

یک آزمون نیکویی برازش نمایی برای داده‌های سانسور شده تصادفی با استفاده از واگرایی L2

یاسر هاشم‌زائی^۱، سیدمهدی امیرجهانشاهی^۲، محمدحسین دهقان^۳

تاریخ دریافت: ۱۴۰۳/۰۹/۱۷

تاریخ پذیرش: ۱۴۰۳/۱۲/۲۳

چکیده:

در این مقاله آزمونی مبتنی بر واگرایی L2 جهت سنجش نمایی بودن توزیع داده‌های سانسور شده تصادفی معرفی شده است. توان آزمون پیشنهادی به روش شبیه‌سازی مونت‌کارلو با سایر آزمون‌های رقیب شامل آزمون‌های کولموگروف-اسمیرنوف، اندرسون-دارلینگ و کرامر-فون-میزس که مبتنی بر تابع توزیع تجربی هستند و آزمون‌های مبتنی بر معیار اطلاع مورد مقایسه قرار گرفته است. نتایج مطالعات شبیه‌سازی این پژوهش نشان داد که آزمون پیشنهادی در این مطالعه، عموماً عملکرد بهتری نسبت به سایر آزمون‌های رقیب دارد.

واژه‌های کلیدی: آزمون نیکویی برازش، سانسور تصادفی، سنجش نمایی بودن، آزمون واگرایی L2.

۱ مقدمه

تحلیل‌های آماری به‌ویژه در آزمون طول عمر مورد استفاده قرار می‌گیرد، از این رو آزمون‌های بسیاری نیز در این زمینه معرفی شده‌اند. در این مقاله آزمون نیکویی برازش واگرایی L2، جهت سنجش نمایی بودن توزیع داده‌های سانسور شده تصادفی معرفی شده است و به کمک شبیه‌سازی با آزمون‌های رقیب مورد مقایسه قرار گرفته است.

سنجش میزان برازندگی یک توزیع بر مجموعه‌ای از داده‌ها، آزمون نیکویی برازش^۴ نام دارد که یک مسئله‌ی بسیار مهم است. در بسیاری از پژوهش‌ها در علوم مختلف توزیع داده‌های مورد پژوهش برای ما حائز اهمیت بوده و معمولاً نامعلوم است. به همین جهت محققین بسیاری به ارائه روش‌های بررسی توزیع‌ها پرداخته‌اند. در بررسی داده‌های طول عمر در بسیاری از تحقیقات به دلایل گوناگون غالباً پژوهشگران با نمونه‌های ناکامل مواجه می‌شوند. به‌عنوان مثال در بازه‌ی تحقیق تعدادی از آزمودنی‌ها در زمان‌های نامشخص سانسور می‌شوند که در این حالت پژوهشگر با داده‌های سانسور شده^۵ مواجه می‌شود. داده‌های سانسور شده به داده‌هایی اشاره دارند که بخشی از اطلاعات آن‌ها به دلیل دلایل مختلف قابل مشاهده نیست یا از دست رفته است. این نوع داده‌ها معمولاً در حوزه‌های آماری و تحلیل داده‌ها رخ می‌دهند، به‌ویژه در زمینه‌هایی مانند پزشکی، اقتصادی و اجتماعی شاهد بروز چنین داده‌هایی هستیم [۱].

۲ آزمون‌های نیکویی برازش برای داده‌های سانسور شده تصادفی

آزمون نیکویی برازش یک روش آماری است که برای ارزیابی میزان برازش یک توزیع خاص به داده‌های مورد نظر بکار می‌رود. بسته به نوع و میزان سانسور، توزیع داده‌ها و سؤال تحقیق، روش‌ها و چالش‌های مختلفی برای انجام آزمون‌های نیکویی برازش برای داده‌های سانسور شده تصادفی وجود دارد. برای داده‌هایی که تصادفی سانسور شده‌اند، چالش این است که برخی از مقادیر داده‌ها به‌طور کامل مشاهده نمی‌شوند، اما فقط تا حدی شناخته شده‌اند یا در یک بازه زمانی مشخص هستند.

پژوهشگران برحسب شرایط داده‌ها روش‌هایی از نیکویی برازش برای داده‌های سانسور شده ارائه می‌کنند. فرضیه نمایی بودن در بسیاری از

^۱دبیر ریاضی آموزش و پرورش زاهدان

^۲عضو هیئت علمی دانشگاه سیستان و بلوچستان (نویسنده مسئول: mjahan@math.usb.ac.ir)

^۳عضو هیئت علمی دانشگاه سیستان و بلوچستان

^۴Goodness of Fit Test

^۵Censored Data

مختلفی نظیر نمایی، رایلی و ... را تجزیه و تحلیل کردند. هم‌چنین کیم^{۱۲} (۲۰۱۲) [۱۹] با استفاده از آزمون نیکویی برازش داده‌های تصادفی سانسور شده را مورد بحث قرار داد. دنیش و اسلم^{۱۳} (۲۰۱۳، ۲۰۱۴) [۸، ۹] از نتایج بیزی برای داده‌ها سانسور تصادفی با توزیع وایبول استفاده کردند. کریشنا^{۱۴} و همکاران (۲۰۱۵) [۲۳] پارامترهای توزیع ماکسول^{۱۵} را تحت سانسور تصادفی برآورد کردند. علاوه بر این کریشنا و همکاران (۲۰۱۸) [۲۲] استنتاج بیزی در توزیع نمایی دو پارامتری با داده‌های سانسور شده تصادفی را مورد مطالعه قرار دادند. امیدی و همکاران (۲۰۱۹) [۲۸]، آزمون نیکویی برازشی بر اساس معیار اطلاع برای داده‌های سانسور شده تصادفی معرفی کردند. کومار و کومار^{۱۶} (۲۰۲۰) [۲۴] پارامترهای توزیع پارتو را تحت سانسور تصادفی برآورد نمودند. حسب‌الله^{۱۷} و همکاران (۲۰۲۳) [۱۵]، توزیع نمایی را تحت سانسور تصادفی برای داده‌های کلینیکی مورد بررسی قرار دادند. در داده‌های طول عمر، ممکن است برخی از رویدادها در زمان تحقیق مشاهده نشده باشند و یا زمان رخداد آن‌ها به صورت دقیق مشخص نباشد. در این صورت، از سانسور تصادفی برای مدل‌سازی این داده‌ها استفاده می‌شود. در سانسور تصادفی، مشاهداتی وجود دارند که زمان رویداد آن‌ها به صورت دقیق مشخص نشده است؛ به عبارت دیگر، در داده‌های سانسور شده تصادفی، برخی از مشاهدات به دلیل عدم دسترسی به اطلاعات کامل، به صورت ناقص یا سانسور شده هستند. در تحلیل داده‌های سانسور شده تصادفی، از روش‌های مختلفی مانند روش‌های ناپارامتری، روش‌های پارامتری استفاده می‌شود [۲۸].

۱۰.۳ تابع درست‌نمایی برای داده‌های سانسور راست تصادفی

فرض کنید C_i یک متغیر تصادفی و $X_1, X_2, \dots, X_n$ یک نمونه‌ی تصادفی از توزیع F باشند. سانسور تصادفی بجای مشاهدات $Y_i = \min(X_i, C_i)$ ، ما مشاهدات را به صورت $X_1, X_2, \dots, X_n$

بنابراین، آزمون‌های استاندارد نیکویی برازش، مانند آزمون کولموگروف-اسمیرنوف یا آزمون مجذور کای، ممکن است برای داده‌های سانسور شده قابل اجرا یا قابل اعتماد نباشد. به همین جهت برخی از روش‌های متداول مبتنی بر اصلاح آزمون‌های استاندارد برای داده‌های سانسور معرفی شدند که در آن‌ها از شبیه‌سازی مونت‌کارلو برای تولید نمونه‌هایی از توزیع هدف و سپس سانسور آن استفاده می‌شود [۵].

