

*** دو روش برای ارزیابی مدل های یادگیری ماشین وجود دارد: **ارزیابی آفلاین**، **ارزیابی آنلاین**

ارزیابی آفلاین مدل را بر اساس معیار هایی که از داده های تاریخی، ایستا و توزیع شده یاد گرفته و ارزیابی شده اند، انداز گیری می کند. معیار هایی مانند دقت و *precision-recall* معمولاً در مرحله آموزش آفلاین استفاده می شوند. تکنیک های ارزیابی آفلاین شامل روش *hold-back* و اعتبار سنجی متقاطع *n-fold* هستند.

ارزیابی آنلاین به ارزیابی معیار ها پس از استقرار مدل اشاره دارد. نکته کلیدی این است که این معیار ها ممکن است با معیار های مورد استفاده برای ارزیابی عملکرد مدل در حالت استقرار زنده متفاوت باشند. ارزیابی آنلاین، به ویژه در عصر دیجیتال، می تواند از تست های چندمتغیره برای درک بهترین مدل های عملکردی پشتیبانی کند. حلقه های بازخورد (*Feedback Loops*) برای اطمینان از عملکرد درست و مورد نظر سیستم ها کلیدی هستند و به درک بهتر مدل در زمینه استفاده از آن کمک می کنند. این بازخوردها می توانند از طریق یک عامل انسانی، یک عامل خودکار هوشمند مبتنی بر زمینه و یا کاربران مدل جمع آوری شوند.

*** چگونه می توانم مدل خود را ارزیابی کنم؟

تعیین معیار های ارزیابی مناسب برای سنجش عملکرد مدل بر اساس هدفی که مدل برای آن طراحی شده است، می باشد. اغلب، داده های استفاده شده از انتخاب الگوریتم مهتر هستند؛ هر چه ویژگی های استفاده شده بهتر باشند، عملکرد مدل نیز بهبود می یابد. برای معیار های ارزیابی از کتابخانه هایی مانند *metrics* در *R* و *scikit-learn* در پایتون می توان استفاده کرد.

• برای ارزیابی مدل های طبقه بندی (Classification)، معیار های رایج عبارتند از:

*** دقت (*Accuracy*) ساده ترین تکنیک برای شناسایی این است که آیا یک مدل پیش بینی های صحیح انجام می دهد یا خیر. این معیار به عنوان درصد پیش بینی های صحیح از کل

پیش بینی های انجام شده محاسبه می شود. $Accuracy = \text{number of correct predictions} / \text{number of total predictions}$

*** ماتریس در هر یختگی (*Confusion Matrix*) طبقه بندی های صحیح و نادرست مدل را تفکیک کرده و آن ها را به برچسب های مناسب نسبت می دهد. در نتیجه، دقت را می توان به شکل زیر باز نویسی کرد $(TP + TN) / (TP + TN + FP + FN)$

*** صحت (*Precision*) بررسی می کند که از بین مواردی که مدل به عنوان مرتبط پیش بینی کرده است، چند مورد واقعاً مرتبط هستند.

*** فراخوانی (*Recall*) بررسی می کند که از بین مواردی که واقعاً مرتبط هستند، مدل چند مورد را به درستی پیش بینی کرده است.

$Precision: (\text{correctly predicted Positive}) / (\text{total predicted Positive}) = TP / TP + FP$

$Recall: (\text{correctly predicted Positive}) / (\text{total correct Positive observation}) = TP / TP + FN$

- حساسیت (*Sensitivity*): این معیار همان فراخوانی (*Recall*) است و نشان می دهد که از بین موارد واقعاً مثبت، مدل چند مورد را به درستی شناسایی کرده است.

- اختصاصیت (*Specificity*): این معیار بررسی می کند که مدل چقدر خوب می تواند موارد منفی را به درستی شناسایی کند.

$Specificity: (\text{correctly predicted Negative}) / (\text{total Negative observation}) = TN / TN + FP$

*** معیار $F\text{-measure} / F1 \text{ score}$: این معیار فراتر از میانگین حسابی رفته و میانگین هارمونیک صحت و فراخوانی را محاسبه می کند که در آن p نشان دهنده دقت (*Precision*) و r نشان دهنده فراخوانی (*Recall*) است. این معیار برای متعادل سازی دقت و فراخوانی در مواردی که توزیع کلاس ها نامتوازن است، مفید است.

$$F = \frac{1}{\frac{1}{2} \left(\frac{1}{p} + \frac{1}{r} \right)} = \frac{2pr}{p+r}$$

