی و زمانی است. لذا مکان و زمان دو پارامتر مهم و تعیین کننده در نحوه تغییرات شوری در رودخانه هستند. علاوه بر این، پارامتر دبی که به گونه‌ای هم با شدت جریان و هم هندسه مقطع رودخانه در تماس است یک پارامتر تأثیرگذار دیگر بر فرایند تغییرات شوری در رودخانه است. می توان مسئله را با کمک پارامترهای اضافی دیگر نظیر ضریب پخش شوری^۳ در رودخانه، سرعت انتقال جریان^۴ یا حتی هندسه حل نمود. اما بایست توجه داشت که بسیاری از پارامترها ممکن است در مکان و زمان‌های مختلف موجود نبوده و درون‌یابی آن‌ها نیز منجر به افزوده شده خطاهای درون‌یابی بر خطای پیش‌بینی گردد. لذا با توجه به مجموعه داده‌ای که در اختیار این تحقیق بوده ۶ پارامتر مؤثر شامل سه پارامتر مکانی (طول جغرافیایی، عرض جغرافیایی و کیلومتر از)

در این روش، مقدار میانگین داده‌ها در یک پنجره زمانی مشخص که شامل چند اپک است، محاسبه شده و به عنوان روند محلی داده در نظر گرفته می‌شود. سپس با کم کردن این میانگین متحرک از مقادیر اصلی سری زمانی، باقیمانده‌ها به دست می‌آیند. با این فرض که باقیمانده‌ها می‌بایست از توزیع آماری نرمال^۱ تبعیت کنند، می‌توان از آزمون آماری ۳ برابر انحراف معیار برای شناسایی مقادیر پرت استفاده کرد. در این آزمون، مقادیری از باقیمانده‌ها که بیش از سه برابر انحراف معیار از میانگین باقیمانده‌ها که غالباً نزدیک صفر است، فاصله دارند، به عنوان مقادیر پرت در نظر گرفته می‌شوند و حذف خواهند شد. این روش با وجود سادگی، کارایی مناسبی دارد.

۲-۲- پارامترهای مؤثر بر انتقال شوری در رودخانه

مدل فیزیکی توصیف کننده تغییرات شوری در طول

^۲ Advection-di

^۳ Diffusion coefficient

^۴ Stream velocity

^۱ Normal distribution

قرار گیرد. شبکه‌های عصبی معمولی معمولاً از ۲ لایه پنهان استفاده می‌کنند در صورتی که شبکه‌های عصبی عمیق می‌توانند تعداد لایه‌های پنهان بیشتری داشته باشند. شبکه عصبی پرسپترون چندلایه^۳ یکی از ساده‌ترین و پرکاربردترین انواع شبکه‌های عصبی عمیق است. این شبکه‌ها در یادگیری ماشینی ارزش بسزایی دارند زیرا می‌توانند روابط غیرخطی در داده‌ها را یاد بگیرند و آن‌ها را به مدل‌های قدرتمندی برای اهدافی نظیر طبقه‌بندی، رگرسیون و تشخیص الگو تبدیل می‌کند. این شبکه از سه بخش اصلی تشکیل شده است: لایه ورودی^۴، لایه‌های پنهان^۵ و لایه خروجی^۶. هسته اصلی این شبکه‌ها نورون‌های مصنوعی هستند که از طریق وزن‌ها و بایاس‌ها با نورون‌های لایه‌های دیگر مرتبط می‌باشند. با استفاده از توابع فعال‌سازی^۷ غیرخطی، شبکه‌های فوق‌قادر به مدل‌سازی مسائل پیچیده و غیرخطی هستند. تابع ضرر^۸ اختلاف بین مقادیر واقعی شوری مشاهده‌شده در دوره آموزش و شوری مدل‌سازی شده را اندازه‌گیری می‌کند. برای این منظور از سنج میانگین مربعات خطا^۹ استفاده شده است. این سنج به خطاهای بزرگ حساس بوده و از طریق یک الگوریتم بهینه‌سازی مناسب کمینه خواهد شد. پارامترهای شبکه عصبی عمیق شامل وزن‌ها و بایاس‌ها هستند که از طریق کمینه‌سازی تابع هدف یافت می‌شوند. علاوه بر این، ابرپارامترها^{۱۰} شامل تعداد لایه‌های مخفی، تعداد نورون‌ها، نرخ یادگیری، اندازه دسته‌ها و تابع فعال‌ساز بایست قبل از آموزش شبکه عصبی تنظیم شوند. انتخاب بهینه این پارامترها بر

به همراه دو پارامتر زمانی (سال و ماه) و همچنین یک پارامتر هیدرولوژیکی (دبی رودخانه) به‌عنوان مشخصه‌های ورودی جهت پیش‌بینی تمرکز شوری به کار گرفته شدند. علت استفاده هم‌زمان از سال و ماه به‌عنوان پارامترهای زمانی آن است که نوسانات شوری از بارش و تبخیر تأثیر می‌پذیرد که خود این پارامترها تابعی از زمان هستند و دارای روند بلندمدت و تغییرات فصلی هستند. لذا دو پارامتر زمانی سال و ماه می‌تواند در پیش‌بینی انواع طول موج‌های بلند و کوتاه پدیده کارآمد باشد. علت استفاده از پارامترهای طول و عرض نیز وابستگی مکانی شدید شوری به شرایط هیدرولوژیکی در یک منطقه است. همچنین ضریب پخش شوری علاوه بر اینکه تابعی از زمان است، تابع مکان نیز بوده و می‌تواند در طول رودخانه به دلیل تغییر شرایط هیدرولوژیکی و ویسکوزیته تغییر نماید. همچنین انتقال شوری خود تابع جریان و کیلومتر از بالادست رودخانه است. لذا در حالت بهینه مجموعه‌ای از داده‌ها اعم از طول، عرض، کیلومتر از، سال، ماه و دبی می‌تواند با توجه به ارتباط با همدیگر و سایر عوامل مؤثر بر شوری تأثیر بسزایی بر دقت پیش‌بینی شوری داشته باشند. در این تحقیق این پارامترها به‌عنوان داده‌های ورودی که مشخصه‌های مؤثر بر تغییرات شوری هستند، لحاظ شده‌اند.

۲-۳- شبکه عصبی عمیق

شبکه‌های عصبی مصنوعی می‌توانند به‌طور مؤثر و با یک تکنیک محاسباتی قوی، تمرکز شوری در رودخانه را پیش‌بینی کنند. برای این منظور از پارامترهای مؤثر بر انتقال شوری در رودخانه به‌عنوان داده‌های تاریخی برای آموزش شبکه عصبی استفاده می‌شود. پس از اتمام آموزش، مدل توسعه‌یافته می‌تواند برای پیش‌بینی تمرکز شوری برای گام‌های زمانی آینده مورد استفاده

³ Multilayer perceptron

⁴ Input layer

⁵ Hidden layers

⁶ Output layer

⁷ Activation function

⁸ Loss function

⁹ Mean square error

¹⁰ Hyper parameters

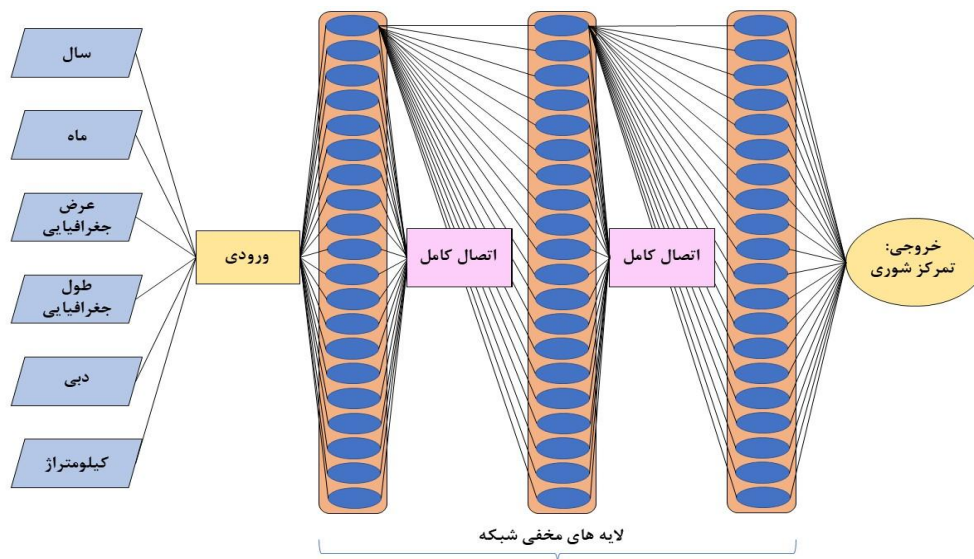
¹ Advection

² Artificial neural networks

پیکربندی‌های مختلفی از شبکه با تعداد لایه‌ها، نورون‌ها، توابع فعال‌ساز متفاوت و مقادیر ابرپارامترهای متنوع مورد ارزیابی قرار گرفتند و مدل‌های مختلف به‌صورت تکراری آموزش داده شد تا پیکربندی‌ای که بهترین معیارهای خطا را در مجموعه داده‌های آموزش، تست و صحت سنجی فراهم می‌کند، شناسایی شود. اگرچه این روش تجربی از نظر محاسباتی زمان‌بر است، اما در عین سادگی تضمین می‌کند که مدل نهایی به‌خوبی با ویژگی‌های داده‌ها تطابق دارد. یکی از چالش‌های مدل‌سازی بواسطه روش‌های مختلف یادگیری ماشین بحث ابتدای مدل ساخته شده به بیش‌برازش می‌باشد. برای ممانعت از بیش‌برازش توسط شبکه عصبی چندلایه ساخته شده اولاً، از تعداد لایه کم و نورون‌های کم استفاده شده است که این موضوع خود می‌تواند ریسک منجر شدن به بیش‌برازش را کاهش دهد. ثانیاً، از پارامتر پایداری‌سازی^۶ برای کاهش ریسک بیش‌برازش استفاده می‌شود. این امر منجر به نرم شدن و کاهش وریانس خروجی‌های شبکه از طریق جریمه کردن تابع هدف با نرخ تعیین شده توسط این پارامتر خواهد شد. برخی از ابرپارامترها که قبلاً به آن‌ها اشاره نشده عبارتند از: تابع فعال‌ساز تانژانت-سیگموئید^۷، حجم دسته آموزش^۸ برابر ۳۲، تعداد اپک آموزش برابر ۱۰۰۰ و نرخ پایداری‌سازی از نوع $L2$ برابر 10^{-3} بوده است.

