



Efficiency of shrinkage pretest estimator for the intercept parameter in simple linear regression model

Z. Mahdizadeh¹, M. Naghizadeh Qomi² *

1. Department of Statistics, University of Mazandaran, Babolsar, Iran.

2. Corresponding Author, Department of Statistics, University of Mazandaran, Babolsar, Iran.

* E-mail: m.naghizadeh@umz.ac.ir

Article Info	ABSTRACT
Article type: Research Article	Introduction Traditionally the classical estimators of unknown parameters of the linear regression model are based exclusively on the sample information, like the maximum likelihood method. Sometimes, in practice the researcher has some prior information about the unknown intercept parameter as a guess that is called non-sample prior information. In this case, the shrinkage and shrinkage preliminary test estimators are introduced by a linear combination of sample and non-sample prior information. Now, if the non-sample prior information is available according to the experimenter's experiences or the results of previous experiments, it can be expressed in the form of a preliminary hypothesis test and the desired parameter can be estimated using the shrinkage pretest estimator. The results show that the closer the guessed value to the real parameter, the better shrinkage pretest estimator performs compared to the maximum likelihood estimator.
Article history: Received: 29 November 2020 Received in revised form: 22 November 2021 Accepted: 10 December 2021 Published online: 3 December 2023	In the past years, in the regression estimation problem, the behavior of the shrinkage and shrinkage pretest estimators under the squared error loss (SEL) and LINEX loss functions was investigated. The squared error loss (SEL) function is popularly used for estimating the unknown parameter in decision theory because of its mathematical and interpretational convenience. Due to the symmetric nature, the use of SEL function may not be appropriate, when positive and negative errors have different consequences. The SEL and LINEX loss functions are symmetric and asymmetric loss functions, respectively, but both are unbounded. Also, The SEL and LINEX loss functions have an infinite maximum value which isn't always appropriate. Sometimes in practice it is necessary to use bounded loss functions to estimate parameters. So we used the reflected normal loss function to estimate parameters, which is the bounded loss function.
Keywords: Non-sample Prior information, Shrinkage pretest estimator, Reflected normal loss function.	Material and methods In this article, the shrinkage and shrinkage pretest estimators are introduced for the intercept parameter of the simple linear regression according to the non-sample prior information and their risk functions and relative efficiencies are derived under the reflected normal loss function. The behavior of shrinkage pretest estimator is compared with respect to the shrinkage and

maximum likelihood estimator using graphical method. The intervals where the shrinkage pretest estimator has less risk compared to the maximum likelihood estimator are presented. The optimum value of the significant level of test is determined using the max-min method. Besides, several methods of finding shrinkage coefficient of the shrinkage pretest estimators are proposed. Then, the application of the proposed estimators is shown using a real data set.

Results and discussion

In comparison between the desired estimators, our findings show that the shrinkage and shrinkage pretest estimator are better than the maximum likelihood estimator when non-sample prior information is close to the real value. Also, the shrinkage pretest estimator with smaller level of significance has higher efficiency. An important issue for the shrinkage pretest estimator is the proper selection of the shrinkage coefficient. So, we applied several methods for finding the shrinkage coefficient to obtain shrinkage pretest estimator.

Conclusion

According to the reported results, the shrinkage pretest estimator has smaller risk than the maximum likelihood estimator in neighborhood of null hypothesis. Therefore, in practice, if the researcher can have prior information of the unknown intercept parameter from previous knowledge or experience that is not far from the real value, the shrinkage pretest estimator would be the best choice as intercept estimator.

The effective intervals where the shrinkage pretest estimator has less risk compared to the maximum likelihood estimator was presented. By increasing the significant level of the test and the sample size, the length of the effective interval decreases. The optimal value of the significant level of the test was obtained using the max-min method. Then, the application of the proposed estimators is shown using a real data set.

How to cite: Z. Mahdizadeh,, M. Naghizadeh Qomi, (2023). Efficiency of shrinkage pretest estimator for the intercept parameter in simple linear regression model. *Mathematical Researches*, 9 (2), 125 – 145.



© The Author(s).

Publisher: Kharazmi University



کارایی برآوردگر پیش‌آزمون انقباضی پارامتر عرض از مبدأ در مدل رگرسیونی خطی ساده

زهره مهدی‌زاده^۱، مهران نقی‌زاده قمی^{۲*}

۱. گروه آمار دانشگاه مازندران، ایران.

۲. گروه آمار دانشگاه مازندران، ایران. رایانامه: m.naghizadeh@umz.ac.ir

اطلاعات مقاله	چکیده
نوع مقاله: مقاله پژوهشی	معمولاً برای برآورد پارامترهای یک مدل رگرسیونی از روش‌های مرسوم برآوردیابی مانند روش ماکسیمم درست‌نمایی استفاده می‌شود. گاهی اوقات محقق دارای اطلاع پیشین درمورد پارامتر عرض از مبدأ به صورت یک حدس است که به آن اطلاع پیشین غیرنمونه‌ای می‌گویند. در این مقاله برآوردگرهای انقباضی و پیش‌آزمون انقباضی برای پارامتر عرض از مبدأ با توجه به اطلاع پیشین غیرنمونه‌ای معرفی و کارایی نسبی آن‌ها تحت تابع زیان نرمال بازتابیده مورد بررسی قرار می‌گیرد. رفتار برآوردگر پیش‌آزمون انقباضی در مقایسه با برآوردگرهای انقباضی و ماکسیمم درست‌نمایی به صورت گرافیکی ارزیابی می‌شود. بازه‌هایی که برآوردگر پیش‌آزمون انقباضی دارای مخاطره کمتری نسبت به برآوردگر ماکسیمم درست‌نمایی است ارائه می‌شود. نتایج نشان می‌دهد هرچه مقدار حدس زده شده به پارامتر واقعی نزدیکتر باشد برآوردگر پیش‌آزمون انقباضی عملکرد بهتری نسبت به برآوردگر ماکسیمم درست‌نمایی دارد. همچنین با استفاده از روش ماکس‌مین مقدار بهینه سطح معنی‌داری آزمون تعیین می‌شود. سپس با استفاده از یک مجموعه داده واقعی، کاربرد برآوردگرهای پیشنهادی نشان داده می‌شود.
تاریخ دریافت: 1399/9/9 تاریخ بازنگری: 1400/9/1 تاریخ پذیرش: 1400/9/10 تاریخ انتشار: 1402/9/12	
واژه‌های کلیدی: اطلاع پیشین، برآوردگر پیش‌آزمون انقباضی، تابع زیان نرمال بازتابیده	

استناد: زهره مهدی‌زاده، مهران نقی‌زاده قمی (1402). کارایی برآوردگر پیش‌آزمون انقباضی پارامتر عرض از مبدأ در مدل رگرسیونی خطی ساده. پژوهش‌های ریاضی، 9 (2)، 125 – 145.



© نویسندگان.

ناشر: دانشگاه خوارزمی

مقدمه

در روش‌های کلاسیک آمار، براساس اطلاعات موجود در نمونه و با روش‌های معمول مانند روش ماکسیمم درستنمایی به برآوردیابی پارامتر نامعلوم می‌پردازند. گاهی در عمل، محقق دارای اطلاعاتی درباره پارامتر نامعلوم به صورت یک حدس یا گمان است که اطلاعات غیرنمونه‌ای نامیده می‌شوند. در این حالت، برآوردگر پیش‌آزمون انقباضی^۱ با ترکیب خطی اطلاعات غیرنمونه‌ای و اطلاعات موجود در نمونه معرفی می‌شود که دارای کارایی نسبی بیشتری نسبت به برآوردگرهای معمول می‌باشد [1].

