



A Nonparametric Test for Stochastic Ordering in Type I Censored Data

Amini-Seresht, E. 

Department of Statistics, Faculty of Science, Bu-Ali Sina University, Hamedan, Iran.

Corresponding author: E. Amini-Seresht, e.amini@basu.ac.ir

Received: 6/8/2025 Revised: 18/11/2025 Accepted and Published Online: 20/11/2025.

Introduction

This paper develops a nonparametric test for the usual stochastic order $F \prec F_0$ under Type I censoring. We extend Banerjee (2008) test statistic to handle incomplete data, using Simpson's rule for weight optimization and bootstrap methods for empirical distribution estimation. The test's power is evaluated via Monte Carlo simulations against the Lehmann alternative model.

Material and Methods

The proposed test statistic $\Delta_n^{(cens)}$ is constructed using only uncensored observations. Weight coefficients ℓ_i are derived via Simpson's rule. The empirical null distribution is approximated using bootstrap resampling with 10,000 iterations. Performance is assessed through extensive simulations comparing power with KS and SW tests under various censoring levels (0%–60%) and sample sizes ($n = 30, 50, 100$).

Results and Discussion

The proposed test demonstrates superior power compared to KS and SW tests, particularly at low censoring levels (10%–20%). For example, at $n = 50$ and 15% censoring, the proposed test achieves 0.75 power, outperforming KS (0.60) and SW (0.62). Power decreases uniformly with increasing censoring, but remains reasonable (e.g., 0.75 at 20% censoring). The test statistic's distribution converges to normality for $n^* > 20$, validated by KS and SW normality tests ($p - \text{value} > 0.05$). In summary, the proposed test offers

a robust and powerful nonparametric method for detecting usual stochastic order in Type I censored data, maintaining competitive performance even in the presence of incomplete information.

Conclusion

The extended Banerjee test provides a reliable nonparametric tool for detecting stochastic ordering in Type I censored data. It maintains acceptable power up to 20% censoring and shows robust performance across various scenarios, making it suitable for practical reliability and survival analysis applications.

Keywords: Usual stochastic ordering, Nonparametric tests, Monte Carlo simulation, Type I censoring.

Mathematics Subject Classification (2020): 62G10, 60E15.



آزمون ناپارامتری برای ترتیب تصادفی معمولی در داده‌های سانسور شده نوع I

ابراهیم امینی سرشت

گروه آمار، دانشکده علوم پایه، دانشگاه بوعلی سینا

نویسنده مسئول: ابراهیم امینی سرشت، e.amini@basu.ac.ir

تاریخ دریافت: ۱۴۰۴/۵/۱۵ تاریخ بازنگری: ۱۴۰۴/۸/۲۷ تاریخ پذیرش و انتشار: ۱۴۰۴/۸/۲۹

چکیده: در این مقاله، یک آزمون ناپارامتری بر پایه داده‌های ناکامل برای بررسی ترتیب تصادفی معمولی با استفاده از تعمیمی از آماره بانرجی برای داده‌های سانسور شده نوع I ارائه می‌شود. این تعمیم با به‌کارگیری ضرایب وزنی مبتنی بر قاعده سیمپسون و روش بوت‌استرپ با ۱۰۰۰۰ تکرار برای تخمین توزیع تجربی آماره پیشنهادی بهینه شده است. توزیع تجربی آماره در حضور سانسور نوع I مطالعه شده و توان آزمون با شبیه‌سازی مونت‌کارلو در مقابل مدل جایگزین لهن ارزیابی گردیده است. **واژه‌های کلیدی:** ترتیب تصادفی، آزمون‌های ناپارامتری، شبیه‌سازی مونت کارلو، سانسور نوع I. **کد موضوع بندی ریاضی (۲۰۲۰):** 60E15, 62G10.

۱ مقدمه

یکی از موضوعات کلیدی در آمار ناپارامتری، طراحی آزمون‌های برازش برای بررسی ساختار ترتیب تصادفی در توزیع‌ها است. یکی از اهداف اصلی در مطالعات قابلیت اطمینان، بررسی وجود رابطه ترتیب تصادفی معمولی بین توزیع داده‌ها و یک توزیع مشخص F است، به عبارت دیگر بررسی وجود رابطه $F \prec F$ به عنوان ترتیب تصادفی معمولی بین توزیع‌های F و F است که در آن F توزیع داده‌های مورد مطالعه می‌باشد. لازم به ذکر است، می‌گوییم متغیر تصادفی X با تابع توزیع F در ترتیب تصادفی معمولی از متغیر تصادفی Y با تابع توزیع G کوچکتر است (یا معادل آن $F \prec G$)، اگر و تنها اگر

$$F(x) \geq G(x), \quad \forall x \in \mathbb{R},$$

©نویسندگان). ناشر انجمن آمار ایران است.

این مقاله با دسترسی آزاد تحت شرایط و ضوابط (CC BY-NC 4.0) توزیع شده است.



یا بطور معادل

$$\bar{F}(x) \leq \bar{G}(x), \quad \forall x \in \mathbb{R},$$

که در آن $\bar{F} = 1 - F$ تابع بقا است. برای مثال، اگر F توزیع زمان شکست یک قطعه جدید و F_0 توزیع قطعه استاندارد باشد، رابطه $F < F_0$ نشان‌دهنده بهبود قابلیت اعتماد است. برای مطالعه انواع ترتیب های تصادفی و کاربردهای آنها، می‌توان به **شیکد و شانتیکومار (۲۰۰۷)** مراجعه کرد.