*** *Logarithmic Loss* (یا *Log-loss*): برای مسائلی استفاده می شود که در آن ها یک احتمال پیوسته پیش بینی می شود، نه یک برچسب کلاس. *Log-loss* یک معیار احتمالی برای سنجش اطمینان از دقت ارائه می دهد و آنتروپی بین توزیع برچسب های واقعی و پیش بینی ها را در نظر می گیرد. برای یک مسئله طبقه بندی دودویی به صورت زیر محاسبه می شود:

$$\text{Log-loss} = -\frac{1}{N} \sum_{i=1}^N y_i \log p_i + (1 - y_i) \log(1 - p_i)$$

*** *AUC* (مساحت زیر منحنی): نرخ مثبت های صحیح (*True Positives*) را در مقابل نرخ مثبت های کاذب (*False Positives*) رسم می کند. این معیار امکان تجسم حساسیت (*Sensitivity*) و اختصاصیت (*Specificity*) طبقه بندی را فراهم می کند و نشان می دهد که چند طبقه بندی مثبت صحیح می توان با تحمل مثبت های کاذب به دست آورد. این منحنی به عنوان *ROC* شناخته می شود مشخصه عملکرد سیستم (*Receiver Operating Characteristic Curve*) یا *ROC* شناخته می شود

• برای ارزیابی مدل های رگرسیون (Regression)، معیار های رایج عبارتند از:

*** 1. *Root-Mean-Squared Error (RMSE)*: این معیار جذر میانگین مربعات خطاها است و میزان اختلاف بین مقادیر پیش بینی شده و مقادیر واقعی را انداز گیری می کند. *RMSE* به خطاهای بزرگ حساس تر است، زیرا خطاها قبل از میانگین گیری به توان ۲ می رسند.

$$RMSE = \sqrt{(1/n) * \sum (y_i - \hat{y}_i)^2}$$

*** 2. *Mean-Squared Error (MSE)*: این معیار میانگین مربعات خطاها است و مشابه *RMSE* عمل می کند، اما جذر گرفته نمی شود. *MSE* نیز به خطاهای بزرگ حساس است.

$$MSE = (1/n) * \sum (y_i - \hat{y}_i)^2$$

*** 3. *Mean-Absolute Error (MAE)*: این معیار میانگین قدر مطلق خطاها است و نسبت به *RMSE* و *MSE* به خطاهای بزرگ حساسیت کمتری دارد، زیرا خطاها به توان ۲ نمی رسند.

$$MAE = (1/n) * \sum |y_i - \hat{y}_i|$$

4.4 (Median Absolute Percentage Error (MAPE))** این معیار میانه درصد خطاها را محاسبه می‌کند و برای داده‌هایی که دارای مقادیر پرت (Outliers) هستند، مناسب است.

$$MAPE = median\left(\left| \frac{y_i - \hat{y}_i}{y_i} \right| \right),$$

5.2 (ضریب تعیین)** این معیار نشان می‌دهد که مدل چقدر از واریانس داده‌ها را توضیح می‌دهد. مقدار R^2 بین 0 و 1 است. 1: مدل تمام واریانس داده‌ها را توضیح می‌دهد. 0: مدل هیچ واریانس را توضیح نمی‌دهد.

$$R^2 = 1 - (SSR/SST) = 1 - \frac{\sum (y_i - \hat{y}_i)^2}{\sum (y_i - \bar{y})^2}$$

******* یک دانشمند داده باتجربه همه داده‌ها را با شک و تردید بررسی می‌کند. داده‌های نامتعادل (Skewed Datasets)، anomalies (ناهنجاری‌ها)، داده‌های نادر (Rare Data)، داده‌های ناسازگار (Inconsistent)، نمونه‌های کلاس نامتعادل (Imbalanced Class Examples)، و outliers (مقادیر پرت) همگی می‌توانند به طور قابل توجهی بر عملکرد مدل تأثیر بگذارند. وجود نمونه‌های بیشتر در یک کلاس نسبت به کلاس دیگر می‌تواند منجر به عملکرد ضعیف مدل شود. تأثیر outliers بزرگ را می‌توان با استفاده از درصد‌های خطا (percentiles of error) کاهش داد. پاکسازی خوب داده‌ها، حذف outliers، و نرمال‌سازی متغیرها از جمله اقدامات مهم برای بهبود کیفیت داده‌ها هستند.

