

رگرسیون ناپارامتری

در این فصل می‌خواهیم در مورد رگرسیون ناپارامتری صحبت کنیم. جفت مشاهدات داده شده مانند $(x_1, Y_1), (x_2, Y_2), \dots, (x_n, Y_n)$ داریم در معادله‌ی زیر به Y متغیر پاسخ می‌گویند که وابسته به متغیر x است.

$$Y_i = r(x_i) + \epsilon_i \quad (1)$$

که به r تابع رگرسیون می‌گویند. ما می‌خواهیم عملکردی را تحت فرضیات ضعیف تخمین بزنیم. که برآوردگر $r(x)$ را به صورت $\hat{r}(x)$ نشان می‌دهیم و همچنین ترجیح می‌دهیم که برآورد ما هموار باشد. و با توجه به این فرض که $\mathbb{V}(\epsilon_i) = \sigma^2$ به x وابسته نیست. ما فرمول بالا به گونه‌ای رفتار می‌کردیم که همبستگی متغیرهای x_i ثابت باشد، اگر به صورت تصادفی باشند داده‌ها به صورت $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ و $r(x)$ به صورت میانگین شرطی Y بر حسب $X = x$ تعریف می‌شود و خواهیم داشت.

$$r(x) = \mathbb{E}(Y|X = x) \quad (2)$$

مثال ۱ داده‌های CMB فراخوانی کنید و دو تا نمودار را رسم کنید. ^۱ محور افقی شکل لحظه‌ی چند قطبی فرکانس نوسانات در زمینه‌ی دما و محور عمودی قدرت نوسانات در هر فرکانس می‌باشد.

^۱ داده‌های تابش زمینه کیهانی که دارای متغیرهای ۱) لحظه‌ی چند قطبی فرکانس نوسانات در زمینه‌ی دما (گشتاور چند قطبی اندازه‌ای از بار) ۲) قدرت نوسانات در هر فرکانس ۳) موج صدا برآورد شده ۴) انحراف معیار برآورد شده می‌باشد

روشی که در این فصل ما در نظر گرفتیم

شیوه ی اول رگرسیون محلی برای فیت کردن توابع و یا سطوح رگرسیون در داده ها

۱ رگرسیون هسته

۲ رگرسیون چند جمله ایی محلی

شیوه ی دوم شیوه ی تنبیه کردن یا مجازات کردن مدل خطی استاندارد (روش حداقل مربعات) در شرایطی ضعیف عمل میکند، در جاهایی که تعداد متغیرها بیشتر از تعداد نمونه هاست که در این شرایط رگرسیون مجازات شده مناسب تر است. با اضافه کردن یک محدودیت در معادله متغیرها را مجازات میکند و از تعداد آن ها میکاهد. نتیجه ی این عمل کاهش مقادیر ضریب به سمت صفر

۱ رگرسیون ریج

۲ رگرسیون لاسو

۳ رگرسیون شبکه ی الاستیک

۱ مروری بر رگرسیون خطی و لجستیک

فرض کنید داده هایی مانند:

$(x_1, Y_1), \dots, (x_n, Y_n)$ $Y_i \in \mathbb{R}$ $x_i = (x_{i1}, \dots, x_{ip})^T \in \mathbb{R}^p$ داریم که و مدل رگرسیون خطی آن برابر است با:

$$Y_i = r(x_i) + \epsilon_i = \sum_{j=1}^p \beta_j x_{ij} + \epsilon_i \quad i = 1, \dots, n \quad (3)$$

که در عبارت فوق $\mathbb{E}(\epsilon_i) = 0$ و $\mathbb{V}(\epsilon_i) = \sigma^2$ برقرار است. ماتریس X به صورت زیر است:

$$X = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \end{bmatrix}$$

بردار $\mathcal{L}$ مجموعه ایی است از ترکیبات خطی ستون های ماتریس X ، با داشتن مقادیر $\beta = (\beta_1, \dots, \beta_p)^T$ ، $\epsilon = (\epsilon_1, \dots, \epsilon_n)^T$ ، $Y = (Y_1, \dots, Y_n)^T$ به صورت زیر نوشته میشود

$$Y = X\beta + \epsilon \quad (4)$$

برآورد حداقل مربعات $\hat{\beta} = (\hat{\beta}_1, \dots, \hat{\beta}_p)$ برداری است که مجموع باقیمانده مربعات خطارا مینیمم میکند

$$RSS = (Y - X\beta)^T(Y - X\beta) = \sum_{i=1}^n \left(Y_i - \sum_{j=1}^p x_{ij}\beta_j \right)^2$$

پس ما به دنبال مینیمم کردن عبارت فوق هستیم.
فرض کنید که $X^T X$ معکوس پذیر باشد، آنگاه برآورد حداقل مربعات به صورت زیر است

$$\hat{\beta} = (X^T X)^{-1} X^T Y \quad (5)$$

و برای برآورد $r(x)$ توسط $x = (x_1, \dots, x_p)^T$ داریم

$$\hat{r}_n(x) = \sum_{j=1}^p \hat{\beta}_j x_j = x^T \hat{\beta}$$

به این ترتیب مقادیر فیت شده $r = (\hat{r}_n(x_1), \dots, \hat{r}_n(x_n))^T$ را میتوان به صورت زیر نوشت:

$$r = X\hat{\beta} = LY \quad (6)$$

$$L = X(X^T X)^{-1} X^T \quad (7)$$

که به ماتریس ایجاد شده ماتریس هت گفته میشود. اعضای بردار $\hat{\epsilon} = Y - r$ را باقیمانه مینامند.