در سال‌های اخیر آزمون‌های ناپارامتری برای بررسی انواع ترتیب های تصادفی مورد توجه بسیاری از پژوهشگران قرار گرفته و در این راستا آزمون‌های ناپارامتری متعددی معرفی شده است. از جمله مطالعات پیشین در این حوزه می‌توان به **چیکاگودار و شاستر (۱۹۷۴)**، **لی و وولوف (۱۹۷۶)**، **کوچر (۱۹۷۹، ۱۹۸۱)** و **پروسچان (۱۹۸۴)** و **چن (۱۹۸۵)** اشاره کرد. به طور خاص، **بانرجی (۲۰۰۸)** یک آماره آزمون ناپارامتری برای بررسی هم‌توزیعی بودن دو متغیر تصادفی در مقابل ترتیب تصادفی معمولی معرفی کرد. این آزمون مبتنی بر یک آماره ی ساده و محاسبه‌پذیر است که به صورت ترکیب خطی از مقادیر $F_0(X_{(i)})$ در داده‌های مرتب‌شده تعریف می‌شود و از قاعده عددی سیمپسون برای تقریب انتگرال استفاده می‌کند. اگرچه آزمون بانرجی در حالت داده‌های کامل از کارایی و سادگی مناسبی برخوردار است، اما در شرایط وجود سانسور به‌طور مستقیم قابل اجرا نیست. با توجه به اهمیت داده‌های سانسور شده در مطالعات تجربی و در تحلیل بقا، توسعه روش‌های ناپارامتری برای داده‌های سانسور شده امری ضروری و مورد توجه است. همچنین لازم به ذکر است که، **حبیبی راد و همکاران (۲۰۱۱)** با تمرکز بر سانسور نوع دوم، معیاری مبتنی بر نسبت درست‌نمایی برای اندازه‌گیری شواهد آماری و تعیین نقطه توقف بهینه (r) در توزیع نمایی توسعه دادند. این رویکرد پارامتری، تعادل بین هزینه و دقت را در آزمایش‌های طول عمر بررسی می‌کند و نشان می‌دهد که در $r = 1$ برای حداکثر شواهد آماری به ازای هر شکست بهینه است. با این حال، در سانسور نوع اول (که زمان سانسور ثابت است و در کاربردهای تجربی رایج‌تر است)، روش‌های ناپارامتری برای بررسی روابط ترتیبی مانند ترتیب تصادفی معمولی ($F < G$) کمتر مورد توجه قرار گرفته‌اند. مقاله حاضر، با تعمیم آزمون ناپارامتری بانرجی به داده‌های سانسور شده نوع اول، این نیاز را برآورده می‌سازد. برخلاف رویکرد پارامتری **حبیبی راد و همکاران (۲۰۱۱)** که بر بهینه‌سازی شواهد آماری در سانسور نوع دوم متمرکز است، روش پیشنهادی در این مقاله ناپارامتری است و همچنین آزادی از پیش‌فرض‌های توزیع را حفظ می‌کند، که بر توان آزمون در شناسایی ترتیب تصادفی معمولی در شرایطی که داده‌ها ناکامل باشند، تأکید دارد.

بنابراین در این مقاله، آزمون بانرجی برای داده‌های سانسور شده نوع I در حالت تک‌نمونه‌ای توسعه داده می‌شود. در آماره ی پیشنهادی، این آزمون تنها بر پایه مشاهدات غیرسانسور شده (شکست‌یافته) تعریف می‌شود و توزیع تجربی آن تحت فرض صفر مورد مطالعه قرار می‌گیرد. علاوه بر این، توان آزمون تحت شرایط مختلف به صورت عددی ارزیابی شده و با حالت داده‌های کامل مقایسه می‌شود. تحلیل‌های ارائه‌شده حاکی از آن است که علیرغم کاهش اطلاعات ناشی از سانسور، آزمون پیشنهادی از عملکرد مناسبی برخوردار بوده و می‌تواند در تحلیل داده‌های ناقص نیز به کار گرفته شود.

در بخش ۲، آماره آزمون پیشنهادی تحت شرایط سانسور داده‌ها مورد بررسی قرار می‌گیرد. در بخش ۳، توزیع آماره پیشنهادی در حضور سانسور نوع I مطالعه می‌شود. در بخش ۴، توان آماره آزمون ارزیابی شده و عملکرد آن در حالت داده‌های کامل و سانسور شده مقایسه می‌گردد. سرانجام، بحث و نتیجه‌گیری در بخش ۵ ارائه شده است.