۳ سانسور تصادفی

سانسور تصادفی^۶ نوع دیگری از راست سانسور است، این سانسور در زمانی اتفاق می‌افتد که مطالعه بعد از مدت C سال پیگیری پایان می‌پذیرد، اما آزمودنی‌ها زمان سانسور یکسانی ندارند (آزمودنی‌ها در زمان‌های متفاوت (تصادفی) وارد آزمایش می‌شوند). غالباً با این شکل از سانسور در آزمایش‌های کلینیکی پزشکی و مطالعه طول عمر حیوانات (که زمان ورود و خروج آزمودنی‌ها به تحقیق، تصادفی است) روبرو می‌شویم.

در عمل ما اغلب در شرایطی قرار داریم که تعدادی آزمایش به صورت تصادفی حذف می‌شود. برای نمونه در آزمایش‌های بالینی یک بیمار ممکن است قبل از اتمام درمان با اختطار قبلی از رده خارج و یا بمیرد. به طور مشابه، در مهندسی قابلیت اطمینان^۷، ممکن است لازم باشد یک آیتم از یک آزمایش قبل از خرابی کامل آن، به علت شکستن یا صرفه‌جویی در وقت و پول، از آزمایش خارج شود. هنگامی که اشیاء به طور تصادفی در زمان‌های مختلف از آزمایش حذف می‌شوند، سانسور تصادفی رخ می‌دهد [۲۲].

تابع بقا سانسور تصادفی توسط گیلبرت^۸ (۱۹۶۲) معرفی شد [۱۴]. این نوع سانسور قبلاً توسط بریسلو و کراولی^۹ (۱۹۷۴) [۶]، کوزیول و گرین^{۱۰} (۱۹۷۶) [۲۱]، مورد بررسی قرار گرفته بود. بعدها، با داده‌های تصادفی سانسور شده گیتانی و الواحدی^{۱۱} (۲۰۰۲) [۱۳]، توزیع‌های

^۶Random Censoring

^۷Reliability

^۸Gilbert

^۹Breslow and Crowley

^{۱۰}Kozioł and Green

^{۱۱}Ghitany and Al-Awadhi

^{۱۲}Kim

^{۱۳}Danish and Aslam

^{۱۴}Krishna

^{۱۵}Maxwell

^{۱۶}Kumaran and Kumar

^{۱۷}Hasaballah

۲.۳ برآورد پارامتر توزیع نمایی تحت داده‌های

سانسور تصادفی

برآورد پارامتر θ در توزیع نمایی به روش ماکسیم درستنمایی تحت سانسور تصادفی به صورت زیر است:

$$\hat{\theta} = \frac{\sum_{i=1}^d x_i + \sum_{i=1}^{n-d} C_i}{d} \quad (2)$$

که در آن C_i زمان سانسور آزمودنی i ام و d تعداد مشاهداتی است که قبل از زمان C_i دچار شکست شده‌اند.

۴ آزمون پیشنهادی

۱.۴ روش فاصله L2

واگرایی L2 که به عنوان فاصله اقلیدسی مجذور یا میانگین مجذور خطا نیز شناخته می‌شود، معیاری برای عدم تشابه بین دو توزیع احتمال است. معمولاً در زمینه‌های مختلف از جمله آمار، یادگیری ماشین و نظریه اطلاع استفاده می‌شود. در این زمینه، واگرایی L2 تفاوت بین دو تابع چگالی احتمال را محاسبه می‌کند. به عنوان مثال، در برآورد آماری، می‌توان از آن برای اندازه‌گیری اختلاف بین توزیع برآورد شده و توزیع واقعی استفاده کرد [۱۷].

واگرایی L2 یکی از معیارهای فاصله بین توزیع‌های احتمالی است که رابطه نزدیکی با نظریه اطلاع دارد. این واگرایی با مربع فاصله اقلیدسی بین توابع چگالی احتمال اندازه‌گیری می‌شود. از طرفی، آنتروپی شانون میزان عدم قطعیت در یک توزیع را نشان می‌دهد و دیورژانس کولبک-لیبلر (KL) به عنوان یک سنجه عدم تقارن، فاصله بین دو توزیع را بر اساس آنتروپی نسبی ارزیابی می‌کند.

ارتباط میان واگرایی L2 و دیورژانس KL در تقریب‌های مرتبه دوم دیده می‌شود، جایی که برای توزیع‌های نزدیک به هم، KL را می‌توان با استفاده از واگرایی L2 تقریب زد. همچنین، در تحلیل توزیع‌های احتمالی و یادگیری ماشین، واگرایی L2 و KL هر دو در تنظیم مدل‌ها و تخمین‌های آماری نقش دارند. در نظریه اطلاع، این معیارها ابزارهایی برای اندازه‌گیری شباهت یا اختلاف بین توزیع‌های مختلف و ارزیابی کارایی تخمین‌گرها محسوب می‌شوند [۳۵].

در نظر می‌گیریم؛ که در آن:

$$Y_i = \begin{cases} X_i, & X_i \leq C_i, \\ C_i, & X_i > C_i, \end{cases}, \quad \delta_i = \begin{cases} 1, & X_i \leq C_i, \\ 0, & X_i > C_i. \end{cases} \quad (1)$$

دیده می‌شوند [۲]. تابع درستنمایی^{۱۸} برای داده‌های سانسور تصادفی به شکل زیر قابل تعریف است:

$$L(\beta) = \prod_{i=1}^n f_{\beta}^{\delta_i}(t_i) S^{1-\delta_i}(t_i)$$

که در آن $S(t)$ تابع بقا و $T_i = (Y_i, \delta_i)$ متغیر تصادفی راست سانسور تصادفی و $\delta_i = I(X_i \leq C_i)$ بوده، Y_i زمان بقا و C_i زمان سانسور فرد i ام و β بیانگر پارامتر سانسور است.

اگر $t_1, t_2, \dots, t_r$ زمان دقیق شکست واحدهای مورد مطالعه و $n-r$ سانسور تصادفی داشته باشیم، آنگاه [۲۰]:

$$L(\beta) = \prod_{i=1}^r f_{\beta}(t_i) \prod_{i=r+1}^n S(C_i)$$

۱.۳ آزمون نمایی بودن جامعه

آزمون‌های نیکویی برازش متعددی جهت سنجش نمایی بودن توزیع داده‌ها معرفی شده‌اند، از جمله این آزمون‌ها می‌توان به آزمون‌های لی لی فورس^{۱۹}، سوست^{۲۰} (۱۹۶۹) [۳۰]، فینکلستین و شفر^{۲۱} (۱۹۷۱) [۱۱]، کولموگروف-اسمیرنوف، کرامر-فون میزس، کویی‌پر، واتسون^{۲۲}، اندرسون-دارلینگ اشاره کرد [۴].

از جمله آزمون‌های نیکویی برازش جهت سنجش نمایی بودن داده‌های سانسور شده تصادفی می‌توان به آزمون‌های مبتنی بر خانواده تابع توزیع تجربی (EDF) توسط کیم (۲۰۱۲) [۱۹] و آزمون‌های مبتنی بر معیار اطلاع توسط امیدی و همکاران (۲۰۱۹) [۲۸] اشاره کرد.

فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه‌ی تصادفی مستقل و هم‌توزیع و نامنفی با تابع چگالی f و تابع توزیع F و میانگین $E(X) = \theta$ باشد. هدف، انجام آزمون فرضیه زیر است:

$$\begin{cases} H_0 : f(x) = f_0(x) \\ H_1 : f(x) \neq f_0(x) \end{cases}$$

در آزمون فرضیه بالا $f_0(x) = \frac{1}{\theta} \exp(-\frac{x}{\theta})$ تابع چگالی نمایی است.