عملکرد و کارایی مدل، صحت پیش‌بینی‌ها و ممانعت از بیش‌برازش^۱ تأثیر بسزایی خواهد داشت. در این مطالعه از شبکه عصبی پرسپترون چندلایه که تعداد و اندازه لایه‌های مخفی آن بهینه‌سازی شده است، استفاده نموده‌ایم. شکل (۱)، معماری شبکه مناسب را نمایش می‌دهد. همان‌گونه از شکل پیداست از سه لایه مخفی هرکدام شامل ۲۰ نورون جهت پیش‌بینی مقادیر تمرکز شوری استفاده شده است. پارامترهای ورودی همان پارامترهای مؤثر بر انتقال شوری هستند که در بخش ۲-۲ معرفی شده‌اند. علاوه بر این، اتصال لایه‌ها به‌صورت کامل برقرار شده است. نرخ یادگیری 0.1 لحاظ شده و از الگوریتم لئونبرگ-مارکوات^۲ جهت کمینه‌سازی تابع ضرر استفاده شده است. تابع ضرر بر اساس میانگین مربعات خطای برازش بر مقادیر تمرکز شوری در دوره زمانی ۱۳۸۰ تا ۱۳۹۵ تعریف شده است. لازم به ذکر است که از داده‌های پس از سال ۱۳۹۵ برای ارزیابی و صحت سنجی روش‌ها استفاده شده است که در بخش‌های آتی به تفصیل پرداخته خواهند شد. علت انتخاب تعداد لایه‌های مخفی آن است که در ارتباط با سایر پارامترها و ابرپارامترها منجر به تحقق بهترین نتیجه‌های دقت در مورد داده‌های تست و راست آزمایی می‌گردند. تحقق بهینه نتیجه‌های دقت به علت چینش پارامترهای مختلف و معماری شبکه براساس آزمون و خطا^۳ ارزیابی شده است. علیرغم اینکه روش‌های کارآمدی نظیر جستجوی تصادفی^۴ یا بهینه‌سازی بایزین^۵ به منظور تنظیم خودکار پارامترها و ابرپارامترها استفاده می‌شوند، در این مطالعه به علت سادگی و قابلیت تفسیر مناسب از رویکرد آزمون و خطا استفاده شد. به عبارت دیگر،

¹ Overfitting² Levenberg-Marquardt algorithm³ Trial and error⁴ Random search⁵ Bayesian optimization⁶ Regularization parameter⁷ Tansig⁸ Batch size



شکل ۳: معماری شبکه عصبی پرسپترون چندلایه به همراه پارامترهای ورودی مؤثر بر انتقال شوری و پارامتر خروجی تمرکز شوری

۲-۴- جنگل تصادفی

جنگل تصادفی، یک الگوریتم یادگیری ماشین گروهی^۱ است که هم برای دسته‌بندی^۲ و هم رگرسیون^۳ استفاده می‌شود. جنگل تصادفی یکی از پرکاربردترین الگوریتم‌های یادگیری ماشین است و شامل تعدادی درخت تصمیم^۴ می‌باشد که در زیرمجموعه‌های مختلف قرار دارد و برای بهبود دقت پیش‌بینی از آن مجموعه داده، میانگین می‌گیرد. عملکرد درخت تصمیم مبتنی بر نظارت^۵ بوده و بدین صورت است که ابتدا بر اساس یک ویژگی داده‌ها را با پرسیدن یک سؤال یا در نظر گرفتن یک شرط تقسیم می‌کند. سپس نوبت به شاخه‌های درخت می‌رسد که رشد کنند. درواقع، هر شاخه نشان‌دهنده یک نتیجه ممکن است (بله/خیر، درست/نادرست یا یک آستانه عددی). درنهایت، در انتهای هر شاخه پیش‌بینی نهایی (یک کلاس برای

طبقه‌بندی، یا یک مقدار برای رگرسیون) وجود دارد که تصمیم نهایی در این قسمت گرفته می‌شود. جهت ساختن درخت تصمیم از الگوریتم‌های مختلفی نظیر CART^۷، C4.5، ID3^۶ و CHID^۸ می‌توان استفاده نمود که اساس کار همگی ارزیابی ویژگی‌ها و تقسیم داده بر اساس ویژگی‌ها است.

جنگل تصادفی به جای تکیه بر یک درخت تصمیم، پیش‌بینی را از هر درخت می‌گیرد و نتیجه نهایی را به‌عنوان خروجی در نظر می‌گیرد [۲۰]. تعداد بیشتر درختان در جنگل منجر به دقت بالاتری می‌شود و از بروز مشکل بیش برآزش جلوگیری می‌کند [۲۱، ۲۲]. در روش جنگل تصادفی اولین مرحله آماده‌سازی داده‌ها برای آموزش مدل است که شامل پاکسازی داده‌ها، حذف مقادیر اشتباه، بازیابی مقادیر از دست رفته و تبدیل متغیرهای طبقه‌بندی شده به متغیرهای عددی است. سپس داده‌ها به مجموعه‌های آموزشی و آزمایشی (تست) تقسیم می‌شوند. مرحله بعد شامل انتخاب

¹ Ensembled

² Classification

³ Regression

⁴ Decision tree

⁵ Supervised

⁶ Iterative Dichotomiser 3

⁷ Classification And Regression Trees

⁸ Chi-square Automatic Interaction Detector

همچنین، قراردادن پارامتر نمونه‌گیری با جایگذاری^۲ برای آموزش هر درخت تصمیم می‌تواند وریانس مدل یا بیش برآزش را کاهش دهد. برخی از ابرپارامترهای تنظیم شده برای مدل به قرار ذیل است: تعداد درخت^۳ ۱۰۰ واحد، عمق ماکزیمم برابر ۱۴، حداقل تعداد نمونه برای تقسیم گره^۴ برابر ۱۰، بیشینه تعداد ویژگی برای انتخاب در هر تقسیم^۵ برابر ۲، ضرایب وزندهی نمونه‌ها برابر $10^{-4} \times 9/14$ و معیار تقسیم گره^۶ برابر رگرسیون کمترین مربعات وضع شده است.

۲-۵- تقویت گرادیان مضاعف

تقویت گرادیان مضاعف یک الگوریتم یادگیری ماشینی مبتنی بر درخت تصمیم است، با این مفهوم که از چارچوب تقویت گرادیان توأم با درخت‌های تصمیم استفاده می‌کند. این الگوریتم، ترکیبی مناسب از تکنیک‌های بهینه‌سازی را برای به دست آوردن نتایج بهتر با استفاده از منابع محاسباتی کمتر در کوتاه‌ترین زمان ممکن به کار می‌گیرد. از این الگوریتم برای طبقه‌بندی و رگرسیون استفاده می‌شود. در این الگوریتم درخت‌های تصمیم به صورت متوالی ایجاد می‌شوند و درختان تصمیم‌گیری کوچک و ساده به وسیله تکنیکی به نام تقویت به یک مدل واحد و قوی‌تر تبدیل می‌شوند، که هر درخت جدید سعی می‌کند خطاهای درخت‌های قبلی را برطرف نماید [۲۳].

تقویت^۷ یک روش گروهی است که از تجمیع مدل‌های ضعیف به مدل قوی‌تر می‌رسد. درواقع روش تقویت گرادیان منجر به کمینه شدن تابع ضرر خواهد شد. این امر توسط افزودن مدل‌های جدید که سبب کاهش گرادیان تابع ضرر خواهند شد، صورت می‌پذیرد.

تصادفی زیرمجموعه‌ای از ویژگی‌ها برای هر درخت در جنگل است. این کار برای کاهش بیش برآزش و افزایش تنوع درختان انجام می‌شود. سپس درخت‌های تصمیم با استفاده از زیرمجموعه‌ای از ویژگی‌ها که به طور تصادفی انتخاب شده‌اند، ساخته می‌شوند. درخت‌های تصمیم با استفاده از یک الگوریتم تقسیم باینری بازگشتی ساخته می‌شوند، که در آن هر گره داخلی داده‌ها را بر اساس یک ویژگی و مقدار آستانه انتخاب شده به دو زیرمجموعه تقسیم می‌کند. هنگامی که تمام درخت‌های تصمیم ساخته شدند، مرحله بعدی پیش‌بینی مجموعه تست است. این کار با جمع‌آوری پیش‌بینی‌ها از تمام درختان جنگل با استفاده از مکانیزم رأی‌گیری انجام می‌شود. در مسائل طبقه‌بندی، کلاسی که توسط اکثریت درختان به عنوان پیش‌بینی نهایی، پیش‌بینی شده است، انتخاب می‌شود. در مسائل رگرسیون، میانگین پیش‌بینی درختان به عنوان پیش‌بینی نهایی انتخاب می‌گردد. در نهایت، عملکرد مدل با استفاده از معیارهای مختلفی که بیانگر صحت هستند، ارزیابی می‌شود. اگر عملکرد مدل رضایت‌بخش باشد، می‌توان از آن برای پیش‌بینی داده‌های جدید استفاده کرد. این مراحل به صورت مکرر، با زیرمجموعه‌های مختلف از ویژگی‌ها و زیرمجموعه‌های مختلف داده، تکرار می‌شوند تا به سطح موردنظر از صحت دست پیدا کنند.

در روش جنگل تصادفی چون از تجمیع درخت‌های تصمیم استفاده می‌شود، خودبه‌خود ریسک ابتلا به بیش برآزش کمتر خواهد شد. با این حال، پارامترهایی نظیر عمق ماکزیمم^۱ که عمق هر درخت تصمیم را محدود می‌سازد و همچنین افزایش مقدار پارامتر کمترین تعداد نمونه‌هایی که در یک برگ مورد نیاز است، می‌تواند منجر به پیش‌بینی‌های نرم‌تر شود.

² Bootstrap

³ N-estimator

⁴ Min-sample-split

⁵ Max-feature

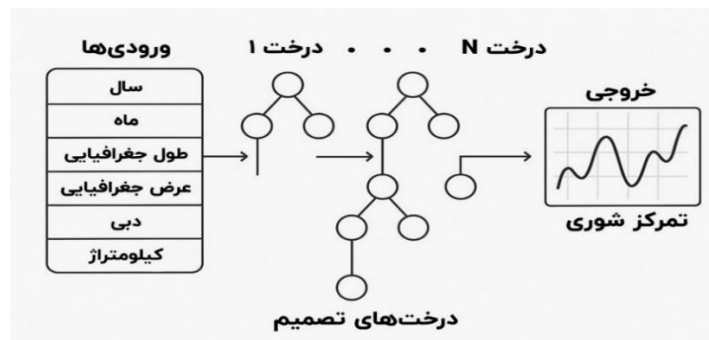
⁶ Criterion

⁷ Boosting

¹ Maximum depth

نتایج آن‌ها را برای کاهش واریانس و جلوگیری از بیش‌برازش ترکیب می‌نماید. درمقابل، الگوریتم تقویت گرادیان مضاعف به روش تقویت عمل می‌کند، یعنی درخت‌ها به صورت متوالی ساخته می‌شوند و هر درخت سعی می‌کند خطاهای باقی‌مانده از درخت قبلی را با استفاده از بهینه‌سازی مبتنی بر گرادیان تصحیح کند. درواقع اگر بخواهیم به شکل بصری ارتباط بین ویژگی‌های ورودی و درخت‌های تصمیم و خروجی را در الگوریتم‌های جنگل تصادفی و تقویت گرادیان مضاعف نمایش دهیم، می‌توان به شکل (مراجعه نمود. همان‌گونه در شکل پیداست، ورودی‌ها شامل پارامترهای مؤثر بر متوسط شوری ماهیانه در رودخانه بوده که شامل پارامترهای زمانی (سال و ماه)، پارامترهای مکانی (طول و عرض جغرافیایی و کیلومتر از بالادست رود) و مهم‌ترین پارامتر هیدرولوژیکی یعنی دبی هستند.

الگوریتم تقویت گرادیان مضاعف ابتدا با یک مدل ساده (یک درخت تصمیم) که به صورت ضعیف پیش‌بینی می‌کند شروع می‌شود. سپس خطاهای این درخت را محاسبه می‌کند، به عبارت دیگر بررسی می‌کند پیش‌بینی‌ها به چه میزان به مقادیر واقعی نزدیک هستند. در ادامه نیز یک درخت جدید برای پیش‌بینی خطاهای باقیمانده آموزش داده می‌شود تا اشتباهات را اصلاح کند. در آخر این فرآیند آن قدر تکرار می‌شود تا همه درختان برای پیش‌بینی نهایی ترکیب شوند و به یک جواب با صحت بالادست یابیم. روش جنگل تصادفی و روش تقویت گرادیان مضاعف هر دو از روش‌های یادگیری تجمعی مبتنی بر درخت تصمیم هستند، اما در نحوه ساخت مدل و بهینه‌سازی با یکدیگر تفاوت‌های اساسی دارند. جنگل تصادفی چندین درخت تصمیم‌گیری را به صورت مستقل و با استفاده از نمونه‌برداری تصادفی از داده‌ها ایجاد کرده و