در یک مدل رگرسیونی خطی ساده برای برآورد ضرایب رگرسیونی معمولاً از روش‌های کلاسیک با استفاده از اطلاعات نمونه ای مانند ماکسیمم درستنمایی استفاده می‌شود. حال اگر اطلاعات پیشین غیرنمونه‌ای با توجه به تجربیات آزمایشگر یا نتایج آزمایش‌های قبلی در اختیار باشد می‌توان این اطلاعات غیرنمونه‌ای را در قالب یک آزمون فرض مقدماتی بیان کرد و به کمک برآوردگر پیش‌آزمون انقباضی به برآورد پارامتر مورد نظر پرداخت. در مساله برآورد رگرسیونی، [5] و [13] رفتار تعدادی از برآوردگرها از جمله برآوردگرهای انقباضی و پیش‌آزمون را برای پارامتر عرض از مبدأ مدل رگرسیونی خطی ساده تحت تابع زیان توان دوم خطا بررسی کردند. [2] کارایی نسبی برآوردگر پیش‌آزمون را تحت تابع زیان لاینکس^۲ مورد مطالعه قرار دادند. تابع زیان توان دوم خطا به علت راحتی در تحلیل و محاسبات ریاضی برای برآورد پارامتر نامعلوم در نظریه تصمیم استفاده می‌شود. از طرفی به دلیل متقارن بودن تابع زیان توان دوم خطا ممکن است، تابع زیان مناسبی نباشد زیرا به بیش برآوردی و کم برآوردی جریمه‌های یکسانی نسبت می‌دهد. تابع زیان‌های توان دوم خطا و لاینکس غیرکراندار هستند. گاهی در عمل زیان ناشی از برآورد پارامترها باید کراندار باشد لذا لازم است که از توابع زیان کراندار استفاده کرد. در این مقاله، به بررسی کارایی نسبی برآوردگر پیش‌آزمون انقباضی پارامتر عرض از مبدأ در مدل رگرسیونی خطی ساده تحت تابع زیان نرمال بازتابیده می‌پردازیم. تابع زیان نرمال بازتابیده که توسط [9] معرفی شد، دارای فرمی به صورت

$$L(D) = k(1 - e^{-\frac{D^2}{2I^2}}) \quad (1)$$

است که در آن $D = d - q$ ، d برآوردگر پارامتر q ، $k > 0$ ماکسیمم تابع زیان و $I > 0$ پارامتر شکل است. تابع زیان (1) کراندار و اساساً برگردانی از تابع چگالی احتمال نرمال از بالا می‌باشد و بر همین اساس تابع زیان نرمال بازتابیده نامیده می‌شود. شکل 1 تابع زیان نرمال بازتابیده را به ازای $k=1$ و مقادیر مختلف $I=1, 5, 10$ نشان می‌دهد. ویژگی‌های تابع زیان (1) عبارتند از: $L(D)$

(۱) تابعی کراندار است.

(2) $L(D)$ به صورت نامتناهی مشتق‌پذیر و در نتیجه پیوسته است.

¹ Shrinkage pretest estimator

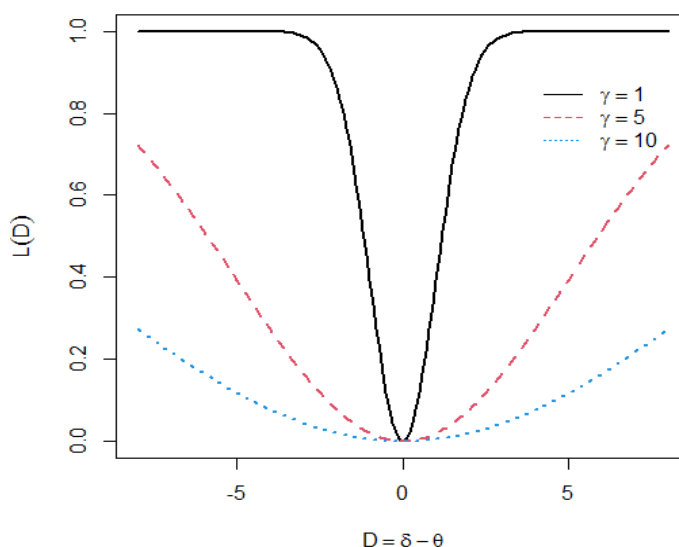
² LINEX loss function

(3) $L(D)$ دارای مینیمم یکتای صفر به ازای $D=0$ است و در بازه $(-\infty, 0)$ یک تابع نزولی و در بازه $(0, \infty)$ یک تابع صعودی است. همچنین $\lim_{D \rightarrow \pm\infty} L(D) = k$.

(4) تابع زیان (1) حالت خاصی از فاصله هلینگر¹ به ازای $k=1$ و $g=2s$ به فرم

$$H^2(f(x|q), f(x|d)) = \frac{1}{2} E_q \left\{ \frac{\sqrt{f(X|q)} - \sqrt{f(X|d)}}{\sqrt{f(X|q)}} \right\}^2,$$

است که در آن $f(x|q)$ و $f(x|d)$ به ترتیب تابع‌های چگالی احتمال توزیع نرمال $N(q, s^2)$ و $N(d, s^2)$ است. تابع زیان (1) توسط [8]، [11] و [7] در مسائل مختلف برآوردیابی مورد استفاده قرار گرفت. از آنجایی که مقدار $k > 0$ هیچ تأثیری در نتایج ندارد بنابراین در این مقاله بدون از دست دادن کلیت مسئله $k=1$ در نظر گرفته می‌شود.



شکل 1: تابع زیان نرمال بازتابیده برای $k=1$.

در این مقاله، در بخش 2، مدل رگرسیونی خطی ساده، برآوردهای انقباضی و پیش‌آزمون انقباضی پارامتر عرض از مبدا معرفی می‌شوند. تابع مخاطره برآوردهای ماکسیمم درست‌نمایی، انقباضی و پیش‌آزمون انقباضی پارامتر عرض از مبدا مدل رگرسیونی در بخش 3 محاسبه می‌شوند. در بخش 4 کارایی نسبی این سه برآوردگر به صورت گرافیکی مورد بررسی قرار می‌گیرد.

¹ Hellinger distance

گیرند. بازه‌های موثر که به ازای آن برآوردگر پیش‌آزمون انقباضی نسبت به برآوردگر ماکسیمم درست‌نمایی دارای مخاطره کمتری است ارائه می‌شود. همچنین مقدار بهینه سطح آزمون معنی‌داری با استفاده از روش ماکس‌مین به دست می‌آید. در بخش ۵ یک مجموعه داده واقعی برای کاربرد برآوردهای انقباضی و پیش‌آزمون انقباضی ارائه خواهد شد. در بخش پایانی به بحث و نتیجه‌گیری پرداخته می‌شود.

۲ معرفی مدل و برآوردگر پیش‌آزمون انقباضی

فرض کنید یک نمونه n تایی از مدل رگرسیونی خطی ساده به صورت

$$Y_i = b_0 + b_1 x_i + e_i \quad i = 1, 2, \dots, n$$

در اختیار داریم که در آن x_i متغیر پیشگو، Y_i متغیر تصادفی پاسخ، b_0 پارامتر عرض از مبدأ، b_1 پارامتر شیب و $e_1, \dots, e_n$ یک نمونه تصادفی n تایی از $N(0, \sigma^2)$ می‌باشد.

مطابق با [6] برآوردهای ماکسیمم درست‌نمایی پارامتر عرض از مبدأ b_0 و شیب b_1 در مدل رگرسیون خطی ساده به صورت $b_0^0 = \bar{y} - b_1^0 \bar{x}$ و $b_1^0 = S_{xy}^{-1} S_{xx}$ هستند که در آن $S_{xx} = \sum_{i=1}^n (x_i - \bar{x})^2$ ، $S_{xy} = \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$.

b_0^0 به عنوان برآوردگر دارای توزیع نرمال با میانگین b_0^0 و واریانس $S^2 H$ است و $\bar{y} = n^{-1} \sum_{i=1}^n y_i$ و $\bar{x} = n^{-1} \sum_{i=1}^n x_i$ است. که در آن $H = \frac{1}{n} + \frac{\bar{x}^2}{S_{xx}}$ می‌باشد، به عبارت دیگر

$$Z = \frac{b_0^0 - b_0}{s \sqrt{H}} : N(0, 1). \quad (2)$$

فرض می‌کنیم محقق دارای اطلاع پیشین (باور یا عقیده) غیرنمونه‌ای در خصوص پارامتر b_0 است که با b_{00} نشان می‌دهیم که از طریق تجربه یا مطالعات گذشته قابل دسترس باشد. بر اساس [10] برآوردگر انقباضی با ترکیب خطی از مقدار حدس زده شده و برآوردگر نمونه‌ای به صورت

$$\hat{b}_0^{SE} = d b_0^0 + (1 - d) b_{00}$$

تعریف می‌شود، که در آن $0 \leq d \leq 1$ ضریب انقباضی است و توسط محقق با توجه به عقیده او نسبت به مقدار حدس زده شده تعیین می‌شود. اگر $d = 0$ تنها از مقدار حدسی و اگر $d = 1$ از برآوردگر نمونه‌ای استفاده می‌شود. اطلاع پیشین غیرنمونه‌ای می‌تواند به کمک فرض صفر $H_0 : b_0 = b_{00}$ در مقابل $H_1 : b_0 \neq b_{00}$ آزمون شود. برآوردگر پیش‌آزمون انقباضی پارامتر b_0 از ترکیب اطلاع پیشین غیرنمونه‌ای (b_{00}) و اطلاعات نمونه‌ای (b_0^0) به صورت

$$\begin{aligned}\hat{b}_0^{SP} &= \begin{cases} db_0^0 + (1-d)b_{00} & F_{1n} < F_{1-a} \\ b_0^0 & F_{1n} > F_{1-a} \end{cases} \\ &= b_0^0 - (1-d)(b_0^0 - b_{00})I(A)\end{aligned}$$

تعریف می‌شود، که در آن $I(A)$ تابع نشانگر مجموعه $A = \{F_{1n} < F_{1-a}\}$ ، F_{1n} آماره آزمون به صورت $F_{1n} = \frac{(b_0^0 - b_{00})^2}{S_H^2}$ ، با $S_n^2 = \frac{1}{(n-2)} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2$ که برآوردگری ناریب برای s^2 است و در آن $\hat{Y}_i = b_0^0 + b_1^0 x_i$ می‌باشد و مقدار چندک مرتبه $F_{1-a} = F_{1n}(1-a)$ درصد توزیع فیشر با درجات آزادی به ترتیب 1 و $\nu = n-2$ است.