۲ آماره آزمون تحت سانسور نوع I

برای آزمون فرضیه $F = F_0 : H_0$ در برابر $F \prec F_0 : H_1$ ، بانرجی (۲۰۰۸) آماره ناپارامتری $\Delta_n = \sum_{i=1}^n \ell_i F_0(X_{(i)})$ را معرفی نمود، که در آن $X_{(i)}$ نشان‌دهنده آماره ترتیبی i -ام، ℓ_i ضرایب وزنی حاصل از قاعده عددی سیمپسون برای تقریب انتگرال و Δ_n ترکیب خطی از مقادیر $F_0(X_{(i)})$ محسوب می‌شود. این آماره تحت فرض صفر دارای خاصیت آزادی از توزیع بوده که امکان اجرای آزمون ناپارامتری بدون نیاز به پیش‌فرض‌های پارامتریک را فراهم می‌سازد. در بسیاری از کاربردهای عملی، به‌ویژه در تحلیل بقا، داده‌ها ممکن است به صورت سانسور شده مشاهده شوند. یکی از متداول‌ترین انواع سانسور، سانسور نوع I است که در آن همه واحدها تا زمان ثابتی τ تحت مشاهده قرار می‌گیرند و داده‌هایی که زمان وقوع آن‌ها از τ بیشتر است، به صورت سانسور شده ثبت می‌شوند. در ادامه، آماره آزمون Δ_n را برای حالت داده‌های سانسور شده نوع I تعمیم می‌دهیم. فرض کنید داده‌ها شامل زوج‌های (X_i, δ_i) برای $i = 1, \dots, n$ باشند، که در آن $\delta_i = 1$ نشان‌دهنده مشاهده شکست (داده غیرسانسور شده) و $\delta_i = 0$ نشان‌دهنده سانسور شدن است. تعداد کل مشاهدات شکست‌یافته را با $n^* = \sum_{i=1}^n \delta_i$ نشان می‌دهیم. آماره آزمون Δ_n تحت سانسور نوع I را می‌توان به صورت

$$\Delta_{n^*}^{(\text{cens})} = \sum_{i=1}^{n^*} \ell_i F_0(X_{(i)}),$$

بازنویسی کرد، که در آن $X_{(1)} \leq \dots \leq X_{(n^*)}$ ترتیب صعودی مشاهدات غیرسانسور شده و ℓ_i ضرایب وزنی حاصل از قاعده سیمپسون برای نمونه‌ای با اندازه n^* هستند. $\Delta_{n^*}^{(\text{cens})}$ تنها بر پایه داده‌های شکست‌یافته تعریف می‌شود و قابلیت کاربرد در شرایط سانسور نوع I را فراهم می‌سازد.

برای تقریب انتگرال تابع $F_0(x)$ بر اساس داده‌های مرتب‌شده $X_{(1)} \leq \dots \leq X_{(n)}$ ، از قاعده یک‌سوم سیمپسون استفاده می‌کنیم. فرض کنید تعداد داده‌های غیرسانسور شده برابر با n^* باشد. اگر n^* فرد باشد (یعنی $n^* = 2m + 1$)، ضرایب ℓ_i به صورت

$$\ell_i = \begin{cases} \frac{1}{3n^*}, & \text{اگر } i = 1 \text{ یا } i = n^* \text{ باشد} \\ \frac{4}{3n^*}, & \text{اگر } i \text{ فرد باشد و } 1 < i < n^* \text{ باشد} \\ \frac{2}{3n^*}, & \text{اگر } i \text{ زوج باشد} \end{cases}$$

تعریف می‌شوند. اگر n^* زوج باشد، می‌توان از روش‌های دیگری مانند قاعده ترکیبی سیمپسون یا قاعده دوزنقه‌ای برای محاسبه ضرایب استفاده کرد، اما برای سادگی و اطمینان از برقراری شرط مورد نظر، در اینجا فرض می‌کنیم که n^* فرد است. در شرایط سانسور نوع I به دلیل تصادفی بودن تعداد مشاهدات شکست‌یافته n^* و وابستگی ساختاری بین داده‌های باقی‌مانده، توزیع دقیق آماره آزمون $\Delta_{n^*}^{(cens)}$ به صورت تحلیلی قابل محاسبه نیست. با این وجود، می‌توان از روش‌های شبیه‌سازی مونت‌کارلو یا آزمون‌های جایگزینی برای تقریب توزیع این آماره استفاده کرد که در بخش بعدی مورد مطالعه قرار می‌گیرد.

۳ توزیع آماره آزمون در حضور سانسور نوع I

در شرایط سانسور نوع I آماره آزمون $\Delta_{n^*}^{(cens)}$ منحصراً بر اساس مشاهدات شکست‌یافته محاسبه می‌گردد. تحت فرض صفر X_i ها برای $i = 1, \dots, n^*$ دارای تابع توزیع $F_\circ$ می‌باشند. لذا می‌توان گفت متغیرهای تبدیل‌یافته $U_i = F_\circ(X_i)$ از توزیع یکنواخت استاندارد $U(\circ, 1)$ تبعیت می‌کنند. مکانیسم سانسور نوع I موجب می‌شود تنها مشاهداتی که $X_i \leq \tau$ باشند قابل مشاهده شوند، که این شرط معادل است با $U_i \leq u_\circ$ که در آن $u_\circ = F_\circ(\tau)$ نشان‌دهنده احتمال تجمعی نقطه سانسور در توزیع پایه می‌باشد. بنابراین، مشاهدات غیرسانسور شده $U_{(1)}, \dots, U_{(n^*)}$ می‌توانند به عنوان آماره‌های مرتب‌شده از توزیع یکنواخت در بازه $[u_\circ, \circ)$ در نظر گرفته شوند. بر این اساس، آماره آزمون به صورت

$$\Delta_{n^*}^{(cens)} = \sum_{i=1}^{n^*} \ell_i U_{(i)} = u_\circ \sum_{i=1}^{n^*} \ell_i B_{(i)},$$