ماتریس هت متقارن است یعنی $L = L^T$

ماتریس هت خودتوان است یعنی $L^2 = L$

که طبق (۷) داریم و هرگاه ماتریس ما خودتوان باشد داریم که رتبه L برابر است با $tr(L)$ یعنی داریم

$$\begin{aligned} rank(X(X^T X)^{-1} X^T) &= tr(X(X^T X)^{-1} X^T) \\ &= tr(X^T X(X^T X)^{-1}) \\ &= tr(I_p) \\ &= p \end{aligned}$$

با توجه به ۱ خواهیم داشت:

$$\hat{r}_n(x) = \ell(x)^T Y = \sum_{i=1}^n \ell_i(x) Y_i \quad (8)$$

که داریم

$$\ell(x)^T = x^T (X^T X)^{-1} X^T$$

و برای برآورد ناریب σ^2 خواهیم داشت

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^n n(Y_i - \hat{r}_n(x_i))^2}{n-p} = \frac{\|\hat{\epsilon}\|^2}{n-p} \quad (9)$$

برای پیدا کردن یک فاصله اطمینان برای $r(x)$ باید تابع های $a(x), b(x)$ به گونه ای باشند که

$$\mathbb{P}(a(x) \leq r(x) \leq b(x)) \geq 1 - \alpha \quad (10)$$

با توجه به (۸) برای به دست آوردن واریانس $\hat{r}_n(x)$ خواهیم داشت:

$$\mathbb{V}(\hat{r}_n(x)) = \sigma^2 \sum_{i=1}^n \ell_i^2(x) = \sigma^2 \|\ell(x)\|^2$$

که پیشنهاد میشود به فرم زیر مورد استفاده قرار بگیرد:

$$I(x) = (a(x), b(x)) = (\hat{r}_n(x) - c\hat{\sigma}\|\ell(x)\|, \hat{r}_n(x) + c\hat{\sigma}\|\ell(x)\|) \quad (11)$$

که c مقداری ثابت است. اگر تابع توزیع $F_{p,n-p}$ داشته باشیم که F تابع توزیع برای p متغیر و با درجه آزادی $n-p$ میباشد، آنگاه با تعریف یک کران بالا مانند $F_{\alpha;p,n-p}$ آنگاه خواهیم داشت:

$$\mathbb{P}(F_{p,n-p} > F_{\alpha;p,n-p}) = \alpha$$

قضیه ۱ در فاصله اطمینان در رابطه ی (۱۱) اگر به جای مقدار c رابطه ی $c = \sqrt{pF_{\alpha;p,n-p}}$ قرار بگیرد آنگاه در رابطه ی (۱۰) هم صدق میکند.

۲ رگرسیون لجستیک

یک مدل آماری رگرسیون برای متغیرهای وابسته دوسویی مانند بیماری یا سلامت، مرگ یا زندگی است. این مدل را می توان به عنوان مدل خطی تعمیم یافته ای که از تابع لجیت (لگاریتم طبیعی نسبت به شانس) ^۲ به عنوان تابع پیوند استفاده می کند و خطایش از توزیع چندجمله ای پیروی می کند، به حساب آورد. منظور از دو سویی بودن، رخ داد یک واقعه تصادفی در دو موقعیت ممکنه است. تفاوت مهم مدل رگرسیون خطی و رگرسیون لجستیک در دو ویژگی رگرسیون لجستیک می تواند دیده شود. اول توزیع شرطی $y|\vec{x}$ یک توزیع برنولی به جای یک توزیع گوسی است چونکه متغیر وابسته دودویی است. دوم مقادیر پیش بینی احتمالاتی است و محدود بین بازه صفر و یک و به کمک تابع توزیع لجستیک بدست می آید رگرسیون لجستیک احتمال خروجی پیش بینی می کند.

$$\text{logit}(p) = \ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 x_{1i} + \dots + \beta_k x_{ki} \quad i = 1, \dots, n^*$$

تفاوت رگرسیون خطی با رگرسیون لجستیک

۱. برای مدل کردن متغیرهایی که مقادیر محدودی به خود میگیرند بهتر از رگرسیون خطی عمل میکند زیرا مدل خطی هر مقداری را در خروجی تولید میکند درحالی که برای چنین متغیرهایی مقادیر محدودی مورد نیاز است.

۲. در رگرسیون خطی مقدار متغیر مورد نظر از ترکیب خطی متغیرهای مستقل بدست می آید در حالیکه در لجستیک رگرسیون از ترکیب خطی تابع logit استفاده میشود.

۳. در رگرسیون خطی پارامترها به روش least squares بدست می آیند در حالیکه این روش برای لجستیک رگرسیون فاقد کارایی بوده و از روش maximum likelihood estimation برای پیدا کردن پارامترها استفاده میشود.

وقتی Y_i ها پیوسته نباشد، ممکن است رگرسیون خطی معمولی مناسب نباشد. برای نمونه فرض کنید $Y_i \in \{0, 1\}$ در این حالت مدلی که استفاده میشود مدل رگرسیون لجستیک است.

$$p_i = p_i(\beta) = \mathbb{P}(Y_i = 1) = \frac{e^{\sum_j \beta_j x_{ij}}}{1 + e^{\sum_j \beta_j x_{ij}}} \quad (12)$$