3 محاسبه تابع مخاطره b_0^0 ، b_0^{SE} و b_0^{SP}

در این بخش، تابع مخاطره b_0^0 ، b_0^{SE} و b_0^{SP} تحت تابع زیان نرمال بازتابیده محاسبه می‌شوند. برای سهولت در محاسبه تابع مخاطره از لم‌های زیر استفاده می‌کنیم.

لم 1. اگر $Z : N(0,1)$ و $S : c_k^2$ مستقل از هم باشند، آنگاه در صورت وجود امید ریاضی $\hat{A} \otimes (0, \infty) : f$ و برای هر $c < 0.5$ رابطه زیر را داریم

$$E[e^{cZ^2} f(Z, S)] = \frac{1}{\sqrt{1-2c}} E\left[\frac{Z}{\sqrt{1-2c}} f\left(\frac{Z}{\sqrt{1-2c}}, S\right)\right]$$

برهان: با استفاده از خاصیت امید ریاضی داریم

$$\begin{aligned}E[e^{cZ^2} f(Z, S)] &= E[E[e^{cZ^2} f(Z, S) | S]] \\ &= E\left[\int_{-\infty}^{+\infty} e^{cz^2} f(z, S) \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} dz\right] \\ &= E\left[\int_{-\infty}^{+\infty} f(z, S) \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}(1-2c)} dz\right] \\ &= E\left[\int_{-\infty}^{+\infty} f(z, S) \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(z\sqrt{1-2c})^2} dz\right].\end{aligned}$$

داریم $J = \frac{1}{\sqrt{1-2c}}$ و ژاکوبی تبدیل $u = z\sqrt{1-2c}$ با تغییر متغیر

$$\begin{aligned}E[e^{cZ^2} f(Z, S)] &= \frac{1}{\sqrt{1-2c}} E\left[\int_{-\infty}^{+\infty} f\left(\frac{u}{\sqrt{1-2c}}, S\right) \frac{1}{\sqrt{2\pi}} e^{-\frac{u^2}{2}} du\right] \\ &= \frac{1}{\sqrt{1-2c}} E\left[f\left(\frac{Z}{\sqrt{1-2c}}, S\right)\right].\end{aligned} \quad W$$

لم 2. [3]. اگر $Z : N(0,1)$ و $S : C_k^2$ مستقل از هم باشند، آنگاه در صورت وجود امید ریاضی $\hat{A} \in \mathcal{A}(0, \infty)$: f و برای هر $\hat{A} \in \mathcal{A}$ رابطه زیر را داریم

$$E[e^{cZ} f(Z, S)] = e^{\frac{c^2}{2}} E[f(Z + c, S)].$$

برای محاسبه تابع مخاطره $b_0^{\%}$ تحت تابع زیان (1) با توجه به هم توزیع بودن $(b_0^{\%} - b_0)^2$ و $Z^2 S^2 H$ و با استفاده از لم 1 داریم

$$\begin{aligned} R(b_0^{\%}, b_0) &= E[L(b_0^{\%}, b_0)] = 1 - E[e^{-\frac{1}{2g^2}(b_0^{\%} - b_0)^2}] \\ &= 1 - E[e^{-\frac{1}{2g^2} Z^2 S^2 H}] \\ &= 1 - E[e^{-aZ^2}] = 1 - \frac{1}{\sqrt{1+2a}} = 1 - (1+2a)^{-\frac{1}{2}}, \end{aligned}$$

که در آن $a = \frac{S^2 H}{2g^2}$. تابع مخاطره $\hat{b}_0^{SP}$ تحت تابع زیان نرمال بازتابیده برابر است با

$$\begin{aligned} R(\hat{b}_0^{SP}, b_0) &= E[1 - e^{-\frac{1}{2g^2}(\hat{b}_0^{SP} - b_0)^2}] \\ &= 1 - E[e^{-\frac{1}{2g^2}(d\hat{b}_0^{\%} + (1-d)b_{00} - b_0)^2} I(F_{1n} < F_{1-a})] \\ &\quad - E[e^{-\frac{1}{2g^2}(b_0^{\%} - b_0)^2} I(F_{1n} > F_{1-a})]. \end{aligned} \quad (3)$$

آماره آزمون F_{1n} با توجه به رابطه (2) به صورت

$$\begin{aligned} F_{1n} &= \frac{[H^{-\frac{1}{2}}(b_0^{\%} - b_{00})S^{-1}]^2}{S/n} \\ &= \frac{[H^{-\frac{1}{2}}(b_0^{\%} - b_0)S^{-1} + H^{-\frac{1}{2}}(b_0 - b_{00})S^{-1}]^2}{S/n} \\ &= \frac{(Z+D)^2}{S/n} \end{aligned}$$

است که در آن $D = H^{-\frac{1}{2}}(b_0 - b_{00})S^{-1}$ و $S = \frac{nS_n^2}{2}$ است. بنابراین اولین امید ریاضی در رابطه (3) به صورت

$$\begin{aligned} E[e^{-\frac{1}{2g^2}(d\hat{b}_0^{\%} + (1-d)b_{00} - b_0)^2} I(F_{1n} < F_{1-a})] &= E[e^{-\frac{1}{2g^2}(ds\sqrt{H}Z - (1-d)DS\sqrt{H})^2} I(F_{1n} < F_{1-a})] \\ &= e^{-(1-d)^2 a D^2} E[e^{c_1 Z^2} f(Z, S)] \end{aligned}$$

محاسبه می‌شود، که در آن $c_1 = -ad^2$ و $f(Z, S) = e^{2ad(1-d)DZ} I(\frac{(Z+D)^2}{S/n} < F_{1-a})$ با استفاده از لم 1 داریم

$$e^{-(1-d)^2 a D^2} E[e^{c_1 Z^2} f(Z, S)] = \frac{e^{-(1-d)^2 a D^2}}{\sqrt{1+2ad^2}} E[e^{\frac{2ad(1-d)D}{\sqrt{1+2ad^2}} Z} I(\frac{(\frac{Z}{\sqrt{1+2ad^2}} + D)^2}{S/n} < F_{1-a})]$$

$$= \frac{e^{-(1-d)^2 a D^2}}{\sqrt{1+2ad^2}} E[f^*(Z, S) e^{c_2 Z}],$$

که در آن $c_2 = \frac{2ad(1-d)D}{\sqrt{1+2ad^2}}$ و $f^*(Z, S) = I(\frac{(Z + \sqrt{1+2ad^2}D)^2}{S/n} < (1+2ad^2)F_{1-a})$ در نتیجه با استفاده از لم 2 داریم

$$E[e^{-\frac{1}{2g^2}(\hat{b}_0^{SE} - b_0)^2} I(F_{1n} < F_{1-a})] = \frac{e^{-(1-d)^2 a D^2}}{\sqrt{1+2ad^2}} e^{\frac{c_2^2}{2}} E[I(\frac{(Z+c_2+\sqrt{1+2ad^2}D)^2}{S/n} < (1+2ad^2)F_{1-a})]$$

$$= \frac{e^{-(1-d)^2 a D^2 + \frac{c_2^2}{2}}}{\sqrt{1+2ad^2}} P(\frac{(Z+c_2+\sqrt{1+2ad^2}D)^2}{S/n} < (1+2ad^2)F_{1-a})$$

$$= \frac{e^{-(1-d)^2 a D^2 + \frac{c_2^2}{2}}}{\sqrt{1+2ad^2}} G_{1n}((1+2ad^2)F_{1-a}; (c_2 + \sqrt{1+2ad^2}D)^2), \quad (4)$$