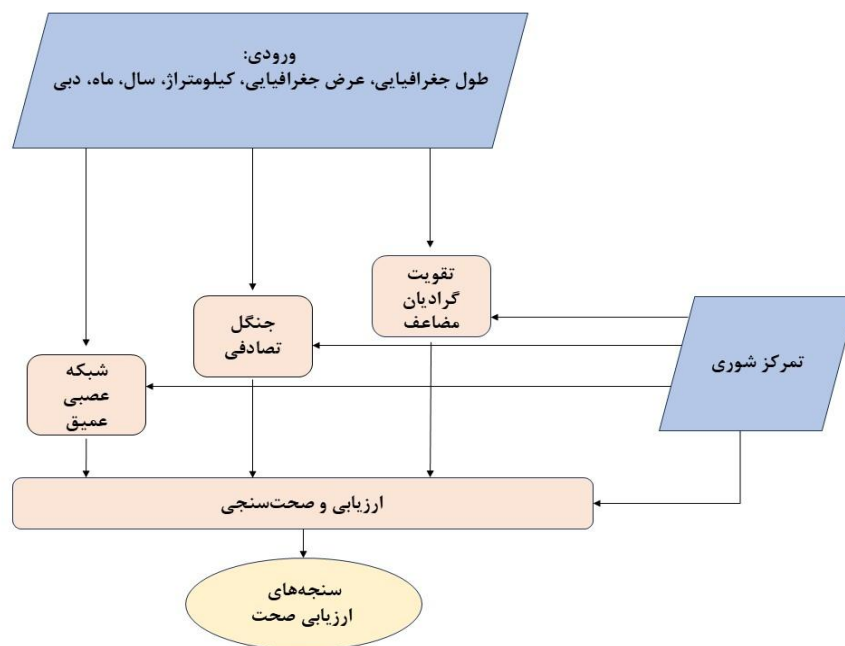
شکل ۴: ارتباط بین ورودی‌ها و خروجی‌ها در الگوریتم‌های جنگل تصادفی و تقویت گرادیان مضاعف

در روش تقویت گرادیان مضاعف نیز چندین ابرپارامتر قادر هستند از بیش‌برازش ممانعت کنند. این ابرپارامترها عبارتند از عمق ماکزیمم درخت‌های تصمیم، پارامتر زیرنمونه که معرف نمونه‌های آموزشی استفاده شده در هر درخت است و همچنین پارامتر

کمترین مجموع وزن نمونه‌ها که دقیقاً مانند پارامتر کمترین تعداد نمونه‌های یک برگ در روش جنگل تصادفی عمل می‌کند. برخی از ابرپارامترهای تنظیم شده برای مدل به قرار ذیل است: ضرایب پایدارسازی $L1$ و $L2$ به ترتیب به صورت $\alpha = 10^{-3}$ و $\lambda = 10^{-3}$ انتخاب شده‌اند. نرخ یادگیری ۰/۱، حداکثر عمق ۱۴ و پارامتر گاما برابر ۱ در نظر گرفته شده است. تعداد

¹ Subsample

می‌شوند. توسط هر روش مدل‌سازی انجام‌شده و صحت مدل‌سازی با ۲۰ درصد داده تست که از مشاهدات جدا شده‌اند و در مدل‌سازی مشارکتی نداشته‌اند ارزیابی می‌گردد. علاوه بر این، عملکرد هر سه مدل برای ۱۸ ماه پس از داده‌های تست بررسی شده و با مشاهداتی واقعی موجود در این ۱۸ ماه صحت سنجی می‌گردند. شکل (۵) طرح کلی پردازش داده‌ها را نشان می‌دهد.



شکل ۵: طرح کلی پردازش داده‌ها

این منظور پیش‌بینی‌های شوری توسط هر سه روش با مشاهدات انجام‌شده در این بازه مقایسه شده و از سنج‌های $RMSE$ و $NRMSE$ برای بررسی صحت نتایج استفاده شد.

سنج R^2 سهم واریانس قابل پیش‌بینی در یک تابع را به‌وسیله متغیر مستقل نمایش می‌دهد. در حقیقت نشان می‌دهد مقادیر بازتولید شده توسط مدل چقدر به مشاهدات نزدیک بوده‌اند. این سنج مقادیری بین صفر و یک اختیار کرده که هر قدر به یک نزدیک‌تر باشد بهتر است. اما برای اینکه از سنج‌های استفاده شود که

درخت ۱۰۰ واحد و پارامتر زیر نمونه ۰/۷ اختیار شده است.

۲-۶- طرح کلی پردازش داده‌ها

داده‌های ورودی که شامل پارامترهای مؤثر بر تمرکز شوری هستند و در بخش ۲-۲ معرفی شدند وارد هر کدام از سه روش مبتنی بر یادگیری ماشین که در بخش‌های ۲-۳، ۲-۴ و ۲-۵ به آن‌ها پرداختیم،

۲-۷- ارزیابی مدل‌ها و راست آزمایی

همان‌طور قبلاً بیان شد، داده‌های تحقیق به دودسته تقسیم شدند. ابتدا، داده‌های مربوط به مدل‌سازی که از ابتدای سال ۱۳۸۰ لغایت پایان ۱۳۹۵ را شامل می‌گردند. از این داده‌ها جهت آموزش شبکه عصبی عمیق، جنگل تصادفی و الگوریتم تقویت گرادیان مضاعف استفاده شد. ارزیابی مدل ساخته‌شده با هر روش توسط سنج‌های R^2 و $RMSE$ بررسی شدند. همچنین داده‌های ۱۳۹۶ تا میانه ۱۳۹۷ به مدت ۱۸ ماه جهت صحت سنجی نتایج به کار گرفته شد. برای

میانگین درصد خطای مطلق ($MAPE$)^۲ یک معیار رایج برای ارزیابی دقت مدل‌های پیش‌بینی است که میزان خطای نسبی بین مقادیر مشاهداتی و پیش‌بینی‌شده را به صورت درصد بیان می‌کند و شامل این مزیت است که خطا را به صورت مستقل از مقیاس داده‌ها بیان می‌کند. این معیار از میانگین قدر مطلق نسبت خطاها به مقادیر واقعی نظیر رابطه (۵) به دست می‌آید:

$$MAPE = \frac{1}{n} \sum_{i=1}^n \left| \frac{y_i - \hat{y}_i}{y_i} \right| \quad \text{رابطه (۵)}$$

همچنین به منظور ارزیابی و مقایسه عملکرد مدل‌ها در بازه‌ی صحت‌سنجی از آزمون آماری دایبولد-ماریانو^۳ (DM) استفاده می‌گردد. هدف از اجرای این آزمون آماری آن است که مدل‌های ساخته‌شده به صورت جفت از لحاظ صحت پیش‌بینی‌ها در بازه صحت‌سنجی مقایسه شود. اگر در یک زمان خاص مانند t خطای مدل اول $e_{1,t}$ باشد و خطای مدل دوم $e_{2,t}$ در نظر گرفته شود و $g(\cdot)$ معرف تابع ضرر^۴ میانگین مربعات خطا (MSE)^۵ باشد که از رابطه (۶) بدست می‌آید:

$$g(e) = \frac{1}{T} \sum_{t=1}^T e_t^2 \quad \text{رابطه (۶)}$$

در اینصورت رابطه (۷) اختلاف تابع ضرر بین دو مدل را بیان کرده که با d_t نشان داده می‌شود:

$$d_t = g(e_{1,t}) - g(e_{2,t}) \quad \text{رابطه (۷)}$$

در آزمون DM فرض اولیه H_0 به صورت زیر در نظر گرفته می‌شود که اختلاف قابل توجهی بین دقت پیش‌بینی‌های ارائه شده توسط دو مدل نیست و فرض ثانویه H_1 بدان صورت است که یک مدل از لحاظ دقت بر مدل دیگر برتری دارد. برای اجرای این آزمون آمارهای به صورت رابطه (۸) تعریف می‌شود:

هم شامل بایاس بوده و هم به خطاهای بزرگ‌تر حساس‌تر باشد از $RMSE$ نیز جهت ارزیابی مدل‌ها استفاده شد. اگر مشاهدات تمرکز شوری را با y_i ، مقادیر بازتولید شده یا پیش‌بینی‌شده توسط مدل را با $\hat{y}_i$ و متوسط n مقدار مشاهده‌شده را با $\bar{y}$ نمایش دهیم، سنج‌های به کاررفته جهت ارزیابی و صحت‌سنجی از روابط (۱)، (۲)، (۳)، (۴) و (۵) محاسبه می‌گردند.

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad \text{رابطه (۱)}$$

سنج $NRMSE$ همان $RMSE$ است که بر میانگین مقادیر مشاهدات تقسیم می‌شود تا همچنان به خطاهای بزرگ حساس باشد و نیز خطاها را به طور نسبی بیان کند یعنی یک سنج مستقل از مقیاس بوده و دقت را برای همه نوع داده با مقادیر مختلف کمتر یا بیشتر نمایش دهد. میانگین قدر مطلق خطا (MAE) میانگین اختلافات مطلق بین مقادیر پیش‌بینی‌شده و مشاهداتی را محاسبه می‌کند و به همه خطاها وزن یکسانی می‌دهد. در مقابل، $RMSE$ با مجذور کردن خطاها، به خطاهای بزرگ‌تر وزن بیشتری می‌دهد و نسبت به داده‌های پرت حساس‌تر است. به طور کلی، MAE معیار ساده‌تری برای تفسیر است، درحالی‌که $RMSE$ در شرایطی که خطاهای بزرگ اهمیت بیشتری دارند، سنج مناسب‌تر است.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad \text{رابطه (۲)}$$

$$NRMSE = \frac{RMSE}{\bar{y}} \quad \text{رابطه (۳)}$$

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad \text{رابطه (۴)}$$

^۲ Mean Absolute Percentage Error

^۳ Diebold-Mariano

^۴ Loss function

^۵ Mean Square Error

^۶ Test statistic

^۱ Mean Absolute Error

نمونه مصنوعی (نمونه‌های بوت‌استرپ) تولید می‌شود. سپس می‌توان یک آماره نظیر انحراف معیار را برای هر نمونه محاسبه نمود. این فرایند امکان تخمین فاصله اطمینان^۵ و عدم قطعیت هر مدل را فراهم می‌سازد.

۳- نتایج عددی

در ابتدا برای اینکه ارتباط بین هر کدام از مشخصه‌های استفاده شده در مدل‌سازی را با تمرکز ماهیانه شوری بررسی کنیم، شکل (۸) ارائه می‌گردد. در این شکل محورهای قائم تمرکز ماهیانه شوری را نشان می‌دهند که میزان آن از ۲۰۰ ppm تا ۱۵۰۰ ppm تغییر می‌کند. همچنین محورهای افقی معرف هر کدام از مشخصه‌های بیان شده در بخش ۲-۲ به‌عنوان پارامترهای مؤثر بر تمرکز شوری هستند. برای بررسی تأثیر ویژگی‌های مشاهده‌شده مختلف بر تمرکز ماهیانه شوری رودخانه، نمودارهای جعبه‌ای برای شش متغیر شامل سال، ماه، طول جغرافیایی، عرض جغرافیایی، دبی جریان و طول رود ترسیم شدند. با توجه به شکل ترسیم‌شده برای متغیر سال می‌توان گفت میانگین شوری و کران بالا و پایین آن در بازه زمانی ۱۳۸۰ تا ۱۳۹۵ روند افزایشی داشته است. میانگین شوری در طول رودخانه در ابتدای بازه مطالعه حدود ۷۴۰ ppm بوده و در اواخر بازه به حدود ۹۰۰ ppm رسیده است. به طور میانگین افزایش ۱۶۰ ppm در یک بازه زمانی ۱۶ ساله وجود داشته که نرخ افزایش تقریبی ۱۰ ppm در سال را نشان می‌دهد. ادامه این روند افزایشی می‌تواند پیامدهای جبران‌ناپذیری را در آینده داشته باشند. نمودارهای ترسیم شده در شکل (۶) نشان می‌دهند که طول و عرض جغرافیایی تأثیر مکانی قابل توجهی بر توزیع شوری دارند. به‌ویژه، روند افزایشی مشخصی در شوری با کاهش طول جغرافیایی مشاهده می‌شود که می‌تواند ناشی از افزایش شوری در امتداد مسیر غربی رودخانه به سمت خلیج فارس و نزدیکی

$$DM = \frac{\bar{d}}{\sqrt{\hat{V}_{(d)} / T}} \quad \text{رابطه (۸)}$$

که در رابطه (۸)، $\bar{d} = \frac{1}{T} \sum_{t=1}^T d_t$ میانگین اختلاف تابع ضرر بین دو مدل و $\hat{V}_{(d)}$ واریانس تخمین زده شده برای d_t ها می‌باشد. برای تخمین این واریانس معمولاً از تخمین ناهمسانی واریانس و خودهمبستگی سازگار^۱ (HAC) استفاده می‌شود (رابطه (۹)).

$$\begin{cases} \hat{V}_{(d)} = \gamma_0 + 2 \sum_{k=1}^{h-1} \gamma_k \\ \gamma_k = \frac{1}{T} \sum_{t=k+1}^T (d_t - \bar{d})(d_{t-k} - \bar{d}) \end{cases} \quad \text{رابطه (۹)}$$

در رابطه (۹)، h تعداد گام‌های پیشروی را برای محاسبه همبستگی نشان می‌دهد.