¹⁸Likelihood Function

¹⁹Lilliefors

²⁰Van-Soest

²¹Finkelstein and Schafer

²²Watson

²³Kullback-Leibler

۴. ویژگی نابرابری مثلثی را دارد:

$$L2(P, Q) + L2(Q, R) \geq L2(P, R)$$

در ادامه به برآورد آماره $L2$ به روش واسیچک جهت سنجش نمایی بودن برای داده‌های کامل پرداخته شده است.

۲.۴ آماره $L2$ تحت سانسور تصادفی برای سنجش نمایی

فرض کنید $x_1, x_2, \dots, x_n$ یک نمونه تحت سانسور تصادفی باشند. هدف انجام آزمون فرضیه زیر است:

$$\begin{cases} H_0 : f(x) = f_0(x) = \frac{1}{\theta} e^{-\frac{1}{\theta}x}, & x > 0, \theta > 0 \\ H_1 : f(x) \neq f_0(x) \end{cases}$$

آماره $L2$ تحت سانسور تصادفی به صورت زیر بیان می‌گردد:

$$L2(f, f_0, C_i) = \int_{-\infty}^{C_i} [f(x) - f_0(x)]^2 dx, \quad C_i \geq 0, \quad i = 1, 2, 3, \dots, n$$

که در آن C_i زمان سانسور مشاهده i ام است.

۳.۴ برآوردگر آماره $L2$ به روش واسیچک تحت سانسور تصادفی

حال جهت برآورد آماره ذکر شده در بالا جهت سنجش نمایی بودن از برآوردگر واسیچک به واسطه‌ی سادگی در محاسبات و دقت بالای آن، استفاده خواهد شد.

$$L2_V^C(f, f_0, C_i) = \frac{1}{n} \sum_{i=1}^d \frac{d}{\gamma m} (x_{(i+m)} - x_{(i-m)}) * \left[\left(\frac{d}{\gamma m} (x_{(i+m)} - x_{(i-m)}) \right)^{-1} - \frac{1}{\theta} e^{-\frac{1}{\theta}x(i)} \right]^2$$

که در آن $\hat{\theta}$ برآورد ماکسیمم درست‌نمایی پارامتر θ تحت سانسور تصادفی است که از رابطه‌ی (۲) به دست آمده است. علاوه بر این m یک مقدار صحیح مثبت کمتر از $\frac{n}{4}$ است. پژوهش‌های متعددی جهت انتخاب مقدار بهینه برای m صورت گرفته است. ویکزورکوفسکی^{۲۴} و گریزگورزسکی^{۲۵} [۳۳] اندازه مقدار بهینه را از رابطه $m = \lceil \sqrt{n} + 0.5 \rceil$ به دست آورده‌اند [۳].

یکی از کاربردهای رایج واگرایی $L2$ در پردازش تصویر و بینایی کامپیوتری است. می‌توان از آن برای اندازه‌گیری تفاوت بین دو تصویر با در نظر گرفتن مقادیر پیکسل به عنوان معیار عدم تشابه بصری استفاده نمود [۷].

مقدار کوچک واگرایی $L2$ نشان‌دهنده شباهت بیشتر بین توزیع‌ها یا نقاط داده است، درحالی‌که مقدار بزرگ واگرایی $L2$ نشان‌دهنده عدم تشابه بیشتر است.

اگر دو مجموعه توزیع احتمال P و Q در فضای متناهی با تکیه‌گاه مشترک، دارای احتمالات p_i و q_i مطابق رابطه (۳) باشند.

$$P = \{p_1, p_2, \dots, p_n\} \quad \text{و} \quad Q = \{q_1, q_2, \dots, q_n\} \quad (۳)$$

با توجه به مثبت بودن مقادیر احتمالات روابط زیر را داریم:

$$p_i \geq 0, \quad q_i \geq 0, \quad \sum_{i=1}^n q_i = 1, \quad \sum_{i=1}^n p_i = 1 \quad (۴)$$

فاصله $L2$ طبق رابطه (۳) به صورت زیر محاسبه می‌شود:

$$L2(p, q) = \sum_{i=1}^n (p_i - q_i)^2 \quad (۵)$$

برای متغیرهای تصادفی پیوسته، این اندازه به شکل زیر تعریف می‌شود:

$$L2(p, q) = \int (p(x) - q(x))^2 dx \quad (۶)$$

اندازه فاصله $L2$ همیشه بین مقادیر صفر و یک است. اگر توزیع احتمالات بین دو توزیع یکسان باشند، مقدار فاصله $L2$ صفر می‌گردد. زمانی که دو توزیع احتمال کاملاً متفاوت باشند، مقدار فاصله $L2$ برابر یک خواهد بود. از ویژگی‌های واگرایی $L2$ می‌توان به موارد زیر اشاره کرد:

۱. هرگاه $P = Q$ باشد این معیار کمترین مقدار یعنی صفر را می‌گیرد.

$$L2(P, Q) = 0 \iff P = Q$$

۲. این معیار همیشه نامنفی است، فقط زمانی مقدار مثبت می‌گیرد که P و Q مخالف هم باشند:

$$L2(P, Q) > 0 \iff P \neq Q$$

۳. متقارن است یعنی:

$$L2(P, Q) = L2(Q, P)$$

²⁴Wieczorkowski

²⁵Grzegorzewski

۵. آماره امیدی و همکاران (۲۰۱۹) [۲۸]

$$\begin{aligned} \hat{T}_{n2}^* &= \sqrt{n} \sum_{i=1}^{n+1} (\hat{U}_i - \hat{U}_{i-1}) (\hat{F}_n(\hat{U}_i)) \ln(\hat{F}_n(\hat{U}_i)) \\ &\quad - \sqrt{n} \sum_{i=1}^{n+1} (\hat{F}_n(\hat{U}_i)) (\hat{U}_i \ln \hat{U}_i - \hat{U}_i - \hat{U}_{i-1} \ln(\hat{U}_{i-1}) + \hat{U}_{i-1}) \end{aligned}$$

که در آماره‌های فوق $\hat{U}_i = \hat{F}_n(Y_i) = \hat{F}_n(\cdot)$ برآوردگر کاپلان-مایر است.

۷ مقایسه توان آزمون‌ها

با استفاده از شبیه‌سازی مونت‌کارلو توان آماره‌ی آزمون $L2_V^C$ با آماره‌های $\hat{T}_{n1}^*$ و $\hat{T}_{n2}^*$ که مبتنی بر معیار اطلاع و آماره‌های مبتنی بر تابع توزیع تجربی (EDF) مورد مقایسه قرار خواهد گرفت.

جهت این ارزیابی با استفاده از شبیه‌سازی مونت-کارلو از هر جایگزین رقیب ۱۰۰۰۰ نمونه با حجم $n = 20, 30, 50$ با نسبت سانسور $40\%, 20\%$ و 50% درصد، در سطح معنی‌داری $\alpha = 0.05$ و $\alpha = 0.1$ شبیه‌سازی شده است و سپس توان هر یک از آزمون‌ها از نسبت تعداد دفعاتی که آماره آزمون از مقادیر بحرانی آن بیشتر است به تعداد کل تکرارها، برآورد خواهد شد. همچنین به بررسی عملکرد آزمون‌های معرفی‌شده جهت سنجش نمایی بودن با تعدادی توزیع جایگزین که توسط کوزیول و گرین (۱۹۷۶) پیشنهاد شده است، پرداخته خواهد شد؛ که به شرح زیر در نظر گرفته شده‌اند:

- توزیع نمایی با پارامتر ۱
- توزیع وایبول $W(\theta)$ ، با پارامتر مقیاس ۱ و پارامتر شکل $\theta = 0.5, 2$
- توزیع گاما $G(\theta)$ ، با پارامتر مقیاس ۱ و پارامتر شکل $\theta = 0.5, 2$
- توزیع لگ نرمال با تابع چگالی، $\frac{1}{\theta x \sqrt{\pi}} \exp\{-(\log x)^2 / (2\theta^2)\}$ ، $\theta = 1$
- توزیع یکنواخت استاندارد

۵ مقادیر بحرانی آماره $L2_V^C$

تعیین توزیع دقیق آماره $L2_V^C$ تحت فرضیه صفر به صورت تحلیلی قابل محاسبه نیست. برای تعیین نقاط بحرانی چندک $(1-\alpha)$ ام آماره $L2_V^C$ از شبیه‌سازی مونت‌کارلو استفاده شده است.