که در آن $G_{i,j}(q; q)$ تابع توزیع تجمعی توزیع F نامرکزی با (i, j) درجه آزادی و پارامتر غیرمرکزی q است. دومین امید ریاضی در رابطه (3) برابر است با

$$E[e^{-\frac{1}{2g^2}(\hat{b}_0^{SE} - b_0)^2} I(F_{1n} > F_{1-a})] = E[e^{-\frac{s^2 H}{2g^2} Z^2} I(F_{1n} > F_{1-a})]$$

$$= E[e^{-a Z^2} I(\frac{(Z+D)^2}{S/n} > F_{1-a})]$$

$$= \frac{1}{\sqrt{1+2a}} E[I(\frac{(\frac{Z}{\sqrt{1+2a}} + D)^2}{S/n} > F_{1-a})]$$

$$= \frac{1}{\sqrt{1+2a}} P(I(\frac{(Z + \sqrt{1+2a}D)^2}{S/n} > (1+2a)F_{1-a}))$$

$$= \frac{1}{\sqrt{1+2a}} (1 - G_{1n}((1+2a)F_{1-a}; (1+2a)D^2)). \quad (5)$$

با جای‌گذاری روابط (4) و (5) در (3) نتیجه می‌شود.

$$R(\hat{b}_0^{SP}, b_0) = 1 - \frac{e^{-(1-d)^2 a D^2 + \frac{c_2^2}{2}}}{\sqrt{1+2ad^2}} G_{1n}((1+2ad^2)F_{1-a}; (c_2 + \sqrt{1+2ad^2}D)^2)$$

$$- \frac{1}{\sqrt{1+2a}} (1 - G_{1n}((1+2a)F_{1-a}; (1+2a)D^2)). \quad (6)$$

تابع مخاطره $\hat{b}_0^{SE}$ با استفاده از لم 1 برابر است با

$$\begin{aligned}
R(\hat{b}_0^{SE}, b_0) &= E \left[1 - e^{-\frac{1}{2g^2}(\hat{b}_0^{SE} - b_0)^2} \right] \\
&= 1 - E \left[e^{-\frac{1}{2g^2}(d\hat{b}_0 + (1-d)b_{00} - b_0)^2} \right] \\
&= 1 - E \left[e^{-\frac{1}{2g^2}(d(\hat{b}_0 - b_0) - (1-d)(b_0 - b_{00}))^2} \right] \\
&= 1 - E \left[e^{-\frac{1}{2g^2}(ds\sqrt{H}Z - (1-d)Ds\sqrt{H})^2} \right] \\
&= 1 - e^{-(1-d)^2 a D^2} E \left[e^{-ad^2 Z^2} e^{(2ad(1-d)D)Z} \right] \\
&= 1 - \frac{1}{\sqrt{1+2ad^2}} e^{-(1-d)^2 a D^2} E \left[e^{\frac{2ad(1-d)D}{\sqrt{1+2ad^2}} Z} \right] \\
&= 1 - \frac{1}{\sqrt{1+2ad^2}} e^{-(1-d)^2 a D^2} M_Z \left(\frac{2ad(1-d)D}{\sqrt{1+2ad^2}} \right) \\
&= 1 - \frac{1}{\sqrt{1+2ad^2}} e^{-(1-d)^2 a D^2 (1 - \frac{2ad^2}{1+2ad^2})} \quad (7)
\end{aligned}$$

که در آن M_Z تابع مولد گشتاور متغیر تصادفی Z است.

4 مقایسه $\hat{b}_0^{SP}$ نسبت به $\hat{b}_0^{SE}$ و b_0^0

در این بخش ابتدا کارایی نسبی $\hat{b}_0^{SP}$ نسبت به $\hat{b}_0^{SE}$ و b_0^0 محاسبه شده و سپس به صورت گرافیکی مورد بررسی بیشتر قرار می‌گیرند.

1.4 کارایی نسبی $\hat{b}_0^{SP}$ نسبت به b_0^0

برای هر مقدار D مخالف صفر از رابطه (6)، تابع مخاطره برآوردگر $\hat{b}_0^{SP}$ را می‌توان به صورت

$$R(\hat{b}_0^{SP}, b_0) = R(b_0^0, b_0) + g(D),$$

نوشت که در آن $g(D)$ عبارت است از

$$\begin{aligned}
g(D) &= \frac{1}{\sqrt{1+2a}} G_{1n}((1+2a)F_{1-a}; (1+2a)D^2) \\
&\quad - \frac{e^{-(1-d)^2 a D^2 + \frac{c_2^2}{2}}}{\sqrt{1+2ad^2}} G_{1n}((1+2ad^2)F_{1-a}; (c_2 + \sqrt{1+2ad^2}D)^2).
\end{aligned}$$

بنابراین کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به برآوردگر b_0^0 به صورت

$$\text{Eff}[\hat{b}_0^{SP}, b_0^0] = \frac{R(b_0^0, b_0)}{R(\hat{b}_0^{SP}, b_0)} = \frac{1 - (1+2a)^{-\frac{1}{2}}}{1 - (1+2a)^{-\frac{1}{2}} + g(D)} \quad (8)$$

است. تحت فرض $H_0: b_0 = b_{00}$ داریم $D=0$ و در نتیجه

$$g(0) = \frac{1}{\sqrt{1+2a}} G_{1n}((1+2a)F_{1-a}; 0) - \frac{1}{\sqrt{1+2ad^2}} G_{1n}((1+2ad^2)F_{1-a}; 0) \leq 0.$$

بنابراین به ازای $D=0$ و $0 \leq d \leq 1$ برآوردگر $\hat{b}_0^{SP}$ نسبت به برآوردگر $b_0^{\%}$ کارایی نسبی معادل یا بیشتری دارد. اگر $\mathbb{R}(|D|) \neq 0$ آنگاه $g(D)$ همگرا به صفر و در نتیجه کارایی نسبی همگرا به یک می‌باشد ($\mathbb{R}(|D|) \neq 0$). بنابراین وقتی $|D|$ به اندازه کافی بزرگ باشد کارایی نسبی تقریباً یک است.

2.4 مقایسه $\hat{b}_0^{SE}$ نسبت به $\hat{b}_0^{SP}$

کارایی نسبی $\hat{b}_0^{SE}$ نسبت به $\hat{b}_0^{SP}$ با استفاده از روابط (7) و (6) به صورت

$$\begin{aligned} \text{Eff}[\hat{b}_0^{SP}, \hat{b}_0^{SE}] &= \frac{R(\hat{b}_0^{SE}, b_0)}{R(\hat{b}_0^{SP}, b_0)} \\ &= \frac{1 - (1+2ad^2)^{-\frac{1}{2}} e^{-\frac{(1-d)^2 ad^2 (1 - \frac{2ad^2}{1+2ad^2})}}{1 - (1+2a)^{-\frac{1}{2}} + g(D)} \end{aligned} \quad (9)$$

است. در همسایگی فرض صفر ($\mathbb{R}(|D|) = 0$)، با توجه به اینکه $1+2ad^2 \leq 1+2a$ و بنابر غیرنزولی بودن تابع توزیع

$$G_{1,n}((1+2ad^2)F_{1-a}; 0) \leq G_{1,n}((1+2a)F_{1-a}; 0)$$

داریم

$$\text{Eff}[\hat{b}_0^{SP}, \hat{b}_0^{SE}; D=0] = \frac{1 - (1+2ad^2)^{-\frac{1}{2}}}{1 - (1+2a)^{-\frac{1}{2}} + g(0)} \leq 1$$

اگر $|D|$ از یک مقداری بزرگتر شود ($\mathbb{R}(|D|) \neq 0$)، فرض صفر به ازای هر $0 < a < 1$ رد می‌شود و در نتیجه $\hat{b}_0^{SP} = b_0^{\%}$ و همینطور $\mathbb{R}(\hat{b}_0^{SE}, b_0) \neq 0$ آنگاه داریم

$$\lim_{|D| \neq 0} \text{Eff}[\hat{b}_0^{SP}, \hat{b}_0^{SE}] = \frac{1}{1 - (1+2a)^{-\frac{1}{2}}} > 1.$$

بنابراین در این حالت کارایی نسبی برآوردگر $b_0^{\%}$ ($\hat{b}_0^{SP} = b_0^{\%}$) نسبت به $\hat{b}_0^{SE}$ بیشتر است.