چنانچه قدرمطلق آماره DM تعریف شده بیش از $Z_{\alpha/2}$ برای توزیع نرمال باشد، فرض H_0 رد خواهد شد، بدین معنا که خطای پیش‌بینی دو مدل باهم متفاوت بوده‌است. این آزمون را در سطح اطمینان^۲ $\alpha = 0.05$ انجام می‌دهیم که نظیر آن $Z_{0.025} = 1.96$ برای توزیع نرمال می‌باشد. چنانچه قدرمطلق آماره محاسبه شده از ۱/۹۶ بیشتر باشد بدین معناست که مدل‌ها از لحاظ دقت متفاوت بوده‌اند و چنانچه $\bar{d}$ مثبت باشد، مدل اول بر مدل دوم برتری دارد.

در ادامه بحث ارزیابی و مقایسه مدل‌ها به بررسی عدم قطعیت^۳ حاصل از هر مدل می‌پردازیم. برای این منظور از روش بوت‌استرپینگ^۴ استفاده شده است. روش فوق یک روش آماری ناپارامتریک برای برآورد عدم قطعیت و توزیع آماره‌ها است که بدون نیاز به فرض توزیع خاصی برای داده‌ها انجام می‌شود. در این روش، با نمونه‌گیری تصادفی و با جایگذاری از داده‌های موجود، تعداد زیادی

¹ Heteroskedasticity and Autocorrelation Consistent

² Significance level

³ Uncertainty

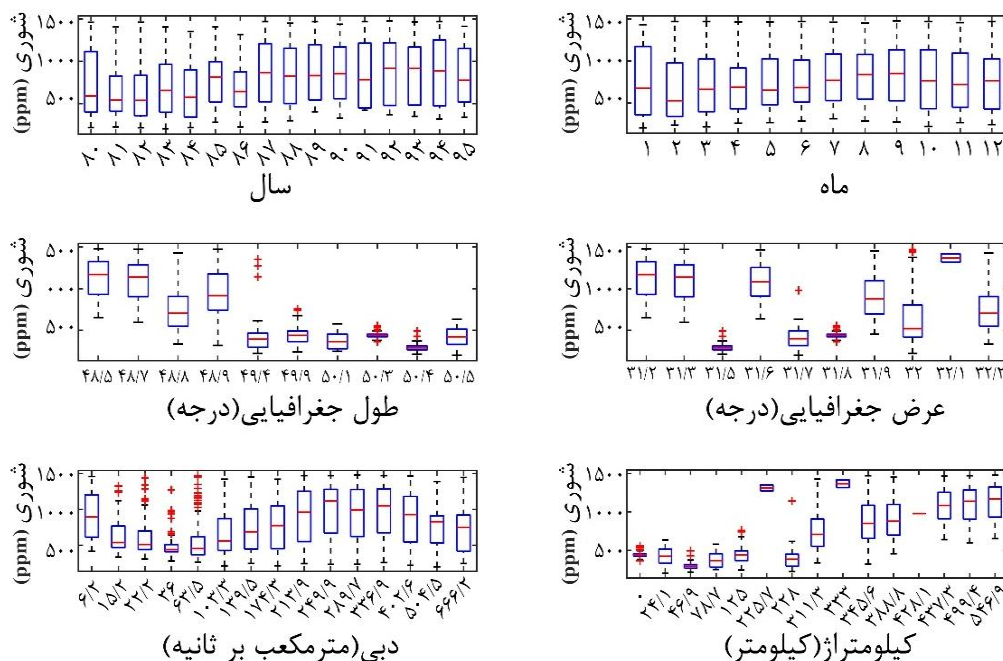
⁴ Bootstrapping

⁵ Confidence interval

با خلیج فارس است پیش می‌رویم، میزان متوسط ماهیانه شوری در اثر اختلاط با آب‌های آزاد افزایش یافته است. تغییرات فصلی (ماه) و بلندمدت (سال) اثرات نسبتاً کمتری نشان می‌دهند، به‌طوری‌که بازه‌های بین چارکی در بیشتر ماه‌ها و سال‌ها همپوشانی دارند؛ این امر حاکی از آن است که تغییرات زمانی در مقایسه با عوامل مکانی و هیدرولوژیکی نقش ثانویه‌ای دارند. در مجموع، این تحلیل چندبعدی بر اهمیت بالای ویژگی‌های مکانی و ویژگی‌های جریان در بازسازی نوسانات شوری تأکید دارد و می‌تواند نشان دهد پارامترهای به کار گرفته‌شده نقش به سزایی در بازتولید متوسط ماهیانه شوری خواهند داشت. این روابط و تعاملات پیچیده بین پارامترها و میزان تمرکز ماهیانه شوری می‌تواند یک عامل ترغیب‌کننده جهت استفاده از تکنیک‌های مدل‌سازی غیرخطی یادگیری ماشین باشد.

مصبب آن باشد. الگوی مشابه اما با نظم کمتر در عرض جغرافیایی نیز دیده می‌شود، به‌طوری‌که مناطق با عرض‌های جغرافیایی (به‌ویژه در حدود ۳۱٫۶ تا ۳۱٫۸ درجه) کاهش محسوسی در شوری نشان می‌دهند که می‌تواند به علت مشاهدات انجام‌شده در ایستگاه‌های سنجش شوری در بخش شرقی رودخانه با عرض کمتر از ۳۲ درجه باشد (به شکل (۲) مراجعه شود).

در مقابل، دبی جریان و طول رود تأثیرات هیدرولوژیکی قابل‌توجهی دارند و رابطه‌ای معکوس با شوری از خود نشان می‌دهند. با افزایش دبی، شوری به‌طور محسوسی کاهش می‌یابد که بیانگر نقش رقیق‌کننده حجم بالای جریان آب و اختلاط بیشتر با آب‌های تازه‌تر در طول مسیر رودخانه است. همچنین از ردیف آخر اشکال پیداست که افزایش کیلومتر از بالادست با میزان تمرکز ماهیانه شوری ارتباط مستقیم دارد یعنی هر چه به سمت مصب (*Estuary*) رودخانه که محل تلاقی آن



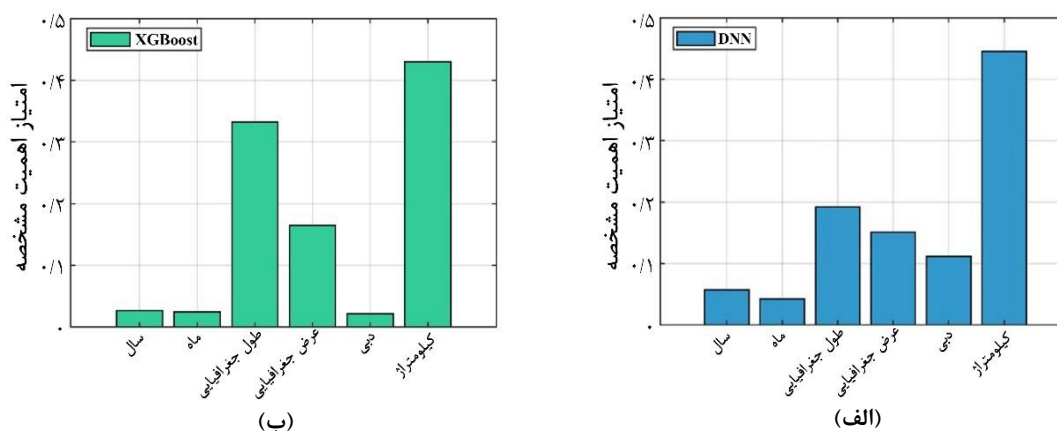
شکل ۶: نمایش تغییرات تمرکز شوری در مقابل ویژگی‌های سال، ماه، طول، عرض، دبی و کیلومتر از بالادست

مهمترین ویژگی‌ها در مدل‌سازی مکانی-زمانی شوری، ویژگی‌های مکانی با همان اهمیت قید شده در قبل هستند. عوامل دیگری وجود دارند که تابع مکان و به خصوص عرض جغرافیای بوده و بر شوری تأثیرگذارند که در این میان می‌توان به بارش، تبخیر، روان‌آب‌های محلی و توپوگرافی اشاره نمود. از آنجایی که رودخانه کارون بسیار طولانی بوده و از چندین استان با ویژگی‌های هواشناسی و شرایط توپوگرافی متفاوت عبور کرده، تغییرات شوری با تغییر عرض جغرافیایی قابل توجیه است.

در شکل (۸) ماتریس همبستگی بین ویژگی‌های مختلف به کاررفته در مدل‌سازی در قالب نقشه رنگی ارائه شده است. اهمیت این شکل در آن است که می‌توان فهمید کدام ویژگی‌ها همسو بوده و دارای همبستگی هستند و در صورت نیاز از تکرار پرهیز کرد. همان‌گونه در شکل پیداست ویژگی‌های طول جغرافیایی و کیلومتر از بالادست همبستگی منفی و قابل توجه دارند، علت این همبستگی آن است که با افزایش کیلومتر از بالادست طول جغرافیایی از شرق به غرب استان خوزستان کاهش می‌یابد. اما از آنجایی که این دو مشخصه مطابق شکل (۷) هر دو دارای مقادیر امتیاز ویژگی بالایی هستند، استفاده از هر دو ویژگی می‌تواند منجر به بهبود کارایی مدل و صحت پیش‌بینی‌ها گردد. بیشترین میزان همبستگی بین دو متغیر کیلومتر از دبی رودخانه است که برابر ۰/۴۷ بوده و خیلی قابل توجه نیست. در مورد سایر ویژگی‌ها همبستگی زیادی مشهود نیست و ارتباط خطی نداشته‌اند.

شکل (۷) اهمیت هر کدام از ویژگی‌های به کار گرفته شده در بازسازی را نشان می‌دهد. برای این منظور از امتیاز ویژگی نرمالایز شده استفاده شده است. برای محاسبه امتیاز ویژگی کفایت یکی از ویژگی‌های ورودی را به‌طور تصادفی تغییر داده و مقادیر آن را جابجا کنیم. سپس با یک روش یادگیری ماشین مثل روش شبکه عصبی عمیق در شکل (۷-الف) یا تقویت گرادیان مضاعف در شکل (۷-ب)، برای داده‌های ورودی جدید که جابجا شده‌اند، پیش‌بینی ارائه دهیم. اگر خطای پیش‌بینی ارائه شده نسبت به خطای به دست آمده از حالت اولیه اختلاف کمی داشته باشد، یعنی اثر این ویژگی بر پدیده تمرکز ماهیانه شوری کمتر است و این منجر به امتیاز ویژگی کمتر می‌گردد و بالعکس اگر اختلاف خطاها زیاد باشد یعنی ویژگی فوق بر پدیده تمرکز ماهیانه شوری بسیار مؤثر بوده است و دارای امتیاز ویژگی بالایی می‌باشد. در حقیقت این شکل نشان می‌دهد که کدام ویژگی‌ها در بازسازی تمرکز ماهیانه شوری تأثیر بیشتری داشته‌اند و اهمیت آن‌ها را اولویت‌بندی می‌کند. همان‌گونه از شکل پیداست ویژگی‌های مکانی اعم از کیلومتر از بالادست، طول جغرافیایی و عرض جغرافیایی به ترتیب بیشترین سهم را در تأثیرگذاری بر تمرکز ماهیانه شوری رودخانه کارون داشته‌اند و اثر ویژگی‌های هیدرولوژیکی و زمانی کمتر بوده است. لازم به ذکر است تأثیر دبی رودخانه نیز در روش مدل‌سازی با شبکه عصبی مصنوعی نسبت به ویژگی‌های زمانی برجسته‌تر بوده است. این موضوع قبلاً نیز در تحلیل شکل (۶) اشاره شده بود. در حقیقت، پارامتر زمان بیشتر در بازسازی نوسانات فصلی پدیده مؤثر بوده و تأثیر به سزایی در بازسازی روندهای اصلی نداشته است. همچنین یک بار هم امتیاز ویژگی‌های ورودی با روش تقویت گرادیان مضاعف محاسبه شد و در شکل (۷-ب) ارائه گردید. همچنان مشاهده می‌گردد

¹ Feature-score



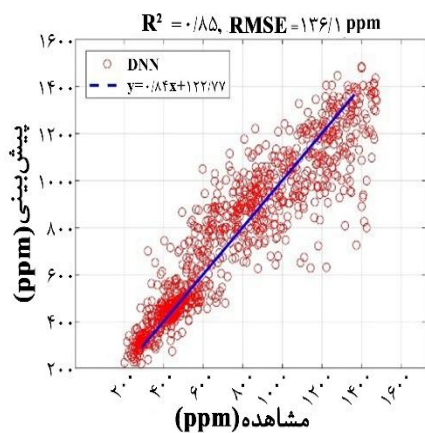
شکل ۷: (الف) امتیاز ویژگی برای مشخصه‌های ورودی به کاررفته در مدل سازی که با مدل شبکه عصبی عمیق محاسبه شده‌اند، (ب) امتیاز ویژگی برای مشخصه‌های ورودی که با مدل تقویت گرادیان مضاعف محاسبه شده‌اند.