میتوانیم با جدا کردن $x_{i1} = 1$ برای همه i ها، در این مدل Y_i دارای توزیع برنولی با میانگین p_i میباشد نسبت درست نمایی پارامتر $\beta = (\beta_1, \dots, \beta_p)^T$ اگر تابع احتمال به صورت $f(y) = p^y(1-p)^{1-y}$ باشد. برابر خواهد بود با:

$$\mathcal{L}(\beta) = \prod_{i=1}^n p_i^{Y_i} (1 - p_i)^{1-Y_i} \quad (13)$$

بیشترین نسبت درست نمایی برآوردگر $\hat{\beta} = (\hat{\beta}_1, \dots, \hat{\beta}_p)$ را از الگوریتم زیر میتوان به دست آورد. بدین صورت است که ابتدا با یک $\hat{\beta} = (\hat{\beta}_1, \dots, \hat{\beta}_p)$ شروع میکنیم و در گذاشتن در فرمول بالا p_i را محاسبه میکنیم سپس با انجام الگوریتم زیر و تکرار آن تا آن جایی که همگرا شود. الگوریتم زیر از سه مرحله تشکیل شده مرحله ی اول مجموعه ی زیر را تعریف میکنیم:

$$Z_i = \log \left(\frac{p_i}{1-p_i} \right) + \frac{Y_i - p_i}{p_i(1-p_i)} \quad i = 1, \dots, n$$

مرحله ی دوم برای برآورد β داریم

$$\hat{\beta} = (X^T W X)^{-1} X^T W Z$$

که در عبارت بالا W ماتریس قطری است که عناصر روی قطر اصلی آن برابر است با: $p_i(1-p_i)$. این بدان معناست که انجام دادن رگرسیون خطی Z روی X .

مرحله ی سوم حساب کردن p_i با استفاده از برآوردگر $\hat{\beta}$

رگرسیون خطی و رگرسیون لجستیک مورد های خاص کلاسی به نام مدل های خطی تعمیم یافته هستند.

۳ هموار کننده ی خطی

همه ی برآورد های ناپارامتری در این فصل هموار کننده ی خطی هستند.

تعریف ۱ برآورگر $\hat{r}_n$ از r یک هموار کننده ی خطی است هرگاه :

$$\forall x, \quad \exists \ell(x) = (\ell_1(x), \dots, \ell_n(x))^T, \quad \hat{r}_n(x) = \sum_{i=1}^n \ell_i(x) Y_i \quad (14)$$

که برای فیت کردن r با توجه به $Y = (Y_1, \dots, Y_n)^T$ خواهیم داشت:

$$r = (\hat{r}_n(x_1), \dots, \hat{r}_n(x_n))^T \quad (15)$$

$$r = LY \quad (16)$$

که در عبارت بالا L یک ماتریسی مربعی $n \times n$ میباشد که سطر i ام توسط $\ell(x_i)^T$ به جود می آید یعنی هر درآیه ی آن به صورت: $L_{ij} = \ell_j(x_i)$ می باشد. سطر i ام نشان دهنده ی وزن هایی است که به هر Y_i برای برآورد $\hat{r}_n(x_i)$ میدهد.

تعریف ۲ ماتریس L ماتریس هموارساز یا هت ماتریس گفته می شود. i امین سطر از L کرنل موثر برای برآورد کردن $r(x_i)$ گفته می شود. مانند (۱) درجه ی آزادی موثر را به صورت زیر تعریف میکنیم:

$$v = tr(L) \quad (17)$$

نکته: وزن های موجود در همه هموارسازها که ما از آنها استفاده می کنیم یک خصوصیت دارند که

$$\forall x, \quad \sum_{i=1}^n \ell_i(x) = 1$$

این نشان می دهد که هموارساز، منحنی را ثابت نگه می دارد. پس اگر $Y_i = c$ برای همه ی i ها آنگاه :

$$\hat{r}_n(x) = c$$

مثال ۲ (رگرسوگرام). فرض کنید که $i = 1, \dots, n$ که $a \leq x_i \leq b$ که در بازه ی (a, b) است را به m بلوک مساوی تقسیم میکنیم که هر بلوک را با $\beta_1, \beta_2, \dots, \beta_m$ نمایش میدهیم. $\hat{r}_n(x)$ را به صورت زیر تعرف میکنیم.

$$\hat{r}_n(x) = \frac{1}{k_j} \sum_{i: x_i \in \beta_j} Y_i, \quad \forall x_i \in \beta_j \quad (18)$$

که k_j ها تعداد نقاط موجود در β_j ها میباشد. به عبارت دیگر برآورد $\hat{r}_n$ تابعی است که توسط میانگین گیری کردن Y_i ها در هر بلوک به دست آمده است. این برآورد، رگرسوگرام گفته می شود. حال تابع $\ell(x)$ را به صورت زیر تعریف میکنیم:

$$\ell_i(x) = \begin{cases} \frac{1}{k_j} & x_i \in \beta_j \\ 0 & o.w \end{cases}$$

که خواهیم داشت

$$\ell(x)^T = (0, 0, \dots, 0, \frac{1}{k_j}, \dots, \frac{1}{k_j}, 0, \dots, 0)$$

برای دیدن شکل ماتریس هموارساز L فرض کنید که $n = 9, m = 3, k_1 = k_2 = k_3 = 3$ سپس برای ماتریس هت خواهیم داشت:

$$L = \frac{1}{3} \times \begin{bmatrix} 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 \end{bmatrix}$$

ماتریس فوق ماتریس مربعی است که $v = tr(L) = m$ درجه ی آزادی موثر میباشد. پهنای هر بلوک به صورت $h = \frac{(b-a)}{m}$ می باشد که نشان می دهد، برآورد به چه میزان هموار است. برای بدست آوردن هت ماتریس، خود صورت سوال گفته که بازه را به سه قسمت مساوی تقسیم کنید و در هر بازه سه نقطه وجود دارد و در کل تعداد نقاط برابر است با ۹ تا پس یعنی سطر اول هت ماتریس سه نقطه ی اول در β_1 وجود دارد پس در سطر اول سه درآیه اول آن یک سوم میشود و همین طور تا آخر.