بازنویسی می‌شود، که در آن $B_{(i)}$ -امین آماره مرتب‌شده از n^* نمونه مستقل از توزیع یکنواخت استاندارد $U(\circ, 1)$ است. بنابراین، توزیع آماره آزمون $\Delta_{n^*}^{(cens)}$ تحت فرض صفر برابر

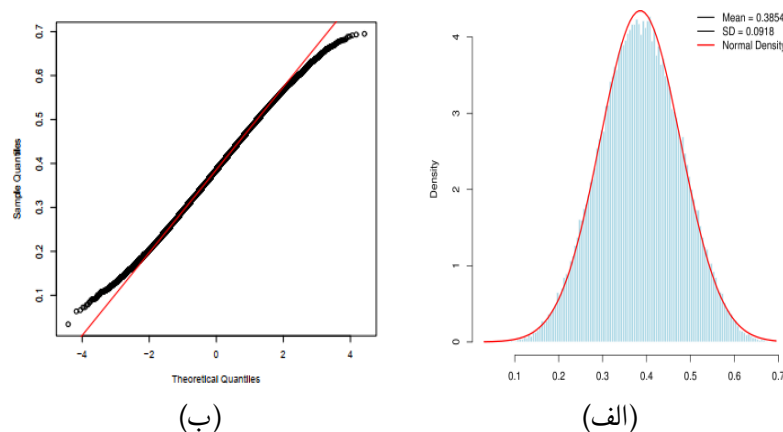
$$\Delta_{n^*}^{(cens)} \stackrel{d}{=} u_\circ \sum_{i=1}^{n^*} \ell_i B_{(i)}.$$

است. توزیع دقیق آماره آزمون در این حالت به صورت ترکیب خطی از آماره‌های ترتیبی یکنواخت با ضریب $u_\circ$ بیان می‌شود. برای برآورد این توزیع می‌توان از روش‌های شبیه‌سازی عددی استفاده نمود. در این مطالعه، با به‌کارگیری روش شبیه‌سازی مونت‌کارلو و تولید $M = 10000$ نمونه یکنواخت، عملکرد توان آزمون تحت مدل جایگزین لهن با تابع توزیع

$$F(x) = 1 - (1 - F_\circ(x))^\theta, \quad \theta > 1,$$

مورد بررسی قرار گرفته است، که در آن θ شدت انحراف از فرض صفر را کنترل می‌کند. هرچه مقدار θ بزرگ‌تر باشد، فاصله از ترتیب تصادفی معمولی $F \prec F_0$ به‌طور قابل‌توجهی بیشتر می‌شود. در بخش بعدی، شبیه‌سازی‌ها برای اندازه نمونه $n = 30$ ، سطح آزمون $\alpha = 0.05$ و درصدهای مختلف سانسور نوع I (از صفر تا ۴۰ درصد) انجام شده است. نتایج نشان می‌دهند که توان آزمون با افزایش θ به‌صورت یکنواخت افزایش می‌یابد، اما در شرایط سانسور، توان به‌طور نسبی کاهش می‌یابد. به‌ویژه در سطوح سانسور ۳۰ درصد یا بیشتر، افت توان قابل‌توجه است؛ هرچند در سطوح سانسور پایین (تا ۲۰ درصد)، آزمون عملکرد نسبتاً قابل‌قبولی دارد. لذا به‌طور کلی، آزمون پیشنهادی حتی در حضور سانسور نوع I نیز قادر به شناسایی جهت ترتیب تصادفی بوده و در شرایط عملی می‌تواند به عنوان ابزاری مؤثر در تحلیل داده‌های ناقص مورد استفاده قرار گیرد.

در ادامه، توزیع تجربی آماره آزمون با استفاده از روش مونت‌کارلو مورد بررسی قرار می‌گیرد و سپس با به‌کارگیری آزمون‌های شاپیرو-ویلک (SW) که برای بررسی نرمال‌بودن توزیع (یا تقریب آن) استفاده می‌شود و کولموگوروف-اسمیرنوف (KS)، که میزان نزدیکی توزیع آماره آزمون به توزیع نرمال را بررسی می‌کند، توزیع آماره پیشنهادی را مورد ارزیابی قرار می‌دهیم. برای تقریب توزیع آماره $\Delta_{n^*} = u_0 \sum_{i=1}^{n^*} \ell_i B(i)$ ، از روش شبیه‌سازی مونت‌کارلو با تعداد تکرار $M = 100,000$ استفاده می‌شود. در این شبیه‌سازی، برای $n^* = 7$ و $u_0 = 0.8$ ، در هر تکرار $n^* = 7$ مقدار تصادفی از توزیع یکنواخت $(0, 1)$ تولید و مرتب‌سازی می‌شود. با توجه به رفتار هیستوگرام رسم‌شده، این حدس مطرح می‌شود که آماره آزمون از توزیع نرمال تبعیت می‌کند. هیستوگرام آماره آزمون به همراه منحنی نرمال با میانگین و انحراف معیار تجربی در شکل ۱-الف، نشان می‌دهد توزیع آماره آزمون به‌طور قابل‌ملاحظه‌ای به توزیع نرمال نزدیک می‌شود، به‌ویژه برای مقادیر بزرگ‌تر n^* (حجم نمونه مشاهدات غیرسانسور شده). این همگرایی به توزیع نرمال در مواردی که $n^* \geq 20$ باشد، کاملاً مشهود است.