3.4 مقایسه برآوردگرها به صورت گرافیکی

در این بخش کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به برآوردگر $b_0^{\%}$ و $\hat{b}_0^{SE}$ به صورت گرافیکی در نرم افزار R ارزیابی می‌شود. در شکل 2 کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به برآوردگر $b_0^{\%}$ با استفاده از رابطه (8) به ازای مقادیر $n=25$ ، $a=0/01, 0/05, 0/1, 0/2$ ، $g=2$ ، $d=0/2, 0/4, 0/6, 0/8$ و $s^2H = \frac{1}{n}$ نسبت به D رسم شده است. از شکل 2 ملاحظه می‌شود که برآوردگر $\hat{b}_0^{SP}$ در همسایگی فرض صفر ($\mathbb{R}(|D|) = 0$) به دلیل نزدیک بودن مقدار اولیه b_{00} به مقدار واقعی b_0 ، مقدار مخاطره برآوردگر $\hat{b}_0^{SP}$ نسبت به $b_0^{\%}$ کمتر است بنابراین در این حالت برآوردگر $\hat{b}_0^{SP}$ دارای کارایی نسبی بیشتری

نسبت به $b_0^{\%}$ است. علاوه بر این مشاهده می‌شود برآوردهای $\hat{b}_0^{SP}$ با سطح معنی‌داری کوچکتر ($\alpha = 0/01$) دارای کارایی نسبی بیشتری نسبت به سایر حالات متناظر هستند زیرا در این حالت فرض صفر با اطمینان بیشتری پذیرفته می‌شود که اعتماد بیشتر آزمایشگر نسبت به مقدار اولیه b_{00} را نشان می‌دهد. همچنین با کوچک شدن ضریب انقباضی (d) یا به عبارتی بزرگ شدن ضریب b_{00} ، کارایی نسبی برآوردهای $\hat{b}_0^{SP}$ افزایش می‌یابد. در شکل ۳ کارایی نسبی برآوردهای $\hat{b}_0^{SP}$ نسبت به d به ازای $D = 0$ برای مقادیر مشخص $n = 10, 20, 30, 40$ ، $a = 0/01, 0/05, 0/1, 0/2$ و $s^2H = \frac{1}{n}$ رسم شده است. از شکل ۳ ملاحظه می‌شود که کارایی نسبی تحت فرض صفر ($D = 0$) نسبت به d ، تابعی نزولی است، زیرا با بزرگ شدن مقدار d ضریب اعتماد برآوردهای $b_0^{\%}$ افزایش می‌یابد، به طوری که وقتی d به یک میل کند ($d \rightarrow 1$)، آنگاه $b_0^{\%}$ به $\hat{b}_0^{SP}$ نزدیک می‌شود. در این حالت کارایی نسبی برآوردهای $\hat{b}_0^{SP}$ نسبت به $b_0^{\%}$ به یک میل می‌کند. با افزایش حجم نمونه (n) کارایی نسبی برآوردهای $\hat{b}_0^{SP}$ نسبت به $b_0^{\%}$ برای $n = 30, 40$ به دلیل اعتماد بیشتر به برآوردهای نمونه‌ای $b_0^{\%}$ ، در مقایسه با $n = 10, 20$ کاهش می‌یابد. همچنین برای $\alpha = 0/01$ کارایی نسبی $\hat{b}_0^{SP}$ نسبت به $b_0^{\%}$ ، بیشتر از سایر برآوردها است. در شکل ۴ کارایی نسبی $\hat{b}_0^{SP}$ نسبت به $\hat{b}_0^{SE}$ با استفاده از رابطه (۹) به ازای مقادیر $n = 20$ ، $d = 0/4$ ، $g = 2$ ، $a = 0/01, 0/05, 0/1, 0/2$ و $s^2H = \frac{1}{n}$ رسم شده است. همانطور که در شکل ۴ ملاحظه می‌شود در همسایگی فرض صفر کارایی نسبی $\hat{b}_0^{SP}$ کمتر از یک است در نتیجه $\hat{b}_0^{SE}$ کاراتر از $\hat{b}_0^{SP}$ می‌باشد. اما وقتی $|D|$ از یک مقداری بزرگتر شود که نشان دهنده‌ی غلط بودن اطلاع پیشین b_{00} است، فرض صفر به ازای هر $0 < \alpha < 1$ رد می‌شود بنابراین $b_0^{\%} = \hat{b}_0^{SP}$ و کارایی نسبی $\hat{b}_0^{SP}$ نسبت به $\hat{b}_0^{SE}$ طبق شکل ۴ بزرگتر از یک و نسبت به $|D|$ صعودی است. با توجه به شکل برای $\alpha = 0$ به دلیل برابر بودن برآوردهای $\hat{b}_0^{SP}$ و $\hat{b}_0^{SE}$ ، کارایی نسبی $\hat{b}_0^{SP}$ نسبت به $\hat{b}_0^{SE}$ روی خط یک منطبق می‌باشد.

در جدول ۱ بازه‌های موثر یعنی دامنه‌ای از مقادیر D در $g = 2$ که به ازای آن برآوردهای $\hat{b}_0^{SP}$ نسبت به برآوردهای $b_0^{\%}$ دارای کارایی نسبی بیشتر از یک است، خلاصه شده است. با توجه به نتایج به دست آمده در جدول ۱، با افزایش α ، طول بازه موثر کاهش می‌یابد. همچنین با افزایش n ، دامنه‌ای از مقادیر D که کارایی نسبی برآوردهای $\hat{b}_0^{SP}$ نسبت به $b_0^{\%}$ بزرگتر از یک است، کاهش می‌یابد. از طرفی با افزایش ضریب انقباضی d ، طول بازه موثر افزایش می‌یابد.

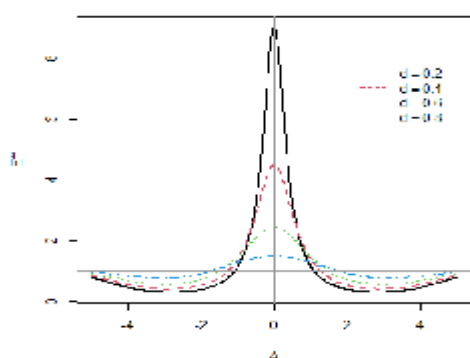
۴.۴ تعیین سطح معنی‌داری بهینه برای برآوردهای $\hat{b}_0^{SP}$

تابع مخاطره برآوردهای $\hat{b}_0^{SP}$ به سطح معنی‌داری پیش‌آزمون (α)، ضریب انقباضی (d)، پارامتر D و پارامتر شکل تابع زیان نرمال بازتابیده (g) بستگی دارد. بدین منظور، کارایی نسبی برآوردهای $\hat{b}_0^{SP}$ نسبت به برآوردهای $b_0^{\%}$ را به عنوان تابعی از α و D در نظر می‌گیریم. بنابراین داریم

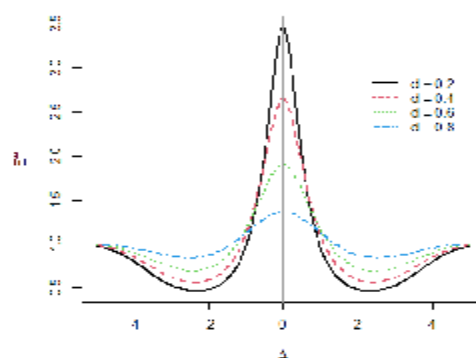
$$\text{Eff}[\hat{b}_0^{\text{SP}}, b_0^0; a, D] = \frac{1 - (1 + 2a)^{-\frac{1}{2}}}{1 - (1 + 2a)^{-\frac{1}{2}} + g(D)}. \quad (10)$$

تابع $\text{Eff}[\hat{b}_0^{\text{SP}}, b_0^0; a, D]$ در $D=0$ به ازای تمام $0 < a < 1$ ، بزرگتر یا مساوی با یک می‌باشد. همچنین اگر $|D|$ از مقدار صفر دور شود $\text{Eff}[\hat{b}_0^{\text{SP}}, b_0^0; a, |D|]$ به طور یکنواخت کاهش می‌یابد و با عبور از خط یک به کمترین مقدار در نقطه $|D|=|D_0|$ می‌رسد. سپس وقتی $|D| \gg |D_0|$ ، کارایی نسبی به طور یکنواخت از نقطه مینیمم به سمت یک افزایش می‌یابد. علاوه بر این برای $D=0$ و $a \in [0, 1]$ داریم

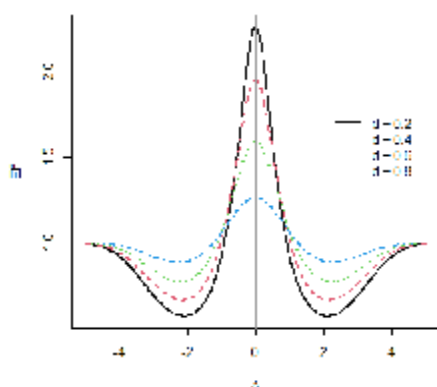
$$\max_a \text{Eff}[\hat{b}_0^{\text{SP}}, b_0^0; a, D=0] = \text{Eff}[\hat{b}_0^{\text{SP}}, b_0^0; a=0, D=0] = \frac{1 - (1 + 2a)^{-\frac{1}{2}}}{1 - (1 + 2ad^2)^{-\frac{1}{2}}} \geq 1.$$



