



Sample Size Determination in Non-Probability Research Based on Modeling Objectives and Predictive Power

Golinari, H.¹, Khorashadizadeh, M.^{1,*}, Mohtashami Borzadaran, Gh. R.²

¹Department of Statistics, University of Birjand, Birjand, Iran.

²Department of Statistics, Ferdowsi University of Mashhad, Mashhad, Iran.

Corresponding author: M. Khorashadizadeh, m.khorashadizadeh@birjand.ac.ir

Received: 12/12/2025 **Revised:** 20/6/2026 **Accepted and Published Online:** 22/6/2026.

Introduction

The absence of a complete sampling frame in many applied studies necessitates the use of non-probability sampling, for which no unified framework for sample size determination exists. This article proposes a novel “Sample Size Determination Based on Modeling Objectives and Predictive Power” (SBMF) framework. Its core innovation is a paradigm shift from population-based to model-based inference, defining the optimal sample size as the minimum size at which a statistical model achieves a target predictive performance.

Material and Methods

The SBF framework is operationalized in seven steps: **Step 1:** The objective is set to prediction/modeling. **Step 2:** “Model Performance Generalization” is selected. **Step 3:** A logistic regression model for purchase intention is defined, with statistical power to detect an age effect ($\delta = 0.2$, $\alpha = 0.05$) as the criterion. **Step 4:** A realistic synthetic reference population of 100,000 social media users is generated. **Step 5:** The core Monte Carlo simulation explicitly models a convenience sampling mechanism with a bias favoring younger users via an unnormalized weight function. For sample sizes from $n = 200$ to 1500, 1,000 replicates are drawn. **Step 6:** The model is fitted to each sample, and power is estimated. **Step 7:** The optimal n^* is the minimum n achieving 0.8 power. A Composite Performance Index (CPI) is developed for systematic comparison. The CPI averages four normalized

sub-criteria: Prediction Accuracy, Statistical Power, Bias Management Capability, and Flexibility.

Results and Discussion

The simulation was conducted under two distinct scenarios to assess the sensitivity of the framework. Under a weak-effect scenario ($\beta_{age} = 0.02$), the effect size was too small for the model to detect reliably, and the target power of 0.8 remained unattainable even at a sample size of $n = 1500$, where the achieved power was only 0.600. In contrast, under a stronger-effect scenario ($\beta_{age} = 0.05$), the SBMF framework successfully identified an optimal sample size of $n^* = 500$, which yielded a statistical power of 0.829. This outcome demonstrates the framework's practical efficiency and its direct dependence on the anticipated effect size. A systematic comparison of methods using the CPI revealed that SBMF achieved the highest overall performance score of 0.829. It was followed by probability sampling with a complete frame (CPI 0.75), the classical power analysis method assuming simple random sampling (CPI 0.72), the qualitative approach of theoretical sufficiency (CPI 0.65), and the rule of ten (CPI 0.55). The superior performance of SBMF is primarily attributable to its unique ability to systematically model and manage the selection bias inherent in non-probability sampling, a feature absent in classical approaches.

Conclusion

The simulation study confirms the efficacy of the SBMF framework for determining a defensible sample size in non-probability research. The model-based inference paradigm and the novel CPI provide a systematic pathway for enhancing the rigor of such studies. The successful application of the framework critically depends on the availability of credible auxiliary data and sound prior knowledge of key model parameters.

Keywords: Non-probability sampling, Sample size determination, Monte Carlo simulation, Model-based inference, Composite performance index.

Mathematics Subject Classification (2020): 62D05, 62F03.



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

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

حجم نمونه در پژوهش‌های غیراحتمالی مبتنی بر اهداف و قدرت پیش‌بینی‌کنندگی

حسن گلی ناری^۱، محمد خراشادی‌زاده^{۱*}، غلامرضا محتشمی برزادران^۲

^۱گروه، آمار دانشگاه بیرجند

^۲گروه، آمار دانشگاه فردوسی مشهد

نویسنده مسئول: محمد خراشادی‌زاده، m.khorashadizadeh@birjand.ac.ir

تاریخ دریافت: ۱۴۰۴/۹/۲۱ تاریخ بازنگری: ۱۴۰۵/۳/۳۰ تاریخ پذیرش و انتشار: ۱۴۰۵/۴/۱

چکیده: تعیین حجم نمونه در نمونه‌گیری غیراحتمالی، به دلیل نبود مبانی ریاضی احتمالی، همواره چالشی اساسی بوده است. این مقاله با معرفی چارچوب نوین تعیین حجم نمونه مبتنی بر اهداف مدل‌سازی و قدرت پیش‌بینی، پارادایم استنتاج را از جمعیت‌مبنا به مدل‌مبنا تغییر می‌دهد. در این رویکرد، حجم نمونه بر اساس حداقل اندازه لازم برای بهینه‌سازی عملکرد مدل، تشخیص روابط معنادار و دستیابی به پیش‌بینی تعمیم‌پذیر تعریف می‌شود. چارچوب پیشنهادی با تلفیق شبیه‌سازی پیشرفته و تحلیل توان آماری، و با اتکا به بسته‌های تخصصی نرم‌افزار R، راهکاری نظام‌مند و عملی برای ارتقای اعتبار و استانداردسازی پژوهش‌های غیراحتمالی ارائه می‌کند. اگرچه مطالعه شبیه‌سازی کارایی این روش را در سناریوی واقع‌گرایانه نشان می‌دهد، اما وابستگی آن به داده‌های کمکی معتبر و دانش پیشین از پارامترهای مدل، محدودیت اصلی کاربرد آن به شمار می‌رود.

واژه‌های کلیدی: نمونه‌گیری غیراحتمالی، تعیین حجم نمونه، شبیه‌سازی مونت کارلو، استنتاج مدل‌مبنا. کد موضوع‌بندی ریاضی (۲۰۲۰): 62F03، 62D05.

۱ مقدمه

تعیین حجم نمونه بهینه همواره به عنوان یکی از ارکان اساسی طراحی پژوهش‌های علمی مطرح بوده است. در روش‌های نمونه‌گیری احتمالی، این امر بر پایه مبانی ریاضی مستحکمی صورت می‌پذیرد که امکان استنتاج آماری از

نمونه به جامعه را فراهم می‌سازد. در این چارچوب، حجم نمونه لازم برای دستیابی به دقت و اطمینان آماری مورد نظر با استفاده از فرمول‌های شناخته‌شده محاسبه می‌شود. با این حال، در بسیاری از موقعیت‌های پژوهشی واقعی، محدودیت‌هایی نظیر عدم دسترسی به چارچوب نمونه‌گیری کامل، ماهیت اکتشافی پژوهش، یا تمرکز بر گروه‌های خاص و کم‌دسترس، پژوهشگران را ناگزیر به استفاده از نمونه‌های غیراحتمالی می‌کند. در این گونه طرح‌ها، مبانی ریاضی نمونه‌گیری احتمالی اعتبار خود را از دست می‌دهند و در نتیجه، منطق تعیین حجم نمونه اساساً دگرگون می‌شود. فقدان یک چارچوب تحلیلی یکپارچه و پذیرفته‌شده برای برآورد حجم نمونه در پژوهش‌های غیراحتمالی، به عنوان یک شکاف روش‌شناختی جدی و چالشی بنیادین در این حوزه خودنمایی می‌کند (الیوت و والیانت، ۲۰۱۷). گزارش انجمن آمریکایی پژوهش افکار عمومی (کوپر و همکاران، ۲۰۱۳) نیز به صراحت بر ضرورت توسعه راه‌حل‌های روش‌مند برای نمونه‌های غیراحتمالی تأکید کرده است.

گسترش چشمگیر پژوهش‌های کیفی، آمیخته و اکتشافی در دهه‌های اخیر که ذاتاً با روش‌های نمونه‌گیری غیراحتمالی از قبیل هدفمند و گلوله‌برفی همخوانی بیشتری دارند (کرسول و پلانز کلارک، ۲۰۱۷)، ضرورت پرداختن به مسئله تعیین حجم نمونه در این نوع پژوهش‌ها را دوچندان ساخته است. این مقاله با هدف غلبه بر این چالش، چارچوب نوینی را با عنوان تعیین حجم نمونه مبتنی بر هدف و قدرت مدل ارائه می‌دهد. نوآوری بنیادین این چارچوب، ایجاد یک تحول پارادایمی از استنتاج جمعیت‌مبنا به استنتاج مدل‌مبنا است که در آن، حجم نمونه بهینه بر اساس بهینه‌سازی عملکرد مدل آماری تعیین می‌شود. مقاله حاضر در پی پاسخگویی به سوالات پژوهشی زیر است: مبانی نظری تعیین حجم نمونه در چارچوب استنتاج مدل‌محور SBMF چیست؟ رویکرد پیشنهادی چگونه می‌تواند شکاف موجود در روش‌های رایج را پر کند؟ اجرای عملی این چارچوب در یک مطالعه شبیه‌سازی‌شده تا چه اندازه می‌تواند به تعیین حجم نمونه‌ای معتبر و بهینه منجر شود؟ و مزایا و محدودیت‌های اصلی آن در مقایسه با روش‌های متعارف کدامند؟

۲ مبانی نظری و مرور ادبیات

درک تفاوت‌های بنیادین بین دو رویکرد نمونه‌گیری احتمالی و غیراحتمالی، پیش‌نیاز هرگونه بحث روش‌شناختی درباره تعیین حجم نمونه است. رویکرد احتمالی بر پایه انتخاب تصادفی و شناخته بودن احتمال انتخاب هر واحد استوار است که امکان استنتاج آماری از نمونه به جامعه را با هدف دستیابی به تعمیم‌پذیری آماری فراهم می‌سازد (کوکران، ۱۹۷۷). در مقابل، رویکرد غیراحتمالی عمدتاً ریشه در پژوهش‌های کیفی و اکتشافی داشته و بر هدفمند بودن انتخاب نمونه برای دستیابی به غنای اطلاعاتی یا کشف نظریه تأکید دارد (گست و همکاران، ۲۰۲۰). در این رویکرد، تعمیم‌پذیری از نوع تحلیلی یا نظری است (فایرستون، ۱۹۹۳). این تقابل فلسفی، مبانی متفاوتی را برای مفهوم حجم نمونه طلب می‌کند.