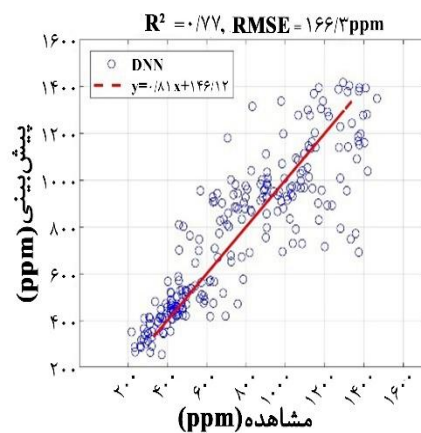
شکل ۸: ماتریس همبستگی ویژگی‌ها

اشکال در متلب ترسیم شده‌اند. در شکل (۹) مقادیر بازسازی شده توسط هر کدام از مدل‌ها با مقادیر مشاهده شده مقایسه شده‌اند و عملکرد مدل از حیث نحوه بازسازی و برازش بر مشاهدات در طول دوره آموزش بررسی شده است.

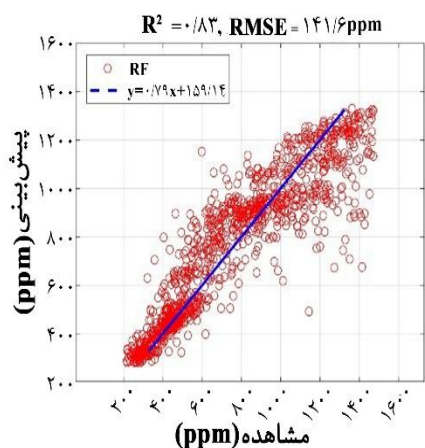
در ادامه، مدل سازی با هر سه روش یادگیری ماشین که در روش تحقیق بیان شد، انجام شده است. لازم به ذکر است که جهت مدل سازی با روش‌های شبکه عصبی عمیق و جنگل تصادفی از نرم افزار متلب و برای مدل سازی به روش تقویت گرادیان مضاعف از نرم افزار پایتون استفاده شده است. برای ارائه نتایج نیز همه



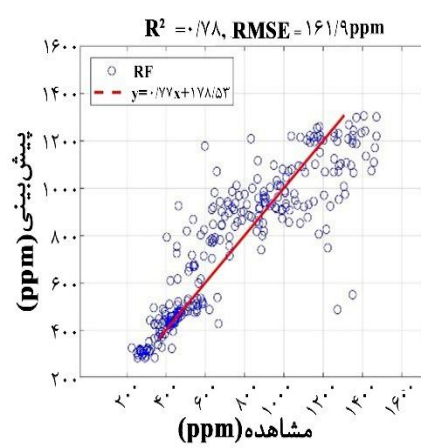
(ب)



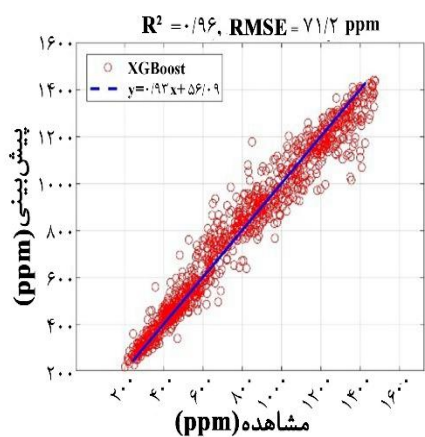
(الف)



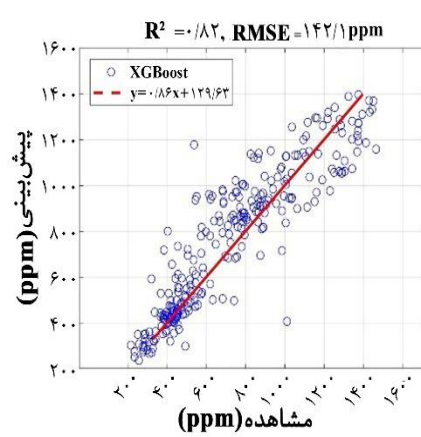
(د)



(ج)



(و)



(ه)

شکل ۹: (الف) مقایسه شوری مدل شده توسط روش شبکه عصبی عمیق و شوری مشاهده شده در مجموعه داده‌های تست و (ب) آموزش، (ج) مقایسه شوری مدل شده توسط روش جنگل تصادفی و شوری مشاهده شده در مجموعه داده‌های تست و (د) آموزش، (ه) مقایسه شوری مدل شده توسط روش تقویت گرادیان مضاعف و شوری مشاهده شده در مجموعه داده‌های تست و (و) آموزش

است.

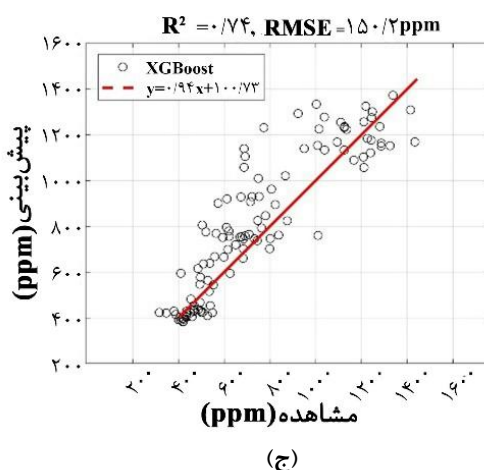
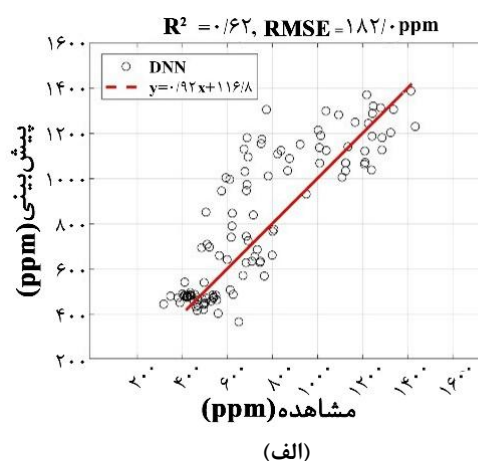
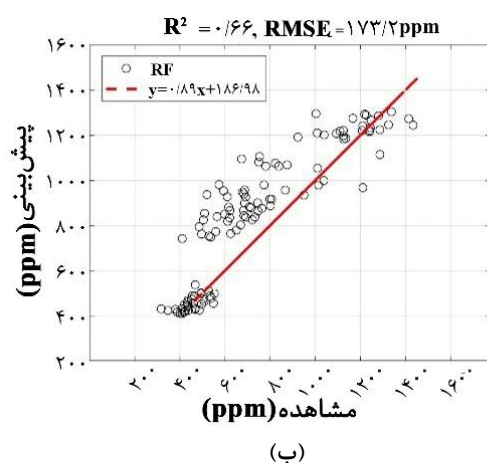
علاوه بر بررسی عملکرد و ارزیابی هر سه مدل در بازه زمانی ۱۳۸۰ تا ۱۳۹۵ که دوره آموزشی مدل‌ها محسوب می‌شد از مدل‌های ساخته‌شده برای پیش‌بینی متوسط ماهیانه مقادیر شوری در ایستگاه‌ها در بازه زمانی ۱۸ ماهه پس از دوران آموزش استفاده شده است. همان‌گونه قبلاً بیان شد این بازه زمانی دوره صحت سنجی نامیده شده است. لذا پیش‌بینی‌های ارائه‌شده در این بازه با مشاهدات موجود مقایسه شده و صحت سنجی با سنج‌های $RMSE$ و $NRMSE$ انجام شد. در شکل (۱۰) پیش‌بینی‌های انجام‌شده در دوره ۱۸ ماهه صحت سنجی با مقادیر مشاهده‌شده مقایسه شده و سنج‌های R^2 و $RMSE$ جهت ارزیابی توانایی هر روش و صحت پیش‌بینی‌های انجام‌شده در بالای اشکال ارائه شده‌اند. ردیف‌های (الف)، (ب) و (ج) به ترتیب مربوط به روش‌های شبکه عصبی عمیق، جنگل تصادفی و تقویت گرادیان مضاعف می‌باشند. نظیر نتایج مشاهده‌شده در مورد داده‌های تست، در دوره صحت سنجی نیز مدل تقویت گرادیان مضاعف برتری محسوسی نسبت به سایر مدل‌ها داشته است. توانایی پیش‌بینی مدل جنگل تصادفی تا حدودی بهتر از مدل شبکه عصبی عمیق بوده است. مقدار $RMSE$ در روش تقویت گرادیان مضاعف نسبت به روش‌های جنگل تصادفی و شبکه عصبی عمیق به ترتیب ۱۳ درصد و ۱۸ درصد بهتر بوده است. همچنین بیشتر شدن سنج R^2 در روش تقویت گرادیان مضاعف نشان‌دهنده ارتباط خطی قوی‌تر پیش‌بینی‌های این روش با مقادیر مشاهداتی بوده و از شکل (۱۰-ج) نیز معلوم است که ترسیم پیش‌بینی‌ها در مقابل مشاهدات بیشتر به خط نیمساز نزدیک بوده است. لذا می‌توان گفت پیش‌بینی‌های ارائه‌شده توسط روش فوق انحراف معیار و بایاس کمتری نسبت به سایر روش‌ها داشته که این موضوع مؤید دقت و صحت بیشتر پیش‌بینی‌هاست. در شکل (۱۱-الف) مقادیر $RMSE$ و در شکل (۱۱-ب)

داده‌های دوره آموزش ۸۰ درصد داده‌ها و داده‌های دوره تست ۲۰ درصد کل داده‌های ورودی در بازه زمانی ۱۳۸۰ تا ۱۳۹۵ بوده‌اند. برای مقایسه هر سه روش بایست پارامترها در هر کدام از روش‌های به کار گرفته‌شده به بهترین صورت ممکن تنظیم شوند و این امر مستلزم آزمون و خطا جهت تعیین بهینه پارامترها و ابرپارامترهای هر روش است. سپس با پارامترهای تعیین‌شده که منجر به بهترین عملکرد هر روش خواهند شد مدل‌سازی انجام‌شده و نتایج هر سه روش باهم مقایسه می‌گردد.

