



Piecewise Regression Model Based on Scale Mixture Normal Distribution

Hashemi, F. 

Department of Statistics, University of Kashan, Kashan, Iran.

Corresponding author: F. Hashemi, farzane.hashemi1367@kashanu.ac.ir

Received: 23/5/2023 Revised: 18/10/2024 Accepted and Published Online: 18/10/2024.

Introduction

In many applied regression problems, the relationship between the response variable and the independent variables often changes suddenly in some places, called breakpoints or change points. Sudden changes in these places cause deviations in the analysis, and Data is analyzed. Although non-linear regression models can be used to examine such data, the importance of breakpoints needs to be addressed, and parameters may be estimated with bias. Hence, piecewise regression models show the importance of breakpoints and appropriate estimation of model parameters. In particular, identifying breakpoints is critical to understanding when and how to change the model. This article obtains breakpoints simultaneously by classifying the data into homogeneous clusters. Therefore, clustering methods can be used. Hybrid models for data clustering are powerful and popular models that estimate parameters and detect breakpoints using the Expectation-Maximum (EM) algorithm and its generalizations. Most of these methods focus on the average and variance changes in the range of two consecutive breakpoints.

Material and Methods

One of the valuable and practical tools to obtain MLE model parameters when the form of the density function is complex is the EM algorithm. This method, introduced by Dempster et al. (1977), is used when the data is only partially observed. The incompleteness of the data may be due to various reasons, such as missing observations or truncated or censored observations. The EM algorithm's convergence speed to obtain the optimal solution may be low compared to other methods. Nevertheless, the simplicity of imple-

mentation and obtaining trusted optimal answers can be the main reason for this method. We will use the EM algorithm to estimate the model parameters, i.e., the vector. To implement this algorithm, it is assumed that there are also hidden observations for the observed data.

Results and Discussion

According to the observations in the tables, the average values of TPR and FPR in the CN-BP model are respectively higher and lower than in the t-BP model. Also, the value of TPR decreases in both models as the sample size increases. With the increase in the sample size, the value of FPR increases in both models. Therefore, the CN-BP model is better than the t-BP model in detecting outliers.

Conclusion

In this article, we proposed a method to solve piecewise regression problems with errors from scale mixture normal (SMN) distribution by using finite mixed models. Finding change points is similar to categorizing similar observations into groups in mixed models. We created regression lines to estimate change points. Also, using the shown simulations, outlier locations are identified using some special modes of the SMN model. In addition, this numerical study, with actual and simulated data sets, shows the performance of the superior model proposed. As a result, the proposed SMN breakpoint (SMN-BP) model usually gives accurate estimates of change points and regression lines. Therefore, the SMN-BP model is suitable for solving piecewise regression problems. Piecemeal regression can be applied to problems with an important change point, such as quality control, medicine, economics, and decision-making decision-making.

Keywords: Break point, EM algorithm, Piecewise regression.

Mathematics Subject Classification (2010): 62J02, 62J05.



©The Author(s). The Publisher is Iranian Statistical Society.

This is an open access article distributed under the terms and conditions of [\(CC BY-NC 4.0\)](https://creativecommons.org/licenses/by-nc/4.0/)

مدل بندی رگرسیون تکه‌ای همراه با خطاهای مخلوط مقیاسی نرمال

فرزانه هاشمی

گروه آمار، دانشگاه کاشان

نویسنده مسئول: farzane.hashemi1367@yahoo.com

تاریخ دریافت: ۱۴۰۲/۳/۲ تاریخ بازنگری: ۱۴۰۳/۷/۲۶ تاریخ پذیرش و انتشار: ۱۴۰۳/۷/۲۷

چکیده: یکی از پرکاربردترین مباحث آماری، مسایل رگرسیونی است. در مسایل رگرسیونی فرض اساسی بر روی خطاها، نرمال بودن آنهاست که این فرض در برخی موارد به سبب وجود ویژگی‌های عدم تقارن یا مکان‌های شکست در داده‌ها برقرار نمی‌باشد. مدل رگرسیون تکه‌ای یکی از راه‌های برون رفت در شرایط نرمال نبودن خطاهاست که به‌طور گسترده در حوزه‌های مختلفی به کار گرفته شده‌اند، که در آن‌ها تشخیص نقطه شکست مهم است و مکان‌های شکست در مدل‌های رگرسیون تکه‌ای برای دانستن زمان و چگونگی تغییر الگوی ساختار داده ضروری است. یکی از مشکلات عمده در این داده‌ها وجود دم سنگینی است که با استفاده از برخی توزیع‌ها که به عنوان تعمیمی از توزیع نرمال هستند این مشکل برطرف شده است. در این مقاله بر اساس توزیع مخلوط مقیاسی نرمال، مدل رگرسیونی تکه‌ای مورد بررسی قرار خواهد گرفت که می‌توان به جای نرمال با به کارگیری تعمیم‌هایی از توزیع نرمال این مشکل را برطرف نمود. همچنین این مدل با مدل رگرسیون تکه‌ای استاندارد که برگرفته از خطاهای نرمال است مورد مقایسه قرار خواهد گرفت.

واژه‌های کلیدی: الگوریتم EM، توزیع مخلوط مقیاسی نرمال، رگرسیون تکه‌ای، نقطه شکست.

کد موضوع بندی ریاضی (۲۰۱۰): 62J05, 62J02.

۱ مقدمه

در بسیاری از مسایل کاربردی رگرسیون، رابطه بین متغیر پاسخ و متغیرهای مستقل اغلب به طور ناگهانی در برخی مکان‌ها تغییر می‌کند که به آن‌ها نقاط شکست یا نقاط تغییر^۱ می‌گویند. تغییرات ناگهانی که در این مکان‌ها اتفاق می‌افتد باعث وجود انحراف در تجزیه و تحلیل داده‌ها می‌شود. اگرچه می‌توان از الگوهای رگرسیون غیرخطی برای بررسی این گونه داده‌ها استفاده کرد، باز هم اهمیت نقاط شکست اغلب نادیده گرفته می‌شود و امکان دارد پارامترها با اربیبی تخمین زده شوند. از این رو، مدل‌های رگرسیون تکه‌ای برای نشان دادن اهمیت نقاط شکست و برآورد مناسب پارامترهای مدل مورد استفاده قرار می‌گیرند. به طور خاص، شناسایی نقاط شکست برای درک زمان و چگونگی تغییر مدل بسیار مهم هستند.

در مقالات، مدل‌های رگرسیونی که شامل نقاط شکست هستند به رگرسیون تکه‌ای^۲ موسوم‌اند. استفاده از رگرسیون تکه‌ای به جای رگرسیون غیر خطی بسیار رایج است زیرا رگرسیون تکه‌ای به طور گسترده در زمینه‌های مختلف مانند پزشکی، اقتصاد سنجی و هواشناسی به کار گرفته شده‌اند. به عنوان مثال در تحقیقات پزشکی، خطر تولد نوزاد مبتلا به سندرم داون برای مادران جوان کم است، اما این خطر به طور قابل توجهی با افزایش سن مادر پس از یک آستانه سنی خاص افزایش می‌یابد (موجنو، ۲۰۰۸). در اقتصاد سنجی، نقاط شکست نشانه‌های مهمی از زمانی هستند که ساختار بازار به دلیل مداخلات دولت در بازارهای مالی دچار تغییر شده‌اند (لوشی و همکاران، ۲۰۱۰). یکی دیگر از کاربردهای رگرسیون تکه‌ای، دسته‌بندی اشیاء به خوشه‌های همگن است. به عنوان مثال، چانگ و همکاران (۲۰۱۵) وابستگی میزان مصرف سوخت (اندازه‌گیری شده با مایل در هر گالن) به وزن ماشین و قدرت آن را مورد بررسی قرار دادند. آن‌ها دو خوشه جدا مشخص کردند که میزان مصرف سوخت به طور ناگهانی پس از فراتر رفتن وزن خودرو و سرعت آن از یک آستانه مشخص تغییر می‌کند.

مسئله اصلی در مدل‌های رگرسیون تکه‌ای، مشخص کردن تعداد بخش‌ها و مرزهای تقسیم‌بندی شده است که این مسئله اغلب به طور جداگانه در تحقیقات رگرسیون تکه‌ای مورد بحث قرار می‌گیرند. به عنوان مثال، سیوپرکا (۲۰۱۱، ۲۰۰۹) ابتدا یک برآوردگر را برای مشخص کردن نقاط شکست، پیشنهاد کرد و سپس برآورد پارامترهای مدل را تحت این نقاط شکست انجام داد. روش‌های متعددی برای تخمین نقاط شکست و برآورد پارامترها در مدل رگرسیون تکه‌ای پیشنهاد شده است که برای اطلاعات بیشتر می‌توان به سیوپرکا (۲۰۰۴)؛ سیوپرکا و داپزل (۲۰۰۸)؛ فرنهاد (۲۰۰۶)؛ چن و همکاران (۲۰۱۱) مراجعه کرد. جلیلی و همکاران (۱۳۹۵) نمودار کنترل جدیدی مبتنی بر رویکرد آزمون خطی تعمیم یافته و رگرسیون تکه‌ای، به منظور پایش پروفایل‌های خطی چندگانه با اثرات متقابل در فاز ۲ ارائه دادند. آنها نشان دادند نمودار کنترل پیشنهادی عملکرد بسیار بهتری نسبت به نمودار کنترل مبتنی بر روش حداقل مربعات خطا دارا بودند. یکی از معایب مربوط به این روش‌ها تشخیص نقطه شکست واحد است. مشکل عمده دیگر مربوط به داده‌های پرت است. برای غلبه بر این مشکل لو و چانگ (۲۰۲۱) مدل رگرسیون

¹Break-points or Change-points²Piecewise regressions

تکه‌ای را بر پایه توزیع تی معرفی کردند و برآورد پارامترهای مدل را بر اساس روش درستنمایی انجام دادند. آنها نشان دادند که این مدل برای پوشش داده پرت مناسب‌تر از مدل بر پایه توزیع نرمال است. اخیراً، **انتبرت و بین (۲۰۲۴)** روشی هوشمند را برای تخمین نقاط شکست معرفی کردند. آنها نشان دادند که نقاط شکست در این روش در یک مرحله بدست می‌آیند. در این مقاله، ما از روشی برای تخمین چندین نقطه شکست و پارامترهای رگرسیون در حالتی که خطای‌های مدل از توزیع مخلوط مقیاسی نرمال^۱ (SMN) پیروی می‌کند ارائه می‌دهیم که به طور همزمان با استفاده از مدل‌های آمیخته بیان شده‌اند. درباره اهمیت استفاده از توزیع SMN می‌توان به خواصی از این توزیع در بعضی حالت‌های خاص آن برای شناسایی نقاط پرت نام برد. در این مقاله به کمک شبیه‌سازی به این خاصیت اشاره شده است. این موضوع از انگیزه‌های معرفی این مدل است. از دیگر موارد انگیزشی و نوآوری می‌توان به برآورد پارامترهای مدل به صورت فرم بسته در برخی از حالت‌های خاص توزیع SMN اشاره کرد.

