

On Bayesian Prediction of k-Record Values via Record Ranked Set Sampling Scheme in a Rayleigh Distribution

Ehsan Golzade Gervi⁽¹⁾¹, Simin mansouri Boroujeni⁽²⁾ and Kazem Fayyaz Heidari ⁽³⁾

^{(1),(2),(3)} Department of Statistics, Payame Noor University, Tehran, Iran.

Received: 8 February 2025

Accepted: 14 April 2026

Published online: 12 June 2026

Abstract:

Introduction

There are some tests where have been done sequentially, and only record values are observed. The record values have been used in applied examinations, such as quality control, hydrology, economics, climatology and etc. This is why many authors attracted to study on it. An observation X_j is an upper record value if its value exceeds all previous observations. Predicting future of k-record values as an influencing factor is one of the important aims in real life. In this article, two sample Bayesian prediction has been investigated via an upper record ranked set sampling scheme in a Rayleigh model. The record ranked set sampling scheme as a substitute method for producing record, has been formally suggested by Salehi and Ahmadi (2014). After a brief introduction to the record ranked set sampling scheme, the first aim of this article is to study the behavior of Bayesian point predictions of k-records from a future sequence based on an upper record ranked set sampling as an informative sample with respect to both squared error symmetric loss function and linear exponential (linex) asymmetric loss function for different values of α and β . We have compared the Bayesian point predictors in terms of the mean square prediction errors. Our second purpose of this article is to study the Bayesian prediction intervals for k-record in this scheme. Evaluation criteria such as the expected lengths and coverage probabilities are employed to consider the performance of this Bayesian prediction intervals.

Material and Methods

In this article, by considering a Rayleigh model, the record ranked set sampling scheme is used to comparing the Bayesian predictions for k-record values. In order to monitor the performance of the Bayesian k-record predictors, a Monte Carlo simulation is used. We consider $m = k = 1, 2, 3$. Also, the values of the a in linex loss function is considered to be -1, 1. The value of λ is considered to be 0.042 in Rayleigh model. We consider different values (1.5, 0.5 and 0) for α and β . We assumed that the case ($\alpha = \beta = 0$) for non-informative prior. Next, an application example (related to wind speed) on a real data set for six years 2015 to 2020 of Brunswick new Province and Bas Caraquej Station for government Canada is presented to compare and show the applicability of the article's results. The given data are obtainable at IOC website. The p-value of the goodness of fit test Anderson-Darling of the Rayleigh distribution, is 0.971

¹Corresponding author

E-mail addresses: (Ehsan Golzade Gervi) e_golzade_g@pnu.ac.ir, (Simin mansouri Boroujeni) s_mansouri@pnu.ac.ir, (Kazem Fayyaz Heidari) fayyaz@pnu.ac.ir

which supports the suitability of the fitting Rayleigh model on real data set. From this data set, it can be shown that $r = (3.05, 6.94, 6.39, 3.61, 6.39, 6.67)^T$, with $n = 6$. Data were analyzed by R software.

Results and discussion

In this article, k-record values predictors under the squared error loss function performs better than the linear exponential loss functions. The Bayesian k-record predictors based on $(\alpha = \beta = 1.5)$ are better than the others in terms of mean square prediction errors. The mean square prediction errors are increasing in m . Moreover, when m increases, the coverage probability become closer to coverage percentage (0.95). The expected lengths under the highest posterior density method are smaller than those of survival method and the coverage probabilities of the highest posterior density method is more than survival method. These Bayesian intervals improve which was expected as the α and β increase jointly. Also, the Bayesian prediction intervals for k-record based on $(\alpha = \beta = 1.5)$ and $(\alpha = \beta = 0.5)$ as compared to the Bayesian prediction intervals based on non-informative prior, performances very well, in terms of coverage probability and expected length. The results in the real data part confirm the results of the simulation part.

Conclusion

This paper contributes to the field of statistical analysis by emphasizing the appropriate sampling scheme to decrease the cost, increase efficiency and introducing the RRSS scheme. In this article, we have considered the behavior of predictors of future k-record via record ranked set sampling scheme in a Rayleigh model. Results showed that the Bayesian predictors of k-record under the squared error loss function performs better than the Bayes predictors under the linear exponential loss function. Also, the performance highest posterior density method are superior to survival method. Finally, using a record ranked set sampling scheme for Bayesian prediction of k-record is recommended if the conditions are met for it.

Keywords: Bayesian point prediction, Bayesian interval prediction, k-Record values, Bayesian predictive density, Record ranked set sampling.





درباره پیشگویی بیزی مقادیر k - رکورد بر اساس طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی در توزیع رایلی

احسان گل‌زاده گروی^(۱)، سیمین منصوری بروجنی^(۲) و کاظم فیاض حیدری^(۳)

(۱)، (۲)، (۳) گروه آمار، دانشگاه پیام نور، تهران، ایران.

تاریخ انتشار: ۱۴۰۵/۳/۲۲

تاریخ پذیرش: ۱۴۰۵/۱/۲۵

تاریخ دریافت: ۱۴۰۳/۱۱/۱۹

چکیده: در این مقاله، با در نظر گرفتن توزیع رایلی، به پیشگویی بیزی مقادیر k - رکورد بالا، زمانی که داده‌ها به شیوه طرح نمونه‌گیری مجموع رتبه‌دار رکوردی جمع‌آوری شده باشد، پرداخته شده است. در ابتدا به پیشگویی‌های نقطه‌ای تحت تابع زیان متقارن توان دوم خطا و تابع زیان نامتقارن لاینکس و محاسبه میانگین توان دوم خطا پیشگویی آنها پرداخته می‌شود. سپس، برای ارزیابی بازه‌های پیشگویی به‌دست‌آمده به روش تابع بقا و روش HPD، متوسط طول بازه و احتمال پوشش ارائه می‌گردد. در پایان، با یک مطالعه شبیه‌سازی و ارائه یک مثال کاربردی از داده‌های واقعی مرتبط با میزان سرعت وزش باد، مقایسه نتایج و عملکرد نشان داده شده است.

واژه‌های کلیدی: پیشگویی بیزی نقطه‌ای، پیشگویی بیزی بازه‌ای، مقادیر k - رکورد، چگالی بیزی پیشگویی، نمونه‌گیری مجموع رتبه‌دار رکوردی.

رده‌بندی ریاضی: 62G30 Primary; 62G25; 62C12

۱ مقدمه

مفهوم پیشگویی^۲ یکی از مفاهیم ریشه‌دار است که تا به امروز دست‌مایه کار بسیاری از آماردانان بوده است. در مسائل پیشگویی، پژوهشگر علاقه‌مند است با استفاده از پیشامد E که به وقوع پیوسته است، پیشامد رخ نداده‌ای مثل F را پیشگویی کند. پیشگویی، در تصمیم‌گیری وضع آینده سیستم‌ها و تحلیل آماری آن نقش ویژه‌ای ایفا می‌کند که به‌صورت استنباط مقادیر متغیرهای تصادفی مجهول آینده یا توابعی از آن متغیرهای تصادفی تعریف می‌شود. به‌طور خاص، پیشگویی و توانایی حدس‌زدن جنبه‌های غیرتعیینی حوادث قبل از تصمیم‌گیری با استفاده از اطلاعات

^۱نویسنده مسئول مقاله

موجود در نمونه از دیدگاه بیزی، به سوژه مورد علاقه‌ای برای محققان تبدیل شده است. دیدگاه بیزی این امکان را به پژوهشگر می‌دهد تا اطلاعات قبلی را در تحلیل مسائل پیش رو وارد کند. از جمله پژوهش‌های ارزشمندی که در مبحث پیشگویی به ثبت رسیده است، می‌توان به تحقیقات در [۱۰، ۱۲، ۱۳، ۱۵، ۱۷، ۱۸، ۲۱، ۲۲، ۲۶، ۳۰، ۳۱] اشاره کرد. در این مطالعه، پیشگویی بیزی دو نمونه‌ای در نظر گرفته شده است. در پیشگویی دو نمونه‌ای، فرض می‌کنیم داده‌های یک نمونه یا یک دنباله از متغیرهای تصادفی را در اختیار داریم و در صدد هستیم تا متغیرهای تصادفی آینده را از یک نمونه یا دنباله دیگر پیشگویی کنیم. در این روش، متغیرهایی که علاقه‌مند به پیشگویی آن‌ها در آینده هستیم مستقل از متغیرهای مشاهده شده است. این مقاله با در نظر گرفتن توزیع رایلی، به مساله پیشگویی بیزی از نوع نقطه‌ای و بازه‌ای برای مقادیر k - رکورد بر اساس داده‌های حاصل از طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی پرداخته است. ما کار خود را با ارائه تعریف مقادیر رکوردی در زیر آغاز می‌کنیم:

تعریف ۱.۱ ([۱۱]). فرض کنید $\{X_i, i \geq 1\}$ یک دنباله نامتناهی از متغیرهای تصادفی مستقل و هم‌توزیع با تابع توزیع تجمعی مشترک $F(x; \theta)$ و تابع چگالی احتمال $f(x; \theta)$ باشد، در این صورت مشاهده‌ی X_j یک رکورد بالا است، هرگاه با احتمال یک، برای همه مقادیر i که $j > i$ برقرار باشد. به عبارت دیگر، برای اینکه یک رکورد جدید رخ دهد باید بتوان آن را با مشاهدات قبلی مقایسه نمود. گوئیم مشاهده X_j یک رکورد بالا است، هرگاه از تمام مشاهدات ما قبل خود بزرگتر باشد. به شیوه‌ای مشابه، رکوردهای پایین نیز قابل تعریف هستند. امروزه داده‌های رکوردی دارای کاربردهای وسیعی به ویژه در اقتصاد، زلزله نگاری، مهندسی، پزشکی، کنترل کیفیت، آب شناسی، استخراج نفت و معادن، قابلیت اطمینان، هوا و فضا، تحلیل داده‌های طول عمر و هواشناسی می‌باشد. برخی از این اطلاعات درباره رکوردها را در منابع کلیدی [۹، ۱۱، ۱۹، ۲۳] می‌توان دید. اکنون فرض کنیم $Y_1, Y_2, \dots$ یک دنباله نامتناهی از متغیرهای تصادفی پیوسته و مستقل از $\{X_i, i \geq 1\}$ باشد. در این صورت مقادیر k - رکورد را می‌توان به عنوان تعمیمی از رکوردها در یک دنباله متغیرهای تصادفی هم‌توزیع و مستقل، مشاهدات متوالی از k امین بزرگ‌ترین مقدار نمونه تعریف کرد. در واقع مقادیر k - رکورد بالا در دنباله‌ای از متغیرهای تصادفی، k امین بزرگ‌ترین مشاهده در یک نمونه‌گیری است، [۳۱]. برای اطلاع بیشتر درباره رکوردها بویژه k - رکوردها، می‌توان به تحقیقات در ۱, 2, 4, 5, 6, 7, 8, 13, 20 نیز مراجعه نمود.

تعریف ۲.۱ ([۱۱]). فرض کنیم $\{Y_i, i \geq 1\}$ و $\{T_{\lambda(k)}, T_{\lambda(k)}, \dots\}$ به ترتیب نشان‌دهنده i امین آماره ترتیبی از یک نمونه تصادفی به حجم m ، دنباله زمان‌های رخداد k - رکورد بالا و مقادیر k - رکورد بالا باشند. با قرار دادن $T_{\lambda(k)} = k$ برای $m \geq 2$ داریم:

$$T_{m,k} = \min\{j : j > T_{m-\lambda(k)}, Y_j > Y_{T_{m-\lambda(k)-k+\lambda(k)}}\}$$

در این صورت دنباله k - رکوردهای بالا عبارتند از: $R_{\lambda(k)} = Y_{\lambda(k)}, R_{m(k)} = Y_{T_{m-\lambda(k)-k+\lambda(k)}}$

لم ۳.۱ ([۱۱]). فرض کنیم $\{X_i, i \geq 1\}$ دنباله‌ای نامتناهی از متغیرهای تصادفی مستقل با تابع توزیع تجمعی F و تابع چگالی f باشند. آنگاه تابع چگالی حاشیه‌ای m امین k - رکورد عبارت است از

$$f_{R_{m(k)}}(u) = \frac{k^m}{(m-1)!} \{-\log(1-F(u))\}^{m-1} \{1-F(u)\}^{k-1} f(u) \quad (۱.۱)$$

در تمام این مقاله، از k - رکورد بالا استفاده می‌نماییم. برای تشریح بیشتر k - رکورد بالا به مثال زیر توجه کنید:

مثال ۴.۱. فرض کنیم یک نمونه تصادفی به حجم $n = ۱۶$ با مقادیر زیر داشته باشیم. این داده‌ها مربوط به زمان خرابی‌های پی در پی در سیستم تهویه مطبوع جت‌بوئینگ ۷۲۰ از ناوگان هواپیمایی ۸۰۴۵ است که توسط پورشان در مرجع [۲۵]، گزارش شده است:

$Y_i : ۱۰۲ - ۲۰۹ - ۱۴ - ۵۷ - ۵۴ - ۳۲ - ۶۷ - ۵۹ - ۱۳۴ - ۱۵۲ - ۲۷ - ۱۴ - ۲۳۰ - ۶۶ - ۶۱ - ۳۴$

الف: در این صورت آماره‌های رکوردی نمونه با استناد به تعاریف ۱.۱ و ۲.۱ به شرح زیر است:

j	۱	۲	۳
زمان رکورد بالا	۱	۲	۱۳
مقدار رکورد بالا	۱۰۲	۲۰۹	۲۳۰

ب: به‌ازای $k = ۲$ ، زمان‌های رخداد m /امین k - رکورد و مقادیر m /امین k - رکورد عبارت است از

m	۱	۲	۳	۴
$T_{m,k}$	۲	۹	۱۰	۱۳
$R_{m(k)}$	۱۰۲	۱۳۴	۱۵۲	۲۰۹

ج: به‌ازای $k = ۳$ ، زمان‌های رخداد m /امین k - رکورد و مقادیر m /امین k - رکورد عبارت است از

m	۱	۲	۳	۴	۵	۶
$T_{m,k}$	۳	۴	۷	۹	۱۰	۱۳
$R_{m(k)}$	۱۴	۵۷	۶۷	۱۰۲	۱۳۴	۱۵۲

مطالعات واقعی متعددی وجود دارد که به دلایل نامعلومی تعدادی از مشاهدات آن ناپدید شده اند و فقط تعداد محدودی از مشاهدات آن در دسترس است که مقادیر آن‌ها از تمام مشاهدات ماقبل خود بزرگتر یا کوچکتر می‌باشند. همین سبب می‌شود که پژوهشگر به همه مشاهدات دسترسی نداشته باشد. بنابراین ناچار است در چنین وضعیتی برای پیش‌گویی مقادیر آینده از همان تعداد محدود مشاهدات (مقادیر رکورد) به‌عنوان مقادیر نمونه‌ای مشاهده شده استفاده کند. به‌کارگیری یک روش آماری مناسب نمونه‌گیری بر اساس طرح‌های مختلف که موجب به حداکثر رساندن دقت خصیصه‌های جامعه، صرفه جویی در زمان و هزینه می‌شود، همواره مورد علاقه آماردانان بوده است. در این مطالعه، با طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی RRSS^۱ که ایده‌ی آن بر مبنای ساختار نمونه مجموعه رتبه‌دار نخستین بار توسط صالحی و احمدی [۲۸]، پیشنهاد شد به اختصار آشنا می‌شویم:

فرض کنیم n دنباله‌ی مستقل از متغیرهای تصادفی در اختیار است. نمونه‌گیری در هر دنباله به نحوی است که تنها رکوردها ثبت می‌شوند. همچنین i امین نمونه‌گیری دنباله‌ای زمانی متوقف می‌شود که i امین رکورد مشاهده شود. برای تحلیل آماری در دنباله i ام، تنها مقدار رکورد i ام نیاز است. به این ساختار، طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی گفته می‌شود. با توجه به توضیحات فوق، مراحل استخراج نمونه‌ای با حجم n به روش نمونه‌گیری مجموعه رتبه‌دار رکوردی با استفاده از دیاگرام زیر دقیق‌تر نشان داده می‌شود:

$$\begin{aligned}
 ۱ : R_{(۱)۱} &\Rightarrow R_{۱,۱} = R_{(۱)۱} \\
 ۲ : R_{(۱)۲} \quad R_{(۲)۲} &\Rightarrow R_{۲,۲} = R_{(۲)۲} \\
 \dots &\Rightarrow \dots \dots \\
 n : R_{(۱)n} \quad R_{(۲)n} \quad R_{(n)n} &\Rightarrow R_{n,n} = R_{(n)n}
 \end{aligned} \tag{۲.۱}$$