مثال ۳ میانگین محلی. $h > 0$ را در نظر بگیرید و مجموعه ایی مانند زیر داریم: n_x تعداد نقاطی است که در β_x وجود دارد. برای هر x که $n_x > 0$ باشد به صورت زیر تعریف میشود.

$$\hat{r}_n(x) = \frac{1}{n_x} \sum_{i \in \beta_x} Y_i.$$

که یک برآورد میانگین محلی از $r(x)$ است. در این مورد $\hat{r}_n(x) = \sum_{i=1}^n Y_i \ell_i(x)$ که تابع $\ell(x)$ به صورت زیر تعریف میشود.

$$\ell_i(x) = \begin{cases} \frac{1}{n_x} & |x_i - x| \leq h \\ 0 & o.w \end{cases}$$

حالا در یک حالت ساده تصور کنید که $h = \frac{1}{9}$, $n_i = \frac{i}{9}$, $n = 9$. که با توجه به موارد بالا ماتریس هت به صورت زیر خواهد شد.

$$L = \begin{bmatrix} \frac{1}{2} & \frac{1}{2} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} \end{bmatrix}$$

سطر اول ماتریس مربوط به مشاهده اول است یعنی وزن y_1 را تعیین میکند. سطر دوم ماتریس مربوط به مشاهده دوم و وزن y_2 را تعیین میکنند همین طور تا آخر. در نهایت با میانگین گیری وزنی برآورد تابع ناپارامتری تعیین میشود، یعنی

$$\hat{r}_n(x) = \frac{1}{n_x} \sum_{i=1}^n y_i \ell_i(x)$$

برای حل مثال داریم ابتدا مقدار $i = 1$ را قرار میدهیم و خواهیم داشت.

$$\beta_1 = \{1, |x_1 - x| \leq \frac{1}{9}\}$$

$$x_1 = \frac{1}{9} \Rightarrow |\frac{1}{9} - x| \leq \frac{1}{9}$$

$$0 \leq x \leq \frac{2}{9}$$

$$x_2 = \frac{2}{9} \Rightarrow |\frac{2}{9} - x| \leq \frac{1}{9}$$

$$\frac{1}{9} \leq x \leq \frac{3}{9}$$

⋮

$$\beta_1 = [0, \frac{2}{9}] \Rightarrow \{1/9, 2/9\} \rightarrow n_x = 2$$

$$\ell_1(x) = \begin{cases} \frac{1}{2} & i = 1, 2 \\ 0 & o.w \end{cases}$$

$$\beta_2 = [\frac{1}{9}, \frac{3}{9}] \Rightarrow \{1/9, 2/9, 3/9\} \rightarrow n_x = 3$$

$$\ell_2(x) = \begin{cases} \frac{1}{3} & i = 1, 2, 3 \\ 0 & o.w \end{cases}$$

⋮

و همین طور تا آخر، که ماتریس L ساخته میشود.

۴ انتخاب پارامتر هموار

در برآورد $r(x)$ به مشاهده X_i وزن بیشتری و به نقاط دور تر وزن کمتری بدهد، معیاری که بتواند تعیین کند کدام مشاهده دور تر است و چه وزنی باید بدهد و کدام مشاهده نزدیک است و چه وزنی باید بگیرد را پهنای باند یا پارامتر هموار ساز میگویند. پس ما نیاز داریم راهی برای انتخاب مناسب ترین h انتخاب کنیم. میانگین مربعات خطا را به صورت زیر تعریف میکنیم:

$$R(h) = \mathbb{E} \left(\frac{1}{n} \sum_{i=1}^n (\hat{r}_n(x_i) - r(x_i))^2 \right) \quad (۱۹)$$

ما باید h را به گونه ایی انتخاب کنیم تا $R(h)$ مینیمم شود، اما $R(h)$ وابسته است به تابع ناشناخته ایی به اسم $r(x)$ ، در عوض ما برآورد تابع ریسک را مینیمم میکنیم. در حدس اول: اولین برآوردی که به ذهن می رسد، میانگین توان دوم خطای مانده ها می باشد، به صورت زیر است:

$$\hat{R}(h) = \frac{1}{n} \sum_{i=1}^n ((Y_i - \hat{r}_n(x_i))^2) \quad (۲۰)$$

این برآورد $R(h)$ را کمتر از مقدار واقعی برآورد می کند، چون از یک مجموعه داده ها دوبار استفاده میشود یعنی یک بار برای برآورد $\hat{r}_n(x_i)$ و بعد از برآورد آن دوباره از همون داده ها برای برآورد $R(h)$ به خاطر همین خطا را افزایش میدهد و برآورد را کمتر از مقدار واقعی نشان میدهد. در واقع مینیمم سازی عبارت $\sum_{i=1}^n (Y_i - \hat{r}_n(x_i))^2$ برای تعیین h باعث می شود که $\hat{r}(x)$ طوری تخمین زده شود که $\hat{R}(h)$ مینیمم شود و این امر باعث کم کردن $R(h)$ می شود. به خاطر این مشکلی که به وجود میآید باید از اعتبارسنجی متقابل استفاده کرد.