در مواجهه با چالش تعیین حجم نمونه در پژوهش‌های غیراحتمالی، رویکردهای مختلفی متداول شده است. نخستین رویکرد که در پژوهش‌های کیفی مبتنی بر نظریه زمینه‌ای کاربرد دارد، بر مبنای کفایت نظری است؛ نقطه‌ای

که داده‌های جدید کمکی به توسعه نظریه نمی‌کنند. این معیار تفسیری، می‌تواند تکرارپذیری پژوهش را تحت تأثیر قرار دهد. مطالعات تجربی نشان می‌دهد دستیابی به این سطح از کفایت ممکن است به ۱۲-۱۵ مصاحبه نیاز داشته باشد (گست و همکاران، ۲۰۰۶). برخی محققان (هنیک و همکاران، ۲۰۱۷) بین کفایت کدگذاری (توقف استخراج کدهای جدید) و کفایت معنایی (درک کامل ابعاد مفاهیم) تمایز قائل می‌شوند که دستیابی به دومی نیازمند حجم نمونه بزرگ‌تری است. رویکرد دوم، قاعده‌های تجربی ساده‌شده‌ای مانند قاعده ده برابر در مدل‌سازی معادلات ساختاری است (اسکرایبر و همکاران، ۲۰۰۶) که فاقد پشتوانه نظری قوی بوده و ممکن است به برآوردهای نادرست منجر شوند. رویکرد سوم، استفاده از تحلیل توان آماری در طرح‌های شبه‌آزمایشی است که با وجود مبانی مستحکم، تفسیر آن در نمونه‌های کاملاً غیراحتمالی به دلیل نقض مفروضات انتخاب تصادفی بحث‌برانگیز است (شونلاو و کوپر، ۲۰۱۷).

مرور ادبیات سه خلأ روش‌شناختی عمده را آشکار می‌سازد: (۱) فقدان چارچوبی نظام‌مند برای تعیین حجم نمونه مبتنی بر اهداف مدل‌سازی در داده‌های غیراحتمالی، (۲) نیاز به روشی عینی و داده‌محور برای تصمیم‌گیری درباره کفایت حجم نمونه در پژوهش‌های کمی‌محور، و (۳) عدم وجود چارچوبی یکپارچه که اهداف پژوهش، ساختار مدل آماری و داده‌های کمکی را هم‌زمان در نظرگیرد. چارچوب پیشنهادی این مقاله با ایجاد گذار از استنتاج جمعیت‌مبنا به استنتاج مدل‌مبنا، پاسخی مستقیم به این نیازهاست.

۳ معرفی چارچوب پیشنهادی SBMF

در پاسخ به خلأهای روش‌شناختی شناسایی‌شده، این بخش به تشریح چارچوب پیشنهادی این پژوهش می‌پردازد. چارچوب SBMF بر یک تحول پارادایمی اساسی استوار است که گذار از استنتاج جمعیت‌مبنا به استنتاج مدل‌مبنا را در پژوهش‌های غیراحتمالی ممکن می‌سازد. در رویکرد نمونه‌گیری احتمالی، هدف غایی، تخمین بی‌طرفانه پارامترهای جامعه θ با حداقل خطای استاندارد $se(\hat{\theta})$ است. این استنتاج مستلزم دانستن توزیع نمونه‌گیری آماره $\hat{\theta}$ است که تنها تحت یک مکانیزم انتخاب تصادفی قابل استنتاج است (سازندال و همکاران، ۱۹۹۲). در شرایط فقدان چارچوب نمونه‌گیری، این مبانی ریاضی اعتبار خود را از دست می‌دهند. چارچوب پیشنهادی SBMF بر این اصل استوار است که در چنین شرایطی، اعتبار پژوهش نه از طریق تعمیم یک برآورد نقطه‌ای، بلکه از طریق توانایی مدل آماری نهایی در شناسایی روابط پایدار و داشتن قدرت پیش‌بینی سنجیده می‌شود. این چرخش، همسو با فلسفه علم معاصر است که بر اهمیت قابلیت پیش‌بینی و تبیین مدل‌ها به عنوان معیار اعتبار علمی تأکید می‌کند (شمونلی، ۲۰۱۰). در این رویکرد جدید، حجم نمونه بهینه n^* حجمی است که در آن مدل آماری M به توانایی کافی برای تشخیص روابط معنادار و مقاوم در برابر عدم قطعیت ناشی از مکانیزم نمونه‌گیری غیراحتمالی دست می‌یابد. این رویکرد، مسئله تعیین حجم نمونه را از حیطه نظریه برآورد کلاسیک به حیطه نظریه تصمیم‌گیری آماری و بهینه‌سازی منتقل می‌کند.

۳.۱ گام‌های اجرایی چارچوب SBMF

اجرای چارچوب SBMF برای تعیین حجم نمونه بهینه در روش‌های نمونه‌گیری غیراحتمالی، در هفت گام الگوریتمی و به هم پیوسته صورت می‌پذیرد. این گام‌ها که پژوهشگر را به صورت نظام‌مند از مرحله طراحی مفهومی تا تعیین حجم نمونه نهایی راهنمایی می‌کنند، به شرح زیر هستند:

گام ۱، تعیین هدف پژوهش: پژوهشگر در این گام، هدف اصلی مطالعه خود را در یکی از سه دسته اصلی مشخص می‌کند. هدف می‌تواند اکتشاف و تولید نظریه باشد که در آن معیار کیفی اشباع نظری ملاک عمل قرار می‌گیرد. هدف دوم، توصیف یک زیرجامعه مشخص است که در آن از فرمول‌های کلاسیک با در نظرگیری تغییرپذیری و دقت مطلوب استفاده می‌شود. هدف سوم و محوری این چارچوب، مدل‌سازی و پیش‌بینی برای جامعه بزرگتر است که منطق استنتاج مدل‌مبنا در آن به کار گرفته می‌شود.

گام ۲، انتخاب استراتژی تعمیم‌پذیری: با توجه به هدف تعیین شده در گام اول، استراتژی مناسب تعمیم‌پذیری انتخاب می‌شود. برای اهداف اکتشافی، استراتژی تعمیم‌پذیری از نوع تعمیم تحلیلی است. برای اهداف توصیفی، استراتژی تعمیم‌پذیری از نوع تعمیم توصیفی با استفاده از برآوردهای ناریب و کران‌های اطمینان است. برای اهداف مدل‌سازی، استراتژی تعمیم‌پذیری از نوع تعمیم عملکرد مدل است. در این حالت، پژوهشگر به دنبال تعمیم پارامترهای مدل نیست، بلکه به دنبال تعمیم عملکرد پیش‌بینی مدل است. هدف، اطمینان از این است که عملکرد کلی مدل، یعنی توانایی آن در پیش‌بینی دقیق نتایج برای داده‌های جدید، به جامعه بزرگتر تعمیم می‌یابد.

گام ۳، تعریف مدل آماری و معیار عملکرد (مختص هدف مدل‌سازی و پیش‌بینی): در این گام، پژوهشگر مدل آماری نهایی خود را به دقت مشخص می‌کند. این مدل بر اساس سوالات تحقیق و تئوری‌های موجود در حوزه موضوعی، و نه بر اساس شناخت کامل از جامعه، طراحی می‌شود. برای مثال، یک مدل رگرسیون لوژستیک برای پیش‌بینی یک نتیجه دودویی یا یک مدل رگرسیون خطی برای پیش‌بینی یک نتیجه پیوسته می‌تواند انتخاب شود. همزمان، معیار کمی عملکرد مدل نیز باید مشخص شود. در مدل‌های پیش‌بینی، این معیار می‌تواند صحت کلی، سطح زیر منحنی^۱ (AUC)، یا R^2 باشد. در مدل‌های علی، معیار عملکرد معمولاً توان آماری برای تشخیص یک اندازه اثر مینیمم مهم از نظر بالینی یا عملی است.

گام ۴، گردآوری داده‌های کمکی و ساخت جامعه مرجع: در شرایط فقدان چارچوب نمونه‌گیری، پژوهشگر باید از داده‌های کمکی برای ساخت یک جامعه مرجع شبیه‌سازی شده استفاده کند. این داده‌های کمکی می‌توانند شامل داده‌های حاصل از پژوهش‌های قبلی با جامعه مشابه، داده‌های پانلی بزرگ‌مقیاس، آمارهای رسمی منتشر شده (نظیر میانگین و انحراف معیار متغیرهای جمعیت‌شناختی)، یا ثبت‌های اداری باشند. جامعه مرجع با توزیع‌های واقع‌بینانه برای متغیرهای کلیدی و با در نظرگیری ساختار همبستگی بین آن‌ها تولید می‌شود. در تعیین ساختار همبستگی بین متغیرها، دو رویکرد کلی قابل اتخاذ است. رویکرد نخست، استفاده از ماتریس همبستگی مستخرج از داده‌های واقعی موجود است که می‌تواند از طریق مطالعات پیشین، داده‌های ثانویه یا پایگاه‌های داده معتبر تأمین شود. رویکرد دوم،

¹Area Under the ROC Curve

تعریف متغیرها به صورت مستقل از یکدیگر است که خود می‌تواند مبتنی بر دو منطق متفاوت باشد: الف) فرض استقلال متغیرها در جامعه واقعی (مانند بسیاری از متغیرهای جمعیت‌شناختی در جوامع بزرگ) و ب) اتخاذ رویکرد محافظه‌کارانه برای ارزیابی خالص توان مدل در تشخیص اثرات مستقل. انتخاب بین این رویکردها باید با توجه به ماهیت متغیرها، ویژگی‌های جامعه هدف و اهداف پژوهش صورت گیرد و به صراحت بیان شود.