در واقع نمونه حاصل از طرح مذکور، بردار n -بعدی، $R = (R_{۱,۱}, R_{۲,۲}, \dots, R_{n,n})^T$ است. در طرح فوق، متغیر $R_{i,(j)}$ ، بر i امین رکورد در دنباله j ام دلالت دارد. درایه‌های بردار R ، مشاهدات حاصل از این طرح هستند که اولاً مستقلند و ثانياً این امکان وجود دارد که مرتب نباشند. اکنون فرض کنیم $R_{i,i}$ ها رکورد i ام در دنباله i ام باشند، آنگاه تابع چگالی احتمال حاشیه‌ای i امین رکورد و تابع چگالی احتمال توام n رکورد به ترتیب به فرم زیر می‌باشند، [۱۱]:

$$f_{i,i}(x) = \frac{\{-\log(1 - F(x))\}^{i-1}}{(i-1)!} f(x) \tag{۳.۱}$$

$$f_R(r; \theta) = \prod_{i=1}^n \frac{\{-\log(1 - F(r_{i,i}; \theta))\}^{i-1}}{(i-1)!} f(r_{i,i}; \theta) \quad , \quad \theta \in \Theta \tag{۴.۱}$$

که در آن‌ها، $r = (r_{۱,۱}, r_{۲,۲}, \dots, r_{n,n})^T$ مقدار مشاهده‌شده $R = (R_{۱,۱}, R_{۲,۲}, \dots, R_{n,n})^T$ فضای Θ و پارامتر است. از زمان ایجاد ایده طرح RRSS، پژوهش‌های ارزشمندی در این حیطه به ثبت رسیده است. صالحی و احمدی برآورد مدل تنش-مقاومت را در طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی برای توزیع نمایی مطالعه کردند [۲۹]. نظریه اطلاع نیز توسط اسکندرزاده و همکاران در این طرح بررسی شده است [۱۴]. در مرجع

^۱Record Ranked Set Sampling

[۲۴]، پاول و توماس طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی همراه را معرفی کردند. صفریان و همکاران برآوردهای بهبود یافته‌ای برای مدل تنش-مقاومت به‌دست آورده‌اند، ([۲۷]). در مرجع [۳]، گل‌زاده‌گروی و همکاران به مقایسه برآوردها و پیشگویی‌های بیزی تجربی بر اساس طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی با طرح نمونه‌گیری معکوس پرداختند. اخیراً استنباط آماری بر اساس رکوردهای بالا در این طرح با استفاده از توزیع پارتو نوع اول، توسط گل‌زاده‌گروی مطالعه شده است، ([۱۶]).

روند نگارش این مقاله بدین شرح است که، پس از ارائه مقدمه و پیشنیازها، در بخش ۲، پیشگویی بیزی نقطه‌ای مقادیر k -رکورد، بر اساس داده‌های حاصل از طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی مورد بحث قرار گرفته است و به دنبال آن، در بخش ۳، به پیشگویی بازه‌ای مقادیر k -رکورد به روش تابع بقا^۱ و روش HPD^۲ پرداخته شده است. در بخش ۴، نتایج به‌دست آمده به کمک شبیه‌سازی را مورد مطالعه قرار داده‌ایم. بخش ۵ نیز به نحوه کاربست پیشگویی‌های به‌دست آمده بر اساس داده‌های واقعی اختصاص داده شده است. در بخش پایانی، نتیجه‌گیری و بحث آورده شده است.

۲ پیشگویی بیزی نقطه‌ای مقادیر k -رکورد

فرض کنید $\{X_i, i \geq 1\}$ دنباله‌ای نامتناهی از متغیرهای تصادفی مستقل و هم‌توزیع دارای توزیع رایلی با تابع چگالی

$$f(x; \lambda) = 2\lambda x e^{-\lambda x^2}, \quad x > 0, \quad \lambda > 0, \quad (1.2)$$

و تابع توزیع تجمعی

$$F(x; \lambda) = 1 - e^{-\lambda x^2}, \quad (2.2)$$

باشد. اگر $L(\lambda|r)$ تابع درستنمایی نمونه‌های استخراج شده بر اساس طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی باشد، با استفاده از روابط (۱.۲) و (۲.۲) و قرار دادن آنها در رابطه (۴.۱) می‌توان نوشت

$$L(\lambda|r) \propto \lambda^N e^{-\lambda t}, \quad (3.2)$$

که در آن $N = \frac{n(n+1)}{2}$ و $t = \sum r_{i,i}^2$ مقدار مشاهده شده $T = \sum R_{i,i}^2$ است. می‌توان نشان داد $R_{i,i}^2$ دارای توزیع گاما با پارامترهای i و λ است. لذا T دارای توزیع گاما با پارامترهای N و λ است. در ادامه می‌خواهیم براساس طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی، بر روی مساله پیشگویی بیزی نقطه‌ای m -امین k -رکورد، $R_{m(k)}$ ، تحت تابع زیان متقارن توان دوم خطا^۳ که به‌صورت $L_1(\lambda, \hat{\lambda}) = (\hat{\lambda} - \lambda)^2$ است، متمرکز شویم.

۱.۲ پیشگویی بیزی نقطه‌ای $R_{m(k)}$ تحت تابع زیان توان دوم خطا

فرض کنیم $\pi(\lambda)$ توزیع پیشین مزدوج گاما، به فرم زیر باشد

$$\Pi(\lambda) = \frac{\beta^\alpha}{\Gamma(\alpha)} \lambda^{\alpha-1} e^{-\beta\lambda}, \quad (4.2)$$

که در آن ابرپارامترهای α و β مقادیری مثبت و معلوم‌اند و $\Gamma(\alpha) = \int_0^\infty \lambda^{\alpha-1} e^{-\beta\lambda} d\lambda$ است. با استفاده از روابط (۳.۲) و (۴.۲)، تابع چگالی پسین λ به فرم زیر به‌دست می‌آید

$$\Pi(\lambda|r) = \frac{(\beta + t)^{N+\alpha}}{\Gamma(N + \alpha)} \lambda^{N+\alpha-1} e^{-\lambda(\beta+t)}. \quad (5.2)$$

¹Survival Function Method

²Highest Posterior Density Method

³Squared Error Loss Function

اکنون فرض کنیم $R_{m(k)}$ ، m امین k -رکورد حاصل از یک دنباله آینده باشد. آنگاه تابع چگالی پیشگوی بیزی^۱ $R_{m(k)}$ ، به شرط r عبارت خواهد بود از

$$f_{R_{m(k)}}^*(u|r) = \int_0^\infty f_{R_{m(k)}}(u|\lambda) \Pi(\lambda|r) d\lambda,$$

که در آن u مقدار مشاهده شده $R_{m(k)}$ می‌باشد. با بکارگیری توزیع رایلی و رابطه (۱.۱) خواهیم داشت

$$f_{R_{m(k)}}(u|\lambda) = \frac{\gamma(k\lambda)^m}{\Gamma(m)} u^{m-1} e^{-k\lambda u}, \quad u > 0, \quad (6.2)$$

در نتیجه طبق تعریف امید ریاضی داریم

$$E(R_{m(k)}^c) = \int_0^\infty u^c f_{R_{m(k)}}(u|\lambda) du = \frac{\Gamma(m + \frac{c}{\gamma})}{\Gamma(m)(k\lambda)^{\frac{c}{\gamma}}}. \quad (7.2)$$

در ادامه با استفاده از روابط (۵.۲) و (۶.۲) و اندکی می‌توان نوشت

$$f_{R_{m(k)}}^*(u|r) = \frac{\gamma}{B(m, N + \alpha)u} p(u)^m (1 - p(u))^{N + \alpha}, \quad (8.2)$$

که در آن $B(m, N + \alpha) = \frac{\Gamma(m)\Gamma(N + \alpha)}{\Gamma(m + N + \alpha)}$ و $p(u) = \frac{k\lambda u}{\beta + t + k\lambda u}$ است.