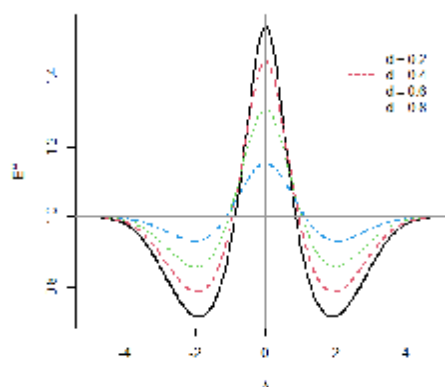
(ب)



(الف)

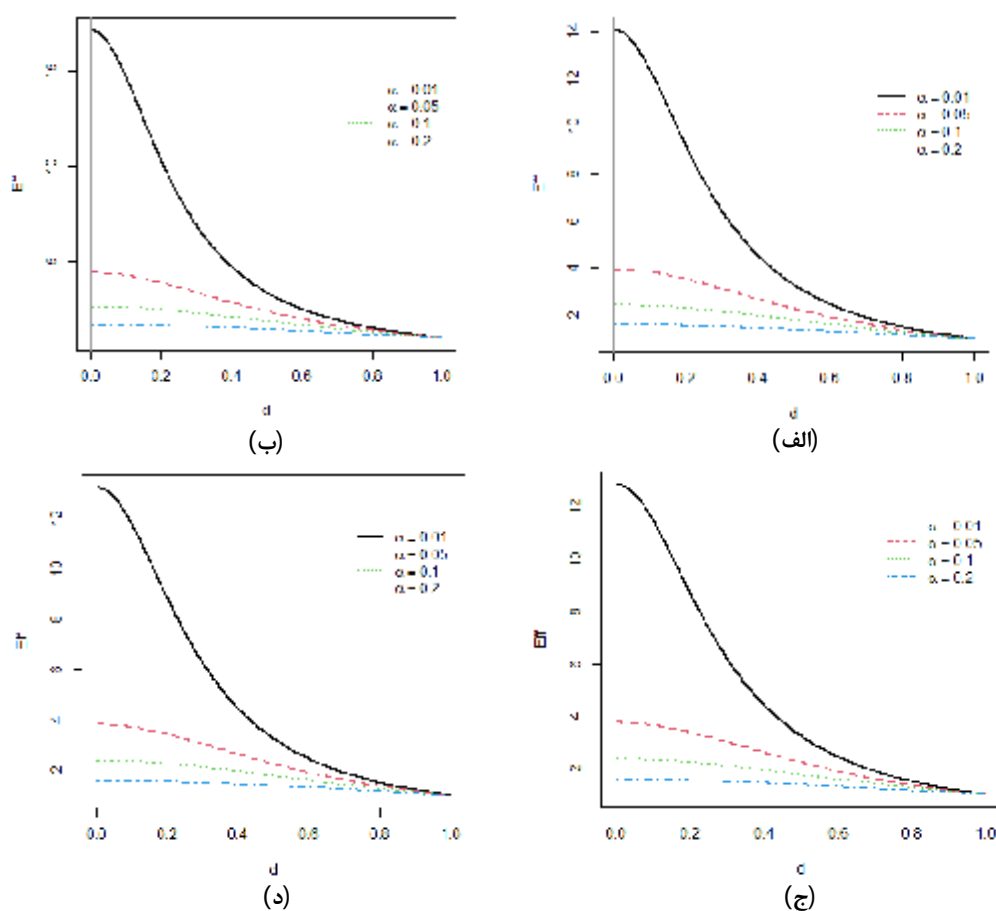


(د)

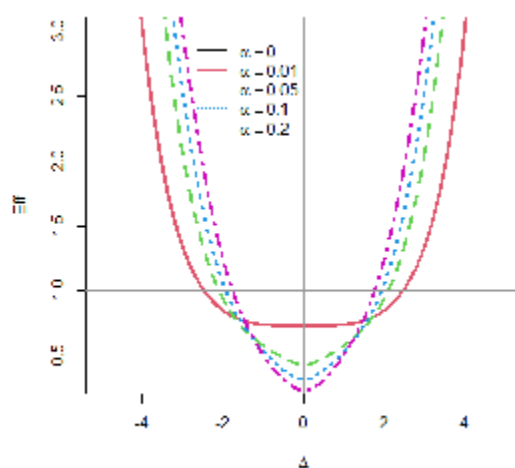


(ج)

شکل ۲: کارایی نسبی برآوردگر $\hat{b}_0^{\text{SP}}$ نسبت به b_0^0 : الف - $a=0/01$ ، ب - $a=0/05$ ، ج - $a=0/1$ ، د - $a=0/2$.



شکل ۳: کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به $b_0^{\%}$: الف - $n=10$ ، ب - $n=20$ ، ج - $n=30$ ، د - $n=40$.



شکل ۴: کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به $\hat{b}_0^{SE}$ به ازای مقادیر $n=20$ ، $d=0/4$ ، $g=2$ و $a=0,0/0,01/0,05/0,1/0,2$ نسبت به D .

جدول ۱: دامنه‌ای از مقادیر $\hat{b}_0^{SP}$ که به ازای آن برآوردگر $\hat{b}_0^{SP}$ کارا تر از b_0^0 است.

a = 0/2	a = 0/1	a = 0/05	a = 0/01	g	n
(-0/8831, 0/8831)	(-0/9502, 0/9502)	(-1/0144, 1/0144)	(-1/1294, 1/1294)	0/2	10
(-0/9782, 0/9782)	(-1/0708, 1/0708)	(-1/1636, 1/1636)	(-1/3453, 1/3453)	0/4	
(-1/0736, 1/0736)	(-1/1978, 1/1978)	(-1/3288, 1/3288)	(-1/6132, 1/6132)	0/6	
(-1/1704, 1/1704)	(-1/3335, 1/3335)	(-1/5152, 1/5152)	(-1/9590, 1/9590)	0/8	
(-0/8644, 0/8644)	(-0/9221, 0/9221)	(-0/9795, 0/9795)	(-1/0927, 1/0927)	0/2	20
(-0/9529, 0/9529)	(-1/0309, 1/0309)	(-1/1112, 1/1112)	(-1/2807, 1/2807)	0/4	
(-1/0404, 1/0404)	(-1/1425, 1/1425)	(-1/2516, 1/2516)	(-1/4991, 1/4991)	0/6	
(-1/1278, 1/1278)	(-1/2586, 1/2586)	(-1/4033, 1/4033)	(-1/7560, 1/7560)	0/8	
(-0/8591, 0/8591)	(-0/9140, 0/9140)	(-0/9691, 0/9691)	(-1/0803, 1/0803)	0/2	30
(-0/9459, 0/9459)	(-1/0195, 1/0195)	(-1/0959, 1/0959)	(-1/2599, 1/2599)	0/4	
(-1/0313, 1/0313)	(-1/1272, 1/1272)	(-1/2299, 1/2299)	(-1/4646, 1/4646)	0/6	
(-1/1162, 1/1162)	(-1/2382, 1/2382)	(-1/3729, 1/3729)	(-1/6998, 1/6998)	0/8	