گام ۵، شبیه‌سازی مونت کارلو: در این گام، با استفاده از جامعه مرجع ساخته‌شده در گام پیشین، نمونه‌های متعددی با حجم‌های مختلف از $n_{\min}$ تا $n_{\max}$ شبیه‌سازی می‌شوند. هدف آن است که فرآیند انتخاب نمونه در دنیای واقعی، که در آن همه اعضای جامعه شانس برابر برای ورود به نمونه ندارند، بازآفرینی شود. برای این منظور، فرض می‌کنیم $f(x)$ توزیع توأم متغیرهای کمکی X در جامعه مرجع باشد. در نمونه‌گیری غیراحتمالی، توزیع نمونه دیگر $f(x)$ نیست، بلکه یک توزیع وزنی به صورت $f^*(x) = w(x) \cdot f(x)$ خواهد بود. در این رابطه، $w(x)$ یک تابع وزن نرم نشده است که میزان انحراف توزیع نمونه از توزیع جامعه را بر اساس مقادیر X تعیین می‌کند. برای تبدیل این توزیع به یک توزیع احتمال خوش‌تعریف، وزن‌ها را نرم می‌کنیم تا مجموع آنها برابر با یک شود. انتخاب فرم تابعی مناسب برای $w(x)$ بر اساس اطلاعات پیشین پژوهشگر درباره ماهیت اریبی در روش نمونه‌گیری غیراحتمالی مورد استفاده صورت می‌گیرد. برای نمونه، اگر طرح نمونه‌گیری از نوع در دسترس باشد، پژوهشگر می‌تواند با تکیه بر دانش تخصصی خود، تابع وزن را به گونه‌ای تعریف کند که احتمال انتخاب را برای زیرگروه‌های خاصی از جامعه (مانند کاربران جوان‌تر) افزایش دهد. پس از محاسبه و نرم‌سازی $w(x)$ برای تمامی واحدهای جامعه مرجع، انتخاب واحدهای نمونه با استفاده از نمونه‌گیری با جایگذاری از توزیع احتمال حاصل صورت می‌پذیرد. این فرآیند نمونه‌گیری وزنی را برای هر حجم نمونه مورد نظر، به تعداد کافی (مثلاً ۱۰۰۰ بار) تکرار می‌کنیم تا توزیع نمونه‌گیری آماره‌های مورد نظر به دقت برآورد شود.

گام ۶، برازش مدل و محاسبه معیار عملکرد: برای هر نمونه شبیه‌سازی شده در هر حجم نمونه، مدل آماری مشخص‌شده در گام ۳ برازش داده می‌شود و معیار عملکرد مربوطه محاسبه می‌گردد. برای مثال، اگر معیار عملکرد، توان آماری برای تشخیص یک پارامتر خاص β_j باشد، فرضیه $H_0: \beta_j = 0$ در مقابل $H_1: \beta_j = \delta$ آزمون می‌شوند، که در آن δ کوچکترین اندازه اثر معنادار است. توان آماری به عنوان نسبت رد کردن صحیح H_0 در تمامی تکرارهای شبیه‌سازی محاسبه می‌شود. در نهایت، منحنی توان (یا هر معیار عملکرد دیگر) بر حسب حجم نمونه رسم می‌شود.

گام ۷، تعیین حجم نمونه بهینه: حجم نمونه بهینه n^* به عنوان کوچکترین حجم نمونه‌ای تعریف می‌شود که در آن معیار عملکرد به آستانه از پیش تعیین‌شده مطلوب می‌رسد. این آستانه می‌تواند یک توان آماری مشخص (مانند ۰/۸ یا ۰/۹) یک دقت پیش‌بینی حداقلی یا یک مقدار آستانه برای سطح زیر منحنی AUC باشد، یعنی $\{n^* = \min\{n \in \mathbb{N} \mid \text{Performance}(n) \geq \gamma\}$ ، که در آن γ سطح آستانه توان مطلوب است.

۳.۲ مبانی نرم‌افزاری

پیاده‌سازی عملی چارچوب SBMF مستلزم استفاده از نرم‌افزارهای آماری قدرتمند و انعطاف‌پذیر است. این پژوهش برای اجرای جامع تمام مراحل شبیه‌سازی و تحلیل آماری، محیط نرم‌افزار R نسخه ۴.۳.۰ را به کار گرفته است. انتخاب R به دلیل قابلیت‌های گسترده این محیط در زمینه شبیه‌سازی، تحلیل‌های آماری پیشرفته و توسعه الگوریتم‌های سفارشی صورت پذیرفته است. برای تولید داده‌های مصنوعی و ساخت جامعه مرجع درگام ۴، از بسته تخصصی simstudy (گروسمن و همکاران، ۲۰۲۳) استفاده شده است. این بسته با ارائه توابع ساختاریافته، امکان تولید داده‌های همبسته با توزیع‌های مشخص آماری را فراهم می‌آورد. برای انجام شبیه‌سازی‌های مونت کارلو درگام ۵ و تحلیل‌های بوت‌استرپ، بسته boot انتخاب شده است. این بسته با ارائه الگوریتم‌های کارآمد، برآوردهای مناسبی از خطای استاندارد و فواصل اطمینان را ممکن می‌سازد (کنتی و ریپلی، ۲۰۲۳). در مرحله تحلیل توان آماری درگام ۶، بسته pwr (شامپلی، ۲۰۲۰) به کار گرفته شده است. این بسته با پیاده‌سازی روابط ریاضی محاسبه توان، دقت لازم برای برآورد حجم نمونه تحت شرایط مختلف اندازه اثر و سطح معنی‌داری را تضمین می‌کند. فرآیند اعتبارسنجی مدل و ارزیابی عملکرد در این گام نیز با بهره‌گیری از بسته جامع caret (کوهن، ۲۰۲۳) انجام پذیرفته است. تمامی مراحل فوق در قالب یک بسته نرم‌افزاری یکپارچه و مستند پیاده‌سازی شده است که شامل توابع سفارشی برای اجرای گام‌های هفت‌گانه چارچوب SBMF می‌باشد. این مجموعه از بسته‌های استاندارد، ضمن اطمینان از دقت محاسباتی، امکان مقایسه‌پذیری نتایج با مطالعات مشابه را فراهم می‌آورد. انتخاب این بسته‌ها بر اساس پایداری، کارایی و قابلیت اطمینان آن‌ها در محیط‌های پژوهشی صورت گرفته است.

۳.۳ پیاده‌سازی گام به گام SBMF

پیش از ورود به جزئیات اجرایی، ساختار کلی این مطالعه موردی به اختصار ترسیم می‌شود. هدف اصلی، مدل‌سازی و پیش‌بینی تمایل به خرید یک محصول جدید در جامعه کاربران یک شبکه اجتماعی پرتعداد است. برای این منظور، یک مدل رگرسیون لوژیستیک به عنوان مدل آماری پایه انتخاب شده است. این انتخاب، بر اساس کاربرد گسترده در بازاریابی برای پیش‌بینی متغیرهای وابسته دودویی صورت گرفته است (هوسمر و همکاران، ۲۰۱۳). استراتژی تعمیم‌پذیری از نوع تعمیم عملکرد مدل و معیار عملکرد نیز توان آماری برای تشخیص اثر متغیر سن تعیین گردیده است. جامعه مرجع، شامل ۱۰۰,۰۰۰ کاربر شبیه‌سازی شده با مشخصات جمعیت‌شناختی واقع‌گرایانه است و نمونه‌گیری غیراحتمالی از نوع در دسترس با اربابی به سمت کاربران جوان‌تر مدل‌سازی می‌شود. هدف نهایی، یافتن کمترین حجم نمونه‌ای است که در آن مدل به توان آماری مطلوب ۰/۸ دست می‌یابد. در ادامه، پیاده‌سازی دقیق هر یک از هفت گام چارچوب SBMF در این مطالعه موردی تشریح می‌گردد.

پیاده‌سازی گام اول: تعیین هدف پژوهش. در این گام، هدف مطالعه مطابق با طبقه‌بندی چارچوب SBMF در دسته سوم، یعنی مدل‌سازی و پیش‌بینی برای جامعه بزرگ‌تر قرار گرفت. به طور مشخص، پژوهش به دنبال تبیین رابطه بین ویژگی‌های کاربران و تمایل به خرید آن‌ها نیست، بلکه هدف، ساخت مدلی با حداکثر توان پیش‌بینی‌کنندگی

برای تمایل به خرید در جامعه هدف است. این تمایز، زیربنای منطقی برای انتخاب‌های بعدی، به‌ویژه استراتژی تعمیم‌پذیری و معیار عملکرد مدل را شکل می‌دهد.

پیاده‌سازی گام دوم: انتخاب استراتژی تعمیم‌پذیری. هم‌راستا با هدف مدل‌سازی و پیش‌بینی، استراتژی تعمیم‌پذیری از نوع تعمیم عملکرد مدل انتخاب گردید. در این استراتژی، بر خلاف تعمیم توصیفی که به دنبال تعمیم یک برآورد نقطه‌ای (مانند میانگین) به جامعه است، هدف آن است که توانایی پیش‌بینی مدل بر روی داده‌های جدید و مشاهده‌نشده، در سطح مطلوبی حفظ شود (شمولی، ۲۰۱۰). این استراتژی، به‌طور خاص برای پژوهش‌های مبتنی بر نمونه‌های غیراحتمالی که فاقد چارچوب نمونه‌گیری هستند، مناسب است، زیرا اعتبار یافته‌ها را نه بر اساس قواعد تصادفی بودن، بلکه بر اساس کارایی پیش‌بینی ارزیابی می‌کند.