در این مقاله بدست آوردن نقاط شکست همزمان با طبقه‌بندی داده‌ها به داخل خوشه‌های همگن انجام می‌شود. بنابراین می‌توان از روش‌های خوشه‌بندی استفاده کرد. مدل‌های آمیخته برای خوشه‌بندی داده‌ها مدل‌هایی قوی و محبوب هستند، که با استفاده از الگوریتم امید-ماکسیم^۲ (EM) و تعمیم‌های آن به برآورد پارامترها و تشخیص نقاط شکست می‌پردازند. برای بررسی بیشتر این مطالب منابع **یانگ (۲۰۱۴)**؛ **یلدیریم (۲۰۱۴)**؛ **کشاورز و هانگ (۲۰۱۴b)**؛ **چانگ و همکاران (۲۰۱۶)**؛ **جعفری و گلعلی‌زاده (۱۴۰۲)** می‌تواند مورد مطالعه قرار گیرند. بیشتر این روش‌ها بر روی تغییرات میانگین و واریانس در محدوده دو نقطه شکست متوالی تمرکز دارند. تفاوت‌های اساسی این مقاله در مقابل مقالات گذشته استفاده از مدل‌های آمیخته برای شناسایی نقاط شکست، برآورد پارامترها با استفاده از الگوریتم EM و بدست آوردن فرم بسته برای برخی پارامترهای مدل، شناسایی خودکار داده‌های پرت و استفاده از یک خانواده از توزیع‌ها برای خطاهای مدل است.

با توجه به مزیت مدل‌های آمیخته برای خوشه‌بندی داده‌ها و تشخیص نقطه تغییر و مرز بین خوشه‌ها، با استفاده از الگوریتم EM، یک الگوریتم قوی برای خوشه‌بندی، تخمین نقاط تغییر و پارامترهای مدل رگرسیونی ارائه می‌دهیم. در بخش ۲ به طور خلاصه توزیع SMN را بررسی و برخی از ویژگی‌های مرتبط با آن را تشریح می‌کنیم. در بخش ۳ مدل پیشنهادی را بیان و برآورد پارامترها و همچنین تخمین تعداد نقاط شکست را ارائه می‌کنیم. دو مطالعه شبیه‌سازی در بخش ۴ و دو مثال واقعی در بخش ۵ برای بررسی اثربخشی روش پیشنهادی انجام شده است.

¹Scale mixture normal distribution

²Expectation and maximization

۲ مروری بر توزیع مخلوط مقیاسی نرمال

متغیر تصادفی X دارای توزیع SMN با پارامتر مکان μ ، پارامتر مقیاس σ^2 و پارامتر ν است، و با نماد $SMN(\mu, \sigma^2, \nu)$ نشان داده می‌شود، هرگاه نمایش تصادفی آن به صورت

$$X \stackrel{d}{=} \mu + W^{-1/2} X_1, \quad X_1 \perp W, \quad (1)$$

باشد، که در آن $\perp$ نماد استقلال دو متغیر، X_1 یک متغیر تصادفی دارای توزیع نرمال با میانگین ۰ و واریانس σ^2 و W یک متغیر تصادفی مثبت با تابع توزیع دلخواه $H(\cdot; \nu)$ و مستقل از X_1 است. با توجه به نمایش تصادفی (۱) به سادگی می‌توان نتیجه گرفت که $X|W = w \sim N(\mu, w^{-1}\sigma^2)$ است. بنابراین می‌توان تابع چگالی بردار تصادفی X را به صورت

$$f_{SMN}(x; \mu, \sigma^2, \nu) = \int_0^\infty \phi(x; \mu, w^{-1}\sigma^2) dH(w; \nu) \quad x \in \mathbb{R},$$

بدست آورد، که در آن $\phi(\cdot; \mu, \sigma^2)$ تابع چگالی نرمال با میانگین μ و واریانس σ^2 است. با توزیع‌های مختلفی که برای W می‌توان در نظر گرفت، موارد خاص از توزیع کلی SMN به دست می‌آیند. در این مقاله روی چند مورد رایج از موارد خاص توزیع SMN متمرکز هستیم.

الف- توزیع نرمال (N): اگر توزیع W تباهیده در نقطه ۱ باشد، یعنی W با احتمال ۱ برابر ۱ باشد، آنگاه توزیع SMN همان توزیع نرمال $N(\mu, \sigma^2)$ است.

ب- توزیع تی (t): اگر W دارای توزیع گاما با پارامتر مقیاس $\nu/2$ و پارامتر شکل $\nu/2$ باشد، آنگاه توزیع SMN همان توزیع تی $t(\mu, \sigma^2, \nu)$ با تابع چگالی

$$f_t(x; \mu, \sigma^2, \nu) = \frac{(\pi\nu)^{-1/2} \Gamma(\frac{\nu+1}{2}) \sigma}{\Gamma(\nu/2) [1 + (x - \mu)/(\sigma^2\nu)]^{(\nu+1)/2}}, \quad (2)$$

است. باید توجه داشت که در توزیع تی اگر $\nu = 1$ باشد، آنگاه توزیع کوشی بدست می‌آید. فرض کنید $W \sim \text{Gamma}(\nu/2, \nu/2)$ و $X \sim t(\mu, \sigma^2, \nu)$ باشد. در این صورت W به شرط $X = x$ دارای توزیع گاما $\text{Gamma}(\nu + 1, \nu + \frac{(x-\mu)^2}{\sigma^2})$ است. به عبارت دیگر داریم:

$$f(w | x) = \frac{w^{\frac{\nu+1}{2}-1}}{\Gamma(\frac{\nu+1}{2})} \left(\frac{\nu + \frac{(x-\mu)^2}{\sigma^2}}{2} \right)^{\frac{\nu+1}{2}} \times \exp \left\{ -\frac{w}{2} \left(\nu + \frac{(x-\mu)^2}{\sigma^2} \right) \right\}.$$

ج- توزیع اسلش^۱ (SL): اگر توزیع W بتا با پارامتر ν و ν_1 باشد، آنگاه توزیع SMN همان توزیع اسلش $X \sim SL(\mu, \sigma^2, \nu)$ با تابع چگالی

$$f_{SL}(x; \mu, \sigma^2, \nu) = \int_0^1 \nu w^\nu \phi(wx; w\mu, \sigma^2) dw \quad x \in \mathbb{R}.$$

است. در صورتی که توزیع W بتا با پارامتر ν و ν_1 باشد و $X \sim SL(\mu, \sigma^2, \nu)$ باشد، آنگاه W به شرط $X = x$ دارای توزیع بریده شده گاما $I_{[\cdot, 1]}(\nu + 1, \nu + \frac{(x-\mu)^2}{\sigma^2}) I_{[\cdot, 1]}$ است، که در آن $I_{[\cdot, 1]}(a) = \int_0^1 t^a dt$ تابع نشانگر است.

د- توزیع نرمال آلوده شده^۲ (CN): اگر توزیع W دارای تابع چگالی به صورت

$$h(w; \nu) = \nu_1 I_{\{\nu_1\}}(w) + (1 - \nu_1) I_{\{1\}}(w), \quad (3)$$

باشد، که در آن $\nu = (\nu_1, \nu_2)$ ، آنگاه توزیع SMN همان توزیع نرمال آلوده شده $X \sim CN(\mu, \sigma^2, \nu)$ با تابع چگالی

$$f_{CN}(x; \mu, \sigma^2, \nu) = \nu_1 \phi(x; \mu, \sigma^2) + (1 - \nu_1) \phi(x; \mu, \nu_2 \sigma^2),$$

است. حال اگر $X \sim CN(\mu, \sigma^2, \nu)$ باشد. آنگاه W به شرط $X = x$ دارای تابع چگالی

$$f(w | x) = \frac{(\nu_1 I_{\{\nu_1\}}(w) + (1 - \nu_1) I_{\{1\}}(w)) \phi(x; \mu, w \sigma^2)}{\nu_1 \phi(x; \mu, \sigma^2) + (1 - \nu_1) \phi(x; \mu, \nu_2 \sigma^2)},$$

است. در این مدل تابع چگالی رابطه (۳) را به صورت

$$h(w; \nu) = \nu_1^{\frac{w-1/\nu_2}{1-1/\nu_2}} (1 - \nu_1)^{\frac{1-w}{1-1/\nu_2}}. \quad (4)$$

بازنوسی می‌کنیم. حال متغیر تصادفی $V = \frac{W-1/\nu_2}{1-1/\nu_2}$ دارای توزیع برنولی با پارامتر ν_2 است. از این روش برای بدست آوردن برآورد پارامترهای مدل به صورت فرم بسته استفاده می‌کنیم.