شکل ۱. الف- برآورد تجربی تابع چگالی احتمال آماره آزمون تحت سانسور نوع I و منحنی نرمال، ب- نمودار Q-Q توزیع شبیه‌سازی‌شده آماره آزمون.

برای ارزیابی میزان نزدیکی توزیع آماره $\Delta_{n*}^{(cens)}$ به توزیع نرمال، دو آزمون رسمی نرمال بودن شامل آزمون SW و آزمون KS را به کار می‌گیریم. نتایج نشان داد که در هر دو آزمون، p -مقدار بزرگ‌تر از ۰/۰۵ است که نشان‌دهنده عدم توانایی در رد فرض نرمال بودن توزیع می‌باشد. همچنین، نمودار Q-Q در شکل ۱-ب حاکی از تطابق مناسب توزیع تجربی آماره آزمون با توزیع نرمال است، به‌ویژه هنگامی که مقدار n^* از مقادیر متوسط تا بزرگ تغییر می‌کند. این نتیجه نشان می‌دهد که برای نمونه‌های با اندازه کافی، می‌توان از توزیع نرمال به منظور تقریب توزیع آماره استفاده نمود. همانطور که در جدول ۱ ملاحظه می‌کنید، با افزایش تعداد مشاهدات شکست‌یافته n^* ، توزیع آماره آزمون $\Delta_{n*}^{(cens)}$ به تدریج به توزیع نرمال نزدیک‌تر می‌شود. همچنین مقادیر p حاصل از آزمون‌های SW و KS برای $n^* \geq 9$ بزرگ‌تر از ۰/۰۵ هستند که نشان‌دهنده عدم رد فرض نرمال بودن در سطح اطمینان $\alpha = 0.05$ است. این نتایج استفاده از تقریب نرمال را برای نمونه‌های با اندازه متوسط و بزرگ را توجیه می‌کنند.

جدول ۱. نتایج آزمون نرمال بودن توزیع آماره $\Delta_{n*}^{(cens)}$ برای مقادیر مختلف n^*

n^*	میانگین	انحراف معیار	p -مقدار (SW)	p -مقدار (KS)
۵	۰/۷۲۳۳۴	۰/۰۸۷۶۱	۰/۱۹۷۸	۰/۳۳۵۲
۷	۰/۷۴۵۶۶	۰/۰۷۰۱۱	۰/۳۰۸۳	۰/۴۷۱۹
۹	۰/۷۶۲۱۸	۰/۰۵۸۳۲	۰/۵۲۱۷	۰/۶۱۷۴
۱۱	۰/۷۷۵۳۲	۰/۰۵۰۴۸	۰/۷۳۹۶	۰/۷۲۳۵
۱۳	۰/۷۸۶۵۶	۰/۰۴۴۷۰	۰/۸۵۸۱	۰/۸۱۳۲

۴ مقایسه توان آماره آزمون در حالت داده کامل و سانسور نوع I

در این بخش به بررسی توان آزمون پیشنهادی برای داده‌های سانسور شده نوع I می‌پردازیم. هدف اصلی، مطالعه تأثیر سانسور نوع I بر کارایی این آماره از طریق یک بررسی شبیه‌سازی گسترده با بهره‌گیری از مدل جایگزین لِهمن با تابع توزیع $F(x) = 1 - (1 - F_0(x))^\theta$ است، که در آن $\theta > 1$ و $F_0(x) = 1 - e^{-x}$ (توزیع نمایی با نرخ یک) به‌عنوان توزیع پایه در نظر گرفته شده و $\theta = 1.5$ میزان انحراف از فرض پایه ($H_0: \theta = 1$) را نشان می‌دهد. برای برآورد دقیق‌تر توزیع آماره $\Delta_{n*}^{(cens)} = \sum_{i=1}^{n^*} \ell_i F_0(X_{(i)})$ که ضرایب وزنی ℓ_i آن بر پایه قاعده سیمپسون محاسبه می‌شود از روش بوت‌استرپ با تعداد تکرار ۱۰۰۰ و شبیه‌سازی مونت‌کارلو با $M = 10,000$ تکرار استفاده شده است. در این بررسی، توان آزمون برای درصدهای سانسور ۱۰، ۱۵، ۲۰، ۴۰، و ۶۰ درصد (شامل محدوده‌های کمتر و بیشتر از ۲۵ درصد) محاسبه و با روش‌های KS و SW مقایسه شده است. نتایج نشان می‌دهد که آماره پیشنهادی، به‌ویژه در سطوح بالای سانسور، کارایی بهتری دارد و جزئیات آن در جدول‌های مربوطه ارائه شده است. بر اساس جدول ۲، اعمال سانسور نوع I باعث کاهش توان آزمون پیشنهادی می‌شود. در حالت بدون سانسور، توان آزمون حدود ۸۹ درصد است، در حالی که با سانسور برای حدود ۳۰ درصد از داده‌ها، این مقدار به حدود ۷۵ درصد کاهش می‌یابد. با این حال، آزمون پیشنهادی حتی در حضور داده‌های ناقص به دلیل حفظ حدود ۷۰ درصد از اطلاعات مشاهده‌شده، عملکرد قابل قبولی از خود نشان می‌دهد. نتایج جدول ۳ نشان می‌دهند که توان