پیاده‌سازی گام سوم: تعریف مدل آماری، معیار عملکرد و طراحی مطالعه شبیه‌سازی. با تثبیت هدف و استراتژی تعمیم‌پذیری، در این گام، ارکان اصلی مطالعه شبیه‌سازی، یعنی مدل آماری، معیار سنجش عملکرد آن، و جامعه مرجع طراحی شدند. برای مدل‌سازی تمایل به خرید، یک مدل رگرسیون لوژستیک به‌عنوان مدل آماری پایه انتخاب گردید. این انتخاب به دلیل جایگاه این مدل در مطالعات بازاریابی برای پیش‌بینی متغیرهای وابسته دودویی صورت پذیرفت (هوسمر و همکاران، ۲۰۱۳). متغیرهای مستقل اولیه این مدل شامل سن، جنسیت و منطقه جغرافیایی (در پنج سطح) تعریف شدند. شایان ذکر است که مدل‌های کامل‌تر تمایل به خرید، عوامل متعدد دیگری نظیر قیمت، کیفیت، وفاداری به برند و تأثیرات اجتماعی را نیز شامل می‌شوند. با این حال، هدف این شبیه‌سازی، ارائه یک تصویر کامل از کلیه عوامل مؤثر بر رفتار خرید نبود، بلکه نمایش عملی کارایی و انعطاف‌پذیری چارچوب SBMF در تعیین حجم نمونه بهینه برای یک مدل آماری مفروض است. بدیهی است که چارچوب پیشنهادی به این مدل خاص محدود نبوده و پژوهشگران می‌توانند مدل‌های بسیار پیچیده‌تر، از جمله مدل‌های تعاملی، غیرخطی و یا چندسطحی را مطابق با نیازهای پژوهشی خود در آن پیاده‌سازی کنند.

هم‌راستا با استراتژی تعمیم‌پذیری انتخاب‌شده، معیار عملکرد مدل، توان آماری برای تشخیص اثر معنادار متغیر سن بر تمایل به خرید تعیین گردید. اندازه اثر مینیمم مهم از نظر علمی $\delta = 0.2$ (بر اساس مقیاس لگاریتم بخت) و سطح معنی‌داری $\alpha = 0.05$ در نظر گرفته شد. توان آماری مطلوب نیز برابر با 0.8 تعیین گردید. این معیار، احتمال تشخیص یک اثر با اندازه مشخص را در صورت وجود آن در جامعه ارزیابی می‌کند و با هدف تعمیم عملکرد مدل همخوانی کامل دارد.

برای اجرای شبیه‌سازی، یک جامعه مرجع مصنوعی شامل ۱۰۰,۰۰۰ کاربر فرضی یک شبکه اجتماعی طراحی گردید. برای واقع‌گرایی بیشتر، توزیع متغیرهای جمعیت‌شناختی در این جامعه به گونه‌ای تعریف شد که می‌تواند از آمارهای منتشرشده پلتفرم‌های اجتماعی واقعی حاصل شده باشد:

- توزیع سن کاربران از یک توزیع نرمال با میانگین ۲۸ سال و انحراف معیار ۸ سال پیروی می‌کند.
- توزیع جنسیتی کاربران به نسبت ۵۵ درصد مرد و ۴۵ درصد زن در نظر گرفته شد.

ضرایب مدل لوژستیک نیز بر اساس مرور نظام‌مند مطالعات پیشین در حوزه رفتار مصرف‌کننده و با هدف تولید نرخ

پایه واقع‌بینانه‌ای برای تمایل به خرید (در محدوده ۲۸ تا ۳۲ درصد) کالبره شدند.

پیاده‌سازی گام چهارم: تولید جامعه مرجع. بر اساس طراحی صورت‌گرفته در گام پیشین، جامعه مرجع با استفاده از بسته simstudy در نرم‌افزار R تولید گردید (گروسمن و همکاران، ۲۰۲۳). برای این منظور، ۱۰۰,۰۰۰ مشاهده با متغیرهای مستقل سن، جنسیت و منطقه جغرافیایی (در پنج سطح) به صورت مستقل از یکدیگر تولید شدند. سپس، متغیر وابسته تمایل به خرید برای هر مشاهده بر اساس مدل لوژیستیک تعریف‌شده در گام سوم شبیه‌سازی گردید. به این ترتیب، جامعه مرجعی ایجاد شد که هم توزیع متغیرهای کمکی و هم رابطه مفروض آن‌ها با متغیر هدف را منعکس می‌سازد.

پیاده‌سازی گام پنجم: شبیه‌سازی مونت کارلو. با استفاده از جامعه مرجع تولیدشده در گام پیشین، فرآیند شبیه‌سازی مونت کارلو برای حجم‌های نمونه مختلف در محدوده ۲۰۰ تا ۱۵۰۰ واحد (با گام‌های ۱۰۰ تایی) اجرا شد. برای هر حجم نمونه، تعداد ۱۰۰۰ تکرار شبیه‌سازی صورت پذیرفت. هسته اصلی این شبیه‌سازی، مدل‌سازی صریح مکانیزم نمونه‌گیری غیراحتمالی^۱ بود. هدف آن بود که فرآیند انتخاب نمونه در دنیای واقعی، که در آن همه اعضای جامعه شانس برابر برای ورود به نمونه ندارند، بازآفرینی شود.

برای این منظور، یک مکانیزم نمونه‌گیری غیراحتمالی از نوع در دسترس با اربیبی سیستماتیک به سمت کاربران جوان‌تر مدل‌سازی گردید. تابع وزن $w(x)$ بر اساس متغیر سن و به صورت یک تابع لوژیستیک معکوس به صورت $i^{-1} = [1 + \exp(0.5 \times (\text{سن}_i - 25))]^{-1}$ انتخاب واحد P تعریف شد، که در آن پارامتر ۰/۵ شدت اربیبی را کنترل می‌کند. نقطه آستانه ۲۵ سال نیز به گونه‌ای انتخاب شد که انحراف از میانگین سنی جامعه (۲۸ سال) به سمت کاربران جوان‌تر اعمال گردد. بر اساس این تابع، یک کاربر ۲۰ ساله تقریباً ۲/۵ برابر یک کاربر ۳۰ ساله شانس انتخاب شدن دارد.

در هر تکرار شبیه‌سازی، ابتدا برای تمامی ۱۰۰,۰۰۰ عضو جامعه مرجع، احتمال انتخاب بر اساس رابطه فوق محاسبه و سپس با تقسیم بر مجموع کل، نرم‌سازی گردید تا یک توزیع احتمال خوش‌تعریف با مجموع احتمالات برابر با یک حاصل شود. در نهایت، انتخاب واحدهای نمونه با استفاده از نمونه‌گیری با جایگذاری از این توزیع احتمال وزنی صورت پذیرفت. این فرآیند برای هر یک از ۱۰۰۰ تکرار در هر حجم نمونه، به طور مستقل تکرار گردید تا تغییرپذیری ناشی از مکانیزم نمونه‌گیری غیراحتمالی در برآورد توان آماری منعکس شود.

پیاده‌سازی گام ششم: برازش مدل و محاسبه معیار عملکرد. برای هر یک از ۱۰۰۰ نمونه شبیه‌سازی‌شده در هر حجم نمونه، مدل رگرسیون لوژیستیک تعریف‌شده در گام سوم برازش داده شد. سپس، معیار عملکرد یعنی توان آماری برای تشخیص اثر معنادار متغیر سن محاسبه گردید. توان آماری برای هر حجم نمونه n ، به‌عنوان نسبت دفعاتی که ضریب متغیر سن در سطح ۰/۰۵ از نظر آماری معنادار تشخیص داده شد، با استفاده از رابطه $\text{Power}(n) = H_1 : \beta_j = \frac{\delta}{1.000} \sum_{i=1}^{1000} I(p_{\text{value},i} < 0.05)$ برآورد گردید، که در آن $I(\cdot)$ تابع نشانگر است و $p_{\text{value},i}$ مقدار احتمال (p-value) متناظر با آزمون معناداری ضریب متغیر سن در i -امین تکرار شبیه‌سازی است. این روش محاسبه توان آماری، رویکردی استاندارد در مطالعات شبیه‌سازی است (شامپلی، ۲۰۲۰).

¹Non-Probability Sampling Mechanism

پیاده‌سازی گام هفتم: تعیین حجم نمونه بهینه. در این گام، منحنی توان آماری بر حسب حجم نمونه رسم گردید. حجم نمونه بهینه n^* به عنوان کوچکترین حجم نمونه‌ای تعریف می‌شود که در آن توان آماری به آستانه مطلوب $\alpha/8$ دست یابد. یعنی $n^* = \min\{n \in \mathbb{N} \mid \text{Power}(n) \geq \alpha/8\}$. این آستانه، سطحی از توان را مشخص می‌کند که در آن احتمال تشخیص یک اثر با اندازه مفروض، در صورت وجود، به سطح قابل قبولی می‌رسد.