¹Slash distribution

²Contaminated-normal distribution

۳ مدل رگرسیون تکه‌ای با خطاهای مخلوط مقیاسی نرمال

در این بخش مدل‌های رگرسیون تکه‌ای براساس توزیع نرمال^۱ (N-BP)، توزیع تی^۲ (t-BP)، توزیع اسلش^۳ (SL-BP) و توزیع نرمال آلوده شده^۴ (CN-BP) در این مقاله مورد بررسی قرار می‌گیرند. در حالت کلی این مدل را با نماد SMN-BP^۵ نمایش می‌دهیم. فرض کنید $(x_1, y_1), \dots, (x_n, y_n)$ یک نمونه تصادفی n تایی از مدل رگرسیونی با متغیر مستقل $x_j = (x_{j1}, \dots, x_{jp})$ و متغیر وابسته $y_j \in \mathbb{R}$ برای $j = 1, \dots, n$ باشد. حال متغیر n -بعدی $x_d^* = (x_{1d}, \dots, x_{nd})^T$ را در نظر می‌گیریم که در آن $d = 1, \dots, p$ است. بر این اساس مشاهدات $(x_1, y_1), \dots, (x_n, y_n)$ را می‌توان به ترتیب افزایش در هر یک از x_d^* قرار داد بطوریکه نقاط شکست را می‌توان با رتبه مشاهدات در مجموعه مرتب شده از x_d^* توصیف کرد. فرض کنید مدل رگرسیون تکه‌ای $c - 1$ نقطه شکست با مقادیر $\tau_1, \dots, \tau_{c-1}$ داشته که $1 \leq \tau_1 < \dots < \tau_{c-1} \leq n - 1$ است. همچنین $\tau_0 = 1$ و $\tau_c = n$ در نظر می‌گیریم. یک مدل رگرسیون تکه‌ای شامل $c - 1$ نقطه شکست را می‌توان به صورت

$$y_{ij} = \beta_{i0} + \beta_{i1}x_{j1} + \dots + \beta_{ip}x_{jp} + \epsilon_{ij},$$

$$\tau_{i-1} \leq y_j < \tau_i, \quad i = 1, \dots, p, \quad j = 1, \dots, n, \quad (5)$$

بیان کرد، بطوری که ϵ_{ij} از هم مستقل و دارای توزیع $SMN(0, \sigma_i, \nu_i)$ هستند. به دلیل اینکه $c - 1$ نقطه تغییر در $n - 1$ مشاهده ابتدایی صورت می‌گیرد، بنابراین $\binom{n-1}{c-1}$ حالت برای $c - 1$ نقطه شکست ایجاد می‌شود. برای هر ترکیب $\tau_1, \dots, \tau_{c-1}$ متغیر تصادفی $\omega_{\tau_1 \dots \tau_{c-1}}$ را به صورت

$$\omega_{\tau_1 \dots \tau_{c-1}} = \begin{cases} 1 & \text{اگر } 1 \leq \tau_1 < \dots < \tau_{c-1} \leq n - 1 \text{ نقاط شکست باشند,} \\ 0 & \text{در غیر این صورت.} \end{cases}$$

تعریف می‌کنیم. حال $\gamma_{\tau_1 \dots \tau_{c-1}} = P(\omega_{\tau_1 \dots \tau_{c-1}} = 1)$ را در نظر می‌گیریم، در نتیجه

$$\sum_{1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1} \gamma_{\tau_1 \dots \tau_{c-1}} = 1.$$

¹Normal Break Point

²Student's t Break Point

³Slash Break Point

⁴Contaminated-normal Break Point

⁵Scale mixture normal break point

حال متغیر Z_{ij} را به صورت

$$\sum_{1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1, \tau_{i-1} \leq j < \tau_i} \dots \sum \omega_{\tau_1, \dots, \tau_{c-1}} \quad 1 \leq i \leq c, \quad 1 \leq j \leq n,$$

تعریف می‌کنیم و برای سایر نقاط $\circ$ در نظر می‌گیریم. بنابراین می‌توان نتیجه گرفت که Z_{ij} یک متغیر برنولی با مقدار ۱ برای حالتی که (x_j, y_j) متعلق به بخش i ام و $\circ$ برای سایر نقاط است. در نتیجه مدل رگرسیون تکه‌ای به یک مدل کلاس‌بندی واضح با متغیرهای مشاهده شده $(x, y) = ((x_1, y_1), \dots, (x_n, y_n))^T$ و متغیرهای غیر مشاهده شده $Z = (Z_1^T, \dots, Z_n^T)^T$ که در آن $Z_k = (Z_{k1}, \dots, Z_{kp})^T$ تبدیل می‌شود. با توجه به مطالب بیان شده شکل تابع درست‌نمایی پیچیده است و بدست آوردن برآورد پارامترها به روش مستقیم کار دشواری است. برای به دست آوردن برآوردهای ماکسیمم درست‌نمایی^۱، (MLE) الگوریتم EM یک روش مناسب است. برای اجرای این الگوریتم روی مدل رگرسیون تکه‌ای با خطاهای مقیاسی نرمال کافی است از نمایش تصادفی

$$\begin{aligned} Y_j | W_j = w_j, Z_{ij} = 1 &\sim N(x_j^T \beta_i, w_j^{-1} \sigma_i^2), \\ W_j | Z_{ij} = 1 &\sim H(\cdot; \nu), \\ Z_j &\sim \text{Multi}(1, \pi_1, \dots, \pi_c) \end{aligned} \quad (۶)$$

استفاده شود، که در آن Multi نماد توزیع چندجمله‌ای و π_i ها وزن‌های آمیخته هستند به طوری که $\sum_{i=1}^c \pi_i = 1$ و مقدار آنها بر اساس سایر پارامترها بدست می‌آید و نیاز به برآورد ندارند.

۳.۱ برآورد ماکسیمم درست‌نمایی با استفاده از الگوریتم EM

یکی از ابزارهای مفید و کاربردی جهت به دست آوردن MLE پارامترهای مدل وقتی که فرم تابع چگالی پیچیده باشد، الگوریتم EM است. این روش که نخستین بار توسط **دمپستر و همکاران** (۱۹۷۷) معرفی شد، زمانی کاربرد دارد که داده‌ها بطور کامل مشاهده نشده باشند. کامل نبودن داده‌ها ممکن است به دلایل مختلف همانند وجود مشاهدات گمشده یا مشاهدات بریده شده یا سانسور شده، باشد. می‌دانیم سرعت همگرایی الگوریتم EM برای بدست آوردن جواب بهینه امکان دارد نسبت به سایر روش‌ها کم باشد. با این وجود، سادگی اجرا و بدست آوردن جواب‌های بهینه مورد اعتماد می‌تواند دلیل اصلی استفاده از این روش باشد. برای برآورد پارامترهای مدل یعنی بردار $\tau = (\tau_1, \dots, \tau_{c-1})$ ، $\sigma^2 = (\sigma_1^2, \dots, \sigma_c^2)$ ، $\beta = (\beta_1, \dots, \beta_c)$ و $\nu = (\nu_1, \dots, \nu_c)$ که بطور کلی با $\theta = (\tau, \beta, \sigma^2, \nu)$ فرض می‌شود برای داده‌های نمایش می‌دهیم از الگوریتم EM استفاده خواهیم کرد. برای اجرای این الگوریتم، فرض می‌شود برای داده‌های مشاهده شده $(x, y) = ((x_1, y_1), \dots, (x_n, y_n))^T$ ، مشاهدات پنهان $W = (W_1, \dots, W_n)^T$ و

^۱Maximum Likelihood Estimator

$\mathbf{Z} = (\mathbf{Z}_1^\top, \dots, \mathbf{Z}_n^\top)^\top$ نیز وجود دارند. در اینصورت لگاریتم تابع درستنمایی داده‌های کامل به صورت

$$\begin{aligned} \ell_c(\boldsymbol{\theta} \mid \mathbf{y}, \mathbf{W}, \mathbf{Z}) = & \sum_{i=1}^c \sum_{j=1}^n Z_{ij} \left\{ \log h(W_j; \nu_i) + \frac{1}{\nu} \log W_j - \log \sigma_i \right. \\ & \left. - \frac{W_j}{\nu \sigma_i^2} (y_j - \mathbf{x}_j^\top \boldsymbol{\beta}_i) \right\} \end{aligned} \quad (۷)$$

خواهد بود. در این صورت الگوریتم EM در دو گام تکرار می‌شود.

الف- گام E: برای هر تکرار k ، تابع Q که به صورت امید ریاضی تابع (۷) تحت شرط مشاهدات و $\hat{\boldsymbol{\theta}}^{(k)}$ محاسبه شده است، را به صورت

$$Q(\boldsymbol{\theta} \mid \hat{\boldsymbol{\theta}}^{(k)}) = \sum_{i=1}^c \sum_{j=1}^n \hat{z}_{ij}^{(k)} \left\{ \hat{\Upsilon}_{ij}^{(k)}(\nu_i) + \frac{1}{\nu} \hat{t}_{ij}^{(k)} - \log \sigma_i - \frac{\hat{w}_{ij}^{(k)}}{\nu \sigma_i^2} (y_j - \mathbf{x}_j^\top \boldsymbol{\beta}_i) \right\}, \quad (۸)$$