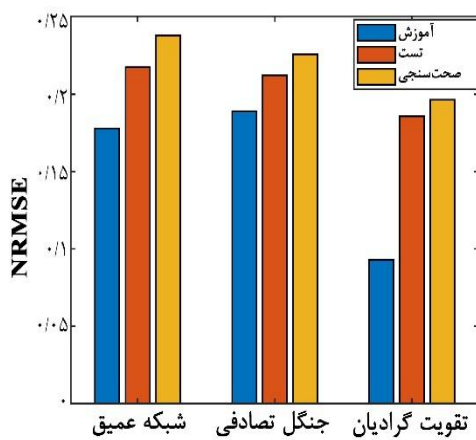
از ردیف‌های (الف)، (ب)، (ج) و (د) در شکل (۹)، مشخص است که عملکرد روش‌های شبکه عصبی عمیق و جنگل تصادفی با پارامترهای مناسب تقریباً شبیه هم بوده ولی توان بازتولید داده‌های آموزشی شبکه عصبی عمیق کمی بیشتر بوده است. سنج R^2 در این بازسازی ۰/۸۵ بوده و مشاهدات دوره آموزشی با $RMSE$ کمی بهتر نسبت به روش جنگل تصادفی بازتولید می‌شوند. اما این برتری صرفاً مربوط به داده‌های آموزش بوده و در مورد داده تست روش جنگل تصادفی کمی بهتر عمل نموده است. در کل مقادیر سنج R^2 در بازسازی داده‌های تست برای دو روش شبکه عصبی عمیق و جنگل تصادفی به ترتیب ۰/۷۷ و ۰/۷۸ بوده و تفاوت فاحشی وجود ندارد و از لحاظ $RMSE$ بازسازی روش جنگل تصادفی کمی برتر بوده است. اما همان‌گونه از اشکال ردیف‌های (ه) و (و) پیداست مدل تقویت گرادیان مضاعف عملکرد بهتری نسبت به دو روش دیگر داشته به‌گونه‌ای که شکل مربوط به بازسازی داده‌های آموزش به خط نیمساز نزدیک بوده و با مقدار ۰/۹۶ برای سنج R^2 و ppm ۷۱ برای $RMSE$ برتری فاحش روش تقویت گرادیان مضاعف را در زمینه بازسازی داده‌های آموزش به رخ می‌کشد. همچنین بایست اقرار داشت که در روش تقویت گرادیان مضاعف سنج‌های R^2 و $RMSE$ نیز برای بازتولید داده‌های تست از سایر روش‌ها بهتر بوده

طولانی بودن دوره ۱۸ ماهه صحت سنجی بایست به موفقیت روش تقویت گرادیان مضاعف در پیش‌بینی بلندمدت مقادیر متوسط ماهیانه شوری اقرار نمود. در ادامه، جدول (۱) ارائه می‌گردد که در آن هر سه مدل ساخته شده از لحاظ تمام سنجه‌های ارزیابی دقت که در بخش ۲-۷ بیان شد، مقایسه می‌گردند. همچنین در جدول (۲) به مقایسه عملکرد مدل‌های ساخته شده از حیث دقت با کمک آزمون آماری DM که در بخش ۲-۷ بیان شد، می‌پردازیم.

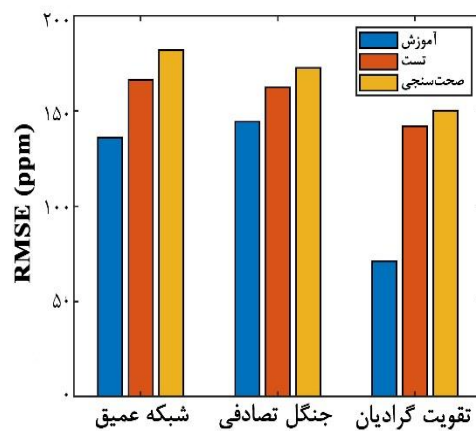
سنجه $NRMSE$ جهت ارزیابی عملکرد و دقت هر سه روش برای مجموعه داده‌های آموزش، تست و صحت سنجی پیش‌بینی‌ها، بررسی شده است. این شکل نیز نشان می‌دهد روش تقویت گرادیان مضاعف برتری خاصی بر سایر روش‌ها داشته است. کوچک بودن سنجه $NRMSE$ در هر سه روش نشان می‌دهد که تمام روش‌ها هم در پیش‌بینی مقادیر کم و هم در پیش‌بینی مقادیر زیاد شوری موفق عمل کرده و البته در روش تقویت گرادیان مضاعف این سنجه که به‌نوعی معرف خطای نسبی است کمتر از ۰/۲ می‌باشد. با توجه به



شکل ۱۰: مقایسه داده‌های متوسط ماهیانه شوری در ۱۸ ماهه زمانی صحت سنجی با مدل شوری ساخته‌شده توسط روش‌های (الف) شبکه عصبی عمیق، (ب) جنگل تصادفی و (ج) تقویت گرادیان مضاعف.



(ب)



(الف)

شکل ۱۱: مقادیر سنجه‌های (الف) $RMSE$ و (ب) $NRMSE$ به کاررفته در صحت سنجی مدل‌های ساخته شده برای بازه زمانی ۱۸ ماه

اگرچه شبکه عصبی عمیق آموزش دیده دارای سنجه R^2 برتری در آموزش در مقایسه با سایر مدل‌هاست، ولی در مورد داده‌های تست و صحت‌سنجی عملکرد ضعیف‌تری داشته است. همچنین بایست اذعان داشت با تنظیم پارامترهای مناسب از بیش‌برازش ممانعت شده است. بنابراین برای توجیه این موضوع بایست گفت در شبکه عصبی چندلایه MLP به خصوص با تعداد نوروها و لایه‌های بهینه، پتانسیل بالایی برای بازسازی رفتارهای پیچیده غیرخطی در طول آموزش وجود دارد که حتی بدون رخ دادن بیش‌برازش می‌تواند به عملکرد بالاتری نسبت به سایر مدل‌ها برای داده‌های آموزش منجر شود. اما بایست توجه داشت به علت دریافت وریانس‌های بالا در ورودی‌ها عملکرد MLP نسبت به روش‌های تجمعی و گروهی مانند جنگل تصادفی و تقویت گرادیان مضاعف که ذاتاً پایدارتر هستند، در مورد داده‌های تست و صحت‌سنجی ضعیف‌تر می‌شود. لذا، روش‌های گروهی می‌توانند حتی علیرغم اینکه در آموزش سنجه‌های ارزیابی ضعیف‌تری داشته باشند در مورد دوره صحت‌سنجی و داده‌های تست عملکرد برتری داشته باشند.

همانگونه که از جدول (۱) پیداست، در مقایسه عملکرد سه مدل یادگیری ماشین شامل شبکه عصبی عمیق، جنگل تصادفی و تقویت گرادیان مضاعف، می‌توان نتیجه گرفت که مدل تقویت گرادیان مضاعف عملکرد بهتری نسبت به سایر مدل‌ها داشته است. این مدل با داشتن کمترین مقادیر MAE ، $RMSE$ و همچنین کمترین مقدار $MAPE$ ، نشان می‌دهد که توانسته است میزان خطای پیش‌بینی را در مقایسه با سایر روش‌ها به میزان قابل توجهی کاهش دهد. همچنین مقدار بالای ضریب تعیین R^2 که برای این مدل در مرحله آموزش محاسبه شده نشان از قدرت بالای تطابق مدل با داده‌های آموزشی دارد. درمقابل، مدل شبکه عصبی عمیق علیرغم عملکرد قابل قبول در آموزش در صحت‌سنجی ضعیف‌تر عمل کرده و بالاترین خطا را شامل می‌گردد. مدل جنگل تصادفی از حیث صحت عملکردی بین دو مدل دیگر داشته است. در مجموع می‌توان ادعا کرد که مدل تقویت گرادیان مضاعف با پایداری و صحت بیشتر گزینه مناسب‌تری برای مسئله پیش‌بینی مکانی-زمانی شوری در رودخانه کارون بوده است.

جدول ۱: مقایسه سنج‌های مختلف ارزیابی عملکرد مدل‌ها در پیش‌بینی ارائه شده برای مجموعه داده‌های آموزش، تست و

صحت‌سنجی

سنجه مدل	MAPE (%)	MAE (ppm)	R^2	NRMSE	RMSE (ppm)		
						آموزش	شبکه
عمیق	۱۳,۳۸	۹۶,۶۲	۰,۸۵	۰,۱۱	۱۳۶,۱۰	تست	عصبی
	۱۵,۸۸	۱۱۷,۸۱	۰,۷۷	۰,۱۳	۱۶۶,۳۰	صحت‌سنجی	عمیق
	۲۰,۲۴	۱۳۶,۱۰	۰,۶۲	۰,۱۶	۱۸۲,۰۴	آموزش	جنگل
تصادفی	۱۳,۸۶	۱۰۲,۷۲	۰,۸۳	۰,۱۱	۱۴۳,۱۸	تست	جنگل
	۱۴,۸۸	۱۰۹,۴۹	۰,۷۸	۰,۱۳	۱۶۱,۹۳	صحت‌سنجی	تصادفی
	۱۹,۲۲	۱۲۷,۲۶	۰,۶۸	۰,۱۵	۱۶۷,۰۸	آموزش	تقویت
مضاعف	۶,۹۳	۴۹,۷۰	۰,۹۶	۰,۰۶	۷۱,۲۵	تست	گرادیان
	۱۵,۲۰	۱۰۲,۱۵	۰,۸۱	۰,۱۱	۱۴۲,۰۶	صحت‌سنجی	مضاعف
	۱۵,۸۶	۱۱۱,۲۶	۰,۷۴	۰,۱۳	۱۵۰,۲۳		

وجود برتری است. همانگونه از جدول ارزیابی می‌شود، مدل‌های جنگل تصادفی و تقویت گرادیان مضاعف بر مدل شبکه عصبی عمیق از حیث دقت برتری دارند. همچنین، مدل تقویت گرادیان مضاعف نیز در سطح اطمینان ۰/۰۵ برتری خود را بر مدل جنگل تصادفی اثبات می‌کند.

در جدول (۲)، نتایج آزمون آماری DM با اعداد ۰ و ۱ ارائه شده‌اند. عدد ۱ به معنای رد شدن فرض اولیه H_0 تعبیر شده که معنای آن اختلاف بین دقت مدل‌های بررسی شده و برتری مدل اول بر مدل دوم است. در مقابل، عدد ۰ به معنای پذیرفتن فرض اولیه H_0 و یکسان بودن عملکرد دو مدل از لحاظ دقت و عدم

جدول ۲: بررسی آماری عملکرد مدل‌های ساخته شده از حیث دقت با آزمون آماری DM

مدل‌ها	H_0	H_1
۱- جنگل تصادفی ۲- شبکه عصبی عمیق	۱	۰
۱- تقویت گرادیان مضاعف ۲- شبکه عصبی عمیق	۱	۰
۱- تقویت گرادیان مضاعف ۲- جنگل تصادفی	۱	۰

انحراف معیار نمونه‌های پیش‌بینی شده مربوط به هر آپک با این روش به عنوان یک تخمین تجربی از عدم قطعیت یا انحراف معیار مدل در نظر گرفته می‌شود. در ادامه برای اینکه بتوان ارزیابی مناسبی از عدم قطعیت

همچنین برای اینکه بتوان ریسک استفاده از هر مدل را برای پیش‌بینی‌ها در نظر گرفت، عدم قطعیت نظیر هر مدل مربوط به هر آپک درون بازه صحت‌سنجی با استفاده از روش بوت استرپینگ بررسی شد. همچنین

هر مدل در طول کل بازه صحت‌سنجی داشت، کران بالا و پایین بدست آمده از این روش و نهایتاً فاصله اطمینان برای پیش‌بینی‌های مدل‌ها در هر یک محاسبه می‌شود. متوسط مقادیر تخمینی کمیت‌های فوق در طول کل بازه صحت‌سنجی در جدول (۳) برای هر مدل ارائه شده‌است. همانگونه از نتایج پیداست، روش تقویت گرادیان مضاعف به طور میانگین کمترین انحراف معیار نمونه‌های بوت استرپ را داشته و نیز