۴ بررسی و تفسیر نتایج شبیه‌سازی

به منظور بررسی دقیق‌تر نحوه عملکرد مدل پیشنهادی در دو سناریوی متفاوت با تغییر ضرایب مدل برای چارچوب پیشنهادی SBMF، عملیات شبیه‌سازی انجام شد. طی شبیه‌سازی‌های انجام شده نتایج بسیار جالبی حاصل شد که در ادامه به بررسی آنها می‌پردازیم. لازم به ذکر است که برای تضمین تکرارپذیری نتایج، در تمام مراحل شبیه‌سازی از مقدار ثابت ۱۲۳ seed استفاده شده است.

۴.۱ تحلیل سناریوی اول با ضرایب ضعیف

نتایج شبیه‌سازی اول با ضرایب $\beta_{\text{سن}} = \alpha/2$ و $\beta_{\text{جست}} = -\alpha/3$ نشان می‌دهد که سیستم حتی با حجم نمونه 1500 واحد نیز نتوانسته است به توان آماری مطلوب $\alpha/8$ دست یابد. این وضعیت حاکی از آن است که اثرات متغیرهای مستقل، به اندازه‌ای ضعیف هستند که مدل قادر به تشخیص روابط معنادار در محدوده حجم نمونه‌های مورد بررسی نیست.

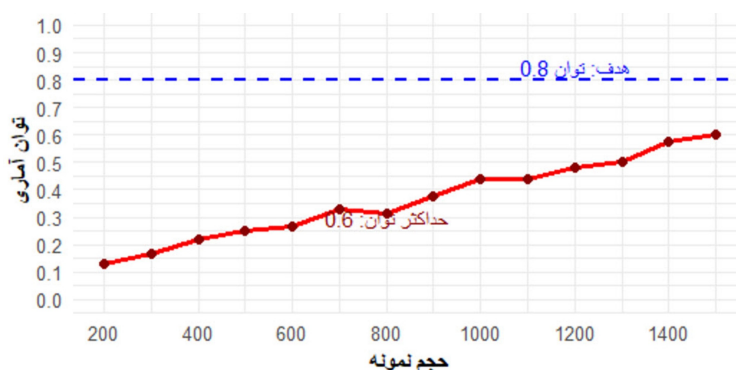
جدول ۱. نتایج شبیه‌سازی سناریوی اول با ضرایب ضعیف ($\beta_{\text{سن}} = \alpha/2$, $\beta_{\text{جست}} = -\alpha/3$)

پارامتر	مقدار
مدت اجرا (دقیقه)	۱۹
حجم جامعه مرجع	۱۰۰,۰۰۰
تکرارهای شبیه‌سازی	۱۰۰۰
حجم نمونه بهینه	دست‌نیافته
توان هدف	$\alpha/8$

همانگونه که در جدول ۲ ملاحظه می‌شود، در سناریوی اول با ضرایب ضعیف، علی‌رغم بررسی حجم نمونه تا 1500 واحد، توان آماری مطلوب $\alpha/8$ حاصل نشد. دلیل اصلی این پدیده را می‌توان در اندازه اثر کوچک متغیرها جستجو کرد. ضریب سن $\alpha/2$ به این معناست که با افزایش هر واحد سن، لگاریتم بخت تمایل به خرید به میزان $\alpha/2$ واحد افزایش می‌یابد که اثر بسیار ناچیزی محسوب می‌شود. این ضعف در قدرت تمایز مدل باعث می‌شود که حتی با حجم نمونه‌های بزرگ، مدل نتواند به درستی اثر متغیرها را تشخیص دهد. منحنی توان شکل ۱ در شبیه‌سازی اول، به وضوح نشان می‌دهد که روند رشد توان بسیار کند است و حتی در حجم نمونه 1500 نیز به آستانه مطلوب

جدول ۲. نتایج تحلیل توان برای حجم‌های نمونه مختلف در سناریوی اول

توان آماری	حجم نمونه (n)
۰/۱۲۸	۲۰۰
۰/۱۶۷	۳۰۰
۰/۲۱۹	۴۰۰
۰/۲۴۹	۵۰۰
۰/۲۶۸	۶۰۰
۰/۳۲۹	۷۰۰
۰/۳۱۴	۸۰۰
۰/۳۷۷	۹۰۰
۰/۴۴۱	۱۰۰۰
۰/۴۳۷	۱۱۰۰
۰/۴۸۱	۱۲۰۰
۰/۵۰۱	۱۳۰۰
۰/۵۷۸	۱۴۰۰
۰/۶۰۰	۱۵۰۰



شکل ۱. منحنی توان بر حسب حجم نمونه SBMF با برآوردهای ضعیف

۸/۰ نمی‌رسد. این موضوع اهمیت برآورد صحیح اندازه اثر در مطالعات پیشین را بیش از پیش آشکار می‌سازد. شیب ملایم منحنی نشان‌دهنده این واقعیت است که افزایش حجم نمونه تأثیر چندانی بر بهبود توان مدل ندارد.

۴.۲ تحلیل سناریوی دوم با ضرایب قوی

با توجه به جدول ۴ پس از افزایش ضرایب به مقادیر $\beta_1 = 0.5$ و $\beta_2 = -0.5$ ، جنسیت، عملکرد مدل به میزان قابل توجهی بهبود یافت. اگرچه این تغییرات جزئی هستند، اما تأثیر محسوسی بر توان آماری مدل داشتند به گونه‌ای که حجم نمونه بهینه ۵۰۰ واحدی با توان ۰/۸۲۹ قابل دستیابی شد. این بهبود قابل توجه ناشی از افزایش قدرت تمایز مدل است. ضریب سن ۰/۵ به معنای اثر قوی‌تر این متغیر بر نتیجه است و مدل می‌تواند با حجم نمونه

جدول ۳. نتایج شبیه‌سازی سناریوی دوم با ضرایب قوی ($\beta_{\text{سن}} = 0/5$, $\beta_{\text{جنسیت}} = -0/5$)

پارامتر	مقدار
مدت اجرا (دقیقه)	۵۷
حجم جامعه مرجع	۱۰۰,۰۰۰
تکرارهای شبیه‌سازی	۱۰۰۰
حجم نمونه بهینه	۵۰۰
توان کسب‌شده	۰/۸۲۹
توان هدف	۰/۸

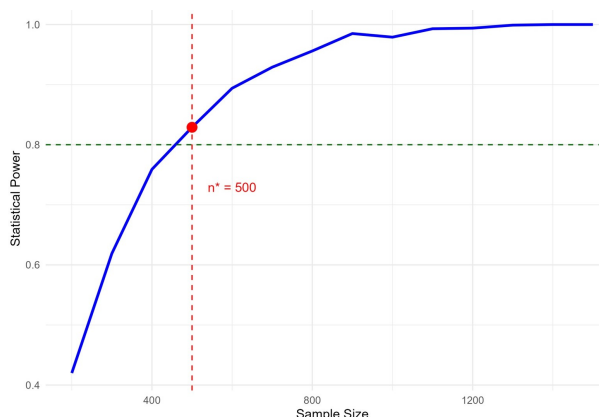
جدول ۴. نتایج تحلیل توان برای حجم‌های نمونه مختلف در سناریوی دوم

حجم نمونه (n)	توان آماری
۲۰۰	۰/۴۲
۳۰۰	۰/۶۱۹
۴۰۰	۰/۷۵۹
۵۰۰	۰/۸۲۹
۶۰۰	۰/۸۹۴
۷۰۰	۰/۹۲۹
۸۰۰	۰/۹۵۶
۹۰۰	۰/۹۸۵
۱۰۰۰	۰/۹۷۹
۱۱۰۰	۰/۹۹۳
۱۲۰۰	۰/۹۹۴
۱۳۰۰	۰/۹۹۹
۱۴۰۰	۱/۰۰۰
۱۵۰۰	۱/۰۰۰

کم‌تری روابط معنادار را تشخیص دهد. افزایش ضریب جنسیت به $0/5$ - نیز سهم بیشتری به این متغیر در پیش‌بینی نتایج اختصاص می‌دهد. شکل ۲ منحنی توان آماری حاصل از شبیه‌سازی دوم را نشان می‌دهد که الگوی سیگموئید مشخصی را دنبال می‌کند. همانگونه که ملاحظه می‌شود، منحنی توان در این مرحله شیب تندتری را نشان می‌دهد که حاکی از کارایی بالاتر مدل در تشخیص اثرات است. حجم نمونه بهینه 500 واحدی در نقطه‌ای شناسایی شده است که توان آماری به $0/829$ می‌رسد و این نقطه با خط چین عمودی در نمودار مشخص شده است. این حجم نمونه نه تنها از نظر آماری معتبر است، بلکه از نظر عملی نیز کاملاً قابل دستیابی می‌باشد. الگوی سیگموئید منحنی به وضوح نشان می‌دهد که پس از نقطه حجم نمونه بهینه، افزایش حجم نمونه بازدهی کمتری در بهبود توان آماری دارد و تأیید می‌کند که مدل در این نقطه به سطح کارایی مطلوب خود رسیده است.

۴.۳ شاخص ترکیبی مقایسه عملکرد

برای مقایسه نظام‌مند چارچوب SBMF با روش‌های متعارف تعیین حجم نمونه، از یک شاخص ترکیبی عملکرد استفاده شده است. این شاخص با هدف فراهم آوردن امکان قضاوت یکپارچه درباره کارایی روش‌های مختلف در



شکل ۲. منحنی توان بر حسب حجم نمونه SBMF با برآوردهای قوی

شرایط فقدان چارچوب نمونه‌گیری احتمالی طراحی شده است. شاخص ترکیبی عملکرد (CPI)^۱ بر اساس چهار زیرمعیار اصلی تعریف می‌شود که هر یک جنبه‌ای از کارایی روش‌های نمونه‌گیری را منعکس می‌کنند. این روش تجمیع خطی با وزن‌های برابر، رویه‌ای استاندارد در ساخت شاخص‌های ترکیبی است (ناردو و همکاران، ۲۰۰۸).

