

آمارگیریهای نمونه ای چند چارچوبی

محمد صالحی^۲

مرجان نورینی^۱

چکیده

چارچوب آماری سنگ بنای یک فرآیند نمونه گیری است. اما گاه چارچوبی که به تنهایی بتواند کلیه واحدهای جامعه مورد مطالعه را پوشش دهد در دسترس نمی باشد. از سوی دیگر ممکن است امکان دستیابی به چارچوبی کامل وجود داشته اما هزینه آمارگیری از آن بالا باشد. در این مواقع در آمارگیریها می توان با استفاده همزمان از چندچارچوب (که لزوماً کامل نبوده و هزینه آمارگیری از آنها کمتر است) کارایی را افزایش و هزینه آمارگیری را کاهش داد. چنین آمارگیریهایی تحت عنوان آمارگیری های چند چارچوبی بکار می روند. مرکز آمار کانادا و ایالات متحده پیشرو در اجرای این طرح می باشند. اهمیت کاربردی این موضوع ما را بر آن داشت تا مقاله حاضر را که مروری بر آمارگیری های نمونه ای چند چارچوبی است، ارائه دهیم.

واژه های کلیدی: چارچوب چندگانه^۳، بررسی های پیچیده^۴، برآورد جک نایف واریانس^۵، برآورد خطی واریانس^۶، روشهای بازنمونه گیری^۷

۱. مقدمه

را کاهش داد. گاهی نیز ممکن است یک چارچوب فهرستی کامل در دسترس باشد اما عملاً با گذشت زمانی نسبتاً طولانی بدلیل بروز تغییرات فراوان در آن، منبعی برای بروز خطاهای غیر نمونه گیری شود. از آنجا که یک فهرست ناحیه ای، چارچوبی کم هزینه، کمتر در معرض تغییرات و سهل الوصول می باشد ترکیب آن با یک چارچوب از اعضای جامعه که احتمالاً ناقص باشد می تواند نتایج مفیدی را حاصل نماید.

شاید بتوان گفت اولین شالوده آمارگیریهای چند چارچوبی در سال (۱۹۴۹) با آمارگیری از فروشگاه های خرده فروشی که مجری آن دفتر سرشماری آمریکا بود گذاشته شد. سپس هارتلی^۱ [۱] در سال (۱۹۶۲) نظریه مقدماتی چارچوبهای چندگانه را توسعه داد. او با فرض اینکه اجتماع چارچوبها، جامعه را پوشش می دهد، واحدهای جامعه را به زیر مجموعه های دو به دو ناسازگار شامل اجتماعها و اشتراکهای چارچوبهای مختلف تقسیم کرد. بعد از وی لاند^۲ [۲] در سال (۱۹۶۸)،

چارچوب آماری یا فهرست واحدهای آمارگیری، اساس و مبنای یک طرح آمارگیری نمونه ای را تشکیل می دهد. گاه ممکن است چارچوبی که کلیه ی واحدهای جامعه ی مورد مطالعه را پوشش دهد در دسترس نباشد، اما امکان دستیابی به پوشش کامل، با تلفیقی از دو یا چند چارچوب فراهم شود. در چنین حالتی به منظور دسترسی به پوشش مناسب از دو چارچوب یا بیشتر بطور همزمان استفاده می شود. گاهی نیز ممکن است یک چارچوب، پوشش کامل را برای جامعه ی مورد مطالعه فراهم کند، اما چارچوب ناقص دیگری موجود باشد که هزینه آمارگیری از چارچوب کامل باشد. در این شرایط به دلیل پایین تر بودن هزینه آمارگیری از این چارچوب، می توان با هزینه ای مشخص و ثابت، از دو چارچوب استفاده کرده و حجم نمونه را بزرگتر و واریانس برآورد

۱- مرکز آمار ایران

۲- دانشکده علوم ریاضی، دانشگاه اصفهان

۳- Multiple Frame

۴- Complex Survey

۵- Jackknife Variance Estimation

۶- Linearization Variance Estimation

۷- Resampling Methods

۸- Hartley

۹- Lund

تحقیقاتی را در زمینه کاربرد و استفاده از طرح چارچوب چندگانه در آمارگیری های مختلف آغاز کرده است.

۲. برآوردگرهای مجموع جامعه

برای سادگی فرض کنید دو چارچوب A و B موجود است که هر دو ناقص و دارای واحدهای مشترک با یکدیگر هستند بطوریکه مجموع آنها رویهم، کل جامعه مورد مطالعه را پوشش می دهد. از چارچوبهای A و B ، سه $(1-2)$ حوزه^۹ دو به دو ناسازگار به دست می آید.

