

در یادگیری نظارت‌شده، دو رویکرد اصلی وجود دارد: * سیستم‌های مبتنی بر *Regression* * سیستم‌های مبتنی بر *Classification*

- **رگرسیون (*Regression*):** برای پیش‌بینی متغیرهای کمی و پیوسته استفاده می‌شود. این روش ضریب و شیب خط را محاسبه می‌کند. در رگرسیون سری زمانی، بعد از زمان نیز اضافه می‌شود. به‌عنوان مثال، از قند خون قبلی بیمار می‌توان برای پیش‌بینی هم‌گلوبین *A1C* استفاده کرد یا از علائم بالینی برای پیش‌بینی بدتر شدن وضعیت بیمار استفاده کرد. یا ریسک عوارض قلبی عروقی مریض را بر اساس عوامل بالینی فورکست کرد. در یک سیستم مبتنی بر رگرسیون، عملکرد با ارزیابی تعداد خطاهای پیش‌بینی اندازه‌گیری می‌شود. عملکرد مدل با معیارهایی مانند **ضریب تعیین (R^2)**، **میانگین خطای مطلق (*MAE*)**، **خطای نسبی مطلق (*RAE*)**، و **ریشه میانگین مربعات خطا (*RMSE*)** ارزیابی می‌شود.

* در **رگرسیون خطی (*linear regression*)**، هدف تنظیم فاصله بین نقاط واقعی و پیش‌بینی‌شده برای حداقل کردن خطاست. خطا به صورت مجموع مربعات تفاضل مقادیر واقعی و پیش‌بینی‌شده محاسبه می‌شود. مدل می‌تواند خطی (درجه یک) یا غیرخطی (درجه دو یا سه) باشد که دقت بیشتری ایجاد می‌کند. شرط برابری واریانس‌ها در بسیاری از داده‌ها برقرار نیست و نویز در مدل‌های خطی می‌تواند بالا باشد. مدل خطی که به صورت $f(x)=ax+b$ تعریف می‌شود، ضرایب a و b به‌گونه‌ای تنظیم می‌شوند که تخمین به‌دقیق‌ترین و نزدیک‌ترین مقدار ممکن برسد. این کار به این صورت انجام می‌شود که مدل در ابتدا دو عدد تصادفی به ضرایب اختصاص می‌دهد و طی فرایند آموزش (*Training*)، مقادیر را به‌روزرسانی می‌کند تا به مقادیر مطلوب برسد. **مولتیپل رگرسیون یا مولتی وریبل** شامل چند متغیر مستقل است، درحالی‌که **مولتی‌وریت آنالیز** به معنای داشتن چندین متغیر وابسته به عنوان پیامد است.

قانون اعداد بزرگ در یادگیری ماشین نشان می‌دهد که نتیجه تعداد زیادی آزمایش مشابه به یک مقدار مشخص همگرا می‌شود. هدف، به حداقل رساندن خطای مدل و ارائه دقیق‌ترین پیش‌بینی‌هاست. این فرایند تکرار می‌شود تا سیستم به بهترین مقادیر برای a و b برسد. به این ترتیب، ماشین از طریق تجربه یاد می‌گیرد و برای استفاده آماده می‌شود. در بسیاری از زمینه‌های یادگیری نظارت‌شده، هدف تنظیم دقیق تابع پیش‌بینی‌کننده همان فرضیه می‌باشد. یادگیری شامل استفاده از الگوریتم‌های ریاضی برای نمایش داده‌های ورودی x در یک حوزه خاص است. به عنوان مثال، x می‌تواند نشان‌دهنده روز باشد و تابع می‌تواند پیش‌بینی مدت زمان انتظار در یک بیمارستان باشد.

دو رویکرد اصلی برای انتخاب ویژگی‌ها وجود دارد:

1. **روش فیلتر (*Filter Method*):** ویژگی‌ها به‌طور مستقل و بدون وابستگی به داده‌ها ارزیابی و قبل از ساخت مدل فیلتر می‌شوند. این روش بر اساس الگوریتمی ثابت عمل می‌کند و وابسته به ویژگی‌های خاص داده‌ها نیست.
2. **روش روکش‌دار (*Wrapper*):** با استفاده از الگوریتم یادگیری ماشین، زیرمجموعه‌ای از داده‌ها بر اساس دقت مدل انتخاب می‌شود. این روش مشخص می‌کند که کدام متغیرها و آستانه‌ها (*Cutoff*) برای بهینه‌سازی مدل استفاده شوند، به‌طوری‌که دقت و تقاطع مدل با خط y بهینه باشد.