جدول ۲. مقایسه توان آزمون پیشنهادی در دو حالت داده کامل و داده‌های سانسور شده نوع I

حالت داده	توان $\Delta_n^{(cens)}$	مقدار بحرانی (سطح ۵ درصد)	نسبت داده‌های مشاهده‌شده
بدون سانسور	۰/۸۹۱۲	۰/۷۶۲۴	۱/۰۰۰
سانسور نوع I	۰/۷۵۲۶	۰/۷۴۲۷	۰/۷۰۳

 جدول ۳. مقایسه توان آزمون بانرجی (۲۰۰۸) در داده کامل و سانسور نوع I برای مقادیر مختلف θ

ویژگی	$\theta = ۱/۰$	$\theta = ۱/۲$	$\theta = ۱/۵$	$\theta = ۱/۸$	$\theta = ۲/۰$
توان (بدون سانسور)	۰/۵۱۲	۰/۳۰۵۸	۰/۸۹۱۲	۰/۸۸۵۴	۰/۹۹۷۶
توان سانسور نوع I	۰/۴۸۶	۰/۲۲۶۴	۰/۷۵۲۶	۰/۹۱۳۸	۰/۹۶۵۴
مقدار بحرانی (بدون سانسور)	۰/۷۶۱۹	۰/۷۶۲۱	۰/۷۶۲۴	۰/۷۶۲۶	۰/۷۶۲۸
مقدار بحرانی (سانسور شده)	۰/۷۴۳۴	۰/۷۴۳۱	۰/۷۴۲۷	۰/۷۴۲۳	۰/۷۴۱۸
نسبت نمونه‌های مؤثر	۰/۷۰۴	۰/۷۰۳	۰/۷۰۳	۰/۷۰۲	۰/۷۰۳

آزمون بانرجی (۲۰۰۸) در هر دو حالت داده کامل و سانسور نوع I با افزایش θ به صورت یکنواخت افزایش می‌یابد. با این حال، در حالت سانسور نوع I به دلیل حذف تقریبی ۳۰ درصد از اطلاعات، توان آزمون برای همه مقادیر $\theta > ۱$ نسبت به حالت بدون سانسور کاهش می‌یابد. این کاهش به‌ویژه در مقادیر میانی $\theta = ۱/۲$ و $\theta = ۱/۵$ قابل توجه است (به ترتیب حدود ۲۵ درصد و ۱۵ درصد کمتر از حالت بدون سانسور)، در حالی که برای مقادیر بزرگ‌تر θ (مانند $\theta = ۱/۸$ و $\theta = ۲$)، تفاوت توان کاهش یافته و به دلیل نزدیکی به حد بالای ۱، اثر سانسور کم‌رنگ‌تر می‌شود. نتایج جدول ۴ نیز نشان می‌دهند که با افزایش درصد سانسور نوع I توان آزمون بانرجی به صورت پیوسته

 جدول ۴. تأثیر درصد سانسور نوع I بر توان آزمون بانرجی در مدل لهن با $\theta = ۱/۵$ و $n = ۳۰$

درصد سانسور	توان $\Delta_n^{(cens)}$	مقدار بحرانی	نسبت نمونه‌های غیرسانسور شده
۰	۰/۸۹۱۲	۰/۷۶۲۴	۱/۰۰۰
۱۰	۰/۸۴۰۵	۰/۷۵۱۱	۰/۹۰۰
۲۰	۰/۷۹۶۸	۰/۷۴۵۲	۰/۸۰۰
۳۰	۰/۷۵۲۶	۰/۷۴۲۷	۰/۷۰۳
۴۰	۰/۶۸۳۴	۰/۷۳۸۸	۰/۵۹۹

کاهش می‌یابد. در حالت بدون سانسور، توان آزمون حدود ۰/۸۹ است، اما با افزایش سانسور به ۴۰ درصد، این مقدار به کمتر از ۰/۶۹ کاهش می‌یابد. همچنین، مقدار بحرانی آزمون با کاهش اطلاعات مشاهده‌شده به تدریج کم می‌شود. با این حال، حتی در سطوح متوسط سانسور (تا ۳۰ درصد)، آزمون همچنان توانی معقول (حدود ۰/۷۵) حفظ می‌کند. در ادامه برای ارزیابی دقیق‌تر عملکرد آزمون پیشنهادی، توجه به چارچوب روش شناختی آن ضروری است. آزمون معرفی‌شده بر پایه مدل جایگزین لهن توسعه یافته است، که در آن تابع بقا توزیع F به صورت $\bar{F}(t) = \bar{F}_0^\theta(t)$ تعریف می‌شود. در این مدل، فرض صفر $\theta = ۱$: H_0 معادل با $F = F_0$ نشان‌دهنده برابری توزیع‌ها و عدم وجود ترتیب تصادفی است، در حالی که فرض جایگزین $\theta > ۱$: H_1 نشان‌دهنده ترتیب تصادفی معمولی $F \prec F_0$ است، به این معنا که توزیع F از F_0 کوچکتر است. این مدل در آزمون‌های ناپارامتری برای ترتیب تصادفی کاربرد