زیرمعیارهای شاخص ترکیبی عملکرد:

۱. دقت پیش‌بینی: این معیار نشان‌دهنده توانایی روش در ارائه برآوردهایی با خطای استاندارد کوچک است. برای روش‌های مبتنی بر شبیه‌سازی (SBMF و نمونه‌گیری احتمالی)، این معیار بر اساس میانگین مربعات خطا (MSE) در ۱۰۰۰ تکرار شبیه‌سازی محاسبه و در دامنه ۰ تا ۱ نرمال‌سازی گردید.

۲. توان آماری: این معیار بیانگر احتمال تشخیص صحیح یک اثر معنادار در سطح معنی‌داری $\alpha = 0.05$ است. برای روش SBMF و روش توان آماری کلاسیک، این معیار مستقیماً از نتایج شبیه‌سازی (جدول ۴) استخراج شده است. برای نمونه‌گیری احتمالی، توان بر اساس فرمول‌های کلاسیک برای حجم نمونه ۱۰۶۷ و اندازه اثر مفروض محاسبه گردید. برای روش اشباع نظری و قاعده ۱۰ برابر، توان بر اساس مطالعات اعتبارسنجی گزارش شده در ادبیات پژوهش برآورد شده است.

۳. قابلیت مدیریت اریبی: این معیار نشان‌دهنده توانایی روش در شناسایی و کنترل اریبی‌های ناشی از مکانیزم نمونه‌گیری غیراحتمالی است. این معیار به صورت کیفی و در سه سطح طبقه‌بندی شده است: سطح ۱ (سیستماتیک) به معنای وجود سازوکار مدون برای شناسایی و تعدیل اریبی، سطح ۲ (مبتنی بر طراحی) به معنای کنترل اریبی از طریق طراحی نمونه‌گیری، سطح ۳ (کیفی) به معنای تشخیص کیفی اریبی، و سطح ۴ (فاقد سازوکار) به معنای عدم توجه به مسئله اریبی. برای کمی‌سازی این معیار، سطوح مذکور به ترتیب

¹Composite Performance Index

به امتیازهای ۱، ۷۵/۰، ۵۰/۰ و ۲۵/۰ تبدیل شده‌اند.

۴. **انعطاف‌پذیری:** این معیار بیانگر توانایی روش در انطباق با شرایط مختلف پژوهشی، انواع متغیرها و مدل‌های آماری گوناگون است. این معیار نیز به صورت کیفی و بر اساس دامنه کاربرد گزارش شده برای هر روش در ادبیات پژوهش، به امتیازهایی در دامنه ۰ تا ۱ تبدیل شده است. روش‌هایی که قابلیت کاربرد در طیف وسیعی از طرح‌های پژوهشی را دارند (مانند SBMF)، امتیاز بالاتری کسب می‌کنند.

روش محاسبه و وزن‌دهی: شاخص ترکیبی عملکرد توسط میانگین وزنی چهار زیرمعیار فوق به صورت

$$CPI = w_1 \times PA + w_2 \times SP + w_3 \times BMC + w_4 \times FL$$

محاسبه می‌شود، که در آن PA ، SP ، BMC و FL به ترتیب امتیازهای دقت پیش‌بینی، توان آماری، قابلیت مدیریت اریبی و انعطاف‌پذیری هستند و w_1 تا w_4 وزن‌های متناظر با آن‌ها می‌باشند. این روش تجمع خطی، رویه استاندارد ساخت شاخص‌های ترکیبی است (ناردو و همکاران، ۲۰۰۸)، که در آن

- PA : امتیاز دقت پیش‌بینی (دامنه ۰ تا ۱)
- SP : امتیاز توان آماری (دامنه ۰ تا ۱)
- BMC : امتیاز قابلیت مدیریت اریبی (دامنه ۰ تا ۱)
- FL : امتیاز انعطاف‌پذیری (دامنه ۰ تا ۱)
- w_1, w_2, w_3, w_4 : وزن‌های مربوط به هر زیرمعیار

برای تعیین وزن‌ها از رویکرد وزن‌دهی برابر $w_1 = w_2 = w_3 = w_4 = ۰/۲۵$ استفاده شده است. دلیل این انتخاب، عدم وجود اولویت نظری مشخص برای ترجیح یک زیرمعیار بر دیگری در ادبیات روش‌شناسی نمونه‌گیری است. این وزن‌دهی برابر، امکان مقایسه متوازن روش‌ها را فراهم می‌آورد و از اریبی احتمالی در جهت‌دهی به نتایج جلوگیری می‌کند. بدیهی است پژوهشگران می‌توانند متناسب با اهداف پژوهشی خود، وزن‌های متفاوتی را برای این زیرمعیارها در نظر بگیرند.

مقادیر شاخص ترکیبی عملکرد برای روش‌های مورد مقایسه: بر اساس روش‌های فوق، مقادیر شاخص ترکیبی عملکرد برای هر یک از روش‌های مورد بررسی در دو سناریوی شبیه‌سازی (ضرایب ضعیف و قوی) به صورت زیر محاسبه شده است:

- برای چارچوب SBMF در سناریوی اول (ضرایب ضعیف): امتیاز $CPI = ۰/۷$ که عمدتاً به دلیل عدم دستیابی به توان آماری مطلوب ($SP = ۰/۶$) و تأثیر آن بر میانگین وزنی است.
- برای چارچوب SBMF در سناریوی دوم (ضرایب قوی): امتیاز $CPI = ۰/۸۲۹$ که با دستیابی به توان آماری مطلوب ($SP = ۰/۸۲۹$) و حفظ امتیازهای بالا در سایر زیرمعیارها حاصل شده است.

- برای نمونه‌گیری احتمالی: امتیاز $CPI = ۰/۷۵$ که با وجود دقت پیش‌بینی مناسب و توان آماری خوب ($SP = ۰/۷۵$)، به دلیل نیاز الزامی به چارچوب نمونه‌گیری (کاهش انعطاف‌پذیری) در این سطح قرار گرفته است.
- برای روش اشباع نظری: امتیاز $CPI = ۰/۶۵$ که عمدتاً ناشی از عدم توانایی در ارائه برآوردهای کمی دقیق است.
- برای قاعده ۱۰ برابر: امتیاز $CPI = ۰/۵۵$ که به دلیل فقدان پشتوانه نظری قوی و عدم توجه به مسئله اریبی، پایین‌ترین امتیاز را کسب کرده است.
- برای توان آماری کلاسیک: امتیاز $CPI = ۰/۷۲$ که با وجود توان آماری مناسب ($SP = ۰/۷۲$)، به دلیل عدم توجه به مکانیزم نمونه‌گیری غیراحتمالی و فقدان سازوکار مدیریت اریبی، امتیاز کمتری نسبت به SBMF در سناریوی دوم دارد.

این مقادیر در جداول ۵ و ۶ و نمودارهای ۳ و ۴ ارائه شده‌اند. شفاف‌سازی روش محاسبه شاخص ترکیبی عملکرد، امکان ارزیابی انتقادی و بازتولید نتایج را برای سایر پژوهشگران فراهم می‌آورد.

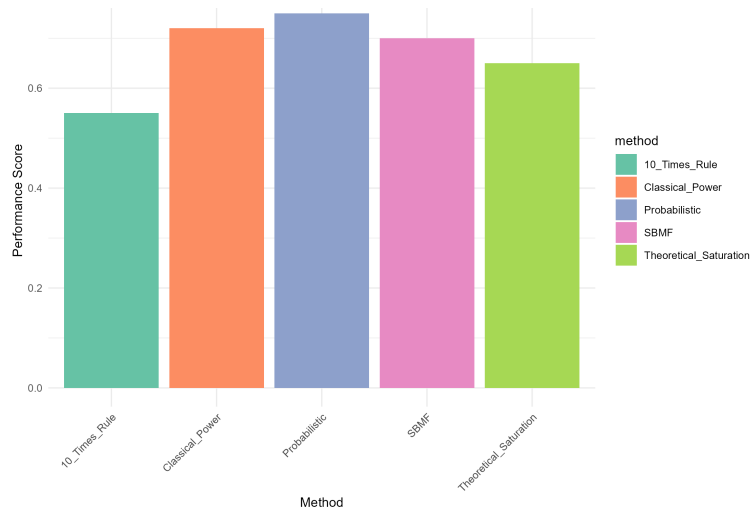
۴.۴ تحلیل مقایسه‌ای روش‌های نمونه‌گیری

امتیاز عملکرد در نمودار مقایسه‌ای بر اساس شاخص‌های متعددی محاسبه شده است. این شاخص‌ها شامل دقت پیش‌بینی، توان آماری، قابلیت مدیریت اریبی، و انعطاف‌پذیری روش در شرایط مختلف می‌باشد. همانطور که در جدول ۵ و شکل ۳ ملاحظه می‌شود، طی شبیه‌سازی سناریوی اول، روش SBMF در رتبه سوم پس از روش‌های احتمالی و کلاسیک قرار گرفته است، که علت اصلی ضعف آن در شناسایی نامطلوب اثرات با اندازه کوچک بود.