برای به دست آوردن چندک توزیع تحت فرض صفر $L2_V^C$ ، تعداد ۱۰۰۰۰ نمونه با اندازه n از توزیع نمایی استاندارد برای مقادیر مختلف n تولید شده است. برای هر نمونه مقدار آماره $L2_V^C$ محاسبه شده و سپس مقادیر تولید شده برای تعیین مقادیر بحرانی در سطح $\alpha = 0.05$ و $\alpha = 0.1$ محاسبه گردید؛ که در جدول ۱ ارائه شده است. H_0 را در سطح معنی‌داری α رد می‌کنیم و H_1 را می‌پذیریم هرگاه $L2_{V,n}^C \geq L2_{V,n,1-\alpha}^C$ باشد، که در آن $L2_{V,n,1-\alpha}^C$ صدک $(1-\alpha)$ ام از آماره $L2_{V,n}^C$ تحت فرضیه H_0 می‌باشد.

۶ معرفی آزمون‌های رقیب

در این بخش، آماره‌های آزمون‌های رقیب نمایی بودن تحت سانسور تصادفی معرفی شده‌اند:

۱. آماره کولموگروف - اسمیرنوف (۱۹۹۳)

$$\hat{D}_n = \sup_{0 < u < 1} |\hat{F}_n(\hat{U}_i) - \hat{U}_i| = \max(D_n^+, D_n^-)$$

۲. آماره کرامر فون - میزس

$$\hat{W}_n^2 = \sum_{i=1}^n \hat{F}_n(\hat{U}_i) (\hat{U}_{i+1} - \hat{U}_i) \{ \hat{F}_n(\hat{U}_i) - (\hat{U}_{i+1} - \hat{U}_i) \} - n/3$$

۳. آماره اندرسون - دارلینگ (۱۹۵۴)

$$\begin{aligned} \hat{A}_n^2 &= n \sum_{i=1}^{n-1} \hat{F}_n(\hat{U}_i)^2 \left(-\ln(1 - \hat{U}_{i+1}) + \ln(\hat{U}_{i+1}) + \ln(1 - \hat{U}_i) - \ln(\hat{U}_i) \right) \\ &\quad - 2n \sum_{i=1}^{n-1} \hat{F}_n(\hat{U}_i) \left(-\ln(1 - \hat{U}_{i+1}) + \ln(\hat{U}_{i+1}) \right) \\ &\quad - n \ln(1 - \hat{U}_n) - n \ln(\hat{U}_n) - n. \end{aligned}$$

۴. آماره امیدی و همکاران (۲۰۱۹) [۲۸]

$$\begin{aligned} \hat{T}_{n1}^* &= \sqrt{n} \sum_{i=1}^{n+1} (\hat{U}_i - \hat{U}_{i-1}) (1 - \hat{F}_n(\hat{U}_i)) \ln(1 - \hat{F}_n(\hat{U}_i)) \\ &\quad - \sqrt{n} \sum_{i=1}^{n+1} (1 - \hat{F}_n(\hat{U}_i)) (\hat{U}_i \ln \hat{U}_i - \hat{U}_i - \hat{U}_{i-1} \ln(\hat{U}_{i-1}) + \hat{U}_{i-1}) \end{aligned}$$

جدول ۱: مقادیر بحرانی آماره $L2_V^C$

$\alpha = 0.1$	$\alpha = 0.05$	نرخ سانسور	n
۱۵/۸۸	۱۹/۶۸	۲۰	۱۰
۲۳/۱۶	۲۸/۳۶	۴۰	
۲۸/۴۵	۳۴/۵۷	۵۰	
۲۴/۵۱	۲۸/۰۴	۲۰	۲۰
۳۵/۰۸	۴۰/۲۱	۴۰	
۴۳/۵۵	۵۰/۰۹	۵۰	
۳۲/۳۸	۳۶/۰۷	۲۰	۳۰
۴۶/۳۹	۵۱/۷۲	۴۰	
۵۷/۴۹	۶۳/۴۶	۵۰	
۴۸/۰۸	۵۲/۲۲	۲۰	۵۰
۶۸/۲۵	۷۴/۶۰	۴۰	
۸۴/۸۵	۹۲/۴۲	۵۰	
۸۶/۰۹	۹۰/۶۵	۲۰	۱۰۰
۱۲۱/۹۸	۱۲۹/۳۶	۴۰	
۱۵۲/۲۱	۱۶۰/۰۳	۵۰	

جدول ۲: مقایسه توان آماره‌های $\widehat{L2}_V^C$, $\widehat{T}_{n2}^*$, $\widehat{T}_{n1}^*$, $\widehat{A}_n^*$, $\widehat{W}_n^*$, $\widehat{D}_n$ در سطح $\alpha = 0.05$ با حجم نمونه $n = 20$