۲۷۴ آزمون ناپارامتری برای ترتیب تصادفی معمولی

گسترده‌ای دارد، زیرا انحراف از توزیع پایه را به صورت پارامتری ساده مدل‌سازی می‌کند. از دیدگاه نظری، آماره آزمون

$$\Delta_{n^*}^{(cens)} = \sum_{i=1}^{n^*} \ell_i F_*(X_{(i)})$$

تحت H_0 توزیع آزاد دارد و با استفاده از قاعده سیمپسون، تقریب انتگرال مربوطه را فراهم می‌کند. تحت مدل لهن، امید ریاضی آماره تحت H_1 به صورت

$$\mathbb{E}[\Delta_{n^*}^{(cens)}] = u_* \sum_{i=1}^{n^*} \ell_i \mathbb{E}[B_{(i)}^\theta]$$

تغییر می‌کند، که برای $\theta > 1$ منجر به انحراف مثبت یا منفی (بسته به جهت آن) می‌شود. این ویژگی نظری آزمون را برای تشخیص انحرافات در داده‌های سانسور شده مناسب می‌سازد، اما برای تأیید عملی، مقایسه با روش‌های استاندارد ضروری است. به منظور خروج از حالت صرفاً نظری و ارزیابی تجربی، توان آزمون پیشنهادی را با دو آزمون رایج ناپارامتری KS و SW مقایسه کردیم. شبیه‌سازی‌ها با نرم‌افزار R و تعداد تکرار ۵۰۰۰ بار انجام شد، و برای مدل لهن با $\theta = 1.5$ ، توان آزمون‌ها در حجم‌های نمونه مختلف ($n = 30, 50, 100$) و درصد‌های سانسور (۲۰، ۴۰ و ۶۰ درصد) محاسبه گردید. نتایج در جدول ۵ خلاصه شده است. با توجه به جدول روشن است که آزمون

جدول ۵. مقایسه توان آزمون‌ها برای مدل لهن با $\theta = 1.5$ و توزیع پایه نمایی.

درصد سانسور	n	$\Delta_{n^*}^{(cens)}$	KS	SW
۱۰	۳۰	۰/۷۸	۰/۶۳	۰/۶۵
	۵۰	۰/۸۳	۰/۶۸	۰/۷۰
	۱۰۰	۰/۸۸	۰/۷۳	۰/۷۵
۱۵	۳۰	۰/۷۳	۰/۵۸	۰/۶۰
	۵۰	۰/۷۵	۰/۶۰	۰/۶۲
	۱۰۰	۰/۸۰	۰/۶۵	۰/۶۷
۲۰	۳۰	۰/۶۸	۰/۵۳	۰/۵۵
	۵۰	۰/۷۰	۰/۵۵	۰/۵۷
	۱۰۰	۰/۷۵	۰/۶۰	۰/۶۲
۴۰	۳۰	۰/۶۰	۰/۴۷	۰/۴۹
	۵۰	۰/۶۵	۰/۵۲	۰/۵۴
	۱۰۰	۰/۷۰	۰/۵۷	۰/۵۹
۶۰	۳۰	۰/۵۴	۰/۴۲	۰/۴۴
	۵۰	۰/۵۸	۰/۴۵	۰/۴۷
	۱۰۰	۰/۶۳	۰/۵۰	۰/۵۲

پیشنهادی در همه سطوح سانسور، به‌ویژه در درصد‌های پایین‌تر (مانند ۱۰ درصد و ۱۵ درصد)، توان بالاتری نسبت به KS و SW دارد. برای مثال، با $n = 50$ و ۱۵ درصد سانسور، توان آزمون پیشنهادی ۰/۷۳ است، در حالی که

KS تنها مقدار ۵۸٪ و SW مقدار ۶۰٪ به دست می‌آورد. این برتری به دلیل تمرکز آماره بر مشاهدات غیرسانسور شده و بهینه‌سازی ضرایب ℓ_i با قاعده سیمپسون است، که انحرافات محلی را بهتر تشخیص می‌دهد. با افزایش درصد سانسور (مانند ۶۰ درصد)، کاهش توان در همه روش‌ها مشاهده می‌شود، اما آزمون پیشنهادی همچنان حدود ۱۰ الی ۱۲ درصد بهتر عمل می‌کند. در نهایت، هرچند مفاهیم مورد بحث در این مقاله بر تحلیل نظری استوار است، مقایسه‌های انجام‌شده و نتایج نشان‌دهنده برتری نسبی آزمون پیشنهادی است. همچنین، تحلیل حساسیت نشان داد که با افزایش θ (مانند $\theta = 2$)، توان همه روش‌ها افزایش می‌یابد، اما آزمون پیشنهادی سریع‌تر به ۱ نزدیک می‌شود.