به دست می‌آوریم، که در آن $\hat{w}_{ij}^{(k)} = E(W_j \mid \hat{\Upsilon}_{ij}^{(k)}(\nu_i)) = E(\log h(W_j; \nu_i) \mid y_j, Z_{ij} = 1, \hat{\boldsymbol{\theta}}^{(k)})$ و $\hat{t}_{ij}^{(k)} = E(\log W_j \mid y_i, Z_{ij} = 1, \hat{\boldsymbol{\theta}}^{(k)})$ و $y_i, Z_{ij} = 1, \hat{\boldsymbol{\theta}}^{(k)}$ به دست می‌آیند. برآورد $\hat{z}_{ij}^{(k)}$ از طریق امید ریاضی شرطی Z_{ij} به شرط y_j و مقدار برآورد پارامترها در مرحله k ام به صورت

$$\begin{aligned} \hat{z}_{ij}^{(k)} = E(Z_{ij} \mid y_j, \hat{\boldsymbol{\theta}}^{(k)}) &= \sum_{\substack{1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1 \\ \tau_{i-1} \leq j < \tau_i}} \dots \sum E(\omega_{\tau_1 \dots \tau_{c-1}} \mid y_j, \hat{\boldsymbol{\theta}}^{(k)}) \\ &= \sum_{\substack{1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1 \\ \tau_{i-1} \leq j < \tau_i}} \dots \sum P(\omega_{\tau_1 \dots \tau_{c-1}} = 1 \mid y_j, \hat{\boldsymbol{\theta}}^{(k)}), \end{aligned}$$

بدست می‌آیند. یک برآورد برای $P(\omega_{\tau_1 \dots \tau_{c-1}} = 1 \mid y_j, \hat{\boldsymbol{\theta}}^{(k)})$ به صورت

$$\begin{aligned} \gamma_{\tau_1 \dots \tau_{c-1}}^* = & \left(\prod_{j=1}^{\tau_1} f_{\text{SMN}}(y_j; \mathbf{x}_j^\top \hat{\boldsymbol{\beta}}_1^{(k)}, \hat{\sigma}_1^{(k)\top}, \hat{\nu}_1^{(k)}) \dots \prod_{j=\tau_{c-1}+1}^n f_{\text{SMN}}(y_j; \mathbf{x}_j^\top \hat{\boldsymbol{\beta}}_c^{(k)}, \hat{\sigma}_c^{(k)\top}, \hat{\nu}_c^{(k)}) \right) \hat{\gamma}_{\tau_1 \dots \tau_{c-1}}^{(k)} / \\ & \sum_{1 \leq g_1 < \dots < g_{c-1} \leq n-1} \dots \sum \prod_{j=1}^{g_1} f_{\text{SMN}}(y_j; \mathbf{x}_j^\top \hat{\boldsymbol{\beta}}_1^{(k)}, \hat{\sigma}_1^{(k)\top}, \hat{\nu}_1^{(k)}) \dots \prod_{j=g_{c-1}+1}^n f_{\text{SMN}}(y_j; \mathbf{x}_j^\top \hat{\boldsymbol{\beta}}_c^{(k)}, \hat{\sigma}_c^{(k)\top}, \hat{\nu}_c^{(k)}) \hat{\gamma}_{g_1 \dots g_{c-1}}^{(k)}, \end{aligned}$$

حاصل می‌شود، که در آن $\hat{\gamma}_{\dots}^{(k)}$ برآورد $\gamma_{\dots}$ تحت $\hat{\boldsymbol{\theta}}^{(k)}$ است. این برآورد تحت $1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1$ است که یکی از حالات نقاط شکست است و صورت کسر بالا از این حالت حاصل می‌شود و مخرج این کسر تمام

$(n-1)$ حالت‌های ممکن برای نقاط شکست است که به صورت جمع در مخرج آمده است. بنابراین برآورد Z_{ij} به صورت

$$\hat{z}_{ij}^{(k)} = \begin{cases} \sum \dots \sum_{\substack{1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1 \\ \tau_{i-1} \leq j < \tau_i}} \gamma_{\tau_1 \dots \tau_{c-1}}^* & \text{اگر } 1 \leq i \leq c, \quad 1 \leq j \leq n, \\ 0 & \text{سایر نقاط} \end{cases}$$

حاصل شده است. حال برآورد پارامترهای $\tau = (\tau_1, \dots, \tau_{c-1})$ را به صورت زیر برزسانی می‌کنیم:

$$\hat{\tau}^{(k+1)} = \arg \max_{1 \leq \tau_1 < \dots < \tau_{c-1} \leq n-1} \gamma_{\tau_1 \dots \tau_{c-1}} | \tau \quad (9)$$

ب- گام CM: در این گام برای هر تکرار k ، پارامتر مجهول θ به صورت $\hat{\theta}^{(k+1)}$ به روزرسانی می‌شود. با ماکسیم‌سازی (۸) نسبت به θ برآوردهای به روز شده به صورت

$$\begin{aligned} \hat{\beta}_i^{(k+1)} &= \left[\sum_{j=1}^n \hat{z}_{ij}^{(k)} \hat{w}_{ij}^{(k)} \mathbf{x}_j \mathbf{x}_j^\top \right]^{-1} \left[\sum_{j=1}^n \hat{z}_{ij}^{(k)} \hat{w}_{ij}^{(k)} y_i \mathbf{x}_j \right], \\ \hat{\sigma}^{2(k+1)} &= \frac{1}{\sum_{j=1}^n \hat{z}_{ij}^{(k)}} \sum_{j=1}^n \hat{w}_{ij}^{(k)} (y_i - \mathbf{x}_i^\top \hat{\beta}^{(k+1)})^2 \end{aligned}$$

به دست می‌آیند. چون فرم بسته‌ای برای برآورد پارامتر ν وجود ندارد، آن را بوسیله ماکسیم‌سازی تابع Q و به صورت $\hat{\nu}^{(k+1)} = \arg \max_{\nu} Q(\nu, \hat{\beta}^{(k+1)}, \hat{\sigma}^{2(k+1)})$ به روز رسانی می‌کنیم.

این فرآیند تا زمانی ادامه دارد که یک شرط همگرایی مانند $|\ell(\hat{\theta}^{(k+1)}) - \ell(\hat{\theta}^{(k)})| < \varepsilon$ برقرار باشد که در آن $\ell(\hat{\theta}^{(k)})$ مقدار تابع درستنمایی در مرحله k ام است.

$$\ell(\hat{\theta}) = \sum_{j=1}^n \log \left(\sum_{i=c}^n f_{\text{SMN}}(y_j; \mathbf{x}_j^\top \hat{\beta}, \hat{\sigma}_i^2, \hat{\nu}) \right).$$

در این طرح برای جلوگیری از عدم همگرایی، از معیار همگرایی ایتکن^۱ با میزان خطا $\varepsilon = 10^{-5}$ ، برای توقف فرایند استفاده شده است.

¹Aickten

۳.۲ برآورد پارامتر ν در مدل CN

با توجه به مطالب ارایه شده در پایان بخش ۲، متغیر W را به V تبدیل می‌کنیم. حال داریم:

$$\begin{aligned}\hat{v}_{ij}^{(k)} &= E(V_j | y_j, Z_{ij} = 1, \hat{\theta}^{(k)}) = P(V_j = 1 | y_j, Z_{ij} = 1, \hat{\theta}^{(k)}) \\ &= \frac{\nu_{1i} \phi(y_j; \mathbf{x}_j^\top \hat{\beta}_i^{(k)}, \hat{\sigma}_i^{(k)})}{f_{\text{CN}}(y_j; \mathbf{x}_j^\top \hat{\beta}_i^{(k)}, \hat{\sigma}_i^{(k)}, \hat{\nu}_i^{(k)})}.\end{aligned}\quad (10)$$

بنابراین برآورد پارامتر ν به صورت

$$\begin{aligned}\hat{\nu}_{1i}^{(k+1)} &= \max \left\{ \nu_{1i}^*, \frac{1}{n_i} \sum_{j=1}^n \hat{z}_{ij}^{(k)} \hat{v}_{ij}^{(k)} \right\}, \\ \hat{\nu}_{2i}^{(k+1)} &= \max \left\{ \nu_{2i}^*, \frac{\sum_{j=1}^n \hat{z}_{ij}^{(k)} (1 - \hat{v}_{ij}^{(k)}) \left(\frac{y_j - \mathbf{x}_j^\top \hat{\beta}_i^{(k)}}{\hat{\sigma}_i^{(k)}} \right)^2}{\sum_{j=1}^n \hat{z}_{ij}^{(k)} (1 - \hat{v}_{ij}^{(k)})} \right\},\end{aligned}$$

حاصل می‌شود، که در آن n_i تعداد مشاهداتی است که در فاصله $[\tau_{i-1}, \tau_i]$ قرار دارند، و $\nu_{1i}^* = 0.01$ و $\nu_{2i}^* = 0.5$.

۳.۳ تشخیص خودکار نقاط پرت

در این بخش تشخیص خودکار نقاط پرت را برای مدل t-BP و CN-BP ارایه می‌دهیم. برای مدل t-BP از فاصله ماحالانوبیس^۱ استفاده می‌کنیم. ابتدا فرض می‌کنیم y_j متعلق به فاصله دو نقطه شکست $(\tau_{i-1}, \tau_i]$ باشد. در این صورت با استفاده از رابطه

$$\delta(y_j; \mathbf{x}_j^\top \hat{\beta}_i^{(k)}, \sigma_i^{(k)}) = \left(\frac{y_j - \mathbf{x}_j^\top \hat{\beta}_i^{(k)}}{\hat{\sigma}_i^{(k)}} \right)^2$$

که مجانباً دارای توزیع خی دو با ۱ درجه آزادی است، نقاط پرت را شناسایی می‌کنیم. در مدل CN-BP ابتدا فرض می‌کنیم y_j متعلق به فاصله دو نقطه شکست $(\tau_{i-1}, \tau_i]$ باشد. سپس مقدار $\hat{v}_{ij}$ را محاسبه می‌کنیم و اگر مقدار $\hat{v}_{ij} > 0.5$ باشد، آنگاه y_j یک نقطه خوب است و در صورتیکه $\hat{v}_{ij} \leq 0.5$ باشد y_j یک نقطه پرت است.

¹Mahalanobis

۳.۴ انتخاب پارامتر اولیه و انتخاب مدل مناسب

در ادامه برخی از فنون اجرای الگوریتم EM بر روی مدل رگرسیون تکه‌ای را مورد بررسی قرار می‌دهیم. این فنون شامل مقادیر اولیه و معیار مناسب بودن مدل است.