قضیه ۲ ترک یک داده نمره اعتبار سنجی متقابل را به صورت زیر تعریف میکنیم:

$$CV = \hat{R}(h) = \frac{1}{n} \sum_{i=1}^n (Y_i - \hat{r}_{(-i)}(x_i))^2 \quad (21)$$

که $\hat{r}_{(-i)}$ برآوردی است که با حذف جفت i ام یعنی دوتایی (x_i, Y_i) . همان طور گفته شد تعریف قبلی در مورد $\hat{r}_{(-i)}$ ناقص است پس خواهیم داشت:

$$\hat{r}_{(-i)}(x) = \sum_{j=1}^n Y_j \ell_{j,(-i)}(x) \quad (22)$$

$$\ell_{j,(-i)}(x) = \begin{cases} 0 & \text{if } j = i \\ \frac{\ell_j(x)}{\sum_{k \neq j} \ell_k(x)} & \text{if } j \neq i \end{cases} \quad (23)$$

به معنای دیگر با قرار دادن وزن x_i به صفر و قرار دادن مجموع وزن های دیگر به یک. برای درک اعتبار سنجی متقابل خواهیم داشت:

$$\begin{aligned} \mathbb{E}(Y_i - \hat{r}_{(-i)}(x_i))^2 &= \mathbb{E}(Y_i - r(x_i) + r(x_i) - \hat{r}_{(-i)}(x_i))^2 \\ &= \sigma^2 + \mathbb{E}(r(x_i) - \hat{r}_{(-i)}(x_i))^2 \\ &\approx \sigma^2 + \mathbb{E}(r(x_i) - \hat{r}_n(x_i))^2 \end{aligned}$$

در رابطه ی بالا $Y_i = r(x_i) + \epsilon$ از این رو خواهیم داشت

$$\mathbb{E}(\hat{R}) \approx R + \sigma^2 = \text{پیش بینی خطر} \quad (24)$$

که در رابطه ی فوق σ^2 میانگین وزنی ϵ میباشد. در نمره اعتبار سنجی متقابل برآورد خطا تاثیر ندارد، نااریب از خطر است. به نظر می رسد که ممکن است ارزیابی $\hat{R}(h)$ وقت گیر باشد زیرا ظاهراً نیاز داریم که پس از رها کردن هر مشاهده، برآوردگر را پس بگیریم. خوشبختانه یک فرمول میانبر برای محاسبات وجود دارد. برای ارزیابی (۲۱) وقت بسیاری نیاز است، چون نیاز به محاسبه مجدد برآوردگر پس از حذف هر مشاهده می باشد.

قضیه ۳ اگر $\hat{r}_n$ یک هموارکننده خطی باشد. آنگاه نمره اعتبار سنجی متقابل $\hat{R}(h)$ را میتوان به صورت زیر نوشت:

$$\hat{R}(h) = \frac{1}{n} \sum_{i=1}^n \left(\frac{Y_i - \hat{r}_n(x_i)}{1 - L_{ii}} \right)^2 \quad (25)$$

که در آن $L_{ii} = \ell_i(x_i)$ عضو i ام قطری هت ماتریس میباشد.

پس پارامتر هموار ساز h با مینیمم کردن $\hat{R}(h)$ قابل انتخاب است. هشدار: شما نمی توانید فرض کنید $\hat{R}(h)$ همیشه از حداقل مشخصی برخوردار باشد. پس باید همیشه $\hat{R}(h)$ را به عنوان تابعی از h ترسیم کنید. یک جایگزین به جای به حداقل رساندن نمره اعتبار سنجی متقابل به طور تقریبی اعتبار متقاطع کلی تعمیم یافته است که هر l_{ii} در معادله (۲۵) قرار میگیرد، با میانگین $n^{-1} \sum_{i=1}^n L_{ii} = v/n$ که $v = \text{tr}(L)$ درجه آزادی موثر میباشد. پس باید رابطه ی زیر را مینیمم کنیم:

$$GCV(h) = \frac{1}{n} \sum_{i=1}^n \left(\frac{Y_i - \hat{r}_n(x_i)}{1 - v/n} \right)^2 \quad (26)$$

پهنای باند میتواند نمره اعتبار سنجی متقابل عمومی را به حداقل برساند. با استفاده از $(1-x)^{-2} \approx 1+2x$ رابطه ی بالا به صورت تقریبی برابر میشود با:

$$GCV(h) \approx \frac{1}{n} \sum_{i=1}^n (Y_i - \hat{r}_n(x_i))^2 + \frac{2v\hat{\sigma}^2}{n} \quad (27)$$

که در رابطه ی بالا $\hat{\sigma}^2 = n^{-1} \sum_{i=1}^n (Y_i - \hat{r}_n(x_i))^2$ میباشد که معادله ی بالا معیاری برای انتخاب متغیر در رگرسیون خطی میباشد. به طور کلی بسیاری از معیار های انتخاب پهنای باند را به صورت زیر نوشت:

$$B(h) = \Xi(n, h) \times \frac{1}{n} \sum_{i=1}^n (Y_i - \hat{r}_n(x_i))^2 \quad (28)$$

در شرایط مناسب انتخاب h باید مینیمم کند $B(h)$ را حال اگر $\hat{h}_0$ مینیمم کند تابع زیان $L(\hat{h}) = n^{-1} \sum_{i=1}^n (\hat{r}_n(x_i) - r(x_i))^2$ و h_0 ریسک را مینیمم کند. سپس همه ی $\hat{h}, \hat{h}_0, h_0$ تمایل دارند به صفر در نرخ $n^{-1/5}$ و همچنین برای برخی ثابت های مثبت مانند $C_1, C_2, \sigma_1, \sigma_2$

$$n(L(\hat{h}) - L(\hat{h}_0)) \rightsquigarrow C_1 \chi_1^2 \quad n^{3/10}(\hat{h} - \hat{h}_0) \rightsquigarrow N(0, \sigma_1^2)$$

$$n(L(h_0) - L(\hat{h}_0)) \rightsquigarrow C_2 \chi_1^2 \quad n^{3/10}(h_0 - \hat{h}_0) \rightsquigarrow N(0, \sigma_2^2)$$

همگرایی $\hat{h}$ برابر است با

$$\frac{\hat{h} - \hat{h}_0}{\hat{h}_0} = O_p \left(\frac{n^{3/10}}{n^{1/5}} \right) = O_p(n^{-1/10})$$