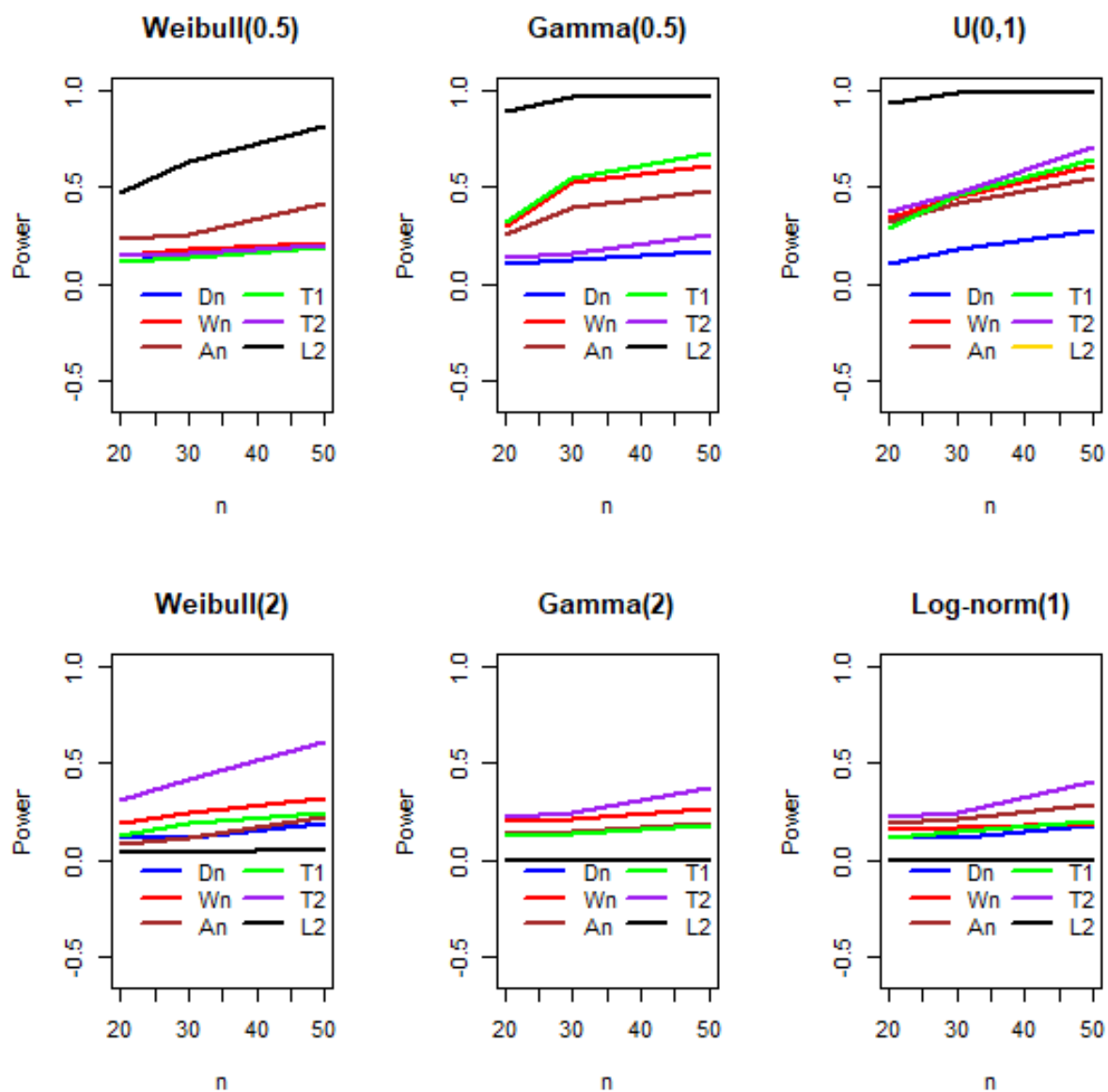
$\widehat{L2}_V^C$	$\widehat{T}_{n2}^*$	$\widehat{T}_{n1}^*$	$\widehat{A}_n^*$	$\widehat{W}_n^*$	$\widehat{D}_n$	نرخ سانسور	توزیع رقیب
۰/۰۵۳	۰/۰۵۱	۰/۰۵۰	۰/۰۵۰	۰/۰۵۰	۰/۰۵۱	۲۰	exponential(۱)
۰/۰۵۲	۰/۰۵۰	۰/۰۵۱	۰/۰۵۱	۰/۰۴۹	۰/۰۴۹	۴۰	
۰/۰۵۱	۰/۰۴۹	۰/۰۵۱	۰/۰۵۰	۰/۰۵۱	۰/۰۵۰	۵۰	
۰/۰۴۸	۰/۰۲۰	۰/۰۱۳	۰/۰۳۵	۰/۰۱۶	۰/۰۱۴	۲۰	Weibull(۰.۵)
۰/۰۴۷	۰/۰۱۵	۰/۰۱۱	۰/۰۲۳	۰/۰۱۵	۰/۰۱۲	۴۰	
۰/۰۴۶	۰/۰۱۳	۰/۰۰۸	۰/۰۲۰	۰/۰۱۱	۰/۰۰۹	۵۰	
۰/۰۹۲	۰/۰۱۸	۰/۰۳۹	۰/۰۳۵	۰/۰۳۷	۰/۰۱۲	۲۰	Gamma(۰.۵)
۰/۰۹۱	۰/۰۱۴	۰/۰۳۲	۰/۰۲۶	۰/۰۳۰	۰/۰۱۰	۴۰	
۰/۰۹۰	۰/۰۱۳	۰/۰۲۸	۰/۰۲۰	۰/۰۲۵	۰/۰۰۹	۵۰	
۰/۰۹۹	۰/۰۴۲	۰/۰۳۵	۰/۰۳۶	۰/۰۳۹	۰/۰۱۸	۲۰	U(۰, ۱)
۰/۰۹۴	۰/۰۳۷	۰/۰۲۹	۰/۰۳۲	۰/۰۳۴	۰/۰۱۰	۴۰	
۰/۰۸۹	۰/۰۳۱	۰/۰۲۳	۰/۰۲۲	۰/۰۲۵	۰/۰۰۸	۵۰	
۰/۰۰۶	۰/۰۳۵	۰/۰۱۷	۰/۰۰۹	۰/۰۲۶	۰/۰۱۳	۲۰	Weibull(۲)
۰/۰۰۵	۰/۰۳۱	۰/۰۱۳	۰/۰۰۸	۰/۰۱۹	۰/۰۱۱	۴۰	
۰/۰۰۴	۰/۰۲۸	۰/۰۱۱	۰/۰۰۷	۰/۰۱۵	۰/۰۱۰	۵۰	
۰/۰۰۵	۰/۰۲۹	۰/۰۱۵	۰/۰۱۷	۰/۰۲۳	۰/۰۱۴	۲۰	Gamma(۲)
۰/۰۰۴	۰/۰۲۲	۰/۰۱۳	۰/۰۱۴	۰/۰۲۰	۰/۰۱۳	۴۰	
۰/۰۰۳	۰/۰۱۷	۰/۰۱۰	۰/۰۱۱	۰/۰۱۴	۰/۰۱۰	۵۰	
۰/۰۰۴	۰/۰۲۹	۰/۰۱۶	۰/۰۲۰	۰/۰۱۷	۰/۰۱۲	۲۰	Log - norm(۱)
۰/۰۰۳	۰/۰۲۲	۰/۰۱۲	۰/۰۱۹	۰/۰۱۶	۰/۰۱۱	۴۰	
۰/۰۰۲	۰/۰۱۹	۰/۰۰۹	۰/۰۱۴	۰/۰۱۴	۰/۰۰۹	۵۰	

جدول ۳: مقایسه توان آماره‌های $\widehat{L2}_V^C$, $\widehat{T}_{n2}^*$, $\widehat{T}_{n1}^*$, $\widehat{A}_n^*$, $\widehat{W}_n^*$, $\widehat{D}_n$ در سطح $\alpha = 0.05$ با حجم نمونه $n = 30$

توزیع رقیب	نرخ سانسور	$\widehat{D}_n$	$\widehat{W}_n^*$	$\widehat{A}_n^*$	$\widehat{T}_{n1}^*$	$\widehat{T}_{n2}^*$	$\widehat{L2}_V^C$
<i>exponential</i> (۱)	۲۰	۰/۵۱	۰/۴۸	۰/۴۹	۰/۵۰	۰/۵۱	۰/۴۸
	۴۰	۰/۴۹	۰/۵۰	۰/۵۱	۰/۴۹	۰/۴۹	۰/۵۰
	۵۰	۰/۴۹	۰/۴۹	۰/۵۰	۰/۴۸	۰/۵۰	۰/۵۱
<i>Weibull</i> (۰/۵)	۲۰	۰/۱۷	۰/۱۹	۰/۳۸	۰/۱۹	۰/۲۲	۰/۶۵
	۴۰	۰/۱۵	۰/۱۸	۰/۲۶	۰/۱۴	۰/۱۶	۰/۶۴
	۵۰	۰/۱۲	۰/۱۵	۰/۲۱	۰/۱۰	۰/۱۵	۰/۶۳
<i>Gamma</i> (۰/۵)	۲۰	۰/۱۴	۰/۵۷	۰/۵۱	۰/۵۸	۰/۱۹	۰/۹۹
	۴۰	۰/۱۳	۰/۵۳	۰/۴۰	۰/۵۵	۰/۱۶	۰/۹۸
	۵۰	۰/۱۱	۰/۴۲	۰/۲۸	۰/۴۵	۰/۱۵	۰/۹۷
<i>U</i> (۰, ۱)	۲۰	۰/۲۵	۰/۵۱	۰/۵۳	۰/۵۲	۰/۵۵	۱
	۴۰	۰/۱۸	۰/۴۵	۰/۴۲	۰/۴۶	۰/۴۷	۰/۹۹
	۵۰	۰/۱۴	۰/۳۴	۰/۲۸	۰/۳۸	۰/۳۹	۰/۹۸
<i>Weibull</i> (۲)	۲۰	۰/۱۶	۰/۳۶	۰/۱۳	۰/۲۸	۰/۵۵	۰/۰۷
	۴۰	۰/۱۲	۰/۲۴	۰/۱۲	۰/۱۹	۰/۴۲	۰/۰۴
	۵۰	۰/۱۱	۰/۱۷	۰/۱۱	۰/۱۵	۰/۳۵	۰/۰۳
<i>Gamma</i> (۲)	۲۰	۰/۱۶	۰/۲۴	۰/۱۹	۰/۱۶	۰/۳۵	۰/۰۰۲
	۴۰	۰/۱۵	۰/۲۱	۰/۱۵	۰/۱۴	۰/۲۵	۰/۰۰۱
	۵۰	۰/۱۲	۰/۱۶	۰/۱۳	۰/۱۲	۰/۱۹	۰/۰۰۱
<i>Log - norm</i> (۱)	۲۰	۰/۱۳	۰/۱۹	۰/۲۳	۰/۱۹	۰/۳۱	۰/۰۰۲
	۴۰	۰/۱۲	۰/۱۷	۰/۲۱	۰/۱۵	۰/۲۵	۰/۰۰۱
	۵۰	۰/۱۱	۰/۱۶	۰/۲۰	۰/۱۴	۰/۲۲	۰/۰۰۱