الف- مقادیر اولیه: به دلیل آنکه مدل‌های آمیخته ساده منجر به تولید توابع لگاریتم درست‌نمایی چند مدی می‌شوند، استفاده از الگوریتم EM بدون در نظر گرفتن مقدار اولیه $(\hat{\theta}^{(0)})$ مناسب ممکن است به جواب‌های بهینه منجر نشود. همچنین مقادیر اولیه مناسب سبب افزایش سرعت همگرایی الگوریتم EM خواهد شد. در نتیجه انتخاب مقدار اولیه مناسب در این مدل‌ها بسیار با اهمیت است. روش‌های متفاوتی برای انتخاب مقادیر اولیه مناسب موجود است. در زیر به یکی از این روش‌ها اشاره می‌شود.

• ابتدا برای نمونه‌ی تحت مطالعه تعداد نقاط شکست یعنی c را مشخص می‌نماییم.

• مقادیر اولیه برای $\hat{\gamma}_{\tau_1 \dots \tau_{c-1}}^{(0)}$ را در هر گروه هم شانس و برابر $\frac{1}{(n-1)}$ در نظر می‌گیریم.

• برای هر i ، بر اساس مقادیر $\hat{\gamma}_{\tau_1 \dots \tau_{c-1}}^{(0)}$ ، مقادیر اولیه $\hat{\sigma}_i^{(0)2} = 1$ برای هر گروه در نظر می‌گیریم. حال با در نظر گرفتن مقادیر داده شده $\hat{\gamma}_{\tau_1 \dots \tau_{c-1}}^{(0)}$ ، $\hat{\sigma}_i^{(0)2} = 1$ ، مقدار $\hat{\beta}_i^{(0)}$ را با استفاده از مرحله CM محاسبه می‌کنیم.

ب- معیار مناسب بودن مدل: برای به دست آوردن تعداد مناسب نقاط شکست c ، باید الگوریتم EM را برای مقادیر متفاوت c اجرا کرد و برای هر c ، یکی از معیارهای مناسب بودن مدل را محاسبه کرده و براساس آن c مناسب را اختیار کرد. به منظور بررسی و مقایسه‌ی مدل‌های آماری مطرح شده در این مقاله، از دو معیار^۱ AIC و^۲ BIC که به صورت زیر محاسبه می‌شوند، استفاده خواهد شد:

$$AIC = 2m - 2\ell_{\max} \quad \text{و} \quad BIC = m \log n - 2\ell_{\max},$$

که در آن m تعداد پارامترهای برآورد شده و $\ell_{\max}$ لگاریتم تابع ماکسیم درست‌نمایی است.

۴ مطالعه شبیه‌سازی

در این شبیه‌سازی ویژگی‌های اریبی (Bias) و میانگین مربعات خطای (MSE) برآوردگرهای ماکسیم درست‌نمایی بر اساس الگوریتم EM را بررسی خواهیم کرد. برای این شبیه‌سازی، به تعداد ۵۰۰ بار نمونه‌هایی به حجم ۱۰۰، ۵۰ و ۲۰۰ از مدل‌های بیان شده با پارامترهای $(\beta_1 = (0, 0.5), \beta_2 = (10, -0.5), \sigma_1 = 1, \sigma_2 = 1.5)$ و

^۱ Akaike information criterion

^۲ Bayesian information criterion

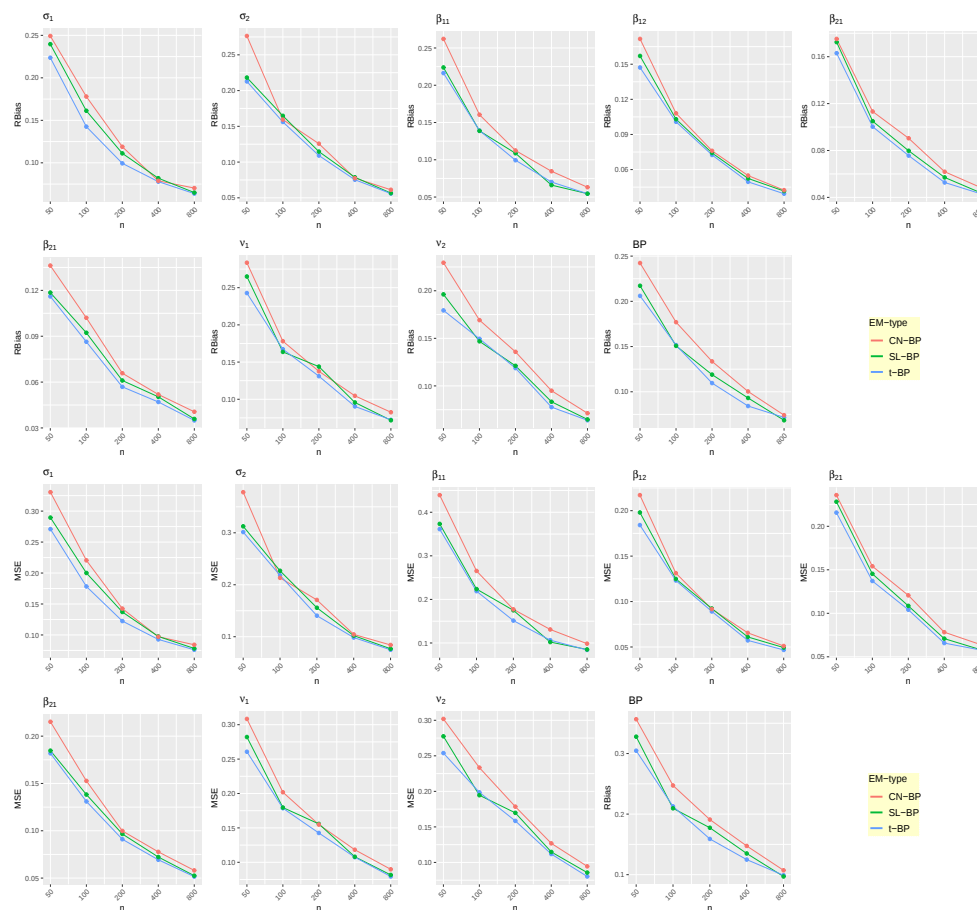
برای مدل t-BP پارامتر $\nu = 3$ ، برای مدل SL-BP پارامتر $\nu = 5$ و برای مدل CN-BP پارامترهای $\nu_1 = 0.4$ و $\nu_2 = 1/2$ تولید شده است. همچنین مقدار نقطه شکست برای همه مدل‌ها به ترتیب ۵۰، ۲۵، ۱۰۰ و ۲۰۰ برای حجم نمونه‌های ۵۰، ۱۰۰، ۲۰۰ و ۴۰۰ است. همچنین متغیر پیشگوی $x = (1, x_{i1})^T$ نیز برای هر n یعنی x_{i1} از فاصله‌ی $(0, 3)$ بطور یکنواخت تولید شده و ثابت در نظر گرفته شده است. در هر تکرار پس از برآورد پارامترها مدل، اریبی و مربع خطای آنها محاسبه شده است. در شکل ۱ میانگین اریبی و مربعات خطای که به صورت زیر محاسبه می‌شوند، گزارش شده است:

$$\text{Bias} = \frac{1}{500} \sum_{i=1}^{500} (\hat{\theta}_i - \theta) \quad \text{و} \quad \text{MSE} = \frac{1}{500} \sum_{i=1}^{500} (\hat{\theta}_i - \theta)^2.$$

با توجه به این شکل مشاهده می‌شود که اریبی و میانگین مربعات خطای برآوردهای ماکسیم درست‌نمایی با افزایش حجم نمونه به سمت صفر میل می‌کنند. بنابراین برآوردهای ML در مدل SMN-BP دارای ویژگی مطلوب مجانبی هستند. همچنین مشاهده شده که تخمین نقاط شکست با افزایش حجم نمونه به خوبی صورت گرفته است. این نقاط با افزایش حجم نمونه به مقدار واقعی نزدیک شده‌اند.

در ادامه، کارایی و عملکرد مدل SMN-BP را از طریق شبیه‌سازی مورد بررسی قرار می‌دهیم. در این شبیه‌سازی به بررسی مقایسه مدل‌ها تحت داده‌های تولید شده از دو فرضیه می‌پردازیم. در فرضیه اول (S_1) خطاهای مدل رگرسیون تکه‌ای از توزیع SMN که در آن $W^{-1} \sim E(0.5)$ (یعنی توزیع نمایی با میانگین ۲) و در فرضیه دوم (S_2) $W^{-1} \sim BS(\alpha, 1)$ (یعنی توزیع بیرنهام-ساندرز (بیرنهام و ساندروز، ۱۹۶۹)) با پارامتر شکل α و پارامتر مقیاس ۱ تولید می‌شوند. این موضوع باعث ایجاد دم سنگینی در داده‌های تولید شده می‌شود. داده‌های تولید شده از مدل‌های بیان شده برای نمونه‌های $n = 50, 100, 200$ با دو نقطه شکست که برای هر نمونه تولید شده‌اند. پارامتر مشترک آنها $\sigma_1 = 1, \sigma_2 = 1.5, \sigma_3 = 2, \beta_1 = (0, 0.5), \beta_2 = (3, -0.5)$ و $\beta_3 = (2, 0.5)$ و متغیر پیشگوی $x = (1, x_{i1})^T$ نیز برای هر n یعنی x_{i1} از فاصله $(0, 5)$ بطور یکنواخت تولید شده و ثابت در نظر گرفته شده است. همچنین پارامتر $\alpha_1 = 3, \alpha_2 = 1, \alpha_3 = 2$ را برای توزیع بیرنهام ساندروز در نظر می‌گیریم. در هر تکرار از شبیه‌سازی مدل‌های CN-BP، t-BP، N-BP و SL-BP بر روی داده‌های تولید شده از دو فرضیه برازش داده می‌شوند و شاخص‌های AIC، BIC، مجموع مربعات باقیمانده^۱ (RSS) و مقادیر نقطه شکست (BP) اندازه‌گیری می‌شود. در انتها میانگین این شاخص‌ها در جدول ۱ گزارش شده است. با دقت در جدول ۱ ملاحظه می‌شود که تاثیر داده‌های پرت به ترتیب در فرضیه S_1 و S_2 در برآوردهای مدل CN-BP و t-BP در مقابل سایر مدل‌ها کمتر است. همچنین مشاهده می‌شود که مدل N-BP بیشترین تاثیر را در وجود داده‌های پرت دارد که نشان از عدم توانایی آن در مدل‌سازی داده‌های پرت دارد. این موضوع نشان می‌دهد که استفاده از سایر حالت‌های خاص مدل SMN-BP یعنی CN-BP و t-BP می‌تواند مناسبتر باشد. این موضوع می‌تواند

¹Residual Sum of Squares



شکل ۱. نتایج شبیه‌سازی برای بررسی اریبی و میانگین مربعات خطای برآوردها.