این سرعت آهسته نشان می دهد که تخمین پهنای باند دشوار است. این سرعت ذاتی است برای مشکل انتخاب پهنای باند از آنجا که این نیز صادق است.

$$\frac{\hat{h}_0 - h_0}{h_0} = O_p \left(\frac{n^{3/10}}{n^{1/5}} \right) = O_p(n^{-1/10})$$

۵ رگرسیون محلی

جال مجدد به رگرسیون محلی ناپارامتری برمیگردیم. فرض کنید x_i ها اسکالر باشد، مدل رگرسیون (۱) را در نظر بگیرید. در این قسمت با در نظر گرفتن برآورد گر ها با توجه به میانگین وزن Y_i ها وزن بیشتری به نقاطی که نزدیک تر است می‌دهد. ما با برآوردگر رگرسیون هسته شروع می‌کنیم.

$$E(Y|X = x) = \int y f(y|x) dy = \int y \frac{f(x, y)}{f(x)}$$

با توجه به عبارت زیر است:

$$\hat{f}(x, y) = \frac{1}{n} \sum_{i=1}^n K_h(x - x_i) K_h(y - y_i), \hat{f}(x) = \frac{1}{n} \sum_{i=1}^n K_h(x - x_i)$$

با جاگذاری در عبارت اول خواهیم دید که

$$\int y \frac{\sum_{i=1}^n K_h(x - x_i) K_h(y - y_i)}{\sum_{j=1}^n K_h(x - x_j)} dy$$

$$\int y K_h(y - y_i) dy = Y_i$$

تعریف ۳ اگر $h > 0$ یک عدد مثبت باشد، که پهنای باند نامیده می‌شود، یک برآورد هسته نادارایا و واتسون است هر گاه:

$$\hat{r}_n(x) = \sum_{i=1}^n \ell_i(x) Y_i \quad (29)$$

که $\ell_i(x)$ به صورت زیر می‌باشد. که در فرمول زیر K یک هسته می‌باشد.

$$\ell_i(x) = \frac{K\left(\frac{x-x_i}{h}\right)}{\sum_{j=1}^n K\left(\frac{x-x_j}{h}\right)} \quad (30)$$

مثال ۴ با فراخوانی داده های CMB چهار رگرسیون مختلف هسته را نمایش می‌دهد، که بر اساس افزایش پهنای باند می‌باشد. هر چه پهنای باند کمتر باشد اتصالات خشن است و اریبی کوچک و واریانس بزرگ و هر چه پهنای باند بزرگتر باشد تخمین ها همار تر اریبی کم و واریانس زیاد تر است.

انتخاب هسته K زیاد مهم نیست. برآوردهای به دست آمده با استفاده هسته های مختلف معمولاً از نظر عددی بسیار مشابه هستند. این مشاهدات توسط محاسبات نظری تأیید شده است که نشان می دهد خطر بسیار غیر حساس است. آنچه اهمیت بیشتری دارد انتخاب است پهنای باند h که میزان صاف کردن را کنترل می کند. ما به پهنای باند اجازه خواهیم داد که به اندازه نمونه بستگی داشته باشد. قضیه زیر چگونگی تأثیر پهنای باند بر برآوردها را نشان می دهد.

قضیه ۴ ریسک در هسته برآودگر ناداریا واتسون به صورت زیر است.

$$R(\hat{r}_n, r) = \overbrace{\frac{h_n^4}{4} \left(\int x^2 K(x) dx \right)^2 \int \left(r''(x) + 2r'(x) \frac{f'(x)}{f(x)} \right)^2 dx}^{\text{مربع ادبی}} \quad (31)$$

$$+ \underbrace{\frac{\sigma^2 \int K^2(x) dx}{nh_n} \int \frac{1}{f(x)} dx}_{\text{واریانس}} + o(nh_n^{-1}) + o(h)n^4$$

به طوری که $h_n \rightarrow 0, nh_n \rightarrow \infty$

آنچه در عبارت بالا مهم است وجود عبارت زیر است:

$$2r'(x) \frac{f'(x)}{f(x)} \quad (32)$$

که در قسمت اریبی وجود دارد.

عبارت بالا را نقشه اریبی مینامیم، زیرا بستگی به نقشه دارد، یعنی به توزیع x_i ها. این بدان معناست که اریبی حساس به موقعیت x_i ها است. می توان نشان داد که برآوردهای هسته همچنین در نزدیکی مرزها اریبی بالایی دارند. در ادامه خواهیم دید میتوان اریبی مرزها را با استفاده از رگرسیون چند جمله ای کاهش داد. اگر عبارات (۳۱) را مشتق بگیریم و برابر صفر کنیم پهنای باند بهینه محاسبه میشود.

$$h_* = \left(\frac{1}{n} \right)^{1/5} \left(\frac{\sigma^2 \int K^2(x) dx \int \frac{dx}{f(x)}}{\int x^2 K^2(x) dx \int \left(r''(x) + 2r'(x) \frac{f'(x)}{f(x)} \right)^2 dx} \right)^{1/5} \quad (33)$$

در نتیجه داریم: $h_* = O(n^{-1/5})$ با قرار دادن h_* در معادله ی (۳۱) خواهیم دید که ریسک با سرعت $O(n^{-4/5})$ در حال کاهش است. در مدل های پارامتری ماکسیم نسبت درستمایی کاهش میدهد ریسک را تا صفر با سرعت $1/n$ پس مناسب است از روش های پارامتری استفاده شود. در عمل نمیتوان پهنای باند را از معادله ی (۳۳) بدست آورد زیرا وابسته به تابع ناشناخته ی r میباشد در عوض از اعتبار سنجی متقابل استفاده میکنیم.