****تنظیم هایپر پارامترها (Hyperparameter Tuning)****

تنظیم هایپر پارامترها به معنای انتخاب مجموعه‌ای از هایپر پارامترهای بهینه برای یک مدل یادگیری ماشین است. مقادیر بهینه هایپر پارامترها دقت پیش‌بینی مدل را به حداکثر می‌رسانند. این فرآیند شامل آموزش مدل، ارزیابی دقت کلی و تنظیم مناسب هایپر پارامترها است. با آزمایش مقادیر مختلف هایپر پارامترها، بهترین مقادیر برای مسئله مشخص می‌شوند که دقت کلی مدل را بهبود می‌بخشند. نمونه‌هایی از هایپر پارامترها در شبکه‌های عصبی (Neural Networks) تعداد و اندازه لایه‌های پنهان، وزن‌ها، نرخ یادگیری (Learning Rate) و غیره و یا درخت‌های تصمیم (Decision Trees) عمق درخت و تعداد برگ‌ها می‌باشند. تنظیم هایپر پارامترها شبیه به آموزش یک مدل یادگیری ماشین است و به عنوان یک مسئله بهینه‌سازی در نظر گرفته می‌شود. رایج‌ترین روش‌های تنظیم هایپر پارامترها 1. جستجوی شبکه‌ای (Grid Search): بررسی تمام ترکیبات ممکن از هایپر پارامترها. 2. جستجوی تصادفی (Random Search): بررسی تصادفی ترکیبات مختلف از هایپر پارامترها می‌باشد. حالت اول بسیار دقیق تر عمل میکند اما حالت دوم سریع تر و با محاسبات کمتر این کار را انجام میدهد.

****آزمون چندمتغیره (Multivariate Testing)**** یک روش بسیار مفید برای تعیین بهترین مدل برای مسئله خاص است. این روش که به عنوان "آزمون فرضیه آماری" نیز شناخته می‌شود، تفاوت بین فرضیه صفر (Null Hypothesis) و فرضیه جایگزین (Alternative Hypothesis) را بررسی می‌کند. - فرضیه صفر: ** مدل جدید تأثیری بر میانگین مقدار معیار عملکرد ندارد. - فرضیه جایگزین: ** مدل جدید میانگین مقدار معیار عملکرد را تغییر می‌دهد. این روش مدل‌های مشابه را مقایسه می‌کند تا بهترین عملکرد را شناسایی کند یا مدل جدید را با مدل قدیمی‌تر مقایسه می‌نماید. معیارهای عملکرد مربوطه مقایسه شده و تصمیم گرفته می‌شود که با کدام مدل ادامه دهند.

اگرچه این فرآیند ساده به نظر می‌رسد، اما چند نکته کلیدی برای توجه وجود دارد.

کدام معیار را برای ارزیابی استفاده کنیم؟ انتخاب معیار مناسب برای ارزیابی مدل به مورد استفاده بستگی دارد. تکرار آزمایش و انجام ارزیابی‌های مکرر می‌تواند احتمال نتایج گمراه‌کننده را کاهش دهد.

همبستگی به معنای علیت نیست. همبستگی بین دو متغیر به این معنا نیست که یکی باعث دیگری می‌شود. در مدل‌های با چندین ویژگی، ممکن است عوامل پنهانی وجود داشته باشند که باعث حرکت همزمان دو متغیر می‌شوند.

چه مقدار تغییر به عنوان تغییر واقعی در نظر گرفته می‌شود؟ تعیین مقدار تغییر مورد نیاز برای رد فرضیه صفر به مورد استفاده بستگی دارد. بهتر است در ابتدای پروژه یک مقدار مشخص تعیین کنید و به آن پایبند باشید.

چند مشاهده مورد نیاز است؟ تعداد مشاهدات مورد نیاز به قدرت آماری مورد نیاز پروژه بستگی دارد. بهتر است این مقدار در ابتدای پروژه تعیین شود.

مدت زمان اجرای آزمون چند متغیره چه قدر باید باشد؟ مدت زمان اجرای آزمون باید به اندازه‌ای باشد که تعداد مشاهدات کافی برای دستیابی به قدرت آماری تعیین‌شده را فراهم کند.

شناسایی تغییرات توزیع (Distribution Drift): پس از استقرار مدل، نظارت بر عملکرد مدل بسیار مهم است. تغییرات در داده‌ها و توسعه سیستم نیازمند تأیید مدل در مقایسه با خط پایه است. اگر تغییر قابل توجهی در معیار اعتبار سنجی مشاهده شود، مدل باید با داده‌های جدید آموزش داده شود.

ثبت تمام تغییرات مدل و یادداشت‌های مربوط به آن‌ها نه تنها به عنوان یک گزارش تغییر برای ذینفعان مفید است، بلکه سابقه‌ای فیزیکی از نحوه تغییر سیستم در طول زمان ارائه می‌دهد.