به دلیل حساسیت دو مدل بیان شده در برابر داده پرت باشد. همچنین مشاهده شده است که با افزایش حجم نمونه تخمین دقیق‌تری از نقاط شکست در مدل SMN-BP صورت می‌پذیرد. بنابراین می‌توان نتیجه گرفت که در مقابل داده‌هایی که دارای دم سنگینی هستند استفاده از مدل‌های CN-BP و t-BP می‌تواند مناسب‌تر باشد. همچنین با توجه به مقادیر بیان شده برای معیار RSS می‌توان نتیجه گرفت که تحت فرضیه S_1 مدل CN-BP نسبت به سایر مدل‌ها بهتر است. در مورد فرضیه S_2 مدل t-BP بهتر است. بنابراین می‌توان نتیجه گرفت که دو مدل t-BP و CN-BP در مقابل داده پرت حساسیت مناسبی دارند. در ادامه برای بررسی شناسایی خودکار داده پرت شبیه‌سازی بعد را طراحی کردیم.

همان طور که نشان داده شد عملکرد مدل‌های CN-BP و t-BP در برابر داده‌های دم سنگین بهتر است.

جدول ۱. نتایج شبیه‌سازی برای عملکرد مدل SMN-BP در برخورد با داده‌های بسیار دم سنگین

فرضیه	نقاط تغیز نقاط	n	معیار	N-BP	t-BP	SL-BP	CN-BP
S1	(۱۵, ۳۲)	۵۰	AIC	۲۹۸,۵	۲۵۵,۱	۲۶۸,۲	۲۵۲,۵
			BIC	۳۰۵,۶	۲۹۰,۷	۲۹۵,۸	۲۸۵,۴
			RSS	۹۷,۵۳	۳۲,۱۴	۵۵,۳۲	۲۲,۸۹
			BP	(۱۰/۱۲, ۴۲/۱۲)	(۱۴/۷۳, ۳۴/۱۷)	(۱۲/۹۸, ۳۶/۰۵)	(۱۵/۱۲, ۳۲/۵۰)
	(۲۵, ۶۷)	۱۰۰	AIC	۶۰۱/۱۷	۵۳۰/۴۲	۵۸۰/۶۴	۵۱۵/۷۷
			BIC	۶۴۰/۱۲	۵۹۰/۳۰	۵۹۷/۴۲	۵۸۸/۹۱
			RSS	۲۱۱/۵۵	۶۹/۱۷	۱۲۰/۱۴	۵۰/۲۳
			BP	(۱۴/۱۴, ۷۵/۱۴)	(۲۲/۴۳, ۶۹/۸۱)	(۲۰/۶۱, ۷۳/۵۵)	(۲۶/۰۸, ۶۶/۸۳)
	(۵۵, ۱۲۵)	۲۰۰	AIC	۱۲۵۰/۰۸	۱۱۰۲/۵۲	۱۱۹۰/۸۶	۱۰۶۷/۴۵
			BIC	۱۳۰۲/۸۵	۱۲۰۲/۴۱	۱۲۴۰/۷۶	۱۱۹۰/۴۵
			RSS	۴۴۰/۸۶	۱۵۰/۴۰	۲۵۰/۷۵	۱۲۸/۶۲
			BP	(۶۳/۸۲, ۱۰۶/۴۱)	(۵۱/۸۶, ۱۲۸/۴۱)	(۶۰/۰۷, ۱۱۵/۶۱)	(۵۳/۹۷, ۱۲۶/۰۸)
S2	(۱۵, ۳۲)	۵۰	AIC	۳۱۷,۹	۲۸۸,۰	۲۹۵,۴	۲۹۰/۶
			BIC	۳۴۹,۰	۳۴۷,۲	۳۵۵,۳	۳۵۰/۶
			RSS	۹۰/۴۰	۱۱,۵۷	۴۳,۸۰	۳۳/۱۵
			BP	(۹/۴۵, ۴۳/۵۲)	(۱۵/۰۹, ۳۲/۱۲)	(۱۳/۵۶, ۳۷/۸۸)	(۱۴/۵۲, ۳۵/۶۳)
	(۲۵, ۶۷)	۱۰۰	AIC	۶۴۰/۰۸	۵۸۰/۱۲	۶۰۸/۳۳	۵۹۹/۰۵
			BIC	۷۸۸/۱۲	۷۱۰/۲۳	۷۲۹/۸۶	۷۲۰/۵۴
			RSS	۱۸۸/۱۷	۵۵/۱۴	۹۰/۵۲	۷۰/۸۶
			BP	(۱۶/۳۳, ۷۳/۱۸)	(۲۵/۱۸, ۶۷/۵۹)	(۲۱/۵۲, ۷۲/۴۷)	(۲۳/۹۲, ۷۱/۳۳)
	(۵۵, ۱۲۵)	۲۰۰	AIC	۱۳۱۰/۳۶	۱۱۸۲/۳۶	۱۲۴۰/۱۷	۱۲۰۵/۱۵
			BIC	۱۵۹۰/۷۷	۱۴۵۰/۱۲	۱۴۹۷/۶۴	۱۴۷۷/۱۰
			RSS	۴۰۰/۶۳	۱۱۰/۲۹	۱۹۲/۴۲	۱۴۹/۳۲
			BP	(۶۵/۴۲, ۱۰۷/۹۲)	(۵۶/۴۴, ۱۲۹/۸۳)	(۶۱/۳۹, ۱۱۲/۸۸)	(۵۰/۶۹, ۱۳۰/۴۱)

بنابراین در ادامه قصد داریم با استفاده از شبیه‌سازی، کارایی مدل‌های بیان شده در شناسایی داده پرت را مورد بررسی قرار دهیم. داده‌های شبیه‌سازی بخش ؟؟ را در نظر می‌گیریم. برای هر مجموعه داده شبیه‌سازی شده، نقاط پرت با اندازه ۲۰ را که به طور یکنواخت در محدوده داده‌های شبیه‌سازی تولید می‌شوند را به داده‌ها اضافه می‌کنیم. در داده‌های شبیه‌سازی شده دو نقطه تغییر وجود دارد و در نتیجه داده‌ها به سه دسته تقسیم می‌شوند. داده‌ها در محدوده اول را با y_{j1} و در محدوده سوم را با y_{j3} نشان می‌دهیم. در هر یک از ۵۰۰ تکرار به داده‌های تولید شده در بازه‌های $(\min(y_{j1}) - 5, \min(y_{j1}))$ و $(\max(y_{j3}), \max(y_{j3}) + 5)$ به ترتیب به دسته اول و سوم داده پرت اضافه می‌کنیم. بنابراین برای هر نمونه تولید شده ۲۰ داده پرت وجود دارد. در جدول ۲ میانگین نسبت نقاط پرتی را که به درستی به عنوان نقاط پرت شناسایی شده‌اند^۱ (TPR) و نسبت نقاط خوب را که به عنوان نقاط پرت^۲ (FPR) مشخص نشده‌اند را برای ارزیابی عملکرد تشخیص نقاط پرت ارائه می‌دهیم. از معیارهای TPR و FPR برای نشان دادن کارایی مدل در شناسایی نقاط پرت استفاده می‌کنیم.

با توجه به جدول ۲ مشاهده می‌شود که میانگین مقادیر TPR و FPR در مدل CN-BP به ترتیب بیشتر و

¹True positive rate

²False positive rate

جدول ۲. میانگین مقادیر TPR و FPR برای ۵۰۰ تکرار برای عملکردهای تشخیص نقاط پرت.

تغییر نقاط	n	۱S		۲S	
		t-BP	CN-BP	t-BP	CN-BP
(۱۵, ۳۲)	۵۰	TPR	۰/۸۷۱	۰/۸۸۲	۰/۸۲۰
		FPR	۰/۱۶۹	۰/۱۷۰	۰/۰۸۲
(۲۵, ۶۷)	۱۰۰	TPR	۰/۸۵۶	۰/۸۶۴	۰/۸۰۲
		FPR	۰/۱۶۷	۰/۱۴۶	۰/۰۵۹
(۵۵, ۱۲۵)	۲۰۰	TPR	۰/۸۴۱	۰/۸۳۱	۰/۷۹۳
		FPR	۰/۱۴۳	۰/۱۳۳	۰/۰۷۳

کمتر از مدل t-BP هستند. در ادامه با افزایش حجم نمونه مقدار TPR در هر دو مدل کاهش پیدا می‌کند. همچنین مشاهده شده است که با افزایش حجم نمونه مقدار FPR در هر دو مدل افزایش پیدا می‌کند. بنابراین مدل CN-BP در شناسایی نقاط پرت بهتر از مدل t-BP است.