$Eff[\hat{b}_0^{SP}, b_0^0; a, D=0]$ به عنوان تابعی از a ، با افزایش a ، کاهش می‌یابد. در حالت کلی برای دو مقدار a_1 و a_2 ، دو منحنی $Eff[\hat{b}_0^{SP}, b_0^0; a_1, |D|]$ و $Eff[\hat{b}_0^{SP}, b_0^0; a_2, |D|]$ در پایین خط یک، یک دیگر را قطع می‌کنند. با افزایش a_1 و a_2 ، مقدار $|D|$ در برخورد دو منحنی کاهش می‌یابد. در نتیجه برای انتخاب یک سطح معنی‌داری بهینه با ماکسیم کارایی نسبی از روش ماکس‌مین استفاده می‌شود. اگر بدانیم $|D|$ طبق جدول ۱ متعلق به بازه موثر است، برآوردگر $\hat{b}_0^{SP}$ انتخاب می‌شود زیرا در این بازه کارایی نسبی بزرگتر از یک می‌باشد. همانطور که در تحلیل کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به b_0^0 بررسی شد، برآوردگر $\hat{b}_0^{SP}$ به طور یکنواخت عملکرد بهتری نسبت به برآوردگر b_0^0 ندارد. همچنین، مقدار D برای محقق نامشخص است. بنابراین یک مقدار برای کمترین کارایی نسبی مانند Eff_0 در نظر می‌گیریم. مجموعه

$$A_a = \{a \mid Eff[\hat{b}_0^{SP}, b_0^0; a, |D|] \geq Eff_0\},$$

را در نظر بگیرید. برآوردگر $\hat{b}_0^{SP}$ که $Eff[\hat{b}_0^{SP}, b_0^0; a, |D|]$ را به ازای هر $a \in A_a$ و $|D|$ ماکزیمم کند انتخاب می‌شود. بنابراین معادله زیر را حل می‌کنیم

$$\max_a \min_{|D|} Eff[\hat{b}_0^{SP}, b_0^0; a, |D|] = Eff_0. \quad (10)$$

بیشترین و کمترین مقدار کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به برآوردگر b_0^0 با حل رابطه (۱۰) نسبت به $|D|$ به دست می‌آید. بیشترین میزان کارایی نسبی (Eff^*) و کمترین میزان کارایی نسبی (Eff_0) و مقدار $|D|$ که برای آن کارایی نسبی، مینیمم می‌شود $|D_0|$ به ازای مقادیر انتخابی a و n ، $d = 0/5$ و $g = 2$ در جدول ۲ خلاصه شده است. برای مثال اگر $d = 0/5$ و $n = 20$ باشد آنگاه برای رسیدن به کارایی نسبی به اندازه حداقل $0/8174$ مقدار بهینه سطح معنی‌داری برابر $a = 0/2$ به ازای $|D_0| = 1/9785$ خواهد بود. از جدول ۲ ملاحظه می‌شود که Eff^* یک تابع نزولی نسبت به a است در حالی که Eff_0 یک تابع صعودی است.

۵.۴ تعیین ضریب انقباضی مناسب

یک مسئله مهم برای برآوردگر $\hat{b}_0^{SP}$ انتخاب مناسب مقدار ضریب انقباضی (d) است. در این بخش دو روش برای انتخاب ضریب انقباضی به صورت زیر ارائه می‌شود.

۱. مقدار (d) از مینیمم کردن تابع مخاطره برآوردگر $\hat{b}_0^{SP}$ در رابطه (۶) نسبت به d از روش‌های عددی حاصل

می‌شود. به طور کلی مقدار D مجهول است بنابراین در این روش برآورد آن به صورت $\hat{D} = H^{-\frac{1}{2}}(b_0^0 - b_{00})S^{-1}$

که در آن $S^2 = S_n$ ، در رابطه (۶) جایگزین می‌شود. مقدار ضریب انقباضی حاصل شده از این روش را با d_1 و

برآوردگر $\hat{b}_0^{SP}$ متناظر با آن را با $\hat{b}_{01}^{SP}$ نشان می‌دهیم.

۲. برای آزمون فرض $H_0: b_0 = b_{00}$ در مقابل $H_1: b_0 \neq b_{00}$ فرض صفر را رد می‌کنیم اگر و تنها اگر $F_{1n} > F_{1-\alpha}$.

مقدار معنی‌داری آزمون به صورت مقدار $G_{1n}(f)$ ، $P = P(F_{1n} > f | H_0) = 1 - G_{1n}(f)$ - به دست می‌آید که f

مقدار مشاهده شده آماره آزمون F_{1n} و $G_{1n}(f)$ تابع توزیع تجمعی توزیع فشر با ۱ و n درجه آزادی است. P

مقدار^۱ این آزمون نشان می‌دهد که چقدر فرض صفر توسط داده‌ها پشتیبانی می‌شود. مقادیر بزرگ P - مقدار در

تایید نزدیک بودن b_{00} به مقدار واقعی b_0 است. بنابراین در این روش که توسط [۱۲] ارائه شد ضریب انقباضی

به صورت

$$d = 1 - (P - \text{مقدار})$$

در نظر گرفته می‌شود. ضریب انقباضی حاصل شده از این روش با d_2 و برآوردگر $\hat{b}_0^{SP}$ متناظر با آن با $\hat{b}_{02}^{SP}$ نشان داده

می‌شود. P - مقدار با افزایش حجم نمونه به صفر میل می‌کند، به عبارت دیگر d_2 به ۱ نزدیک می‌شود.

همچنین می‌توان ریشه دوم P - مقدار که توسط [۴] ارائه شد را برای پشتیبانی قوی‌تر از b_{00} پیشنهاد کرد. بنابراین ضریب

انقباضی به صورت $d = 1 - \sqrt{P - \text{مقدار}}$ در نظر گرفته می‌شود. ضریب انقباضی تعیین شده از این روش با d_3 و برآوردگر

$\hat{b}_0^{SP}$ متناظر با آن با $\hat{b}_{03}^{SP}$ نشان داده می‌شود.

^۵ P-value

جدول 2: ماکسیمم و مینیمم کارایی نسبی برآوردگر $\hat{b}_0^{SP}$ نسبت به b_0^0 برای $d=0/5$ و $g=2$

n					
۵۰	۴۰	۳۰	۲۰	۱۰	a
2/2142	2/2226	2/2369	2/2670	2/3693	Eff *
0/6368	0/6342	0/6298	0/62036	0/5859	Eff ₀ 0/05
2/3445	2/3595	2/3856	2/4422	2/6601	D ₀
1/7511	1/756774	1/7664	1/7867	1/8579	Eff * ۰/۱۰
0/7227	0/7208	0/7175	0/7105	0/6853	Eff ₀
2/1385	2/1484	2/1657	2/2032	2/3446	D ₀
1/5145	1/5184	1/5252	1/5394	1/5899	Eff * 0/15
0/7808	0/7793	0/7767	0/7711	0/7513	Eff ₀
2/0213	2/0287	2/0415	2/0690	2/1722	D ₀
1/3691	1/3720	1/3768	1/3871	1/4238	Eff * 0/20
0/8252	0/8239	0/8218	0/8174	0/8014	Eff ₀
1/9417	1/9474	1/9573	1/9785	2/0574	D ₀
1/2012	1/2027	1/2054	1/2110	1/2313	Eff * ۰/۳۰
0/8898	0/8890	0/8877	0/8847	0/8743	Eff ₀
1/8377	1/8413	1/8475	1/8608	1/9098	D ₀

5 مثال عددی

در این بخش یک مثال عددی برای تشریح برآوردگرهای پیشنهادی ارائه می‌شود. داده‌ها مربوط به سیستم جلو برنده یک موتور راکت از [6] است که وابستگی مقاومت برش به عنوان متغیر پاسخ به زمان هفتگی کارکرد سیستم به عنوان متغیر

پیشگو بررسی شد. ۲۰ مشاهده روی مقاومت برشی و زمان هفتگی کارکرد سیستم جمع‌آوری شد و در جدول ۳ نشان داده شده است. نتایج بررسی شده پیش فرض‌های مناسبت مدل را تایید می‌کند. برآوردهای ماکسیمم درستنمایی برای پارامترهای شیب و عرض از مبدا به ترتیب $b_1^0 = -37/15$ و $b_0^0 = 2627/82$ به دست آمد.

برآوردهای $\hat{b}_{01}^{SP}$ ، $\hat{b}_{02}^{SP}$ و $\hat{b}_{03}^{SP}$ به ترتیب به ازای d_1 ، d_2 و d_3 و $\hat{b}_0^{SE}$ به ازای d که طبق روش اول تعیین ضریب انقباضی از مینیمم کردن تابع مخاطره $\hat{b}_0^{SE}$ حاصل می‌شود، برای چهار مقدار اولیه b_{00} در $g=2$ محاسبه شده‌اند. برای مثال در سطر اول از جدول ۴ برآورد b_0 را وقتی که مقدار حدس زده شده $b_{00}=2200$ است در نظر می‌گیریم، برای برآورد $\hat{b}_0^{SE}$ مقدار $d=0/98$ از مینیمم کردن رابطه (۷) در $\hat{D}=9/68$ به دست آمد. بنابراین برآورد انقباضی $\hat{b}_0^{SE}$ به صورت

$$\hat{b}_0^{SE} = 0/98(2627/82) + (1 - 0/98)(2200) = 2623/351,$$

است. در حالتی که مقدار حدس زده شده $b_{00}=2600$ است برای برآورد $\hat{b}_0^{SE}$ مقدار $d=0/461$ از مینیمم کردن رابطه (۷) در $\hat{D}=0/62$ به دست آمد. بنابراین برآورد انقباضی $\hat{b}_0^{SE}$ در مقدار اولیه $b_{00}=2600$ به صورت

$$\hat{b}_0^{SE} = 0/461(2627/82) + (1 - 0/461)(2200) = 2612/826,$$

است. برای آزمون فرض $H_0: b_0 = 2200$ در مقابل $H_1: b_0 \neq 2200$ در سطح معنی‌داری $\alpha = 0/05$ مقدار آماره آزمون به صورت

$$f = \frac{(b_0^0 - b_{00})^2}{S_n^2 H} = 93/75,$$

جدول ۳: داده‌های مثال عددی

مشاهده	مقاومت برشی	سن عامل (هفته)	مشاهده	مقاومت برشی	سن عامل
i	y_i	x_i	i	y_i	x_i
1	2158/70	15/50	11	2156/20	13/00
2	1678/15	23/75	12	2399/55	3/75
3	2316/00	8/00	13	1779/82	25/00
4	2061/30	17/00	14	2336/75	9/75
5	2207/50	5/5	15	1765/30	22/00
6	1708/30	19/00	16	2053/50	18/00
7	1784/70	24/00	17	2414/40	6/00
8	2575/00	2/50	18	2200/50	12/50
9	2357/90	7/50	19	2654/20	2/00
10	2265/70	11/00	20	1753/70	21/50