حوزه a : شامل واحدهایی است که فقط در چارچوب A می باشند. $a = A \cap B^c$

حوزه b : شامل واحدهایی است که فقط در چارچوب B می باشند. $b = A^c \cap B$

حوزه ab : شامل واحدهایی است که در هر دو چارچوب $ab = A \cap B$ می باشند. (C) نشان دهنده مکمل مجموعه می باشد

N_A و N_B تعداد واحدها در چارچوبهای A و B ، N_a و N_b و N_{ab} تعداد واحدهای موجود در حوزه های a و b و ab می باشد.

دو نمونه مستقل s_A و s_B بر اساس طرحهای نمونه گیری احتمالی $PB(s_A)$ و $PB(s_B)$ به اندازه n_A و n_B از دو چارچوب فوق گرفته می شود. بطوریکه احتمال شمول نمونه حاصل از چارچوب A عبارتست از: $\pi_i^A = P\{i \in s_A\}$ و احتمال شمول نمونه حاصل از چارچوب B برابر است با: $\pi_i^B = P\{i \in s_B\}$.

در مجموع $\tilde{n}_A$ واحد اولیه نمونه گیری^{۱۰} از چارچوب A و $\tilde{n}_B$ واحد اولیه نمونه گیری از چارچوب B انتخاب می شود بطوریکه $\tilde{n} = \tilde{n}_A + \tilde{n}_B$ واحد اولیه نمونه گیری داریم. بعضی از این واحدها در نمونه ممکن است مشابه باشند یعنی به هر دو نمونه $s_A \cap s_B$ تعلق داشته باشند.

بر اساس نمونه های مستقل فوق، $\tilde{n}_a^A$ و $\tilde{n}_{ab}^A$ ، تعداد واحد های اولیه نمونه گیری حاصل از چارچوب A می باشند که به ترتیب در زیر مجموعه های a و ab قرار دارند. به همین ترتیب $\tilde{n}_b^B$ و $\tilde{n}_{ab}^B$ نیز تعداد واحدهای اولیه نمونه گیری حاصل از چارچوب B می باشند که به ترتیب در زیر مجموعه های a و ab هستند. با فرض اینکه Y_a ، Y_b ، Y_{ab} ، μ_a ، μ_b و μ_{ab} به ترتیب مجموع و میانگین جامعه در حوزه های a و b و ab باشند، داریم:

$$Y = Y_a + Y_{ab} + Y_b \quad (1)$$

فولر و بورمیستر^۱ [۳] در سال (۱۹۷۲)، و گل^۲ [۴] در سال (۱۹۷۵)، فورد و بوسکر^۳ [۵] در سال (۱۹۷۶)، بانکیئر^۴ [۶] در سال (۱۹۸۶)، اسکینر^۵ [۷] در سال (۱۹۹۱)، اسکینر و رائو^۸ [۸] در سال (۱۹۹۶) و کالتون و اندرسون^۹ [۹] در سال (۱۹۸۶) مقالاتی را در این زمینه ارائه نموده اند.

اخیراً استفاده از نمونه گیریهای دو چارچوبی، که یکی از چارچوبها را فهرست شماره تلفن افراد و دیگری را یک چارچوب ناحیه ای تشکیل می دهد، در برخی کشورها مرسوم شده است. در ایالات متحده، بررسی های اولیه جهت استفاده از دو چارچوب در طرحهای نمونه گیری کشاورزی، از اواسط دهه ۱۹۵۰ و جایگزینی نمونه گیری مبتنی بر این روش در بسیاری از نمونه گیریهای کشاورزی صورت گرفت. بدین ترتیب می توان مشاهده کرد که استفاده توأم از دو چارچوب در طرحهای آماری کشاورزی از سالها پیش در ایالات متحده باب شده است. لیکن به دلیل نیاز به تکنولوژی پیشرفته برای ایجاد یک چارچوب ناحیه ای، تاکنون تعداد اندکی از کشورها از این روش بهره جسته اند. سازمان خواربار و کشاورزی وابسته به سازمان ملل متحد^{۱۱} FAO به دلیل مزایایی که نمونه گیریهای دو چارچوبی به همراه دارد، برای سرشماری های آتی کشاورزی، استفاده از این روش را به سایر کشورها توصیه کرده است. البته استفاده از طرح دو چارچوبی در طرحهای آماری کشاورزی مستلزم وجود چارچوب ناحیه ای و فهرستی است. برای این منظور، تهیه یک چارچوب ناحیه ای کشاورزی مستلزم صرف هزینه بسیار و نیروی کارشناسی بسیار متبحر و تکنولوژی پیشرفته در زمینه سیستم های ماهواره ای و کامپیوتری است.