* رگرسیون لجستیک (*logistic regression*): احتمال وقوع را بین صفر و یک قرار می‌دهد و داده‌ها را به دو دسته طبقه‌بندی می‌کند. اسم آن رگرشن می‌باشد ولی در واقع برای تقسیم‌بندی استفاده می‌شود دقیقاً چون لگاریتم کار می‌کند بین صفر و یک می‌باشد و مثلاً بیان می‌کند که اگر احتمال کمتر از نیم باشد به سمت صفر می‌رود و اگر احتمال بیشتر از نیم باشد به سمت یک می‌رود و می‌تواند در دو گروه صفر و یک طبقه‌بندی کند.

- طبقه‌بندی (*Classification*): برای دسته‌بندی کردن متغیرهای کیفی استفاده می‌شود. به‌عنوان مثال، در مورد هم‌گلوبین *A1C*، مقدار ۷.۸ درصد به‌عنوان دیابت در نظر گرفته می‌شود چون بالاتر از کات اف تعیین شده می‌باشد. همچنین برای تشخیص وضعیت‌های مختلف مانند پیامدهای دو یا چند حالت مثل بهبودی و عدم بهبودی، دیابتی بودن یا نبودن، مرگ و زندگی می‌توان از این مورد استفاده کرد یا اینکه بر اساس پروفایل سلامت تشخیص بدیم آیا رتینوپاتی دیابتی وجود دارد یا ندارد. مناسب‌ترین روش درمانی چیست مثلاً ادامه درمان است صبر کردن هست یا تغییر کردن دار. دسته‌بندی فرآیندی است که در آن یک یا چند آیت (x) به یکی از دسته‌های از پیش تعریف‌شده و گسسته (یا برچسب‌ها y) تخصیص داده می‌شود. الگوریتم‌های دسته‌بندی از ویژگی‌ها (یا صفات) برای تعیین چگونگی تخصیص یک آیت به حداقل یک کلاس دودویی (دو دسته) یا دو یا چند کلاس استفاده می‌کنند. داده‌های آموزشی به صورت یک بردار ورودی ارائه می‌شوند. این بردار شامل برچسب‌هایی برای ویژگی‌های آیت است. به عنوان مثال، یک آیت ممکن است به این صورت تعریف شود: ("*HbA1c, 7.8%*", "دیابتی") به شکل $\{x_training, output\}$

*** الگوریتم‌هایی که فرایند دسته‌بندی را انجام می‌دهند شامل درخت‌های تصمیم‌گیری (*Decision Tree*)، جنگل تصادفی (*Random Forest*)، بیز ساده (*Naive Bayesian*)، رگرسیون لجستیک، نزدیک‌ترین همسایگان (*kNN*)، ماشین بردار پشتیبان (*SVM*)، و غیره هستند.

درخت تصمیم (*Decision Tree*):

درخت تصمیم یک نمودار درختی است که با تقسیم داده‌های بزرگ به بخش‌های کوچک‌تر و تعیین قوانین *if-then*، تصمیم‌گیری را انجام می‌دهد. ساختار درخت شامل ریشه به‌عنوان گره اولیه، شاخه‌ها برای گسترش درخت، و برگ‌ها به‌عنوان گره‌های نهایی است که نتیجه تصمیم‌گیری را مشخص می‌کنند. یکی از مشکلات درخت تصمیم، افزایش گره‌ها و شاخه‌ها است که احتمال بیش‌از‌اش (*Overfitting*) و پیچیدگی محاسبات را افزایش می‌دهد. برای رفع این مشکل از فرایند هرس (*Pruning*) استفاده می‌شود. سه الگوریتم اصلی برای ساخت درخت تصمیم وجود دارد که تفاوت آن‌ها در معیار هزینه‌ای است که برای انتخاب گره‌ها یا ویژگی‌ها استفاده می‌شود. یکی از ویژگی‌ها این است که اولویت‌ها و روابط در تصمیم‌گیری به وضوح مشخص می‌باشند و خروجی به شکل مجموعه‌ای از قوانین بیان می‌شود. در نهایت، این درخت‌ها فرایند تصمیم‌گیری را به‌صورت قوانین شفاف و قابل‌اجرا نمایش می‌دهند.