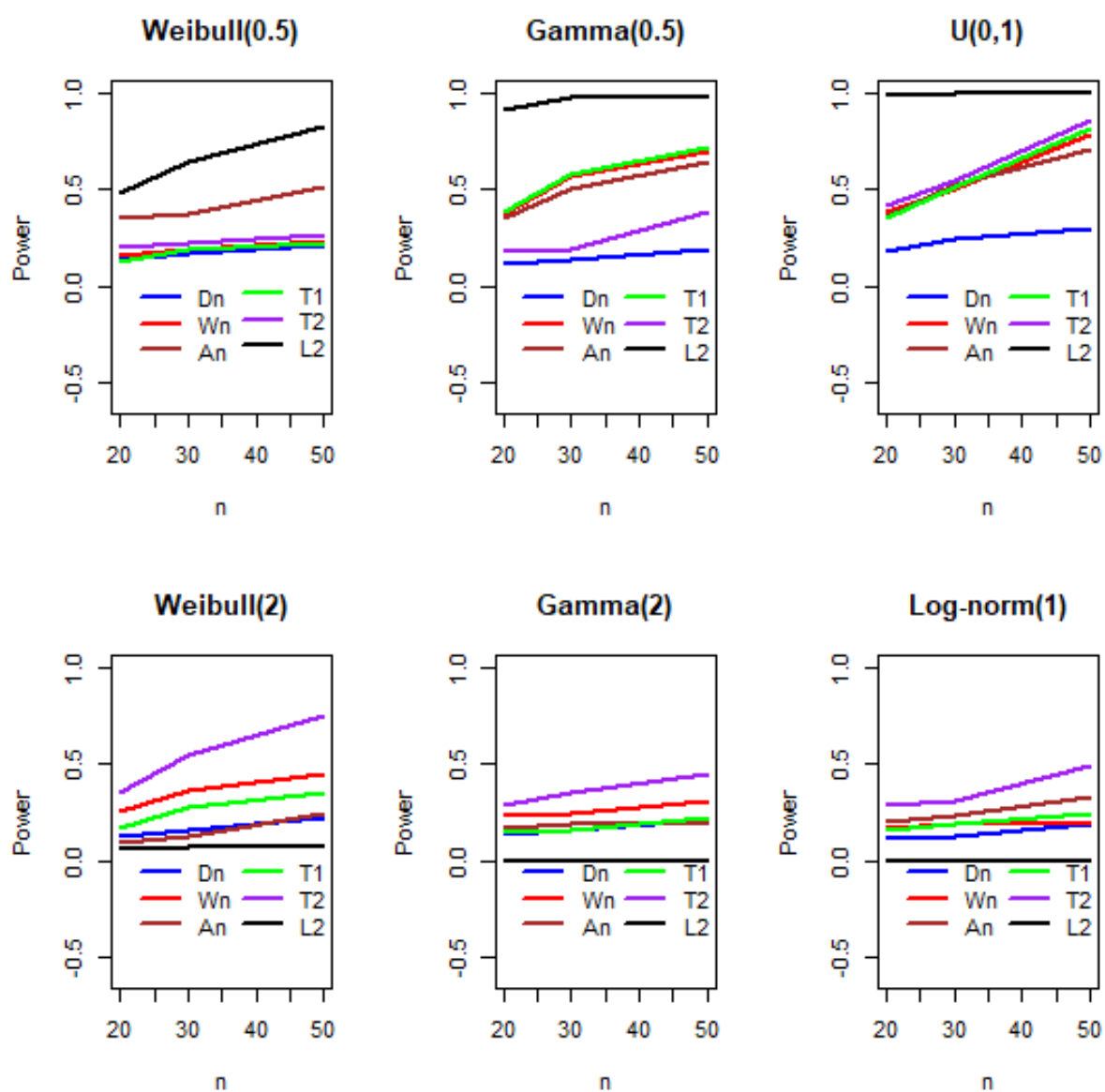
جدول ۴: مقایسه توان آماره‌های $\widehat{L2}_V^C$, $\widehat{T}_{n2}^*$, $\widehat{T}_{n1}^*$, $\widehat{A}_n^*$, $\widehat{W}_n^*$, $\widehat{D}_n$ در سطح $\alpha = 0.05$ با حجم نمونه $n = 50$

توزیع رقیب	نرخ سانسور	$\widehat{D}_n$	$\widehat{W}_n^*$	$\widehat{A}_n^*$	$\widehat{T}_{n1}^*$	$\widehat{T}_{n2}^*$	$\widehat{L2}_V^C$
<i>exponential</i> (۱)	۲۰	۰/۵۱	۰/۵۰	۰/۴۹	۰/۵۱	۰/۵۰	۰/۵۲
	۴۰	۰/۵۰	۰/۵۰	۰/۵۱	۰/۴۹	۰/۴۹	۰/۴۹
	۵۰	۰/۴۹	۰/۴۹	۰/۵۰	۰/۵۱	۰/۵۰	۰/۵۲
<i>Weibull</i> (۰/۵)	۲۰	۰/۲۱	۰/۲۳	۰/۵۲	۰/۲۲	۰/۲۷	۰/۸۳
	۴۰	۰/۱۹	۰/۲۱	۰/۴۲	۰/۱۹	۰/۲۰	۰/۸۲
	۵۰	۰/۱۶	۰/۱۸	۰/۳۵	۰/۱۴	۰/۱۹	۰/۸۲
<i>Gamma</i> (۰/۵)	۲۰	۰/۱۹	۰/۷۰	۰/۶۵	۰/۷۲	۰/۳۹	۰/۹۹
	۴۰	۰/۱۷	۰/۶۱	۰/۴۸	۰/۶۸	۰/۲۶	۰/۹۸
	۵۰	۰/۱۵	۰/۵۲	۰/۴۱	۰/۵۲	۰/۲۰	۰/۹۸
<i>U</i> (۰, ۱)	۲۰	۰/۳۰	۰/۷۹	۰/۷۱	۰/۸۲	۰/۸۶	۱
	۴۰	۰/۲۸	۰/۶۱	۰/۵۵	۰/۶۵	۰/۷۱	۱
	۵۰	۰/۲۴	۰/۴۵	۰/۳۸	۰/۴۸	۰/۶۱	۱
<i>Weibull</i> (۲)	۲۰	۰/۲۲	۰/۴۵	۰/۲۵	۰/۳۵	۰/۷۵	۰/۰۷
	۴۰	۰/۱۹	۰/۳۲	۰/۲۲	۰/۲۴	۰/۶۱	۰/۰۶
	۵۰	۰/۱۴	۰/۲۴	۰/۲۰	۰/۲۱	۰/۴۵	۰/۰۵
<i>Gamma</i> (۲)	۲۰	۰/۲۱	۰/۳۱	۰/۲۰	۰/۲۲	۰/۴۵	۰/۰۰۱
	۴۰	۰/۱۸	۰/۲۷	۰/۱۹	۰/۱۸	۰/۳۸	۰
	۵۰	۰/۱۶	۰/۲۱	۰/۱۸	۰/۱۴	۰/۳۰	۰
<i>Log - norm</i> (۱)	۲۰	۰/۱۹	۰/۲۰	۰/۳۳	۰/۲۴	۰/۴۹	۰
	۴۰	۰/۱۸	۰/۱۹	۰/۲۹	۰/۲۰	۰/۴۱	۰
	۵۰	۰/۱۴	۰/۱۶	۰/۲۵	۰/۱۸	۰/۳۰	۰