جدول 4: برآورد پارامترهای موردنظر در داده‌های مثال

$\hat{b}_{03}^{SP}$	$\hat{b}_{02}^{SP}$	$\hat{b}_{01}^{SP}$	$\hat{b}_0^{SE}$	b_0^0	b_{00}
2627/82	2627/82	2627/82	2623/351	2627/82	2200
2607/436	2612/885	2613/468	2612/826	2627/82	2600
2652/793	2636/46	2627/824	2644/071	2627/82	2700
2627/82	2627/82	2627/82	2637/875	2627/82	2800

با توجه به این که مقدار آماره آزمون $f = 93/75$ از $F_{0.95} = 4/413$ بیشتر است در نتیجه فرض صفر رد می‌شود، بنابراین برآورد $\hat{b}_0^{SP}$ برابر $2627/82$ است. در حالتی که مقدار حدس زده شده $b_{00} = 2600$ است از آنجایی که $f = 0/396$ از $F_{0.95} = 4/413$ کوچکتر است، لذا فرض صفر پذیرفته می‌شود. مقدارهای $d_1 = 0/4841$ از مینیمم کردن تابع مخاطره در $\hat{D} = 0/62$ و مقدار $d_2 = 0/4631$ و $d_3 = 0/2673$ از روش P -مقدار محاسبه می‌شود. بنابراین برآوردهای $\hat{b}_0^{SP}$ به صورت

$$\hat{b}_{01}^{SP} = 0/4841(2627/82) + (1 - 0/4841)(2600) = 2613/468,$$

$$\hat{b}_{02}^{SP} = 0/4631(2627/82) + (1 - 0/4631)(2600) = 2612/885,$$

$$\hat{b}_{03}^{SP} = 0/2673(2627/82) + (1 - 0/2673)(2600) = 2607/436,$$

به دست می‌آیند. سایر برآوردهای ارائه شده در جدول 4 به طور مشابه حاصل می‌شوند. همچنین در این مثال طبق روش ماکس‌مین اگر $n = 20$ و $d = 0/48$ در نظر گرفته شود آنگاه برای رسیدن به کارایی نسبی $\hat{b}_0^{SP}$ نسبت به b_0^0 حداقل $0/6998$ مقدار بهینه سطح معنی‌داری آزمون $a = 0/1$ می‌باشد. در نتیجه فرض $H_0: b_0 = 2600$ در مقابل $H_1: b_0 \neq 2600$ در سطح معنی‌داری $a = 0/1$ آزمون می‌کنیم. از آنجایی که $F_{1,n} = 0/396 < F_{0.90} = 3/006$,

فرض صفر رد نمی‌شود و برآورد $\hat{b}_0^{SP}$ برابر است با

$$\hat{b}_0^{SP} = 0/48(2627/82) + (1 - 0/48)(2600) = 2613/35.$$

6 بحث و نتیجه‌گیری

در این مقاله کارایی نسبی برآوردگرهای انقباضی و پیش‌آزمون انقباضی برای پارامتر عرض از مبدا مدل رگرسیونی خطی ساده تحت تابع زیان نرمال بازتابیده مورد بررسی قرار گرفت. ابتدا تابع زیان، معرفی و سپس مخاطره برآوردگرهای ماکسیمم درست‌نمایی، پیش‌آزمون انقباضی و انقباضی تحت تابع زیان نرمال بازتابیده به دست آمد. با مقایسه برآوردگرهای ماکسیمم

درست‌نمایی و پیش‌آزمون انقباضی می‌توان نتیجه گرفت که هرچه مقدار اطلاع غیرنمونه‌ای (b_{00}) به پارامتر واقعی نزدیک‌تر باشد، کارایی نسبی برآوردگر پیش‌آزمون انقباضی نسبت به برآوردگر ماکسیمم درست‌نمایی بیشتر می‌شود. همچنین برآوردگر پیش‌آزمون انقباضی به طور یکنواخت عملکرد بهتری نسبت به برآوردگر انقباضی ندارد به طوری که در همسایگی $D=0$ کارایی نسبی برآوردگر پیش‌آزمون انقباضی نسبت به برآوردگر انقباضی کمتر از یک است. اما وقتی $|D|$ از یک مقداری بزرگتر شود برآوردگر پیش‌آزمون انقباضی کارایی بیشتری نسبت به برآوردگر انقباضی دارد. بازه‌های موثر یا به عبارتی دامنه‌ای از مقادیر D که به ازای آن برآوردگر پیش‌آزمون انقباضی نسبت به برآوردگر ماکسیمم درست‌نمایی دارای مخاطره کمتری است ارائه شد. با افزایش سطح معنی‌داری آزمون و حجم نمونه، طول بازه موثر کاهش می‌یابد. همچنین مقدار اندازه آزمون بهینه با استفاده از روش ماکس‌مین به دست آمد. در یک مثال واقعی بر اساس چند مقدار انتخابی برای ضریب انقباضی، برآوردهای انقباضی و پیش‌آزمون انقباضی محاسبه شدند.

تقدیر و تشکر

نویسندگان، از سردبیر و داوران محترم مقاله برای ارزیابی و بهبود کیفیت آن کمال تشکر را دارند.

References

1. کیاپور، آ.، برآوردیابی انقباضی کلاسیک و بیزی در توزیع رایلی با استفاده از حدس نقطه‌ای براساس داده‌های سانسور شده، مجله پژوهش‌های ریاضی، جلد 4، شماره 1 (1397)، 63-74.
2. Z. Hoque, S. Hossain, Improved Estimation in Regression with Varying Penalty, Journal of Statistical Theory and Practice, **6** (2012) 260-273.
3. Z. Hoque, S. Khan, J. Wesolowski, Performance of Preliminary Test Estimator Under Linex Loss Function, Communications in Statistics-Theory and Methods, **38** (2009) 252-261.
4. M. Jabbari Nooghabi, Shrinkage Estimation of $P(Y < X)$ in the Exponential Distribution Mixing with Exponential Distribution, Communications in Statistics-Theory and Methods, **45** (2016) 1477-1486.
5. S. Khan, Z. Hoque, A.K.M.E. Saleh, Estimation of the Intercept Parameter for Linear Regression Model with Uncertain Prior Information, Statistical Papers, **46** (3) (2005) 379-395.
6. D. C. Montgomery, E.A. Peck, Vining G.G., Introduction to Linear Regression Analysis, 5th Edition, Wiley (2012).
7. M. Naghizadeh Qomi, A. Kiapour, Improving the MLE of the Location Parameter of a Normal Population under Reflected Normal Loss, 47th Annual Iranian Mathematics Conference, Kharazmi University, Karaj, Iran (2016).
8. M. Naghizadeh Qomi, N. Nematollahi, A. Parsian, Estimation After Selection Under Reflected Normal Loss Function, Communications in Statistics Theory and Methods, **41** (2012) 1040-1051.

9. F. Spring, The Reflected Normal Loss Function, Canadian Journal of Statistics, **21** (1) (1993) 321-330.
10. J.R. Thompson, Some Shrunk Techniques for Estimating the Mean, Journal of the American Statistician Association, **63** (1968) 113-122.
11. M. Towhidi , J. Behboodian, Estimation of a Location Parameter with a Reflected Normal Loss function, Iranian Journal of Science and Technology, **25** (2001) 183-190.
12. S.K. Tse, Tso G., Shrinkage Estimation of Reliability Distribution Lifetimes, Communications in Statistics - Simulation and Computation, **25** (2) (1996) 415-430.
13. H. Waldl, Some Comments on Statistical Papers 46, 270-395, Statistical Papers, **51** (2010) 241-246.