جنگل تصادفی (Random Forest) :

جنگل تصادفی به جای یک درخت تصمیم، از مجموعه ای از درخت ها تشکیل شده که هر کدام متغیری خاص را نشان می دهند. این روش به عنوان یک تکنیک تجمعی (Ensemble) شناخته می شود، زیرا با ترکیب پیش بینی های آن ها دقت مدل را افزایش و واریانس را کاهش می دهد. در رگرسیون، میانگین پیش بینی ها و در طبقه بندی، دسته ای که اکثریت آرا را دارد نتیجه نهایی را تعیین می کنند. این روش نسبت به درخت تصمیم دقت بیشتری دارد، کمتر دچار بیش برآش می شود، با داده های ناقص بهتر عمل می کند و تخمین های بهتری ارائه می دهد. در واقع، هر درخت در جنگل تصادفی به طور مستقل عمل کرده و پیش بینی خاص خود را ارائه می دهد، اما نتیجه نهایی بر اساس ترکیب این پیش بینی ها مشخص می شود. جنگل تصادفی رویکردی قوی تر و مناسب برای مسائل پیچیده طبقه بندی و رگرسیون با متغیرهای متنوع است.

مدل های یادگیری تجمعی (Ensemble Learning)

مدل های ترکیبی (Ensemble Methods) برای بهبود دقت پیش بینی از ترکیب چند مدل استفاده می کنند و به طور همزمان از مزایای مدل های مختلف بهره برداری کنند و در نتیجه عملکرد بهتری نسبت به مدل های تک بعدی داشته باشند. شامل دو روش اصلی *Boosting* و *Bagging* هستند.

- Boosting:** در این روش، مدل ها به صورت ترتیبی آموزش داده می شوند و هر مدل جدید سعی می کند خطاهای مدل قبلی را اصلاح کند. در هر مرحله، مدل جدید نقاطی که مدل قبلی قادر به پیش بینی صحیح آن ها نبوده را شناسایی کرده و سعی می کند این نقاط را درست پیش بینی کند. این فرآیند باعث افزایش تدریجی دقت و کاهش بیش برآش (*Overfitting*) می شود. ویژگی بارز *Boosting*، یادگیری مدل جدید از خطاهای مدل قبلی است. یکی از تکنیک های پیشرفته در این حوزه *Gradient Boosting* است که از درخت های تصمیم گیری و پیش بینی کننده های ضعیف برای کاهش خطا استفاده می کند. نسخه پیشرفته تر آن *XGBoost*، با اضافه کردن منظم سازی (*Regularization*)، قادر است از **واریانس بالا** جلوگیری کرده و دقت پیش بینی را بهبود بخشد. عملکرد بهتری در داده های پیچیده دارد.
- Bagging:** هدف این روش کاهش واریانس و جلوگیری از بیش برآش است. مدل ها به طور مستقل و موازی آموزش داده می شوند و پیش بینی نهایی از میانگین یا ترکیب نتایج آن ها به دست می آید. این روش از نمونه گیری بوت استرپ برای آموزش مدل ها استفاده می کند، که به طور تصادفی داده ها را برای آموزش مدل ها انتخاب می کند. باعث کاهش رفتارهای غیرعادی در پیش بینی می شود. زیرا پیش بینی ها از مدل های مختلف با داده های متفاوت میانگین گیری می شوند. *Random Forest* یکی از معروفترین الگوریتم های *Bagging* است که از ترکیب درخت های تصمیم بهره می گیرد.

ماشین بردار پشتیبان (SVM) (Support Vector Machine) :