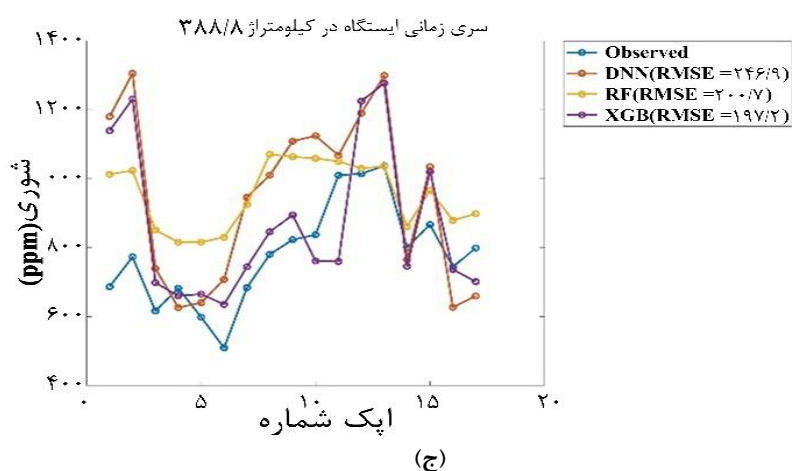
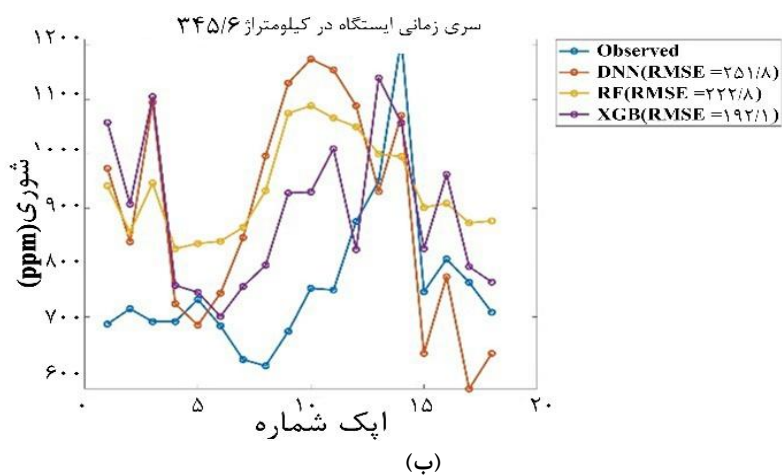
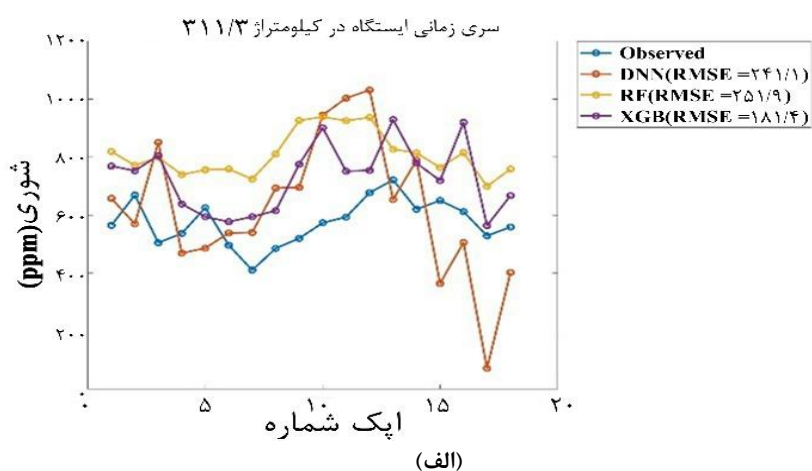
شامل کمترین فاصله اطمینان برآورد شده است که حاکی از عدم قطعیت مناسب این روش در مقایسه با سایر روش‌هاست. همچنین، روش‌های شبکه عصبی عمیق و جنگل تصادفی از حیث پارامترهای تخمین زده شده برای عدم قطعیت تفاوت زیادی نداشته و می‌توان گفت علیرغم سایر ارزیابی‌های مربوط به سنج‌های دقت، عملکرد روش شبکه عصبی عمیق کمی نسبت به جنگل تصادفی از حیث عدم قطعیت بهتر بوده است.

جدول ۳: متوسط کران‌های بالا، پایین، فاصله اطمینان و انحراف معیار نمونه‌های بوت استرپ برای ارزیابی عدم قطعیت هر مدل

مدل	متوسط کران پایین	متوسط کران بالا	متوسط فاصله اطمینان	متوسط انحراف معیار نمونه‌ها
شبکه عصبی عمیق	۵۳۵/۳	۱۱۱۱/۱	۵۷۵/۸	۱۳۶/۲
جنگل تصادفی	۵۷۴/۴	۱۱۸۸/۳	۶۱۳/۹	۱۴۳/۲
تقویت گرادیان مضاعف	۶۵۲/۹	۹۴۷/۱	۲۹۴/۲	۷۰/۶

در شکل (۱۲) به ارائه سری‌های زمانی حاصل از پیش‌بینی توسط روش‌های مختلف در ایستگاه‌های واقع در کیلومتر ۳۱۱/۳ (الف)، ۳۴۵/۶ (ب) و ۳۸۸/۸ (ج) می‌پردازیم. علت اینکه سری‌های زمانی پیش‌بینی‌شده برای این ایستگاه‌ها ارائه شده‌اند، آن است که این ایستگاه‌ها دارای بیشترین انحراف معیار مشاهدات بوده و می‌توانند صحت پیش‌بینی‌های مدل‌ها را به چالش بکشند. همان‌طور از شکل (دیده می‌شود در ایستگاه‌های واقع در کیلومتر ۳۱۱/۳، بیشترین انحراف معیار و پراکندگی برای شوری ثبت‌شده مشاهده شده است. لذا به دنبال مقایسه مدل‌ها از لحاظ توان پیش‌بینی این نوسانات هستیم. لازم به ذکر است از آنجایی که داده‌های مطالعه میانگین ماهیانه شوری در ایستگاه‌ها هستند، نوسانات خیلی زیاد در همه ایستگاه‌ها وجود ندارد. همچنین مقادیر شوری مشاهداتی در دوره صحت‌سنجی با رنگ آبی در شکل نمایش داده شده‌اند. اولاً بایست توجه داشت تمام

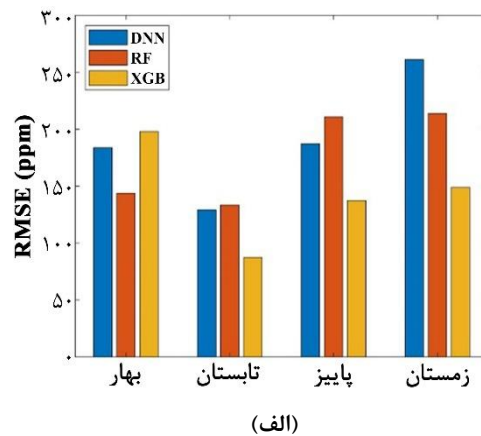
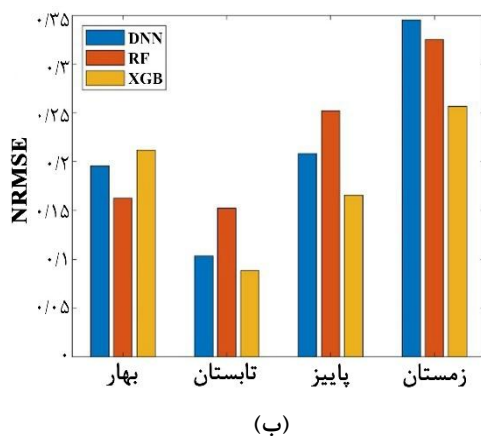
روش‌های به کار گرفته‌شده تا حدودی روندهای کاهشی و افزایشی را درست پیش‌بینی کرده‌اند. پیش‌بینی‌های ارائه شده توسط روش جنگل تصادفی نرم‌تر بوده و نسبت به سایر روش‌ها انحراف معیار کمتری داشته‌اند. روش شبکه عصبی عمیق در ماه‌های آخر دچار میرایی و کاهش شده به گونه‌ای که می‌توان گفت پیش‌بینی‌های انجام‌شده در دو ماه آخر در کیلومتر ۳۴۵/۶ و پیش‌بینی‌های چهار ماه آخر در ایستگاه واقع در کیلومتر ۳۱۱/۳ معتبر نیست. پیش‌بینی‌های روش تقویت گرادیان مضاعف هم به مقادیر مشاهداتی نزدیک‌تر بوده و هم روندهای افزایشی و کاهشی توسط این روش با دقت بالاتری پیش‌بینی شده‌اند. در سمت راست اشکال نیز سنج‌های $RMSE$ مربوط به دقت پیش‌بینی‌های هر روش برای هر ایستگاه نمایش داده شده است که همچنان بر عملکرد بهتر روش تقویت گرادیان مضاعف در پیش‌بینی‌های ارائه شده برای ایستگاه‌های فوق دلالت دارد.



شکل ۱۲: سری‌های زمانی متوسط ماهیانه شوری و مقادیر پیش‌بینی شده توسط مدل‌ها در ۱۸ یک صحت سنجی از آغاز ۱۳۹۶ تا میانه ۱۳۹۷ در: (الف) ایستگاه واقع در کیلومتر ۳۱۱/۳، (ب) ایستگاه واقع در کیلومتر ۳۴۵/۶ و (ج) ایستگاه واقع در کیلومتر ۳۸۸/۸

بهتر بوده و این برتری در فصول پاییز و زمستان نیز قابل رؤیت است. در فصل پاییز شرایط کمی متفاوت بوده و روش جنگل تصادفی ضعیف‌ترین عملکرد را از حیث دقت دارد. در زمستان نیز به ترتیب اولویت عملکرد بهتر با مدل تقویت گرادیان مضاعف، جنگل تصادفی و در آخر شبکه عصبی عمیق است. کاهش دقت شبکه عصبی عمیق در این فصل را می‌توان به تغییرات غیرخطی و نوسانات شوری رودخانه در اثر تغییرات بارش مربوط دانست. به طور کلی، مدل تقویت گرادیان مضاعف صحت و قابلیت اعتماد بیشتری را در مقایسه با سایر مدل‌ها در مقابل تغییرات زمانی و فصلی نشان می‌دهد. از شکل (۱۳-ب) مشهود است که کم یا زیاد شدن خطاها به بازه تغییرات شوری وابسته نبوده و سنجه $NRMSE$ که مستقل از مقیاس است عملکرد مشابهی با سنجه $RMSE$ را در فصول مختلف نشان می‌دهد.

همچنین به منظور بررسی عملکرد مدل در دوره‌های زمانی مختلف اعم از فصول خشک و تر نمودارهای مربوط به سنجه‌های $RMSE$ و $NRMSE$ مربوط به پیش‌بینی‌های مدل‌های ساخته شده در شکل (۱۳) نمایش داده شده است. شکل (۱۳-الف) معرف سنجه $RMSE$ برای پیش‌بینی‌های ارائه شده در بازه زمانی صحت سنجی به تفکیک فصول مختلف و مدل‌های به کار گرفته شده است. در شکل (۱۳-ب) نیز نظیر نمودارهای فوق برای سنجه $NRMSE$ ترسیم شده است. شکل فوق عملکرد مدل‌ها را از حیث صحت در فصول مختلف نشان می‌دهد. در بهار مدل جنگل تصادفی کمترین خطا را داشته و نسبت به سایر مدل‌ها عملکرد بهتری دارد. در تابستان به طور کلی دقت پیش‌بینی‌های همه مدل‌ها بهتر از سایر فصول است و این موضوع را می‌توان به کمتر بودن وریانس و تغییرات شوری در این فصل مربوط دانست. در تابستان عملکرد مدل تقویت گرادیان مضاعف نسبت به سایر مدل‌ها



شکل ۱۳: بررسی سنجه‌های (الف) $RMSE$ و (ب) $NRMSE$ مربوط به دقت مدل‌های ساخته شده برای پیش‌بینی‌های ارائه شده در فصول مختلف در دوره زمانی صحت سنجی

کشاورزی، صنعتی و شرب در جنوب غرب کشور ایفا می‌کند. داده‌های بلندمدت، نشان از روند افزایشی شوری در آب این رودخانه دارند؛ روندی که می‌تواند

۴- بحث و نتیجه‌گیری
رودخانه کارون به‌عنوان یکی از حیاتی‌ترین منابع آبی ایران، نقشی کلیدی در تأمین آب موردنیاز بخش‌های

مشاهده نشد، اما کاهش تدریجی دقت مدل شبکه عصبی در بازه‌های زمانی طولانی‌تر نشان داد که جنگل تصادفی در بلندمدت عملکردی باثبات‌تری دارد. این موضوع را می‌توان به کم بودن ورودی‌های شبکه عصبی ارتباط داد؛ زیرا صرفاً داده‌های ورودی مکانی، زمانی و دبی رودخانه بودند.