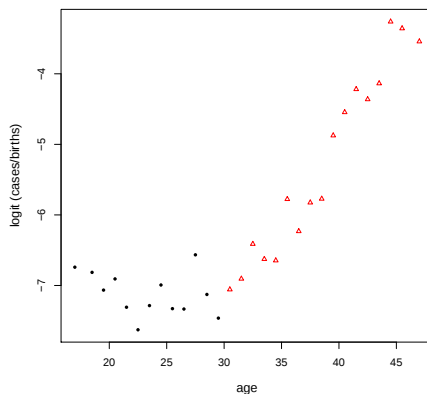
۵ مثال‌های کاربردی

مثال ۱. خطر ابتلا به سندروم داون برای مادران جوان‌تر کم است، اما با افزایش سن مادر این خطر به یک‌باره افزایش می‌یابد. مهم است که بدانیم این خطر چه زمانی و چگونه با سن مادر تغییر می‌کند. در بسته segmented در نرم افزار R داده‌هایی برای بررسی رابطه بین سن مادر و خطر ابتلاء به سندروم داون وجود دارد. این داده‌ها شامل ۳۰ نمونه هستند که در آنها سن مادر و همچنین آزمایشاتی مربوط به سندروم داون در مادر باردار انجام شده است. نمودارهای پراکندگی در شکل ۲ یک تغییر ناگهانی در الگوی خطر ابتلا به سندروم داون را در جنین در حدود ۳۰ سالگی مادر را نشان می‌دهد. با در نظر گرفتن یک نقطه شکست در این داده‌ها، p —مقدار برای آزمون اندرسون دارلینگ^۱ به ترتیب برای دو دسته برابر ۰/۸۵۸ و ۰/۲۳۶ است. این مقادیر نشان دهنده نرمال بودن داده‌ها هستند که فرض استفاده از توزیع SMN را توجیه می‌کنند. نتایج برازش مدل‌های مختلف شامل مقدار تابع درستنمایی و معیارهای مناسبت توزیع یعنی AIC و BIC در جدول ۳ آورده شده است. برای نشان دادن مناسب بودن مدل پیشنهادی برای این داده‌ها از یک مدل رگرسیون غیر خطی نیز استفاده می‌کنیم. مدل غیر خطی بیان شده به صورت

$$y_j = \frac{\beta_0}{1 + \beta_1 \exp\{-\beta_2 x_j\}} + \varepsilon_j \quad (۱۱)$$

در نظر گرفته شده است. خطاهای مدل (۱۱) از توزیع SMN پیروی می‌کنند. نتایج مربوط به مقدار تابع درستنمایی، AIC و BIC نیز در جدول ۳ آورده شده است. همان‌طور که مشاهده می‌شود در تمامی حالت‌های خاص توزیع SMN، مدل رگرسیون تکه‌ای مقادیر AIC و BIC کمتری دارد. این موضوع نشان می‌دهد مدل رگرسیون تکه‌ای بهتر از مدل

^۱Anderson-Darling test



شکل ۲. نمودار پراکندگی سن در مقابل تابع لوجیت موارد تقسیم بر تولد در داده های سندروم داون.

غیر خطی بر این داده‌ها برازش شده است.

جدول ۳. مقادیر تابع درستنمایی و معیارهای AIC و BIC برای داده‌های سندروم داون.

CN-BP		SL-BP		t-BP		N-BP		معیار
تکه‌ای	غیرخطی	تکه‌ای	غیرخطی	تکه‌ای	غیرخطی	تکه‌ای	غیرخطی	
-۵۰/۴۰۹	-۲۷/۲۹۰	-۶۲/۷۵۱	-۳۵/۴۸۳	-۵۶/۴۳۹	-۳۱/۸۷۹	-۷۱/۳۳۲	-۴۵/۴۱۸	$-\ell_{\max}$
۱۱۲/۸۵۸	۷۲/۵۸۰	۱۳۵/۵۰۲	۸۶/۹۶۶	۱۲۲/۸۷۸	۷۹/۷۵۸	۱۵۰/۶۶۴	۱۰۴/۸۳۶	AIC
۱۲۱/۲۶۵	۸۵/۱۹۰	۱۴۲/۵۰۸	۹۸/۱۷۵	۱۲۹/۸۸۴	۹۰/۹۶۷	۱۵۶/۲۶۸	۱۱۴/۶۴۴	BIC

حال با توجه به مطالب بیان شده، نتایج مربوط به برآورد پارامترهای مدل رگرسیون تکه‌ای در جدول ۴ آورده شده است. نتایج برازش مدل‌های مختلف شامل مقدار برآورد شده‌ی پارامترها همراه معیار RSS در جدول ۴ آورده شده است. با توجه به مقادیر AIC و BIC در جدول ۳ دیده می‌شود که مدل CN-BP نسبت به سایر مدل‌ها به داده‌ها بهتر برازش داده شده است. نتایج به دست آمده در مورد تمامی مدل‌ها (به جدول ۴ مراجعه کنید) نشان می‌دهد که خطر ابتلاء جنین به سندروم داون به طور چشمگیری پس از ۲۸/۵ سالگی سن مادر افزایش می‌یابد. حال با توجه به معیار BIC بهترین مدل CN-BP است، که در آن برآورد شیب خط مدل رگرسیون تکه‌ای از ۰/۰۰۹- به ۰/۳۰۸+ می‌رسد. این تغییر شیب ناگهانی در سن ۲۸/۷ سالگی مادران باردار است. همچنین نتیجه دیگر این است که خطر ابتلاء به سندروم داون به طور قابل توجهی با سن مادر قبل از ۲۸/۷ سالگی در مدل CN-BP مرتبط نیست، زیرا شیب خط رگرسیونی در این بازه زیاد نمی‌باشد.

مثال ۲. تجزیه و تحلیل روند تغییرات آب و هوا به ویژه برای بررسی مشکل گرمایش جهانی بسیار مهم است. چندین تحقیق (کارل و همکاران، ۲۰۰۰؛ من، ۲۰۰۵) به روش رگرسیونی برای مطالعه تغییرات روند دما صورت گرفته

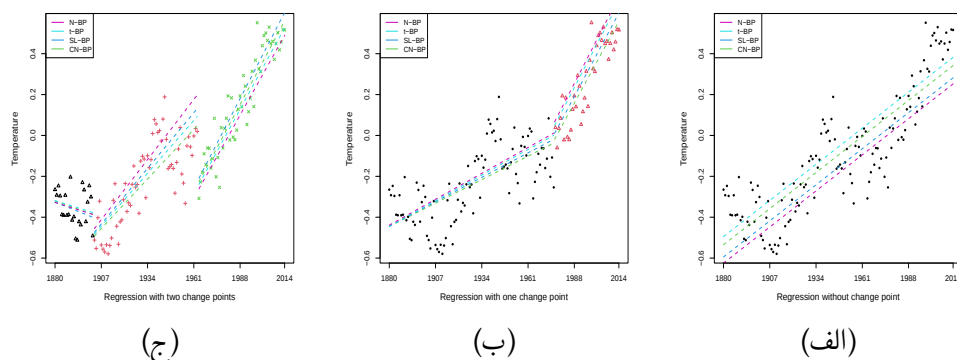
جدول ۴. برآورد پارامترها و برآورد مقدار نقطه تغییر برای داده‌های سندروم داون .

پارامتر	N-BP	t-BP	SL-BP	CN-BP
β_{11}	-۶/۷۰۱	-۶/۷۰۱	-۶/۷۰۱	-۶/۷۰۱
β_{12}	-۰/۰۰۲	-۰/۰۱۱	-۰/۰۱۵	-۰/۰۰۹
β_{21}	-۱۴/۷۹۷	-۱۴/۰۱۸	-۱۵/۱۲۸	-۱۴/۵۲۹
β_{22}	۰/۲۴۷	۰/۲۵۵	۰/۲۱۵	۰/۳۰۸
σ_1	۰/۹۲۷	۱/۰۸۲	۱/۲۳۴	۱/۰۱۸
σ_2	۱/۲۸۶	۱/۱۷۳	۰/۹۹۶	۱/۱۹۷
ν	—	۴۸۲۳	۲۹۵۷	(۰/۴۲۹, ۱/۳۹۸)
BP	۱۲/۴۱	۱۲/۰۲	۱۲/۲۴	۱۲/۱۵
RSS	۲/۵۵۱	۲/۱۵۸	۲/۳۸۷	۲/۲۰۱