از ویژگی‌های تابع چگالی احتمال پیشگوی بیزی، $f_{R_{m(k)}}^*(u|r)$ ، مستقل بودن آن از پارامتر λ است. این ویژگی مثبت، تابع چگالی احتمال پیشگوی بیزی را به یک معیار مناسب برای استنباط آماری در خصوص مقادیر k -رکورد آینده تبدیل کرده است. پیشگویی بیزی $R_{m(k)}$ تحت تابع زیان متقارن توان دوم خطا، $\hat{R}_{m(k)}^S$ ، میانگین تابع چگالی احتمال پیشگوی بیزی است که به فرم

$$\hat{R}_{m(k)}^S = E(R_{m(k)}|r) = \int_0^\infty u f_{R_{m(k)}}^*(u|r) du = \sqrt{\frac{\beta + t}{k}} \frac{\Gamma(m + \frac{1}{\gamma})\Gamma(N + \alpha - \frac{1}{\gamma})}{\Gamma(m)\Gamma(N + \alpha)}, \quad (9.2)$$

محاسبه می‌شود. میانگین توان دوم خطا پیشگویی^۲ (MSPE) برای $\hat{R}_{m(k)}^S$ ، به کمک روابط (۷.۲) و (۹.۲) با اندکی محاسبات برابر است با

$$\begin{aligned} E(\hat{R}_{m(k)}^S - R_{m(k)})^2 &= E(\hat{R}_{m(k)}^S)^2 + E(R_{m(k)})^2 - 2E(\hat{R}_{m(k)}^S)E(R_{m(k)}) \\ &= \frac{\Gamma^2(m + \frac{1}{\gamma})\Gamma^2(N + \alpha - \frac{1}{\gamma})}{k\Gamma^2(m)\Gamma^2(N + \alpha)} (\beta + E(T)) \\ &\quad + \frac{m}{k\lambda} - \frac{2\Gamma^2(m + \frac{1}{\gamma})\Gamma(N + \alpha - \frac{1}{\gamma})}{k\lambda^{\frac{1}{\gamma}}\Gamma^2(m)\Gamma(N + \alpha)} E\sqrt{\beta + t} \end{aligned}$$

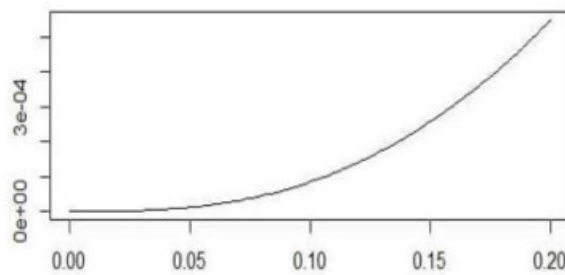
با بهره‌گیری از بسط دوجمله‌ای، می‌توان آن را به صورت زیر دوباره نویسی کرد

$$\begin{aligned} E(\hat{R}_{m(k)}^S - R_{m(k)})^2 &= \frac{1}{k\lambda} \left(\frac{\Gamma^2(m + \frac{1}{\gamma})\Gamma^2(N + \alpha - \frac{1}{\gamma})}{k\Gamma^2(m)\Gamma^2(N + \alpha)} \right) (\beta\lambda + N) \\ &\quad + m - \frac{2(\beta\lambda)^{\frac{1}{\gamma}}\Gamma^2(m + \frac{1}{\gamma})\Gamma(N + \alpha - \frac{1}{\gamma})}{k\lambda^{\frac{1}{\gamma}}\Gamma^2(m)\Gamma(\alpha)\Gamma(N + \alpha)} \sum_{l=0}^{\infty} \binom{\frac{1}{\gamma}}{l} \frac{\Gamma(\alpha + l)}{\beta^l}. \quad (10.2) \end{aligned}$$

این تابع با فرض ثابت نگه داشتن سایر پارامترها، تابعی یکنوا صعودی از m است. شکل ۱ را ببینید.

¹Bayesian Predictive Density Function

²Mean square prediction error



شکل ۱: نمودار میانگین توان دوم خطا پیشگویی $\hat{R}_{m(k)}^S$ نسبت به m برای $\lambda = \alpha = \beta = 2$.

۲.۲ پیشگویی بیزی نقطه‌ای $R_{m(k)}$ تحت تابع زیان لاینکس

تابع زیان نامتقارن لاینکس^۱ به فرم $L_2(\lambda, \hat{\lambda}) = b(e^{a(\hat{\lambda}-\lambda)} - a(\hat{\lambda} - \lambda) - 1)^2$ به فرم $a \neq 0$ شکل تابع زیان را مشخص می‌کند و $b > 0$ نیز برای مقیاس تابع زیان به کار می‌رود. در این بخش، پیشگویی بیزی نقطه‌ای $R_{m(k)}$ تحت تابع زیان لاینکس، $\hat{R}_{m(k)}^L$ ، با استفاده از داده‌های حاصل از طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی ارائه می‌شود. پیش‌گویی بیزی نقطه‌ای برای $\hat{R}_{m(k)}^L$ به فرم $-\frac{1}{a} \ln E(e^{-aR_{m(k)}}|r)$ است. با انجام اندکی محاسبات می‌توان دید

$$\hat{R}_{m(k)}^L = -\frac{1}{a} \ln \int_0^\infty e^{-au} f_{R_{m(k)}}^*(u|r) du = -\frac{1}{a} \ln \sum_{i=0}^\infty \frac{\left(-a\sqrt{\frac{\beta+t}{k}}\right)^i}{i!} \frac{\Gamma(m+\frac{1}{\gamma})\Gamma(N+\alpha-\frac{1}{\gamma})}{\Gamma(m)\Gamma(N+\alpha)}.$$

میانگین توان دوم خطا پیشگویی برای $\hat{R}_{m(k)}^L$ ، فرم بسته‌ای ندارد. برای حل آن باید به سراغ روش‌های عددی رفت.

۳ پیشگویی بیزی بازه‌ای مقادیر k -رکورد

پیشگویی‌های نقطه‌ای، اطلاعات زیادی درباره دقت پیشگویی ارائه نمی‌دهند. این نوع از پیشگویی‌ها یا صد درصد درست هستند و یا صد درصد نادرست. لذا در بسیاری از موارد، به‌دست آوردن یک بازه پیشگویی که احتمال قرار گرفتن متغیر تصادفی قابل پیشگویی در آن زیاد است کفایت می‌کند. در این بخش، قصد ایجاد یک بازه پیشگویی $(1-\gamma)$ درصدی برای $R_{m(k)}$ به روش تابع بقا و روش HPD با استفاده از داده‌های حاصل از طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی را داریم.

۱.۳ پیشگویی بازه‌ای به روش تابع بقا

برای ایجاد یک بازه پیشگویی بیزی برای $R_{m(k)}$ به روش تابع بقا، به کمک تابع بقا پیشگوی، $F_{R_{m(k)}}^*(u|r)$ ، داریم

$$F_{R_{m(k)}}^*(u|r) = P(R_{m(k)} > u) = \int_0^\infty \frac{2}{B(m, N+\alpha)y} p(y)^m (1-p(y))^{N+\alpha} dy$$

حدود بازه پیشگویی بیزی $(1-\gamma)$ درصدی، از حل دو معادله زیر حاصل می‌شود

$$F_{R_{m(k)}}^*(L(r)) = 1 - \frac{\gamma}{2}, \quad F_{R_{m(k)}}^*(U(r)) = \frac{\gamma}{2} \quad (1.3)$$

¹Linex Loss Function

که در آن‌ها $L(r)$ و $U(r)$ به ترتیب حدود پائین و بالای بازه هستند. از رابطه (۱.۳) داریم

$$\int_0^{L(r)} \frac{2}{B(m, N + \alpha)y} p(y)^m (1 - p(y))^{N + \alpha} dy = \frac{\gamma}{2}$$