برتری مدل تقویت گرادیان مضاعف نه تنها در بازسازی دقیق‌تر داده‌های آموزش بلکه در تطابق بهتر با مشاهدات واقعی در دوره صحت‌سنجی مشهود بود. به‌ویژه در ایستگاه‌هایی که نوسانات مقادیر شوری بیشتر بوده‌اند، پیش‌بینی‌های این مدل در مقایسه با سایر مدل‌ها شباهت بیشتری به واقعیت داشتند. این امر نشان از توانایی این مدل در درک روابط غیرخطی پیچیده میان متغیرهای ورودی و خروجی دارد که از دیدگاه‌های علمی و کاربردی حائز اهمیت فراوان است. همچنین از حیث عدم قطعیت نیز مدل تقویت گرادیان مضاعف بهترین عملکرد را داشت.

درمجموع، نتایج این مطالعه حاکی از آن است که استفاده از روش‌های نوین یادگیری ماشین، به‌ویژه الگوریتم تقویت گرادیان مضاعف، می‌تواند ابزاری کارآمد برای مدل‌سازی و پیش‌بینی شوری در رودخانه‌های مهمی مانند کارون باشد. بهره‌گیری از این روش‌ها می‌تواند به تصمیم‌گیران و تدوینگران سیاست‌های آبی و مقابله با روندهای نگران‌کننده شوری کمک شایانی نماید. برای بهبود مطالعاتی که در آینده در این حیطه انجام خواهد شد، پیشنهاداتی ارائه می‌گردد. از جمله افزایش تعداد پارامترها و در نظر گرفتن عواملی نظیر دما، بارش، روان‌آب، پارامترهای کیفی و تبخیر در آموزش مدل‌ها می‌تواند تأثیر بسزایی داشته و حتی منجر به افزایش بازه زمانی برای پیش‌بینی بلند مدت کارآمد گردد. علاوه بر این، توسعه مدل‌هایی مبتنی بر فیزیک که می‌توانند با در نظر گرفتن معادلات فیزیکی کمبود داده‌های میدانی را تاحدی جبران نمایند، پیشنهاد می‌گردد.

تبعات اجتماعی و اقتصادی جدی برای جوامع محلی و نیز امنیت غذایی منطقه به همراه داشته باشد. در این پژوهش، باهدف مدل‌سازی دقیق و پیش‌بینی‌شده متوسط ماهیانه شوری در طول رودخانه کارون، از سه روش پرکاربرد یادگیری ماشین شامل شبکه عصبی عمیق، جنگل تصادفی و تقویت گرادیان مضاعف بهره گرفته شد. پارامترهای ورودی مدل‌ها شامل طول جغرافیایی، عرض جغرافیایی، کیلومتر از مبدأ رودخانه، سال، ماه و دبی جریان بودند. تحلیل اهمیت این پارامترها نشان داد که عوامل مکانی نسبت به پارامترهای زمانی تأثیر بیشتری بر غلظت شوری داشته‌اند و نیز پارامتر دبی به‌عنوان یک متغیر هیدرولوژیکی نقش پررنگ‌تری نسبت به سال و ماه ایفا نموده است.

آموزش مدل‌ها با استفاده از داده‌های مربوط به بازه زمانی ۱۳۸۰ تا ۱۳۹۵ انجام شد. در این میان، ۲۰ درصد داده‌ها به‌صورت تصادفی به‌عنوان داده‌های تست انتخاب‌شده و در فرآیند آموزش شرکت نداشتند. ارزیابی عملکرد مدل‌ها در بازسازی داده‌های تست نشان داد که مدل تقویت گرادیان مضاعف بهترین برازش را با داده‌های واقعی داشته است. مدل جنگل تصادفی نیز در اکثر موارد نسبت به شبکه عصبی عمیق عملکرد بهتری ارائه داد، هرچند این برتری در تمام ایستگاه‌های اندازه‌گیری مشاهده نشد و در ضمن از حیث عدم قطعیت مدل شبکه عصبی عملکرد بهتری داشت.

در گام بعد، از هر سه مدل برای پیش‌بینی مقادیر شوری در بازه زمانی ۱۸ ماهه از ابتدای سال ۱۳۹۶ تا میانه سال ۱۳۹۷ استفاده شد. مقادیر پیش‌بینی‌شده با داده‌های واقعی ثبت‌شده در همین بازه زمانی مورد صحت‌سنجی قرار گرفتند. نتایج این مرحله نیز مؤید برتری مدل تقویت گرادیان مضاعف در پیش‌بینی تمرکز شوری بود. در این میان، گرچه تفاوت معناداری از نظر آماری بین مدل جنگل تصادفی و شبکه عصبی عمیق

مراجع

- [1] E. Berger, O. Frör, and R. B. Schäfer, "Salinity impacts on river ecosystem processes: a critical mini-review," *Philosophical Transactions of the Royal Society B*, vol. 374, no. 1764, p. 20180010, 2019.
- [2] B. Biemond, H. E. de Swart, and H. A. Dijkstra, "Quantification of salt transports due to exchange flow and tidal flow in estuaries," *Journal of Geophysical Research: Oceans*, vol. 129, no. 11, p. e2024JC021294, 2024.
- [3] M. Halimi and Z. M. Isa, "Advection-Diffusion Equation with Spatially Dependent Coefficients for Instantaneous Pollutant Injection in a River," *Frontiers in Water and Environment*, vol. 5, no. 1, pp. 1-10, 2024.
- [4] R. Szymkiewicz, "A simplified approach for simulating pollutant transport in small rivers with dead zones using convolution," *Journal of Hydrology and Hydromechanics*, vol. 72, no. 4, pp. 538-546, 2024.
- [5] J. Xiong, J. Shen, and Q. Qin, "Exchange flow and material transport along the salinity gradient of a long estuary," *Journal of Geophysical Research: Oceans*, vol. 126, no. 5, p. e2021JC017185, 2021.
- [6] J. M. Hunter et al., "Framework for developing hybrid process-driven, artificial neural network and regression models for salinity prediction in river systems," *HESS*, vol. 22, no. 5, pp. 2987-3006, 2018.
- [7] J. Ubah, L. Orakwe, K. Ogbu, J. Awu, I. Ahaneku, and E. Chukwuma, "Forecasting water quality parameters using artificial neural network for irrigation purposes," *Sci. Rep.*, vol. 11, no. 1, p. 24438, 2021.
- [8] W. Chen, W. Liu, W. Huang, and H. Liu, "Prediction of salinity variations in a tidal estuary using artificial neural network and three-dimensional hydrodynamic models," *Computational Water, Energy, and Environmental Engineering*, vol. 6, no. 01, p. 107, 2017.
- [9] R. Zheng, Z. Sun, J. Jiao, Q. Ma, and L. Zhao, "Salinity Prediction Based on Improved LSTM Model in the Qiantang Estuary, China," *Journal of Marine Science and Engineering*, vol. 12, no. 8, p. 1339, 2024.
- [10] W. Huang and S. Foo, "Neural network modeling of salinity variation in Apalachicola River," *Water Res.*, vol. 36, no. 1, pp. 356-362, 2002.
- [11] A. M. Melesse et al., "River water salinity prediction using hybrid machine learning models," *Water*, vol. 12, no. 10, p. 2951, 2020.
- [12] J. Hu, B. Liu, and S. Peng, "Forecasting salinity time series using RF and ELM approaches coupled with decomposition techniques," *Stoch. Environ. Res. Risk Assess.*, vol. 33, pp. 1117-1135, 2019.
- [13] M. Kulisz, J. Kujawska, Z. Aubakirova, and E. Wojtas, "Prediction of river salinity with artificial neural networks," in *Journal of Physics: Conference Series*, 2023, vol. 2676, no. 1: IOP Publishing, p. 012004.
- [14] K. M. R. Rasha, "Salinity Prediction at the Bhairab River in the South-Western Part of Bangladesh Using Artificial Neural Network," *Nature Environment and Pollution Technology*, vol. 21, no. 3, pp. 1431-1438, 2022.
- [15] P. Duc, "Harnessing Lstm and Xgboost Algorithms for Salinity Prediction in Mekong Delta's Main Rivers," Available at SSRN 4965853.
- [16] M. Karbasi et al., "Multi-step ahead forecasting of electrical conductivity in rivers by using a hybrid Convolutional Neural Network-Long Short-Term Memory (CNN-LSTM) model enhanced by Boruta-XGBoost feature selection algorithm," *Sci. Rep.*, vol. 14, no. 1, p. 15051, 2024.
- [17] M. Niazkar et al., "Applications of XGBoost in water resources engineering: A

- systematic literature review (Dec 2018–May 2023)," *Environmental Modelling & Software*, vol. 174, p. 105971, 2024.
- [18] R. Dehghani and D. Abbaspour, "A Comparison of Statistical and Intelligent Methods for Estimating River Salinity: A Case Study of Simineroud, Western Azerbaijan, Iran," (in Fa), *International Bulletin of Water Resources and Development*, vol. 2, no. 3, p. 182, 2014.
- [19] S. Kanani, G. Asadollahfardi, and A. Ghanbari, "Application of artificial neural network to predict total dissolved solid in Achechay River basin," *World Applied Sciences Journal*, vol. 4, no. 5, pp. 646-654, 2008.
- [20] T. Hastie, R. Tibshirani, and J. Friedman, "The elements of statistical learning. Springer series in statistics," New York, NY, USA, 2001.
- [21] T. K. Ho, "Random decision forests," in *Proceedings of 3rd international conference on document analysis and recognition*, 1995, vol. 1: IEEE, pp. 278-282.
- [22] T. K. Ho, "The random subspace method for constructing decision forests," *IEEE transactions on pattern analysis and machine intelligence*, vol. 20, no. 8, pp. 832-844, 1998.
- [23] T. Chen and C. Guestrin, "Xgboost: A scalable tree boosting system," in *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pp. 785-794, 2016.



Comparison and evaluation of machine learning methods for salinity prediction in the Karun River

Hani Mahbuby ^{1*}, Mohammad Pirayesh ², Yahya Jamour ³

1- Assistant professor in Faculty of Civil, Water and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

2- MSc student of Hydrography, Faculty of Civil, Water and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

3 - Associate professor in Faculty of Civil, Water and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

Abstract

Most regions of Iran are characterized as a hot and arid climate with low precipitation. The primary water resources in the country are groundwater and river water. In recent years, due to little rainfall and drought conditions, rivers have received more attention. In this context, river water salinity is considered as one of the most critical water quality parameters, requiring special attention.

In this study, the long-term variations in the salinity of the Karun River were assessed. Furthermore, the monthly average salinity of the Karun River was modeled from 2001 (1380) to 2016 (1395) using various machine learning methods, including deep neural networks (DNN), random forests (RF), and extreme gradient boosting (XGBoost). Subsequently, the average salinity was forecasted for an 18-month period extending to mid-2018 (1397). The predicted values were compared with the observed measurements during this 18-month validation period, and the performance and accuracy of the different modeling approaches were evaluated and compared.

The results demonstrated that, first, the monthly average salinity increased at a rate of approximately 10 ppm/year throughout the study period, posing a potential threat to the regional ecosystem. The DNN and RF models showed comparable accuracy and could predict the average salinity of the river with the accuracy of 170 _ 180 ppm when their respective hyperparameters were optimally tuned. However, the RF performed slightly better in long-term forecasting.

Among the tested methods, XGBoost outperformed the others, achieving a prediction error of approximately 150 ppm. Compared to RF and DNN, this represents a relative error reduction of about 13% and 18%, respectively.

Key words : River salinity, Deep neural network, Random forest, Extreme gradient boosting.