مثال ۵ نمره اعتبار سنجی متقابل را بر روی داده های CMB در ازای تاثیر درجه ی آزادی انجام دهید. بهترین پارامتر هموار آن است که این نمره را کاهش مینیمم کند.

۶ رگسیون چند جمله ایی

برآوردگرهای هسته از اریبی مرزی رنج می برند. این مشکلات را می توان با استفاده از رگسیون هسته به نام رگسیون چند جمله ای محلی کاهش داد. برای شروع باید $a = \hat{r}_n(x)$ را به گونه ایی انتخاب کرد که عبارت زیر را مینیمم کند. $\sum_{i=1}^n (Y_i - a)^2$ تابع ثابت $\hat{r}_n(x) = \bar{Y}$ تخمین خوبی از $r(x)$ نیست. اکنون با معرفی تابع وزن $w_i(x) = K(\frac{x_i - x}{h})$ حال باید $a = \hat{r}_n(x)$ را به گونه ایی انتخاب کنیم، تا مقدار وزنی مربعها را به حداقل برساند.

$$\sum_{i=1}^n w_i(x)(Y_i - a)^2 \quad (34)$$

با محاسبات خواهیم داشت:

$$\hat{r}_n(x) = \frac{\sum_{i=1}^n w_i(x)Y_i}{\sum_{i=1}^n w_i(x)}$$

این یک برآوردگر ثابت محلی است، از حداقل مربعات دارای وزن محلی بدست آمده. که نشان میدهد به جای یک رگسیون ثابت باید از یک رگسیون چند جمله ایی استفاده کنیم. بگذارید x مقداری ثابت باشد که می خواهیم $r(x)$ را تخمین بزنیم. برای همسایگی x ارزشی به اسم u که به صورت چند جمله ایی به صورت زیر تعریف میشود.

$$P_x(u; a) = a_0 + a_1(u - x) + \frac{a_2}{2!}(u - x)^2 + \dots + \frac{a_p}{p!}(u - x)^p \quad (35)$$

میتوان به صورت تقریبی برآورد $r(u)$ را در همسایگی x به صورت زیر نوشت

$$r(u) \approx P_x(u; a). \quad (36)$$

ما با برآورد $a = (a_0, \dots, a_p)^T$ با انتخاب $\hat{a} = (\hat{a}_0, \dots, \hat{a}_p)^T$ برای مینیمم کردن مجموع مربعات وزن های محلی

$$\sum_{i=1}^n w_i(x)(Y_i - P_x(X_i; a))^2 \quad (37)$$

برآورد a چون به x وابسته است پس مناسب تر است به صورت $\hat{a}(x) = (\hat{a}_0(x), \dots, \hat{a}_p(x))^T$ نوشته شود. اگر بخواهیم این وابستگی را واضح تر بیان کنیم به صورت $\hat{r}_n(u) = P_x(u; a)$ در حالت خاص اگر $u = x$ باشد آنگاه

$$\hat{r}_n(x) = P_x(x; \hat{a}) = \hat{a}_0(x) \quad (38)$$

اگرچه که $\hat{r}_n(x)$ فقط وابسته $\hat{a}_0(x)$ میباشد اما این بدین معنی نیست که به یک مقدار ثابت فیت میشود. با قرار دادن $p = 0$ برآوردگر هسته را باز میگرداند. مقدار مخصوص که $p = 1$ است رگسیون محلی خطی نامیده میشود. برای پیدا کردن $\hat{a}(x)$ ابتدا داریم:

$$X_x = \begin{bmatrix} 1 & x_1 - x & \dots & \frac{(x_1 - x)^p}{p!} \\ 1 & x_2 - x & \dots & \frac{(x_2 - x)^p}{p!} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_n - x & \dots & \frac{(x_n - x)^p}{p!} \end{bmatrix} \quad (39)$$

ما در اینجا W_x را ماتریس قطری $n \times n$ مینامیم که درایه ی (i, i) آن از $w_i(x)$ تشکیل شده باشد. پس میتوانیم بنویسیم:

$$(Y - X_x a)^T W_x (Y - X_x a) \quad (40)$$

حال با مینیم کردن عبارت وزن حداقل مربعات بدست می آید.

$$\hat{a}(x) = (X_x^T W_x X_x)^{-1} X_x^T W_x Y \quad (41)$$

به طور خاص اگر $\hat{r}_n(x) = \hat{a}_0(x)$ آنگاه فقط سطر اول معادله ی (40) باید استفاده شود.

قضیه ۵ برآورد رگسیون چند جمله ایی محلی به صورت زیر میباشد:

$$\hat{r}_n(x) = \sum_{i=1}^n \ell_i(x) Y_i \quad (42)$$

که داریم $\ell(x)^T = (\ell_1(x), \dots, \ell_n(x))$ آنگاه

$$\ell(x)^T = e_1^T (X_x^T W_x X_x)^{-1} X_x^T W_x$$

که $e_1 = (1, 0, \dots, 0)^T$ میباشد. میانگین و واریانس برآورد هم به صورت زیر میباشد.

$$\mathbb{E}(\hat{r}_n(x)) = \sum_{i=1}^n \ell_i(x) r(x_i)$$

$$\mathbb{V}(\hat{r}_n(x)) = \sigma^2 \sum_{i=1}^n \ell_i(x)^2 = \sigma^2 \|\ell(x)\|^2$$

ک بار دیگر، تخمین ما یک سطح صاف خطی است و ما می توانیم با به حداقل رساندن فرمول اعتبار سنجی متقابل که در قضیه ۳ ارائه شده است، پهنای باند را انتخاب کنیم.