با استفاده از تغییر متغیر $p(y) = v$ و محاسباتی نه چندان دشوار، رابطه اخیر را می‌توان به فرم

$$\int_0^{(1 + \frac{\beta+t}{kL(r)})^{-1}} \frac{1}{B(m, N + \alpha)} v^{m-1} (1 - v)^{N + \alpha - 1} dv = \frac{\gamma}{2} \quad (2.3)$$

دوباره نویسی کرد. با استفاده از رابطه (۲.۳) خواهیم داشت، $L(r) = \sqrt{\frac{\beta+t}{k} \frac{B_{1-\frac{\gamma}{2}}(m, N + \alpha)}{1 - B_{1-\frac{\gamma}{2}}(m, N + \alpha)}}$ به‌طور مشابه

نیز داریم: $U(r) = \sqrt{\frac{\beta+t}{k} \frac{B_{\frac{\gamma}{2}}(m, N + \alpha)}{1 - B_{\frac{\gamma}{2}}(m, N + \alpha)}}$. مقصود از نماد $B_{\gamma}(m, N + \alpha)$ ، چندک مرتبه γ ام توزیع بتا با پارامترهای m و $N + \alpha$ است. در نتیجه بازه پیشگویی بیزی برای $R_{m(k)}$ عبارت است از

$$\left[\sqrt{\frac{1}{k} \frac{B_{1-\frac{\gamma}{2}}(m, N + \alpha)}{1 - B_{1-\frac{\gamma}{2}}(m, N + \alpha)}}, \sqrt{\frac{1}{k} \frac{B_{\frac{\gamma}{2}}(m, N + \alpha)}{1 - B_{\frac{\gamma}{2}}(m, N + \alpha)}} \right]. \quad (3.3)$$

طول این بازه تصادفی تابعی از T است. متوسط طول این بازه برابر است با

$$\left(\sqrt{\frac{1}{k} \frac{B_{\frac{\gamma}{2}}(m, N + \alpha)}{1 - B_{\frac{\gamma}{2}}(m, N + \alpha)}} - \sqrt{\frac{1}{k} \frac{B_{1-\frac{\gamma}{2}}(m, N + \alpha)}{1 - B_{1-\frac{\gamma}{2}}(m, N + \alpha)}} \right) \sum_{l=0}^{\infty} \left(\frac{1}{l} \right) \frac{\Gamma(\alpha + l)}{\Gamma(\alpha) \beta^{l - \frac{1}{2}}}.$$

۲.۳ پیشگویی بازه‌ای به روش HPD

شکل ۲، نمودار تابع چگالی احتمال پیشگوی بیزی $f_{R_{m(k)}}^*(u|r)$ و هیستوگرام مقادیر تولید شده از آن بر اساس نمونه‌های شبیه‌سازی شده را نشان می‌دهد. همان‌طور که ملاحظه می‌شود، تابع چگالی احتمال پیشگوی بیزی تک مدی است. لذا به منظور ایجاد یک بازه پیشگویی بیزی به روش HPD (بازه دارای کوتاهترین متوسط طول) برای $R_{m(k)}$ تحت داده‌های حاصل از طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی، نیاز به حل همزمان دو معادله زیر داریم

$$\int_{\xi_2}^{\xi_1} f_{R_{m(k)}}^*(u|r) du = 1 - \gamma, \quad f_{R_{m(k)}}^*(\xi_1|r) = f_{R_{m(k)}}^*(\xi_2|r), \quad 0 < \gamma < 1$$

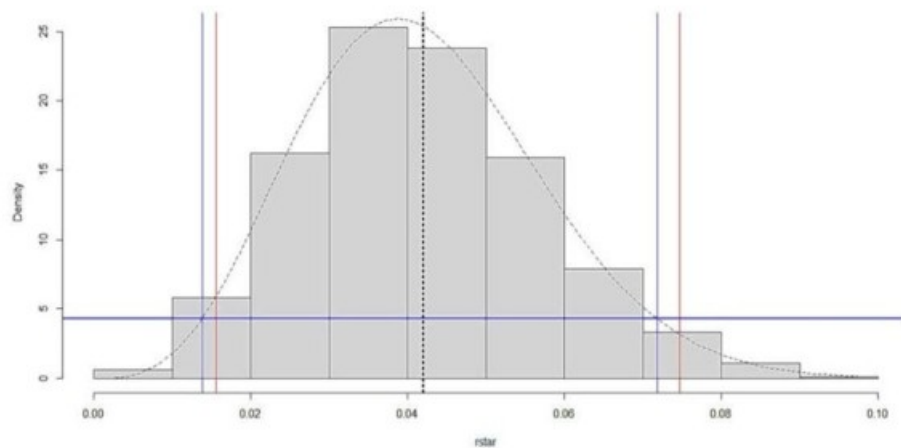
یا به‌طور معادل

$$\int_{\xi_2}^{\xi_1} \frac{2}{B(m, N + \alpha)u} p(u)^m (1 - p(u))^{N + \alpha} du = 1 - \gamma, \quad \left(\frac{\xi_1}{\xi_2} \right)^{2m-1} = \left(\frac{\beta + t + k\xi_1^2}{\beta + t + k\xi_2^2} \right)^{m+N+\alpha} \quad (4.3)$$

به (ξ_1, ξ_2) یک بازه پیشگویی بیزی $(1 - \gamma)$ درصدی برای $R_{m(k)}$ به روش HPD می‌گویند.

۴ مطالعه شبیه‌سازی

در این بخش، ارزیابی عملکرد طرح معرفی شده در پیشگویی بیزی مقادیر k -رکورد آینده با استفاده از شبیه‌سازی مورد مطالعه قرار می‌گیرد. با این هدف، برای ابرپارامترهای (α, β) در رابطه (۴.۲)، از مقادیر $(1/5, 1/5)$ ، $(0.5, 0.5)$ و



شکل ۲: منحنی تابع چگالی احتمال پیشگوی بیزی $f_{R_{m(k)}}^*(u|r)$ و هیستوگرام مقادیر تولید شده از آن بر اساس نمونه‌های شبیه سازی شده با ۱۰۰۰ بار تکرار

(α, β) برای حالت‌های آگاهی‌بخش و ناآگاهی‌بخش استفاده شده است. با استفاده از نرم‌افزار R، بر اساس یک نمونه تصادفی شبیه سازی شده با ۱۰۰۰ بار تکرار، پارامتر مجهول λ در مدل رایلی، تحت تابع زیان توان دوم خطا، برابر با ۰/۴۲ تخمین زده شده است. در شکل ۲، خط‌چین عمودی مشکی، برآورد بیز پارامتر λ تحت توابع زیان توان دوم خطا و لاینکس است. این برآوردها چون نزدیک هم هستند، بر روی یک خط افتاده اند. معیارهای زیر برآورد میانگین توان دوم خطا پیشگویی تحت توابع زیان توان دوم خطا و لاینکس در ۱۰۰۰ بار تکرار را نمایش می‌دهد:

$$MSPE_1(\hat{Y}_i, Y_i) = \frac{1}{1000} \sum_{i=1}^{1000} L_1(\lambda, \hat{\lambda}) \quad , \quad MSPE_2(\hat{Y}_i, Y_i) = \frac{1}{1000} \sum_{i=1}^{1000} L_2(\lambda, \hat{\lambda})$$

کوچک بودن معیار MSPE بیانگر دقت بیشتر و عملکرد بهتر است. به‌ازای مقادیر $k = 1, 2, 3$ ، $m = 1, 2, 3$ و $a = -1, 1$ نتایج شبیه سازی مربوط به مقادیر پیشگویی $R_{m(k)}$ و میانگین توان دوم خطا پیشگویی متناظر آن تحت توابع زیان توان دوم خطا و لاینکس در جدول ۱ گزارش شده است. دیده می‌شود، با فرض ثابت نگه داشتن سایر پارامترها و افزایش پارامتر m ، برآورد میانگین توان دوم خطا پیشگویی نیز افزایش می‌یابد. این نتیجه در شکل ۱ برای $\alpha = \beta = 2$ در قالب نمودار نیز آورده شده است. همچنین در جدول ۱، با استفاده از معیار MSPE می‌توان دید که، برآورد m امین k -رکورد آینده $R_{m(k)}$ تحت تابع زیان توان دوم خطا، کمتر از برآورد آن تحت تابع زیان لاینکس و در نتیجه بهتر از آن است.