طرحهای دو چارچوبی در حالتی که هر دو چارچوب ناقصند نیز بکار می روند. آمارگیری وسیع ملی کودکان و جوانان زیر نظر اداره آمار کانادا برای تهیه اطلاعاتی در رابطه با مطالعه صفات و کارآزمودگی کودکان کانادایی از سه چارچوب ناقص مجزای متداخل استفاده می کند.

اخیراً در ایالات متحده در آمارگیری های مربوط به افراد بی خانمان از طرحهای چند چارچوبی استفاده می شود. مرکز آمار ایران نیز

۱- Fuller and Burmeister

۲- Vogel

۳- Ford and Boseker

۴- Bankier

۵- Skinner

۶- Rao

۷- Kalton and Anderson

۸- Food and Agriculture Organization

۹- domain

۱۰- Primery sampling unit (psu)

هدف برآورد Y است.

$$\hat{Y}_{ab}(\theta) = \theta \hat{Y}_{ab}^A + (1-\theta) \hat{Y}_{ab}^B$$

$$\hat{N}_{ab}(\theta) = \theta \hat{N}_{ab}^A + (1-\theta) \hat{N}_{ab}^B$$

$$\hat{Y}(\theta) = \hat{Y}_a^A + \hat{Y}_b^B + \hat{Y}_{ab}(\theta)$$

که مقدار بهینه θ واریانس $\hat{Y}_{ab}(\theta)$ را می نیمم می کند.

فرض می کنیم که D ، ماتریس واریانس-کواریانس بردار $\Delta = [\Delta Y \quad \Delta N]^T$ نشان دهنده ترانهاد بردار می باشد و نیز تعریف می کنیم $c = \text{cou}\{\Delta, \hat{Y}(\theta)\}$

۱.۲ برآوردگر هارتلی و فولر - بورمیستر

هارتلی [۱] در سال (۱۹۶۲) و فولر و بورمیستر [۳] در سال (۱۹۷۲) برآوردگرهای دو چارچوبی زیر را برای برآورد مجموع پیشنهاد کردند:

$$\begin{aligned}\hat{Y}_H &= \hat{Y}(\theta) + \theta_H \Delta Y \\ \hat{Y}_{FB} &= \hat{Y}(\theta) + \beta^T \Delta\end{aligned}\quad (4)$$

که در صورت ناویژه بودن ماتریس D مقادیر بهینه $\hat{Y}_{FB}$ و $\hat{Y}_H$ به ترتیب واریانس $\beta = -D^{-1}c$ و $\theta_H = -c_1 / D_{11}$ را می نیمم می کند. اما در عمل چون واریانس-کواریانس ها در D و c مجهول بوده و باید از روی نمونه برآورد شوند، به همین خاطر هر دوی $\hat{Y}_{FB}$ و $\hat{Y}_H$ به $\hat{\beta}$ و $\hat{\theta}_H$ و $\hat{D}$ بستگی داشته و برآوردگرهای $\hat{Y}_{FB}(\hat{\beta}_{FB})$ و $\hat{Y}_H(\hat{\theta}_H)$ بطور کلی توابع خطی از y نخواهند بود و وزنهای هر متغیر پاسخ باید بطور جداگانه محاسبه شوند. برای مثال اگر $\hat{Y}_{H1}$ و $\hat{Y}_{H2}$ و $\hat{Y}_{H3}$ برآوردگرهای هارتلی برای تعداد کل بیمارانی که دچار تنگی نفس در گروههای سنی ۱۶-۰ و ۴۵-۱۷ و ۴۵ به بالا باشد آن گاه $\hat{Y}_{H1} + \hat{Y}_{H2} + \hat{Y}_{H3}$ لزوماً مساوی برآورد هارتلی برای تعداد کل بیماران دچار تنگی نفس در جامعه نخواهد بود. در نتیجه فولر و بورمیستر بر آن شدند برآوردگری ارائه دهند که از یک مجموعه وزنهای یکسان برای همه متغیرهای پاسخ استفاده کند. آنها ابتدا حالت خاص انتخاب نمونه های تصادفی ساده از هر دو چارچوب را بررسی کردند. در این حالت برآوردگر زیر را پیشنهاد کردند:

$$\begin{aligned}\tilde{Y}_{FB,srs} &= (N_A - \hat{N}_{ab,srs}) \bar{y}_{a,srs} + \hat{N}_{ab,srs} \bar{y}_{ab,srs} \\ &+ (N_B - \hat{N}_{ab,srs}) \bar{y}_{b,srs}\end{aligned}\quad (5)$$