جدول ۵. مقایسه عملکرد روش‌های نمونه‌گیری در سناریوی اول

روش	حجم نمونه	امتیاز عملکرد CPI	روش مواجهه با اریبی	مبانی نظری
SBMF	دست‌نیافته	۰/۷	سیستماتیک	مدل‌مبنا
نمونه‌گیری احتمالی	۱۰۶۷	۰/۷۵	مبتنی بر طراحی	جمعیت‌مبنا
کفایت نظری	دست‌نیافته	۰/۶۵	کافی	نظریه زمینه‌ای
قاعده ۱۰ برابر	۴۰۰	۰/۵۵	ندارد	تجربی
توان آماری کلاسیک	۷۸۲	۰/۷۲	ندارد	جمعیتی

جدول ۶ و شکل ۴ عملکرد پنج روش نمونه‌گیری را در شبیه‌سازی دوم با ضرایب قوی‌تر، مقایسه می‌کند. چارچوب SBMF ضمن صعود به رتبه اول در تمامی سناریوها به عملکرد بهینه‌تری دست یافته است که نشان‌دهنده پایداری و برتری این چارچوب در شرایط مختلف پژوهشی است و این امر نشان‌دهنده اهمیت و توانایی این چارچوب در بهره‌برداری از اطلاعات موجود است. برای درک بهتر مزیت‌های رقابتی چارچوب پیشنهادی، جدول ۷ به مقایسه سیستماتیک ویژگی‌های کلیدی SBMF با روش‌های متعارف تعیین حجم نمونه می‌پردازد. این مقایسه بر اساس



شکل ۳. نمودار مقایسه امتیاز عملکرد روش‌های نمونه‌گیری با برآورد ضعیف چارچوب SBMF

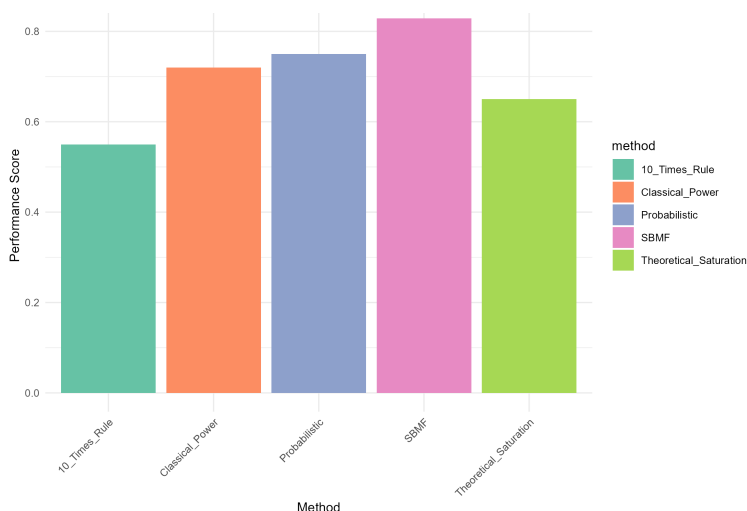
جدول ۶. مقایسه عملکرد روش‌های نمونه‌گیری در سناریوی دوم

روش	حجم نمونه	امتیاز عملکرد CPI	روش مواجهه با ارببی	مبانی نظری
SBMF	۵۰۰	۰/۸۲۹	سیستماتیک	مدل‌مبنا
نمونه‌گیری احتمالی	۱۰۶۷	۰/۷۵	مبتنی بر طراحی	جمعیت‌مبنا
کفایت نظری	دست‌نیافته	۰/۶۵	کیفی	نظریه زمینه‌ای
قاعده ۱۰ برابر	۴۰۰	۰/۵۵	ندارد	تجربی
توان آماری کلاسیک	۷۸۲	۰/۷۲	ندارد	جمعیت‌مبنا

معیارهای روش‌شناختی اساسی صورت گرفته است.

جدول ۷. مقایسه چارچوب SBMF با روش‌های متعارف تعیین حجم نمونه در شرایط فقدان چارچوب نمونه‌گیری احتمالی

روش	مبانی نظری	نیاز به چارچوب نمونه‌گیری	روش مواجهه با ارببی	حجم نمونه پیشنهادی
نمونه‌گیری احتمالی	استنتاج جمعیت‌مبنا	الزامی	از طریق طراحی	محاسباتی
کفایت نظری	نظریه زمینه‌ای	غیرضروری	کیفی	کیفی
قاعده ۱۰ برابر	تجربی	غیرضروری	ندارد	ثابت
توان آماری کلاسیک	استنتاج جمعیت‌مبنا	الزامی	ندارد	محاسباتی
SBMF	استنتاج مدل‌محور	غیرضروری	سیستماتیک	بهینه‌شده



شکل ۴. نمودار مقایسه امتیاز عملکرد روش‌های نمونه‌گیری با برآورد قوی چارچوب SBMF

۴.۵ جمع‌بندی مطالعه موردی

این مطالعه شبیه‌سازی، کارایی چارچوب SBMF را در تعیین حجم نمونه بهینه برای پژوهش‌های مبتنی بر نمونه‌گیری غیراحتمالی به صورت قانع‌کننده‌ای نشان می‌دهد. پیاده‌سازی موفق هفت گام چارچوب در دو سناریوی متفاوت، توانایی آن را در مدیریت عدم قطعیت ناشی از مکانیزم نمونه‌گیری غیراحتمالی تأیید می‌کند. نتایج به وضوح نشان داد که دقت در برآورد پارامترهای مدل و اندازه اثر مفروض، نقشی تعیین‌کننده در محاسبه حجم نمونه بهینه ایفا می‌کند. در سناریوی اول با ضرایب ضعیف، حتی با حجم نمونه ۱۵۰۰ واحد نیز توان مطلوب ۰/۸ حاصل نشد، اما در سناریوی دوم با ضرایب قوی‌تر، حجم نمونه بهینه ۵۰۰ واحد با توان ۰/۸۲۹ شناسایی گردید. این یافته، برتری رویکرد استنتاج مدل‌مبنای SBMF را نسبت به پارادایم سنتی جمعیت‌مبنا به اثبات می‌رساند: در شرایطی که روش‌های کلاسیک حجم نمونه‌های غیرعملی پیشنهاد می‌دهند، SBMF با تمرکز بر بهینه‌سازی عملکرد مدل، راهکاری عملی و قابل دفاع ارائه می‌دهد.

تکرارپذیری کامل نتایج از طریق ثبت تمام پارامترهای شبیه‌سازی، زمان اجرا و تنظیم مقدار اولیه تصادفی تضمین شده است. این شفافیت روش‌شناختی، امکان بازتولید دقیق نتایج را برای جامعه پژوهشی فراهم می‌آورد. در نهایت، این مطالعه بر اهمیت برآورد واقع‌بینانه از پارامترهای مدل در مرحله طراحی پژوهش تأکید می‌کند. کیفیت خروجی چارچوب SBMF به دقت این برآوردهای اولیه وابسته است و پژوهشگران باید با استفاده از منابع معتبر (مطالعات پیشین، فراتحلیل‌ها، داده‌های ثانویه و دانش خبرگان) این برآوردها را انجام دهند. جدول ۸ مثال‌های کاربردی از این چارچوب در حوزه‌های پژوهشی مختلف ارائه می‌دهد.

اجرای موفق چارچوب SBMF در هر یک از مثال‌های فوق، مستلزم چهار پیش‌نیاز روش‌شناختی است: (۱)

جدول ۸. مثال‌های دیگری از کاربرد چارچوب SBMF در حوزه‌های پژوهشی مختلف

حوزه	متغیر وابسته	متغیرهای مستقل (نمونه)	مدل آماری	مکانیزم نمونه‌گیری
بازاریابی و مصرف‌کننده	تمایل به خرید محصول جدید	سن، جنسیت، منطقه جغرافیایی، درآمد	رگرسیون لوژیستیک	نمونه‌گیری در دسترس از کاربران فعال یک پلتفرم
سلامت و پزشکی	ابتلا به دیابت نوع ۲ (بله/خیر)	سن، شاخص توده بدنی (BMI)، سابقه خانوادگی، فعالیت فیزیکی	رگرسیون لوژیستیک	نمونه‌گیری از مراجعین کلینیک‌های تخصصی
کشاورزی	بازده محصول گندم (تن در هکتار)	نوع کود، میزان آبیاری، دمای متوسط، نوع خاک	رگرسیون خطی چندگانه	نمونه‌گیری هدفمند از مزارع داوطلب
آموزش	نمره پایانی دانش‌آموزان	ساعات مطالعه، وضعیت اجتماعی-اقتصادی، حضور در کلاس	رگرسیون خطی یا چندسطحی	نمونه‌گیری خوشه‌ای از مدارس خاص

دسترسی به داده‌های کمکی معتبر برای ساخت جامعه مرجع شبیه‌سازی‌شده، (۲) برآورد اولیه از پارامترهای مدل (اندازه اثرهای مورد انتظار) از طریق مطالعات مقدماتی، فراتحلیل یا مشورت با خبرگان، (۳) تعیین عینی آستانه اثر مینیمم مهم از نظر علمی یا عملی، و (۴) تعریف عملیاتی شفاف از معیار عملکرد مدل (توان آماری، دقت پیش‌بینی یا شاخص‌های برازش). رعایت این چهار شرط، امکان بهره‌گیری مؤثر از چارچوب SBMF را در حوزه‌های کاربردی مختلف فراهم می‌سازد.