جدول ۱: پیشگویی $R_{m(k)}$ و MSPE متناظرش (در داخل پرانتز) تحت توابع زیان توان دوم خطا و لاینکس

$\alpha = \beta = 0$			$\alpha = \beta = 0.5$			$\alpha = \beta = 1.5$			k	m
Sel	Linex($a = -1$)	Linex($a = 1$)	Sel	Linex($a = -1$)	Linex($a = 1$)	Sel	Linex($a = -1$)	Linex($a = 1$)		
۰/۳۹۵(۰/۷۱۶۹)	۰/۳۹۳(۰/۷۱۷۲)	۰/۳۹۱(۰/۷۱۷۶)	۰/۳۹۵(۰/۷۱۶۸)	۰/۳۹۱(۰/۷۱۷۱)	۰/۳۹۳(۰/۷۱۷۵)	۰/۳۹۶(۰/۷۱۶۷)	۰/۳۹۴(۰/۷۱۷۰)	۰/۳۹۲(۰/۷۱۷۴)	۱	۱
۰/۱۹۷(۰/۱۷۹۳)	۰/۹۸۱(۱/۶۶۶۱)	۰/۵۳۶(۱/۶۶۸۳)	۰/۱۹۸(۰/۱۷۹۱)	۰/۹۰۷(۱/۵۳۶۸)	۰/۵۰۷۸(۱/۶۳۱۴)	۰/۱۹۸(۰/۱۷۹۱)	۰/۷۲۹۶(۱/۵۱۸۵)	۰/۴۱۷۷(۱/۵۶۰۰)	۲	۱
۰/۱۳۱(۰/۰۷۹۷)	۰/۹۸۲(۰/۶۶۲۳)	۰/۸۵۵(۱/۶۸۱۴)	۰/۱۳۲(۰/۰۷۹۶)	۰/۵۴۰(۰/۶۶۱۴)	۰/۸۲۴۶(۱/۵۸۰۴)	۰/۱۳۲(۰/۰۷۹۵)	۰/۹۲۳۱(۰/۴۱۸۲)	۰/۵۱۸۸(۱/۴۶۰۱)	۳	۱
۰/۸۳۸(۱/۳۵۱۲)	۰/۸۳۴(۱/۴۵۲۲)	۰/۸۲۹(۱/۶۶۳۴)	۰/۸۳۹(۱/۲۵۱۰)	۰/۸۳۴(۱/۳۵۲۱)	۰/۸۳۰(۱/۵۵۳۳)	۰/۸۴۰(۱/۲۵۰۸)	۰/۸۳۵(۱/۲۵۱۹)	۰/۸۳۰(۱/۵۵۳۱)	۱	۲
۰/۴۱۹(۰/۳۸۷۹)	۰/۴۱۸(۱/۷۶۷۹)	۰/۴۱۷(۱/۷۷۸۳)	۰/۴۱۹(۰/۳۸۷۷)	۰/۴۱۸(۱/۶۸۷۸)	۰/۴۱۷(۱/۶۸۸۰)	۰/۴۲۰(۰/۳۸۷۳)	۰/۵۴۰(۱/۶۸۷۶)	۰/۴۱۸(۱/۶۸۷۹)	۲	۲
۰/۲۷۹(۰/۱۷۶۶)	۰/۳۱۱(۰/۹۵۸۳)	۰/۵۱۵(۱/۶۶۹۴)	۰/۲۷۹(۰/۱۷۲۴)	۰/۴۴۱(۰/۹۵۵۰)	۰/۲۷۵(۱/۶۵۵۲)	۰/۲۸۰(۰/۱۷۲۲)	۰/۲۹۴۹(۰/۸۶۳۶)	۰/۲۰۱(۱/۴۸۴۵)	۳	۲
۰/۱۲۸۴(۱/۳۵۱۷)	۰/۱۲۷۷(۱/۵۴۲۹)	۰/۱۲۷۰(۱/۷۶۵۰)	۰/۱۲۸۵(۱/۳۵۰۶)	۰/۱۲۷۷(۱/۳۵۲۷)	۰/۱۲۷۰(۱/۶۵۴۹)	۰/۱۲۸۵(۱/۳۵۰۳)	۰/۱۲۷۸(۱/۳۵۲۵)	۰/۱۲۷۱(۱/۶۵۴۷)	۱	۳
۰/۶۴۲(۰/۵۸۷۹)	۰/۶۴۰(۱/۷۸۸۵)	۰/۶۳۸(۱/۸۸۸۲)	۰/۶۴۲(۰/۵۸۷۶)	۰/۶۴۰(۱/۷۸۷۸)	۰/۶۳۹(۱/۸۵۸۱)	۰/۶۴۳(۰/۵۸۷۴)	۰/۶۴۱(۱/۷۷۷۳)	۰/۶۴۰(۱/۷۸۸۰)	۲	۳
۰/۴۲۸(۰/۲۶۱۹)	۰/۴۲۷(۱/۰۶۱۲)	۰/۴۲۶(۱/۶۶۷۳)	۰/۴۲۸(۰/۲۶۱۳)	۰/۴۲۷(۰/۹۶۱۲)	۰/۴۲۷(۱/۶۶۱۳)	۰/۴۲۹(۰/۲۶۱۰)	۰/۴۲۸(۰/۹۶۱۱)	۰/۴۲۷(۱/۵۶۱۲)	۳	۳

جدول ۲ برآورد مقدار احتمال پوشش به همراه برآورد میانگین طول بازه پیشگویی بیزی با ضریب پیشگویی ۰/۹۵ برای $R_{m(k)}$ را نشان می‌دهد. برآورد مقدار احتمال پوشش، نسبت تعداد دفعاتی است که در ۱۰۰۰ بار تکرار، مقادیر واقعی $R_{m(k)}$ آینده در بازه‌های پیشگویی مربوطه قرار گیرند و به‌صورت $CP = \frac{1}{1000} \sum_{r=1}^{1000} I(L_r < Y_i < U_r)$ تعریف می‌شود که در آن مقصود از $I(\cdot)$ ، تابع نشانگر است. هم چنین معیار $EL = \frac{1}{1000} \sum_{r=1}^{1000} (L_r - U_r)$

میانگین طول بازه‌های پیشگویی است. نتایج جدول ۲ بیانگر آن است که، طول بازه‌های پیشگویی به‌دست آمده به روش HPD به‌طور متوسط کوتاه‌تر از طول بازه‌های پیشگویی به‌دست آمده به روش تابع بقا است. احتمال پوشش نیز در بسیاری از حالات تفاوت چندانی با ضریب پیش‌گویی ۰/۹۵ ندارد و از این رویکرد نیز خوب عمل کرده است. در شکل ۲، خطوط عمودی قرمز بازه پیشگویی تحت روش تابع بقا با دم‌های برابر است. خطوط عمودی آبی نیز بازه اطمینان تحت روش HPD و خط افقی آبی، چگالی در دو نقطه است که چون نزدیک به هم هستند روی یک خط قرار گرفته است. شکل ۲ نتایج فوق را تایید می‌کند.