SVM یک الگوریتم طبقه بندی خطی باینری غیر احتمالاتی است که برای مسائل طبقه بندی و رگرسیون به کار می رود. این الگوریتم با یافتن ابرصفحه ای که بهترین تفکیک را بین دو کلاس دارد عمل می کند. بردارهای پشتیبان، نقاط داده نزدیک به ابرصفحه هستند که اگر حذف شوند، موقعیت ابرصفحه تغییر می کند. هرچه مقدار حاشیه یا فاصله نقاط داده از ابرصفحه بیشتر باشد، اطمینان بیشتری وجود دارد که داده ها به درستی طبقه بندی شده اند. هدف *SVM* بیشینه کردن حاشیه یا فاصله نقاط داده از ابرصفحه است تا اطمینان بیشتری از طبقه بندی صحیح داشته باشند. این الگوریتم از "هستساز" برای نگاشت داده ها به فضاهای با ابعاد بالاتر و تسهیل طبقه بندی استفاده می کند. داده ها با استفاده از ترنند هسته به صورت تکراری به ابعاد بالاتری نگاشت می شوند تا یک ابرصفحه برای طبقه بندی آن ها تشکیل شود. این روش در مجموعه داده های کوچک و در مسائلی مانند تحلیل متن، تحلیل احساسات، تشخیص تصویر و شناسایی دست نوشته های عددی عملکرد خوبی دارد. اما با افزایش ابعاد داده ها، قابلیت درک مدل کاهش یافته و زمان آموزش افزایش می یابد. همچنین در مواجهه با داده های پرنویز عملکرد کمتری دارد.

نایو بیز (Naïve Bayes):

نایو بیز الگوریتمی برای محاسبه احتمال وقوع یک رویداد بر اساس وقوع رویدادهای دیگر است. این روش با فرض استقلال متغیرها، که در واقعیت کمتر دیده می شود، عمل می کند و برای مجموعه داده های با ابعاد بالا کاربرد دارد. این الگوریتم از فرمول زیر استفاده می کند: $P(B/A) = (P(A/B) * P(B)) / P(A)$ با این فرمول می توان احتمال تعلق یک نمونه به یک کلاس را تعیین کرد؛ مثلاً احتمال خطر بیماری با توجه به وضعیت سلامتی بیمار محاسبه می شود و کلاسی که بیشترین احتمال را دارد، انتخاب می گردد. نایو بیز ساده، اما بسیار کارآمد است. مدل های بایس ساده به راحتی ساخته می شوند و برای مجموعه داده های بزرگ بسیار مفید هستند. علی رغم سادگی، بایس ساده به خاطر عملکرد بهتر نسبت به روش های طبقه بندی پیچیده تر شناخته شده است.

نزدیک ترین همسایه (K-NN (K-Nearest Neighbour):

نباید روش *kNN* را با خوشه بندی *k-means* اشتباه گرفت. روش *kNN* یک شیء ناشناخته را با برچسب اکثریت نزدیکترین *k* همسایگان طبقه بندی می کند. در واقع یک تکنیک غیر پارامتری برای طبقه بندی و رگرسیون است که به جای یادگیری مدل، داده های آموزشی را ذخیره کرده و طبقه بندی نمونه جدید را بر اساس یادگیری با قیاس انجام می دهد. از آنجا که مدلی یاد گرفته نمی شود، *kNN* به عنوان یک یادگیرنده تنبل شناخته می شود. فاصله بین شیء و همسایه ها با استفاده از فاصله اقلیدسی محاسبه می شود. در واقع الگوریتم فاصله بین هر دو نقطه را محاسبه کند. نزدیکترین همسایگان را بر اساس این فواصل زوجی پیدا کند. برای تعیین برچسب کلاس، بر اساس لیست نزدیکترین همسایگان رای گیری اکثریت انجام دهد. پیش بینی در زمان درخواست انجام می شود. در مسائل رگرسیون، میانگین یا میانه *k*-مشابهترین نمونه ها استفاده می شود. در مسائل طبقه بندی، کلاسی که بیشترین فراوانی را در میان *k*-مشابهترین نمونه ها دارد، انتخاب می شود.

از معایب آن می توان به موارد زیر اشاره کرد. 1) نیاز به محاسبات سنگین برای هر پیش بینی چون برای طبقه بندی یک شیء، فاصله همه همسایگان در مجموعه داده آموزشی باید محاسبه شود 2) عدم سازگاری با داده های با ابعاد بالا چون هر متغیر پیش بین به عنوان یک بعد از فضای ورودی در نظر گرفته می شود. افزایش ابعاد به صورت نمایی حجم فضای ورودی را افزایش می دهد. 3) عدم توانایی در مواجهه با داده های ناقص اشاره کرد، زیرا فاصله بین بردارها در صورت وجود داده های ناقص نمی تواند محاسبه شود.

نتیجه گیری: مدل های مختلف هر کدام نقاط قوت و ضعف خودشان را دارند و شرایط خاصی کاربرد دارند و انتخاب بهترین مدل به نوع داده ها و اهداف تحقیق بستگی دارد.