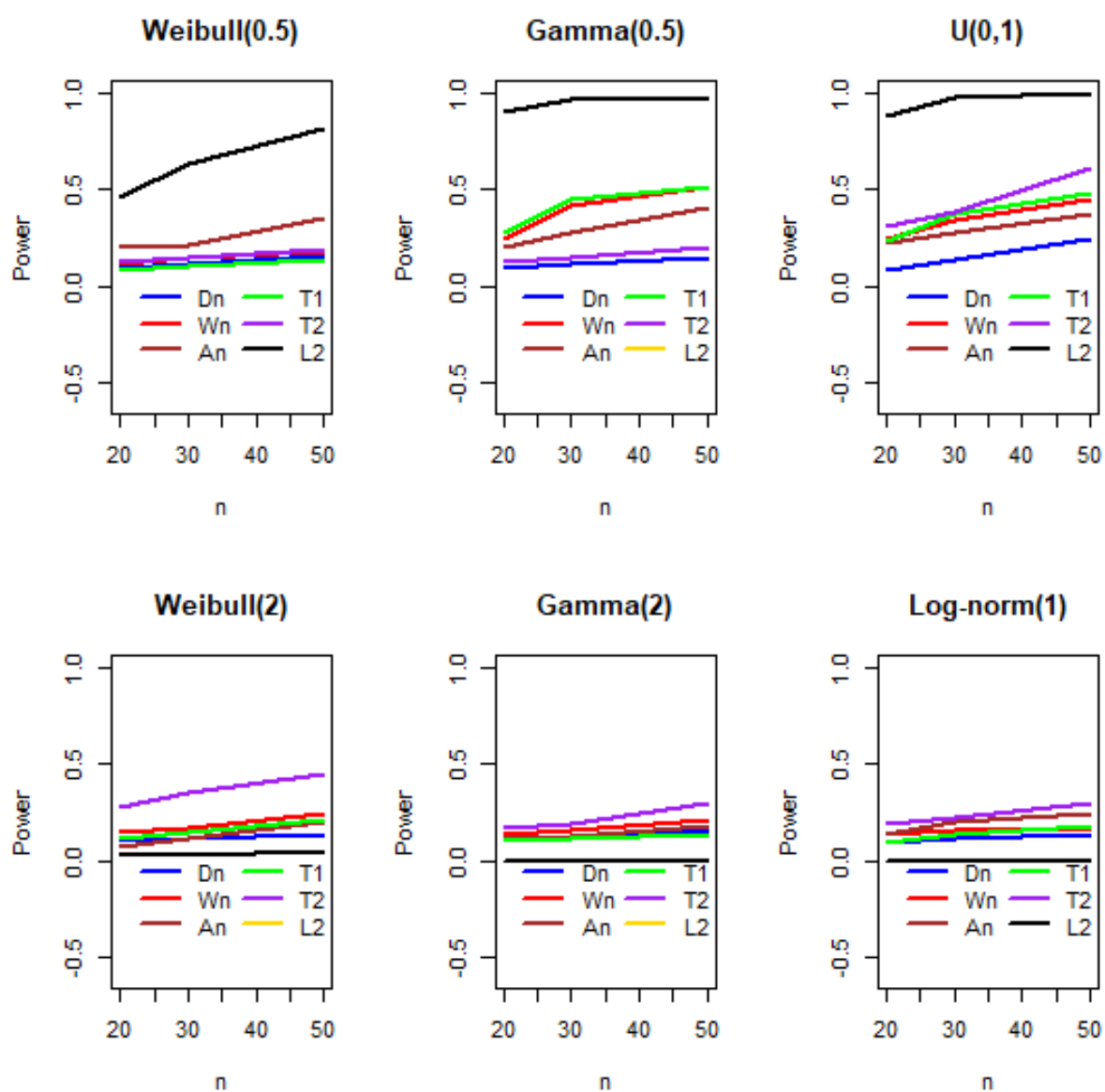
شکل ۱: مقایسه مقادیر توان آزمون پیشنهادی و آزمون‌های رقیب برای توزیع‌های جایگزین با نرخ سانسور ۲۰ درصد و حجم نمونه

$$n = 20, 30, 50$$



شکل ۲: مقایسه مقادیر توان آزمون پیشنهادی و آزمون‌های رقیب برای توزیع‌های جایگزین با نرخ سانسور ۴۰ درصد و حجم نمونه

$$n = 20, 30, 50$$



شکل ۳: مقایسه مقادیر توان آزمون پیشنهادی و آزمون‌های رقیب برای توزیع‌های جایگزین با نرخ سانسور ۵۰ درصد و حجم نمونه

$$n = 20, 30, 50$$

جدول ۵: مقایسه نتایج شبیه سازی آزمون های مختلف نمایی بودن برای داده های سانسور تصادفی

توزیع جایگزین	$Weibull(0.5)$	$Gamma(0.5)$	$Weibull(2)$	$Gamma(2)$	$U(0, 1)$	$Log - norm(1)$
بهترین آزمون	$\widehat{L2}_V^C$	$\widehat{L2}_V^C$	$\widehat{T}_{n2}^*$	$\widehat{T}_{n2}^*$	$\widehat{L2}_V^C$	$\widehat{T}_{n2}^*$

با توجه به مقادیر جداول زیر همان طور که در شکل های (۱)، (۲)، (۳) مشخص است، پرتوان ترین آزمون نسبت به توزیع های وایبول (۰/۵)، گاما (۰/۵)، یکنواخت استاندارد، آزمون واگرایی $L2$ است. اما برای توزیع های وایبول (۲)، گاما (۲)، لگ نرمال (۱)، آزمون واگرایی $\widehat{T}_{n2}^*$ دارای عملکرد مطلوبی نسبت سایر آزمون هاست.

۸ نتیجه گیری

نیازمند مفروضات بیشتری هستند که ممکن است در شرایط سانسور شده نقض شوند. برآوردگر ($L2$) با استفاده از تفاوت های مربعی بین توابع چگالی، دقت بیشتری در بازنمایی تفاوت های ساختاری بین توزیع ها ارائه می دهد. کاربرد این آزمون در مدل های بقا و تحلیل های مربوط به داده های پزشکی، ارزش آن را به عنوان ابزاری مطمئن در آزمون فرضیه افزایش می دهد. نتایج این پژوهش نشان داد که برای توزیع های وایبول، ($Gamma(0.5)$) و یکنواخت استاندارد، آزمون ($L2$) توان بالاتری نسبت به رقبای خود دارد. این نتیجه گیری به پژوهشگران کمک می کند تا در مواجهه با چالش های ناشی از سانسور تصادفی، از معیاری استفاده کنند که هم از نظر نظری و هم از نظر عملی بهینه باشد. در نهایت، یافته های ارائه شده در مقاله، اهمیت استفاده از واگرایی ($L2$) به عنوان معیاری نوین و کاربردی در تحلیل داده های سانسور شده را تأیید می کند.

در این مقاله یک آماره آزمون جدید برای سنجش نمایی بودن داده ها، هنگامی که داده ها به طور تصادفی سانسور می شوند، پیشنهاد شد. آزمون واگرایی ($L2$) به عنوان یک برآوردگر جدید، در تشخیص نمایی بودن داده های سانسور شده تصادفی توانمندی فوق العاده ای از خود نشان داد. این آزمون نسبت به آزمون های مبتنی بر تابع توزیع تجربی مانند کولموگروف اسمیرنوف، اندرسون دارلینگ و کرامر فون میزس عملکرد بالاتری از نظر توان آماری دارد.