جدول ۲: برآورد میانگین طول بازه پیشگویی بیزی و احتمال پوشش ($R_{m(k)}$ در داخل پرانتز) $\alpha = \beta = 0$

$\alpha = \beta = 0$		$\alpha = \beta = 0.5$		$\alpha = \beta = 1.5$		k	m
روش HPD	روش تابع بقا	روش HPD	روش تابع بقا	روش HPD	روش تابع بقا		
۰/۰۷۴۴(۰/۹۷۶)	۰/۰۷۸۵(۰/۹۷۹)	۰/۰۷۴۵(۰/۹۸۳)	۰/۰۷۸۵(۰/۹۸۹)	۰/۰۷۴۷(۰/۹۸۵)	۰/۰۷۸۷(۰/۹۶۱)	۱	۱
۰/۰۳۷۲(۰/۹۹۵)	۰/۰۳۹۳(۰/۹۸۴)	۰/۰۳۷۲(۰/۹۹۴)	۰/۰۳۹۲(۰/۹۹۹)	۰/۰۳۷۵(۰/۹۶۸)	۰/۰۳۹۵(۰/۹۰۶)	۲	
۰/۰۲۴۷(۰/۹۷۷)	۰/۰۲۶۱(۰/۹۶۲)	۰/۰۲۴۸(۰/۹۹۵)	۰/۰۲۶۱(۰/۹۹۳)	۰/۰۲۵۱(۰/۹۲۹)	۰/۰۲۶۴(۰/۹۵۵)	۳	
۰/۰۱۱۴۹(۰/۹۲۳)	۰/۰۱۱۷۹(۰/۹۷۱)	۰/۰۱۱۴۹(۰/۹۷۵)	۰/۰۱۱۷۹(۰/۹۹۵)	۰/۰۱۱۵۱(۰/۹۹۵)	۰/۰۱۱۸۰(۰/۹۷۹)	۱	۲
۰/۰۵۷۳(۰/۹۵۴)	۰/۰۵۸۹(۰/۹۰۴)	۰/۰۵۷۴(۰/۹۹۵)	۰/۰۵۸۹(۰/۹۹۹)	۰/۰۵۷۶(۰/۹۹۲)	۰/۰۵۹۱(۰/۹۱۲)	۲	
۰/۰۳۸۳(۰/۹۳۲)	۰/۰۳۹۴(۰/۹۰۴)	۰/۰۳۸۳(۰/۹۹۳)	۰/۰۳۹۳(۰/۹۷۶)	۰/۰۳۸۴(۰/۹۸۱)	۰/۰۳۹۴(۰/۹۴۳)	۳	
۰/۰۱۴۴۱(۰/۹۴۲)	۰/۰۱۴۶۹(۰/۹۶۲)	۰/۰۱۴۴۰(۰/۹۸۵)	۰/۰۱۴۶۹(۰/۹۶۹)	۰/۰۱۴۴۲(۰/۹۷۹)	۰/۰۱۴۷۱(۰/۹۸۷)	۱	۳
۰/۰۷۲۱(۰/۹۵۴)	۰/۰۷۳۴(۰/۹۰۶)	۰/۰۷۲۰(۰/۹۹۹)	۰/۰۷۳۴(۰/۹۵۷)	۰/۰۷۲۱(۰/۹۹۲)	۰/۰۷۳۶(۰/۹۱۱)	۲	
۰/۰۴۸۰(۰/۹۳۵)	۰/۰۴۸۹(۰/۹۹۴)	۰/۰۴۸۰(۰/۱۰۰)	۰/۰۴۸۹(۰/۹۱۳)	۰/۰۴۸۱(۰/۹۷۹)	۰/۰۴۹۱(۰/۹۹۱)	۳	

۵ مثال واقعی مرتبط با میزان سرعت وزش باد

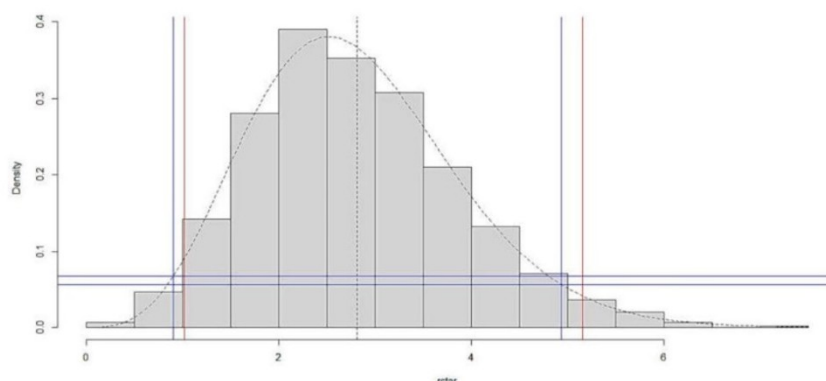
درایه‌های بردار مشاهدات $r = (3.5, 6.94, 6.39, 3.61, 6.39, 6.67)^T$ ، رکوردهای به‌دست آمده بر اساس دیگرام RRSS در رابطه (۲.۱) و داده‌های واقعی مرتبط با میزان سرعت وزش باد کشور کانادا، استان Brunswick new و ایستگاه Bas caraquej بین سال‌های ۲۰۱۵ تا ۲۰۲۰ میلادی می‌باشد که از سایت کمپسیون بین المللی اقیانوس شناسی IOC استخراج و برای ارزیابی مورد استفاده قرار گرفته است. در بررسی آزمون نیکویی برازش توزیع رایلی بر این داده‌ها، p -مقدار آزمون اندرسون-دارلینگ برابر ۰/۹۷۱ محاسبه شده است. بیشتر بودن این مقدار احتمال از مقدار خطای نوع اول ($\gamma = 0.05$) بیانگر مناسب بودن برازش توزیع رایلی بر روی این داده‌ها است. شکل ۳ نمودارهای احتمال $P-P$ و احتمال تجمع $Q-Q$ را نشان می‌دهد. در بررسی شکل ۳، روند خطی دیده می‌شود و این یعنی توزیع رایلی به خوبی بر روی داده‌های میزان سرعت وزش باد برازش یافته است. پارامتر مجهول λ در مدل رایلی تحت توابع زیان توان دوم خطا و لاینکس نیز تخمین زده شده است. بر پایه این داده‌ها، بازه‌های پیشگویی به‌دست آمده به روش تابع بقا و روش HPD به ترتیب $(1.0145, 5.1569)$ و $(0.9016, 4.9316)$ است. پر واضح است که طول بازه پیشگویی به‌دست آمده به روش HPD کوتاه‌تر از طول بازه پیشگویی به روش تابع بقا است. مقدار MSPE تحت توابع زیان توان دوم خطا و لاینکس نیز محاسبه شده است. نتایج نشان داد که برآوردگر بیزی تحت تابع زیان توان دوم خطا بهتر از برآوردگر بیزی تحت تابع زیان لاینکس است. نتایج بر اساس داده‌های واقعی به‌صورت گرافیکی در شکل ۴ و جدول ۳ دیده می‌شود. نتایج به‌دست آمده در بخش داده‌های واقعی، نتایج حاصل از مطالعه شبیه‌سازی را تایید می‌کند.

جدول ۳: مقادیر به‌دست آمده بر اساس داده‌های واقعی مرتبط با میزان سرعت وزش باد

$MSPE_{(a=-1)}^L$	$MSPE_{(a=1)}^L$	$MSPE^{(S)}$	طول بازه (تابع بقا)	طول بازه (HPD)	p - مقدار	$\widehat{\lambda^{(L)}}$	$\widehat{\lambda^{(S)}}$
۴/۰۵۵	۳/۱۱۳	۲/۳۹۳	۴/۱۴۲۴	۴/۰۳	۰/۰۸۶	۲/۸۰۸	۲/۷۹۶



شکل ۳: نمودارهای احتمال P-P و Q-Q بر اساس داده‌های واقعی مرتبط با میزان سرعت وزش باد



شکل ۴: منحنی تابع چگالی احتمال بیزی پیشگوی و هیستوگرام مقادیر تولید شده از آن بر اساس داده‌های واقعی

۶ نتیجه‌گیری

یکی از مسائل مهم در استنباط آماری، به‌دست آوردن اطلاعات از طریق فرایندی است که بتوان با هزینه کمتر و کارایی بیشتر، وضع سیستم‌های آینده جامعه را که با گذشت زمان در حال تغییر هستند پیشگویی و تحلیل کرد. افزایش تعداد داده‌های نمونه همواره راهی است که می‌تواند به آماردانان کمک کند تا شواهد آماری را در جهت استنباط‌های دقیق‌تر در جامعه هدف، افزایش دهد. به هر حال در بسیاری از آزمایش‌های کاربردی میسر نیست که تعداد داده‌های نمونه را افزایش دهیم و عمدتاً تصمیم‌گیری بر اساس تعداد محدود مشاهدات صورت می‌گیرد. در این صورت محقق باید سعی کند تا آنجا که مقدور است اطلاعات دیگری درباره توزیع نمونه‌گیری مشاهدات خود و یا استفاده از اطلاعات قبلی (روش بیزی) راجع به متغیرهای تصادفی آینده به‌دست آورد و یا طرح نمونه‌گیری مناسبی را برگزیند تا بتواند با این روش و با درصد خطای کمتری تصمیم‌گیری کند. در این مقاله طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی مورد استفاده قرار گرفت. با استفاده از داده‌های رکوردی حاصل از این طرح و در نظر گرفتن مدل رایلی، پیشگویی بیزی مقادیر k - رکورد آینده به کمک شبیه‌سازی و داده‌های واقعی مرتبط با میزان سرعت وزش باد ارزیابی شد. بررسی‌ها نشان داد نتایج پیشگویی مقادیر m - امین k - رکورد حاصل از یک دنباله آینده بر اساس تابع زیان توان دوم خطا دقیق‌تر از نتایج به‌دست آمده بر اساس تابع زیان لاینکس است. هم‌چنین طول بازه‌های پیشگویی به‌دست آمده برای مقادیر m - امین k - رکورد حاصل از یک دنباله آینده به روش HPD کوتاه‌تر و دارای احتمال پوشش بیشتر نسبت به بازه‌های پیشگویی به‌دست آمده بر اساس تابع بقا است. در بررسی این نوع پژوهش از دیدگاه آمار بیزی، خوب است که طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی به‌کار گرفته شود.