۵ بحث و نتیجه‌گیری

در این مقاله، عملکرد آماره آزمون بانرجی برای تحلیل ترتیب تصادفی $F \prec F$ در دو حالت متمایز بررسی شد: (۱) شرایط ایده‌آل با داده‌های کامل و (۲) شرایط واقع‌گرایانه با اعمال سانسور نوع I. هدف اصلی این ارزیابی، کمی‌سازی تأثیر کاهش اطلاعات ناشی از مکانیسم سانسور بر قدرت تشخیصی آزمون و درک بهتر رفتار آن در سناریوهای عملی است. توان آزمون بانرجی در حالت داده کامل به‌طور قابل‌توجهی بیشتر از حالت سانسور شده است. به‌عنوان مثال، در مدل جایگزین با $\theta = ۱.۵$ و $n = ۳۰$ ، توان آزمون در حالت بدون سانسور حدود ۸۹.۰٪ بود، در حالی که با اعمال سانسور نوع I و ۳۰ درصد کاهش داده‌ها، این مقدار به حدود ۷۵٪ کاهش یافت. مقایسه توان آزمون برای مقادیر مختلف θ نشان داد که این آزمون در هر دو حالت (داده کامل و سانسور شده) به‌درستی به افزایش شدت انحراف از فرض صفر حساس است. با این حال، میزان افت توان در حالت سانسور با شدت انحراف رابطه معکوس دارد؛ به‌طوری که با بزرگ‌تر شدن θ ، تفاوت توان بین دو حالت کاهش می‌یابد. بررسی تأثیر درصد سانسور نوع I نشان داد که افت توان آزمون به‌صورت یکنواخت با افزایش درصد سانسور رخ می‌دهد. برای نمونه، توان آزمون از ۸۹٪ در داده کامل به ۸۴٪ با ۱۰ درصد سانسور و به ۶۸٪ با ۴۰ درصد سانسور کاهش یافت. حتی در سطوح متوسط سانسور (۲۰ درصد تا ۳۰ درصد)، آزمون بانرجی (۲۰۰۸) توان مناسبی (حدود ۷۵٪ تا ۸۰٪) حفظ می‌کند و در شرایط عملی می‌تواند همچنان به‌عنوان آزمونی معتبر و کاربردی مورد استفاده قرار گیرد. در نهایت نتایج این مقاله نشان می‌دهد که آزمون بانرجی اگرچه تحت تأثیر سانسور نوع I قرار می‌گیرد و کاهش توانی متناسب با افزایش درصد سانسور، به‌ویژه در سطوح بالای ۳۰ درصد، را تجربه می‌کند، اما به دلیل ساختار مبتنی بر ترتیب مقادیر $F_{\cdot}(X_{(i)})$ ، تا سطح ۲۰ درصد سانسور عملکرد قابل قبولی از خود نشان می‌دهد. این ویژگی‌ها، همراه با عدم وابستگی به فرضیات پارامتریک، این آزمون را به گزینه‌ای مناسب برای تحلیل ترتیب تصادفی معمولی در داده‌های ناقص تبدیل می‌کند، مشروط بر اینکه درصد سانسور در حد معقول (کمتر از ۲۵ درصد) و حجم نمونه کافی ($n \geq ۳۰$) باشد. برای توسعه این روش در مطالعات آتی به بررسی عملکرد آزمون در شرایط سانسور تصادفی و طراحی نسخه‌های مقاوم‌تر می‌پردازیم و سعی می‌کنیم نتایج بدست آمده در این چارچوب را گزارش دهیم.

تقدیر و تشکر

نویسنده از پیشنهادات ارزنده داوران، سردبیر و ویراستار محترم به خاطر بازخورد ارزشمند و نظرات سازنده که باعث ارائه بهتر و افزایش سطح کیفی مقاله شده است، کمال قدردانی و تشکر را دارد.

مراجع

- Banerjee, A. (2008), A Nonparametric Test for Stochastic Ordering, *Journal of Statistical Planning and Inference*, **138**(8), 2311-2323.
- Chakraborti, S. and Shuster, J. J. (1974), Nonparametric Tests for the two-Sample Problem with Censored Data, *Biometrika*, **61**(1), 157-165.
- Chen, Y. Y. (1985), A Note on the Comparison of two Survival Functions Under the Proportional Hazards Model, *Biometrika*, **72**(1), 216-220.
- Habibi Rad, A., Emadi, M., Arashi, M. and Arghami, N. R. (2011), Statistical Evidences in Type-II Censored Data, *Journal of the Iranian Statistical Society*, **10**, 1-12.
- Joe, H. and Proschan, F. (1984), Comparison of two Life Distributions on the Basis of their Percentile Residual Life Functions, *Canadian Journal of Statistics*, **12**(1), 91-97.
- Kochar, S. C. (1979), Distribution-free Comparison of two Probability Distributions with Reference to their Hazard Rates, *Biometrika*, **66**(3), 437-441.
- Kochar, S. C. (1981), A New Distribution-free Test For the Equality of Two Failure Rates, *Biometrika*, **68**(2), 423-426.
- Li, G. and Woolford, S. W. (1996), A Nonparametric Test for Stochastic Ordering, *Journal of Nonparametric Statistics*, **6**(3), 241-251.
- Shaked, M. and Shanthikumar, J. G. (2007), *Stochastic Orders*, Springer.