نتایج شبیه سازی مونت کارلو نشان داد که در شرایط داده های سانسور شده، واگرایی ($L2$) در تشخیص صحیح انحراف از نمایی بودن، دقیق تر عمل می کند. دلیل این بهبود عملکرد، حساسیت بالاتر این معیار به تغییرات جزئی در انتهای توزیع است که در موارد سانسور، اهمیت ویژه ای دارد. از سوی دیگر، آزمون های مبتنی بر معیار اطلاع معمولاً

مراجع

- [۱] اسدی، مجید. (۱۳۹۲). آشنایی با نظریه قابلیت اعتماد، چاپ دوم، مرکز نشر دانشگاهی.
- [۲] باوفا، سمیه. پاییز (۱۳۹۰). رفتار مجانبی برآوردگر ناپارامتری تابع رگرسیون در مدل سانسور تصادفی با داده‌های α -آمیزنده. پایان‌نامه کارشناسی ارشد. دانشگاه فردوسی مشهد.
- [۳] پاک‌گوهر، علیرضا. (۱۳۹۷). ویژگی‌هایی از اندازه‌لین-وانگ برای تغییر سانسور نوع یک، دکتری تخصصی، دانشگاه فردوسی مشهد، پردیس بین‌الملل.
- [۴] شکوری، مرضیه. پاییز (۱۳۸۹). مروری بر آزمون‌های نیکویی برازش برای توزیع نمایی. پایان‌نامه کارشناسی ارشد. دانشگاه فردوسی مشهد.
- [5] Barr D.R., and Davidson. T. (1973), A Kolmogorov-Smirnov test for censored samples, *Technometrics*, **15**(4), 739-57.
- [6] Breslow, N. and Crowley, J. (1974), A large Sample Study of the lifetable and product limit estimates under random censorship, *Ann. Statist.*, **2**(3), 437-453
- [7] Cover, T.M. (1999), *Elements of information theory*, New York, John Wiley Sons.
- [8] Danish, M.Y., and Aslam, M. (2013), Bayesian estimation for randomly censored generalized exponential distribution under asymmetric loss functions, *J. Appl. Stat.*, **40**(5), 1106-1119.
- [9] Danish, M.Y., Aslam, M. (2014), Bayesian inference for the randomly censored Weibull distribution, *J.Stat. Comput. Simul.*, **84**(1), 215-230.
- [10] Efron, B. (1967). The Two Sample Problems with Censored Data, *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, **4**, 831-853.
- [11] Finkelstein, J.M., and Schafer, R.E. (1971), Improved goodness-of-fit tests, *Biometrika*, **58**(3), 641-645.
- [12] Ghasabi, M., Deypir, M. (2021). Using optimized statistical distances to confront distributed denial of service attacks in software defined networks, *Intelligent Data Analysis*, **25**(1), 155-176.
- [13] Ghitany, M.E. and Al-Awadhi, S. (2002). Maximum likelihood estimation of Burr XII distribution parameters under random censoring, *J. Appl. Stat.*, **29**, 955-965.
- [14] Gilbert, E. (1962). Random subdivisions of space into crystals, *The Annals of mathematical statistics*, **33**(3), 958-972.
- [15] Hasaballah, M.M., Balogun, O.S., Bakr, M.E. (2023). Investigation of exponential distribution utilizing randomly censored data under balanced loss functions and its application to clinical data, *Symmetry*, **15**(10), 1854.
- [16] Hellinger, E. (1909), Neue begründung der theorie quadratischer formen von unendlichvielen veränderlichen, *Journal für die reine und angewandte Mathematik*, **136**, 210-271.
- [17] Kailath, T. (1967). The divergence and Bhattacharyya distance measures in signal selection, *IEEE transactions on communication technology*, **15**(1), 52-60.
- [18] Kaplan, E.L., and Meier, P. (1958). Nonparametric estimation from incomplete observations, *Journal of the American statistical association*, **53**(282), 457-81.
- [19] Kim, N. (2012). Testing Exponentiality based on EDF Statistics for Randomly Censored Data when the Scale Parameter is Unknown, *The Korean Journal of Applied Statistics*, **25**, 311-319.

- [20] Klein, J.P., Moeschberger, M.L. (2003). *Survival analysis: techniques for censored and truncated data*, New York, Springer.
- [21] Koziol, J.A., and Green, S.B. (1976). A Cramer-von Mises Statistic for Randomly Censored Data, *Biometrika*, **63**, 465-474.
- [22] Krishna, H., and Goel, N. (2018). Classical and Bayesian inference in two parameter exponential distribution with randomly censored data, *Computational Statistics*, **33**, 249-275.
- [23] Krishna, H., Vivekanand, and Kumar, K. (2015). Estimation in Maxwell distribution with randomly censored data, *J. Stat. Comput. Simul.*, **85(17)**, 3560-3578.
- [24] Kumar, K., and Kumar, I. (2020), Parameter estimation for inverse Pareto distribution with randomly censored life time data, *Int. J. Agric. Stat. Sci.*, **16(1)**, 419-430.
- [25] Le Cam, L.M., and Yang, G.L. (2000). *Asymptotics in statistics: some basic concepts*, Springer Science Business Media.
- [26] Lilliefors, H.W. (1969). On the Kolmogorov-Smirnov test for the exponential distribution with mean unknown, *Journal of the American Statistical Association*, **64(325)**, 387-9.
- [27] Meier, P. (1975). Estimation of a Distribution Function from Incomplete Observations, *Journal of Applied Probability*, **12**, 67-87.
- [28] Omid, F., Habibirad, A., and Fakoor, V. (2019). Goodness of fit tests based on information criterion for randomly censored data, *Journal of Statistical Research of Iran*, **16**, 507-533.
- [29] Shannon, C.E. (1948). A Mathematical Theory of Communication, *The Bell System Technical Journal*, **27**, 379-432.
- [30] Soest, J.V. (1969), Some goodness of fit tests for exponential distributions, *Statistica Neerlandica*, **23(1)**, 41-51.
- [31] Tang, J., Cheng, Y., and Zhou, C. (2009). Sketch-based SIP flooding detection using Hellinger distance, *In GLOBECOM 2009-2009 IEEE Global Telecommunications Conference*, 1-6.
- [32] Vasicek, O. (1976). A test for normality based on sample entropy, *ournal of the Royal Statistical Society Series B: Statistical Methodology*, **38(1)**, 54-59.
- [33] Wiecezorkowski, R., and Grzegorzewski, P. (1999). Entropy estimators—improvements and comparisons, *Communications in Statistics-Simulation and Computation*, **28(2)**, 541-67.
- [34] McIlmurray, M.B. and Turkie, W. (1987). Controlled trial of gamma linolenic acid in Duke's C colorectal cancer, *BMJ*, **294**, 1260.
- [35] Yang, X., Yang, X., Yang, J., Ming, Q., Wang, W., Tian, Q., and Yan, J. (2021). Learning high-precision bounding box for rotated object detection via kullback-leibler divergence, *Advances in Neural Information Processing Systems*, **34**, 18381-18394.

A Goodness-of-fit test of Exponentiality for random censored data using L2 divergence

Yasser Hashemzaie¹, Seyyed Mahdi Amir Jahanshahi², Mohammad Hossein Dehghan³

Abstract:

In this paper, the L2 divergence test is introduced to assess the exponentiality of random censored data. The power of this test is compared using Monte Carlo simulation with other competing tests, including the Kolmogorov-Smirnov, Anderson-Darling, and Cramer-von-Mises tests based on the empirical distribution function and some other tests based on the information criterion. The results of the simulation studies of this research showed that the proposed test generally performs better than its other competing tests.

Keywords: Goodness of fit test, random censoring, exponential function test, L2 divergence test.

¹Mathematics teacher, Education and Training in Zahedan

²Associate Professor, Dept.of statistics, Faculty of Math, University of Sistan and Baluchestan, Daneshgah Ave., Zahedan, Iran.

³Assistant Professor, Dept.of statistics, Faculty of Math, University of Sistan and Baluchestan, Daneshgah Ave., Zahedan, Iran.