مراجع

- [۱] افهمی، ب.، مددی، م.، رضاپور، م.، ۱۳۹۴. آزمون نیکویی برازش بر مبنای آنتروپی رکورد از توزیع پارتو تعمیم یافته، مجله علوم آماری، ۹ (۱)، صص. ۴۳-۶۰.
- [۲] پاک، ع.، جعفری، ع. ا.، محمودی، م. ر.، ۱۴۰۱. استنباط درباره شاخص عملکرد طول عمر در توزیع نمایی دو پارامتری: داده‌های سانسور شده و رکورد، پژوهش‌های ریاضی، ۸ (۴)، صص. ۱۹-۴۳.
- [۳] گل‌زاده گروی، ا.، نصیری، پ.، صالحی، م.، ۱۴۰۰. مقایسه برآوردها و پیشگویی‌های بیزی تجربی بر اساس طرح نمونه‌گیری معکوس با طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی، مجله علوم آماری، ۱۵ (۱)، صص. ۱۹۳-۲۱۸.
- [۴] گل‌زاده گروی، ا.، ۱۴۰۲. برآورد پارامتر مدل در تغییرات اقلیمی با بکارگیری داده‌های رکوردی حاصل از طرح نمونه‌گیری مجموعه رتبه‌دار رکوردی، نشریه پژوهش‌های اقلیم شناسی، ۱۴۰۲ (۵۳)، صص. ۵۱-۵۶.
- [5] Ahmadi, J., Doostparast, M., and Parsian, A., 2005. Estimation and Prediction in a two-parameter exponential distribution based on k-record values under LINEX loss function. *Comm. Statist. Theory Methods.*, 34 , pp. 795-805.
- [6] Ahmadi, J., Jafari Jozani, M., Marchand, E., and Parsian, A., 2009. Prediction of k-records from a general class of distributions under balanced type loss functions. *Metrika.*, 70 , pp. 19-33.
- [7] Ahmadi, J., MirMostafaei, S.M.T.K., and Balakrishnan, N., 2010. Bayesian Prediction of order statistics based on k-record values from exponential. *Statistics*, 34, pp. 795-805.
- [8] Ahmadi, J., Mirmostafaei, S.M.T.K., and Balakrishnan, N., 2012. Bayesian prediction of k-record values based on progressively censored data from exponential distribution. *Journal of Statistical computation and simulation*, 82, pp. 51-62.
- [9] Ahsanullah, M., 1995. *Record Statistics*. Nova Science Publishers Inc, Commack, New York.
- [10] Aitchison, J., and Dunsmore, I.R., 1975. *Statistical prediction analysis*. Cambridge University Press, Cambridge.
- [11] Arnold, B.C., Balakrishnan, N., and Nagaraja, H.N., 1998. *Records*, John Wiley & Sons, New York.
- [12] Asgharzadeh, A., Valiollahi, R., Fallah, A., and Nadarajah, S., 2018. Bayesian inference and prediction for normal distribution based on records. *Statistics*, 78, pp. 15-36.
- [13] Doostparast, M., and Doostparast, M., 2018. Prediction of corrossions in Gas and Oil pipelines based on the theory of records. *arXiv preprint arXiv:1801.00959*.
- [14] Eskandarzadeh, M., Tahmasebi, S., and Afshari, M., 2016. Information measures for Record ranked set samples. *Ciencia e Natura.*, 38 , pp. 554-563.
- [15] Ghafouri, S., Habibirad, A., Yousefzadeh, F., and Pakyari, R., 2022. Non-Bayesian Estimation and Prediction under Weibull Interval Censored Data. *Journal of Statistical Research of Iran JSRI*, 17 pp. 171-190.
- [16] Golzade Gervi, E., 2024. Inference for the Pareto Type-I distribution using upper record ranked set sampling scheme. *Int. J. Nonlinear Anal. Appl.*, 15 , pp. 115-123.

- [17] Golzade Gervi, E., Nasiri, P., and Salehi, M., 2019. Comparision of record tanked set sampling and ordinary records in prediction of future record statistics from an exponential distribution. *Statistical Research and Training Center*, 16, pp. 73-99.
- [18] Golzade Gervi, E., Nasiri, P., and Salehi, M., 2021. An overview of Bayesian prediction of future record statistics using upper record ranked set sampling scheme. *Int. J. Nonlinear Anal. Appl.*, 12, pp. 493-507.
- [19] Gulati, S., and Padgett, W.J., 2003. *Parametric and Non-Parametric Inference from Record-breaking Data*. Lecture Notes in Statistics., 172 Springer, New York.
- [20] Klimczak, M., 2006. Prediction of k th record, *Statist. Probab. Lett.*, 76, pp. 117-127.
- [21] Mojiri, A., Waghei, Y., Nili Sani, H.R., and Mohtashami Borzadaran, G.R., 2023. Non-stationary spatial autoregressive modeling for the prediction of lattice data. *Communications in Statistics-Simulation and Computation*, 52, pp. 5714-5726.
- [22] Moradi, M., Jabbari Nooghabi, M., and Rounaghi, M.M., 2021. Investigation of fractal market hypothesis and forecasting time series stock returns for Tehran Stock Exchange and London Stock Exchange. *International Journal of Finance Economics*, 26, pp. 662-678.
- [23] Nevzorov, V., 2001. Records. Mathematical Theory. *Translation of Mathematical Monographs*, 194, American Mathematical Society, Providence, RI.
- [24] Paul, J., and Thomas, P.Y., 2017. Concomitant Record ranked set sampling. *Communications in Statistics -Theory and Methods*, 46, pp. 9518-9540.
- [25] Porschan, F., 1963. Theoretical explanation of observed decreasing failure rate. *Technometrics*, 5, pp. 375-383.
- [26] Saadati Nik, A., Asgharzadeh, A., and Raqab, M.Z., 2021. Estimation and prediction for a new Pareto-type distribution under progressive type-II censoring. *Mathematics and Computers in Simulation*, 190, pp. 508-530.
- [27] Safaryian, A., Arashi, M., and Arabi, R., 2019. Improved Estimators for stress-strength Reliability using Record ranked set sampling scheme. *Communications in statistics-simulation and computation*, 33. doi:10.1080/03610918.2018.1468451.
- [28] Salehi, M., and Ahmadi, J., 2014. Record ranked set sampling scheme. *Metron*, 72, pp. 351-365.
- [29] Salehi, M., and Ahmadi, J., 2015. Estimation of stress-strength Reliability using Record ranked set sampling scheme from the Exponential distribution. *Filomat.*, 29, pp. 1149-1162.
- [30] Salehi, M., Ahmadi, J., and Balakrishnan, N., 2015. Prediction of order Statistics and Record values based on ordered Ranked set sampling. *Statistical Computation and Simulation*, 85, pp. 77-88.
- [31] Valiollahi, R., Asgharzadeh, A., and HKT, Ng., 2022. Prediction for Lindley distribution based on type-II right censored samples. *Journal of the Iranian Statistical Society*, 16, pp. 1-19.