است و بسیاری از این محققان گزارش کردند که قابل توجه‌ترین نقطه تغییر در گرمایش زمین در اواسط دهه ۱۹۷۰ بوده است. در این مثال داده‌های مربوط به گرمایش زمین را که توسط مرکز هادلی^۱ جمع‌آوری شده‌اند را مورد بررسی قرار می‌دهیم. این داده‌ها در بعضی نقاط به صورت گمشده بودند که از آنها صرف نظر شده است. این دماها در بازه زمانی ۱۹۸۰-۲۰۱۴ مورد بررسی قرار گرفته‌اند که میانگین این دماها در سال به عنوان مبنا در نظر گرفته شده است. شکل ۳ نتایج به دست آمده توسط مدل‌های مختلف SMN-BP را با استفاده از مدل‌هایی شامل ۰، ۱، ۲ و نقطه تغییر نشان می‌دهد. برازش داده‌ها با مدل SMN-BP، تغییر قابل توجهی را در سال ۱۹۷۶ نشان می‌دهد که با نتایج مطالعات گذشته مطابقت دارد. با توجه به جدول ۵ در تمامی مدل‌ها، مدلی با ۲ نقطه تغییر به دلیل اینکه BIC کمتری دارد بهتر برازش شده است. حال در بین مدل‌هایی با ۲ نقطه تغییر مدل CN-BP را ترجیح می‌دهیم زیرا هم BIC کمتری دارد و علاوه بر آن مقدار RSS کمتری دارد. این موضوع نشان می‌دهد که مدل بیان شده در برابر سایر مدل‌های رقیب از مجموع مربع باقیمانده کمتری برخوردار است. از دلایل دیگر رد مدل با سه نقطه تغییر می‌توان به این موضوع اشاره کرد که در دسته سوم فقط ۱۰ داده قرار گرفته است که این موضوع نیز می‌تواند باعث اریبی برآوردها در این دسته شده باشد. در مدل با ۲ نقطه تغییر قبل از سال ۱۹۰۲، هیچ ارتباط معنی‌داری بین گرمایش جهانی و زمان وجود نداشت زیرا برآورد شیب خط $\beta_{11} = ۰/۰۰۴۳$ است، که این موضوع می‌تواند به دلیل پیشرفت نچندان زیاد اکثر کشورها در صنعتی شدن باشد. در بازه زمانی ۱۹۰۲ تا ۱۹۶۴ برآورد شیب خط $\beta_{21} = ۰/۰۲۱$ است که یک شیب ملایم برای افزایش دمای جهانی است که از دلایل آن صنعتی شدن برخی از کشورها و تولید گازهای گلخانه‌ای زیادی در سطح جهان است که موجب افزایش دمای کره زمین شده است. اما در بازه زمانی ۱۹۶۳ تا ۲۰۱۴ این شیب به شدت افزایش پیدا کرده است که مقدار آن $\beta_{31} = ۰/۰۴۵$ است که دلیل آن می‌تواند افزایش خارج از قاعده مراکز صنعتی در تمامی کشورهای دنیا باشد. این مراکز صنعتی با تولید گازهای گلخانه‌ای زیاد باعث بالا رفتن بی‌رویه دمای سطح زمین شده‌اند. برای تکمیل و صحت نتایج بدست آمده آزمون اندرسون-دارلینگ برای نرمال بودن داده‌ها انجام شده است. با توجه به دو نقطه شکست و سه دسته داده، p -مقدار برای آزمون اندرسون-دارلینگ به ترتیب برای سه دسته برابر ۰/۴۲۵، ۰/۳۲۸ و ۰/۷۸۲ است. این مقادیر نشان دهنده نرمال بودن داده‌ها هستند که

¹The Met Office Hadley Center

فرض استفاده از توزیع SMN را توجیه می‌کنند.



شکل ۳. کاربرد SMN-BP برای داده‌های گرمایش سطح زمین.

جدول ۵. مقادیر بدست آمده برای شاخص‌های BIC و RSS در داده‌های مربوط به گرمایش زمین.

تعداد نقاط تغییر	شاخص	N-BP	t-BP	SL-BP	CN-BP
۰	BIC	۱۲۸۲۸۴۶	۹۹۸۵۲۳	۱۰۵۹۴۰۱	۹۸۹۴۲۷
	RSS	۳۰۴۹	۲۸۶۰	۲۹۶۸	۲۵۶۰
۱	BIC	۱۱۰۹۵۰۶	۹۵۶۴۸۶	۱۰۰۲۴۸۰	۹۲۱۷۵۰
	RSS	۲۸۸۵	۲۴۷۶	۲۶۸۵	۲۳۰۸
۲	BIC	۱۰۹۶۶۰۴	۹۰۹۹۶۷	۹۷۸۱۷۵	۸۹۰۷۰۹
	RSS	۲۴۵۲	۲۰۱۷	۲۲۲۵	۱۹۸۵

۶ بحث و نتیجه‌گیری

در این مقاله، با استفاده از مدل‌های آمیخته متناهی روشی را برای حل مسائل رگرسیون تکه‌ای با خطاهایی از توزیع SMN پیشنهاد کردیم. پیدا کردن نقاط تغییر فرآیندی مشابه دسته‌بندی داده‌ها به گروه‌هایی از مشاهدات مشابه در مدل‌های آمیخته است. با در نظر گرفتن نیرومندی مدل‌های آمیخته در خوشه‌بندی داده‌ها، یک مدل SMN-BP برای تخمین نقاط تغییر و پارامترهای رگرسیونی ایجاد کردیم. همچنین با استفاده از شبیه‌سازی نشان داده شد که شناسایی نقاط پرت با استفاده از برخی حالت‌های خاص مدل SMN-BP به خوبی صورت می‌پذیرد. علاوه بر این چندین بررسی عددی با مجموعه داده‌های واقعی و شبیه‌سازی شده کارایی و برتری مدل پیشنهادی را نشان می‌دهند. در نتیجه، مدل پیشنهادی SMN-BP معمولاً برآوردهای دقیقی از نقاط تغییر و پارامترهای رگرسیون را ارائه می‌دهد.

از این رو، مدل SMN-BP یک رویکرد مطلوب برای حل مشکلات رگرسیون تکه‌ای است. رگرسیون تکه‌ای را می‌توان در مسائل مختلفی که در آنها نقطه تغییر مهم است مانند کنترل کیفیت، پزشکی، اقتصاد و پردازش سیگنال اعمال کرد. در بیشتر موارد این داده‌ها دارای کجی و چولگی هستند. در تحقیقات آینده، سعی می‌کنیم از مدل‌ها دارای چولگی و کجی برای برآورد نقطه تغییر قویتر برای مجموعه‌های داده با نقاط پرت استفاده کنیم. همچنین قصد داریم از خوشه‌بندی فازی برای برآورد پارامترهای مدل رگرسیون تکه‌ای و نقاط تغییر با داده‌هایی دارای کجی استفاده کنیم.

تقدیر و تشکر

نویسنده بر خود لازم می‌داند؛ مراتب سپاس خود را از سردبیر محترم، داوران گرامی و ویراستار مجله که با نظرات ارزشمند خود موجبات ارتقای این مقاله را فراهم کردند، اعلام نماید.

مراجع

جلیلی، م.، بشیری، م.، منطقی، م. و توفیق، ع. (۱۳۹۵)، توسعه رویکردی مبتنی بر رگرسیون تکه‌ای برای پایش پروفایل‌های خطی چندگانه با اثرات متقابل در فاز، مهندسی و مدیریت کیفیت، ۶، ۲۳۷-۲۴۹.

جعفری، ف.، و گلعلی‌زاده، م. (۱۴۰۲)، مدل‌بندی رگرسیونی چندکی آمیخته تاوانیده دوگانه از طریق رویکرد درستمایی، مجله علوم آماری، ۱۷، ۳۸۹-۴۰۵.

Birnbaum, Z. W., and Saunders, S. C. (1969), A New Family of Life Distributions, *Journal of Applied Probability*, **6** (2), 319-327.

Chang, S. T. and Lu, K. P. (2016), Change-Point Detection for Shifts in Control Charts Using EM Change- Point Algorithms, *Quality and Reliability Engineering International*, **32**, 889-900.

Chang, S. T, Lu, K. P. and Yang, M. S. (2015), Fuzzy Change-Point Algorithms for Regression Models, *IEEE Transactions on Fuzzy Systems*, **23** (6), 2343-2357.

Chen, C. W. S., Chan, J. S., Gerlach K. R., and Hsieh, W. Y. L. (2011), A Comparison of Estimators for Regression Models with Change Points, *Statistics and Computing*, **21**, 395-414.

- Ciuperca, G. (2009), The M-Estimation in a Multi-Phase Random Nonlinear Model, *Statistics & Probability Letters*, **75**, 573–580.
- Ciuperca, G. (2011), A General Criterion to Determine the Number of Change-Points, *Statistics & Probability Letters*, **81**, 1267–1275.
- Ciuperca, G. (2004), Maximum Likelihood Estimator in a Two-Phase Nonlinear Regression Model, *Statistics & Decisions*, **22**, 335–349.
- Ciuperca, G. and Dapzol, N. (2008), Maximum Likelihood Estimator in a Multi-Phase Random Regression Model, *Statistics*, **42**, 363–381
- Dempster, A. P., Laird N. M., and Rubin D. B. (1977), Maximum Likelihood From Incomplete Data via the EM-Algorithm, *Journal of the Royal Statistical Society. Series B (methodological)*, **39**, 1–38.
- Fearnhead, P. (2006), Exact and Efficient Bayesian Inference for Multiple Change-Point Problems, *Statistics and Computing*, **16**, 203–213.
- Julious, S. A. (2001). Inference and Estimation in a Change Point Regression Problem, *Journal of the Royal Statistical Society Series D*, **50**, 51–61.
- Karl, T.R., Knight, R.W. and Baker, B. (2000), The Record Breaking Global Temperatures of 1997 and 1998: Evidence for an Increase in the Rate of Global Warming?, *Geophysical Research Letters*, **27** (5), 719–722
- Keshavarz, M. and Huang, B. (2014a), Bayesian and Expectation Maximization Methods for Multivariate Change-Point Detection, *Computers & Chemical Engineering*, **60**, 339–353.
- Keshavarz M. and Huang, B. (2014b), Expectation Maximization Method for Multivariate Change-Point Detection in Presence of Unknown and Changing Covariance, *Computers & Chemical Engineering*, **69**, 128–146.
- Loschi, R. H., Pontel, J. G. and Cruz F. R. B. (2010), Multiple Change-Point Analysis for Linear Regression models, *Chilean Journal of Statistics*, **1**, 93–112.

- Lu, K. P., and Chang, S. T. (2021), Robust Algorithms for Change-Point Regressions Using the t-Distribution, *Mathematics*, **9**, 2394.
- Menne, J. M. (2005), Abrupt Global Temperature Change and the Instrumental Record, *In: Record, 18th Conference on Climate Variability and Change*.
- Muggeo, V. M. R. (2008), Segmented: an R package to Fit Regression Models with Broken-Line Relationships, *R News*, **8**, 20–25.
- von Ottenbreit, M., and De Bin, R. (2024), Automatic Piecewise Linear Regression, *Computational Statistics*, 1-41.
- Yang, F. (2014), Robust Mean Change-Point Detecting Through Laplace Linear Regression Using EM Algorithm, *Journal of Applied Mathematics*, 2014 (1): 856350.
- Yildirim S., Singh S.S., and Doucet A. (2014), An Online Expectation-Maximization Algorithm for Change-Point Models, *Journal of Computational and Graphical Statistics*, **22**(4), 906–26.