۵ بحث و نتیجه‌گیری

نتایج شبیه‌سازی، کارایی چارچوب SBMF را در تعیین حجم نمونه بهینه تأیید می‌کند. موفقیت این چارچوب در دستیابی به توان آماری ۸۲۹/۰ با حجم نمونه ۵۰۰ واحد، نشان‌دهنده مؤثر بودن گذار از استنتاج جمعیت‌مبنا به استنتاج مدل‌مبنا در شرایط فقدان نمونه‌گیری احتمالی است. این یافته با دیدگاه شموئلی (۲۰۱۰) در مورد تفاوت میان اهداف تبیین و پیش‌بینی همخوانی دارد. همچنین مرور نظام‌مند رحمان (۲۰۲۳) نیز بر ضرورت ارائه راه‌حل‌های ساختاریافته برای این حوزه تأکید کرده است. توانایی چارچوب در مدیریت ارببی سیستماتیک نمونه‌گیری در دسترس از طریق شبیه‌سازی، نقطه قوت بارز آن در مقایسه با روش‌های کلاسیک است که فاقد چنین مکانیزمی هستند. این ویژگی با پژوهش‌های اخیر در زمینه روش‌های دوگانه‌قوی (چن، ۲۰۲۰) و چالش‌های روش‌شناختی مطرح‌شده توسط پفرمن (پفرمن، ۲۰۱۵) همسو است. مزیت اصلی SBMF نسبت به روش‌های رایج، رویکرد یکپارچه و منعطف آن است. این چارچوب بر خلاف قواعد سرانگشتی (اسکرایبر و همکاران، ۲۰۰۶)، محاسبات را بر اساس اهداف پژوهش سفارشی می‌سازد. در مقایسه با کفایت نظری (چارمز، ۲۰۱۴)، راهکاری کمّی برای پژوهش‌های مدل‌محور ارائه می‌دهد و نسبت به توان آماری کلاسیک (شامپلی، ۲۰۲۰)، مکانیزم نمونه‌گیری غیراحتمالی را در شبیه‌سازی وارد می‌کند.

کدهای R برای اجرای چارچوب SBMF و مدل‌سازی نمونه‌گیری غیراحتمالی و محاسبه توان از طریق گیت‌هاب

با آدرس <https://github.com/hassangolinari1360/nonprobability-sample-size> قابل دسترس است.

محدودیت‌ها و پیشنهادها

این چارچوب در شرایط فقدان داده‌های کمکی معتبر یا دانش پیشین از پارامترها با چالش مواجه است و برای تخمین شاخص‌های آمارهای رسمی که مستلزم نمونه‌گیری احتمالی هستند، طراحی نشده است. پژوهش‌های آتی می‌توانند به طراحی بسته نرم‌افزاری کاربرپسند در محیط R، اعتبارسنجی میدانی در حوزه‌های سلامت و علوم اجتماعی، توسعه مبانی نظری تعمیم‌پذیری عملکرد مدل با وجود ارببی (چن و همکاران، ۲۰۲۵) و توسعه نسخه‌هایی برای نمونه‌گیری از جمعیت‌های کم‌دسترس (هرسکوویک، ۲۰۲۵) و گلوله‌برفی (گوجریال و همکاران، ۲۰۲۵) بپردازند. در نهایت، پژوهش‌های آینده می‌توانند نسخه‌های دیگری از این چارچوب برای طرح‌های نمونه‌گیری غیراحتمالی پیچیده‌تر، مانند نمونه‌گیری‌های آمیخته‌مد توسعه دهند. این مسیرهای پژوهشی می‌توانند نقش این چارچوب را به عنوان یک استاندارد روش‌شناختی در پژوهش‌های غیراحتمالی تثبیت کنند.

تقدیر و تشکر

این مقاله حاصل پژوهشی در حوزه روش‌شناسی آمار کاربردی است. نویسندگان بر خود لازم می‌دانند از داوران محترم مجله علوم آماری که با دقت نظر و ارائه پیشنهادهای سازنده، نقش بسزایی در ارتقای کیفیت این اثر ایفا کردند، و همچنین از اعضای هیئت تحریریه و ویراستاران مجله برای مدیریت شایسته فرایند داوری و انتشار، صمیمانه سپاسگزاری کنند. افزون بر این، نویسندگان مراتب قدردانی ویژه خود را تقدیم جناب آقای دکتر فرهاد مهران (مؤسسه بین‌المللی نیروی کار) می‌دارند که ایده اولیه این پژوهش برآمده از نگاه ایشان است و راهنمایی‌ها و همراهی ارزشمندشان در تمام مراحل توسعه و نگارش این مقاله پابرجا بود.

مراجع

AAPOR (2013), *Report of the AAPOR Task Force on Non-Probability Sampling*, Prepared by Couper, M. P., Dever, J. A., and Gile, K. J.

Canty, A. and Ripley, B. D. (2023), *boot: Bootstrap Functions*, R package version 1.3-28, <https://CRAN.R-project.org/package=boot>.

Castro-Martín, L., Rueda, M. and Ferri-García, R. (2025), Evaluation of Available

Techniques and Their Combinations to Address Selection Bias in Nonprobability Surveys, *AStA Advances in Statistical Analysis*, **109**(2), 1–35.

Champely, S. (2020), *pwr: Basic Functions for Power Analysis*, R package version 1.3-0, <https://CRAN.R-project.org/package=pwr>.

Charmaz, K. (2014), *Constructing Grounded Theory* (2nd ed.), Sage Publications, London.

Chen, Y. (2020), Doubly Robust Inference for Non-Probability Samples, *Journal of Survey Statistics and Methodology*, **8**(3), 485–510.

Chen, Y., Li, X. and Wang, Q. (2025), Doubly Robust Estimation for Non-Probability Samples with Heterogeneity, *Journal of Computational and Applied Mathematics*, **457**, 116283.

Cochran, W. G. (1977), *Sampling Techniques* (3rd ed.), John Wiley & Sons, New York.

Creswell, J. W. and Plano Clark, V. L. (2017), *Designing and Conducting Mixed Methods Research* (3rd ed.), Sage Publications, Thousand Oaks.

Elliott, M. R. and Valliant, R. (2017), Inference for Nonprobability Samples, *Statistical Science*, **32**(2), 249–264.

Firestone, W. A. (1993), Alternative Arguments for Generalizing From Data as Applied to Qualitative Research, *Educational Researcher*, **22**(4), 16–23.

Fox, J. and Weisberg, S. (2019), *An R Companion to Applied Regression* (3rd ed.), Sage Publications, Thousand Oaks.

Grossman, J., Waring, S. and Jackson, L. (2023), *simstudy: Simulation of Study Data*, R package version 0.8.0, <https://cran.r-project.org/package=simstudy>.

- Guest, G., Bunce, A. and Johnson, L. (2006), How Many Interviews Are Enough? An Experiment with Data Saturation and Variability, *Field Methods*, **18**(1), 59–82.
- Guest, G., Namey, E. and Chen, M. (2020), A Simple Method to Assess and Report Thematic Saturation in Qualitative Research, *PLOS ONE*, **15**(5), e0232076.
- Gujral, S., Singh, P. and Sharma, N. (2025), The Double-Edged Sword of Consecutive and Snowball Sampling: Practical Utility Versus Methodological Compromise, *Indian Journal of Psychological Medicine*, **48**(1), 81-84. doi: 10.1177/02537176251405469.
- Hennink, M. M., Kaiser, B. N. and Marconi, V. C. (2017), Code Saturation Versus Meaning Saturation: How Many Interviews Are Enough?, *Qualitative Health Research*, **27**(4), 591–608.
- Herskovic, L. (2025), *Sampling Techniques for Hard-to-Reach Populations: Literature Review*, Food and Agriculture Organization, Rome.
- Hosmer, D. W., Lemeshow, S. and Sturdivant, R. X. (2013), *Applied Logistic Regression* (3rd ed.), John Wiley & Sons, New Jersey.
- Kuhn, M. (2023), *caret: Classification and Regression Training*, R Package Version 6.0-94, <https://cran.r-project.org/package=caret>.
- Nardo, M., Saisana, M., Saltelli, A., Tarantola, S., Hoffman, A. and Giovannini, E. (2008), *Handbook on Constructing Composite Indicators: Methodology and User Guide*, OECD Publishing, Paris.
- Pargent, F., Koch, T. K., Kleine, A. K., Lerner, E. and Gaube, S. (2024), A Tutorial on Tailored Simulation-Based Sample-Size Planning for Experimental Designs with Generalized Linear Mixed Models, *Advances in Methods and Practices in Psychological Science*, **7**(4), 1–18.

- Pfeffermann, D. (2015), Methodological Issues and Challenges in the Production of Official Statistics: 24th Annual Morris Hansen Lecture, *Journal of Survey Statistics and Methodology*, **3**(4), 425–467.
- Rahman, M. M. (2023), Sample Size Determination for Survey Research and Non-Probability Sampling Techniques: A Review and Set of Recommendations, *Journal of Entrepreneurship, Business and Economics*, **11**(1), 42–62.
- Särndal, C.-E., Swensson, B. and Wretman, J. (1992), *Model Assisted Survey Sampling*, Springer-Verlag, New York.
- Schonlau, M. and Couper, M. P. (2017), Options for Conducting Web Surveys, *Statistical Science*, **32**(2), 279–292.
- Schreiber, J. B., Nora, A., Stage, F. K., Barlow, E. A. and King, J. (2006), Reporting Structural Equation Modeling and Confirmatory Factor Analysis Results: A Review, *The Journal of Educational Research*, **99**(6), 323–338.
- Shmueli, G. (2010), To Explain or to Predict?, *Statistical Science*, **25**(3), 289–310.
- Valliant, R. (2024), Sample Design Using Models, *Survey Methodology*, **50**(2), 149–183.
- Vehovar, V., Toepoel, V. and Steinmetz, S. (2016), Non-Probability Sampling, *The SAGE Handbook of Survey Methodology*, C. Wolf, D. Joye, T. W. Smith, and Y. Fu (Eds.), Sage Publications, London, 329–345.
- Venables, W. N. and Ripley, B. D. (2002), *Modern Applied Statistics with S* (4th ed.), Springer, New York.
- Xie, Y. (2014), *knitr: A Comprehensive Tool for Reproducible Research in R*, Chapman and Hall/CRC, Boca Raton.
- Zhu, H. (2021), *kableExtra: Construct Complex Table with 'kable' and Pipe Syntax*, R package version 1.3.4, <https://cran.r-project.org/package=kableExtra>.