مثال ۶ با فراخوانی داده های LIDAR که در فصل چهار با آن ها روبرو شده بودیم. چهار نمودار به دست آمده که نمودار سمت چپ بالا داده ها میباشد با فیت کردن توسط رگسیون خطی محلی با $h \approx 37$ که با اعتبار سنجی متقابل بدست آمده. نمودار بالا سمت راست مقدار باقی مانده ها هست، نمودار پایین سمت چپ برآورد $\sigma(x)$ میباشد و نمودار پایین سمت راست یک فاصله اطمینان ۹۵ درصدی میباشد.

قضیه ۶ هموارسازی خطی محلی: اگر $p = 1, \hat{r}_n(x) = \sum_{i=1}^n \ell_i(x) Y_i$ که

$$\ell_i(x) = \frac{b_i(x)}{\sum_{i=1}^n b_i(x)}$$

$$b_i(x) = K \left(\frac{x_i - x}{h} \right) (S_{n,2}(x) - (x_i - x) S_{n,1}(x)) \quad (43)$$

که داریم:

$$S_{n,j}(x) = \sum_{i=1}^n K \left(\frac{x_i - x}{h} \right) (x_i - x)^j, j = 1, 2.$$

مثال ۷ با فراخوانی داده ی CMB برای $p = 0, p = 1$ رگسیون محلی را رسم میکنیم. که نمودار پایینی زوم شده به کران سمت چپ توجه داشته باشید برای $p = 0$ فیت کردن در نزدیکی کران ها ضعیف است که از اریب بودن کران ها است.

مثال ۸ تابع دوپلر به صورت زیر تعریف میشود:

$$r(x) = \sqrt{x(1-x)} \sin \left(\frac{2.1\pi}{x + .05} \right), 0 \leq x \leq 1 \quad (44)$$

این تابع برای برآورد بسیار دشوار میباشد اما یک مورد مناسب برای تست روش رگسیون ناپارامتری است. در فضایی ناهمگن میباشد یعنی همواری (مشتق دوم) آن نسبت به x متفاوت است.

نمودار های آن رسم شده است. که نمودار بالا سمت چپ نمودار خود تابع میباشد سمت راست تعداد ۱۰۰۰ داده شبیه سازی شده است از فرمول $Y_i = r(i/n) + \sigma \epsilon_i$ با توجه به مقادیر زیر $\sigma = 1, \epsilon_i \sim N(0, 1)$ نمودار پایین سمت چپ: نمره اعتبار متقابل سنجی اثر گذاری درجه ی آزادی که با توجه به این نمودار مشخص شد که مینیمم آن حدودا در نقطه ی ۱۶۶ اتفاق می افتد با پهنای باند ۰.۰۰۵ و نمودار پایین سمت راست که نمودار فیت شده ی آن میباشد. خط فیت شده چون دارای درجه آزادی موثر بالایی هست پس از هوش بالایی برخوردار میباشد، زیرا تلاش بر آن است که در نوسانات سریع در نزدیکی صفر بهترین فیت را داشته باشیم اگر از هوار کننده بیشتر استفاده میکردیم قسمت راست فیت شده میتواند مناسب تر باشد هر چند ساختار در نزدیکی صفر از بین میرفت. این مشکلی است که در عملکرد های ناهمگن فضایی وجود دارد که در فصل نه بحث میشود.

در قضیه ی زیر رفتار نمونه بزرگ را در ریسک رگسیون خطی محلی نشان میدهد و چرا رگسیون خطی محلی از رگسیون هسته مناسب تر است.

قضیه ۷ اگر

$$Y_i = r(x_i) + \sigma(X_i)\epsilon_i \quad \text{for } i = 1, \dots, n \quad a \leq X_i \leq b$$

را داشته باشیم و $X_1, \dots, X_n$ یک نمونه باشد که دارای توزیع با تابع چگالی f باشد و شرط های زیر وجود داشته باشد

$$f(x) > 0 \quad ۱.$$

$$f, r'', \sigma^2 \text{ در نزدیکی } x \text{ پیوسته باشد.} \quad ۲.$$

$$h_n \rightarrow 0, nh_n \rightarrow \infty \quad ۳.$$

و اگر $x \in (a, b)$ باشد، و داده های $X_1, \dots, X_n$ معلوم باشد آنگاه برآوردگر خطی محلی و برآوردگر هسته دارای واریانس

$$\frac{\sigma^2(x)}{f(x)nh_n} \int K^2(u)du + o_p\left(\frac{1}{nh_n}\right) \quad (۴۵)$$

و اریبی برآوردگر هسته نادارایا و واتسون به صورت

$$h_n^2 \left(\frac{1}{2} r''(x) + \frac{r'(x)f'(x)}{f(x)} \right) \int u^2 K(u)du + o_p(h^2) \quad (۴۶)$$

میباشد، در حالی که اریبی برآوردگر خطی محلی به صورت تقریبی برابر

$$h_n^2 \frac{1}{2} r''(x) \int u^2 K(u)du + o_p(h^2) \quad (۴۷)$$

میباشد. بنابراین برآوردگر خطی محلی بدون نقشه اریبی است. (۳۲) در نقاط مرزی a, b برآوردگر هسته نادارایا واتسون تقریبی دارای اریبی با نظم h_n است در حالی که برآوردگر خطی محلی دارای اریبی با نظم h_n^2 میباشد، این بدان معناست که برآوردگر خطی محلی اریبی مرز هارا از بین میبرد.

نتیجه ی بالا به طور کلی بیشتر برای محلی چند جمله ایی با درجه ی p میباشد. که با گرفتن p به طرز عجیب غریبی نقشه ی اریبی را در مرز هارا بدون افزایش واریانس کاهش میدهد