بطوریکه:

$$\begin{aligned}\bar{y}_{ab,srs} &= w \bar{y}_{ab}^A + (1-w) \bar{y}_{ab}^B \\ w &= \frac{\tilde{n}_{ab}^A (1-f_B)}{\tilde{n}_{ab}^A (1-f_B) + \tilde{n}_{ab}^B (1-f_A)}\end{aligned}$$

چندین برآوردگر نقطه ای تحت عنوان برآوردگر چارچوب دوگانه برای برآورد Y پیشنهاد شده که همگی آنها به فرم $\hat{Y} = \hat{Y}_a^A + \hat{Y}_{ab} + \hat{Y}_b^B$ می باشند. هر یک از این برآوردگرها بسته به اینکه اطلاعات حاصل از دو نمونه برای برآورد Y چگونه با هم ترکیب می شوند با هم تفاوتی دارند.

مسئله برآورد بطور خاص به اطلاع و آگاهی ما از N_A ، N_B ، N_{ab} و N_b بستگی دارد که بطور کلی سه حالت عمده وجود دارد:

$$(1) N_{ab}, N_B, N_A \text{ معلوم}$$

$$(2) N_B \text{ و } N_A \text{ معلوم اما } N_{ab} \text{ نامعلوم}$$

$$(3) N_{ab}, N_B, N_A \text{ نامعلوم}$$

برای سادگی در این مقاله روی حالت ۲ متمرکز شده و برآوردگرها را در این حالت ارائه می دهیم.

فرض می کنیم w_i^A و w_i^B وزنهای هایک [۱۰] برای طرحهای نمونه گیری مورد استفاده در چارچوبهای A و B باشند:

$$\begin{aligned}w_i^A &= N_A [\pi_i^A \sum_{j \in S_A} (\frac{1}{\pi_j^A})]^{-1} \\ w_i^B &= N_B [\pi_i^B \sum_{j \in S_B} (\frac{1}{\pi_j^B})]^{-1}\end{aligned}\quad (2)$$

دو متغیر نشانگر زیر را برای هر یک از چارچوبهای A و B در نظر می گیریم:

$$\delta_i^A = \begin{cases} 1 & \text{واحد } I \text{ متعلق به چارچوب } A \text{ باشد} \\ 0 & \text{در غیر این صورت} \end{cases}$$

$$\delta_i^B = \begin{cases} 1 & \text{واحد } I \text{ متعلق به چارچوب } B \text{ باشد} \\ 0 & \text{در غیر این صورت} \end{cases}$$

در نتیجه برآوردگر مجموع در سه حوزه a و b و ab به قرار زیر می باشند.

$$\begin{aligned}\hat{Y}_a^A &= \sum_{i \in S_A} w_i^A (1-\delta_i^B) y_i, \quad \hat{Y}_{ab}^A = \sum_{i \in S_A} w_i^A \delta_i^B y_i \\ \hat{Y}_b^B &= \sum_{i \in S_B} w_i^B (1-\delta_i^A) y_i, \quad \hat{Y}_{ab}^B = \sum_{i \in S_B} w_i^B \delta_i^A y_i\end{aligned}$$

برآوردگر اندازه ی هر یک از حوزه ها نیز به صورتی مشابه با قرار دادن $y_i = 1$ در تعاریف $\hat{Y}_a^A$ و $\hat{Y}_{ab}^A$ و $\hat{Y}_b^B$ و $\hat{Y}_{ab}^B$ بدست می آیند.

برای سادگی مقایسه برآوردگرهای آتی، تعاریف زیر را در نظر می گیریم:

$$\begin{aligned}\Delta Y &= \hat{Y}_{ab}^A - \hat{Y}_{ab}^B \\ \Delta N &= \hat{N}_{ab}^A - \hat{N}_{ab}^B\end{aligned}\quad (3)$$

۲.۲ برآوردگرهای چارچوب منفرد

بانکیر [۶] در سال (۱۹۸۶)، کالتون و اندرسون [۹] در سال (۱۹۸۶) و اسکینر [۷] در سال (۱۹۹۱) برای برآورد مجموع جامعه با در نظر گرفتن مشاهدات بگونه ای که آنها از یک چارچوب منفرد با وزنهای اصلاح شده برای مشاهدات در حوزه متداخل ab نمونه گیری شده اند، برآوردی ارائه دادند. وزنهای اصلاح شده برای برآوردگرهای تک چارچوبی اسکینر و کالتون و اندرسون به تشخیص واحدهای مشابه در نمونه ها نیازی ندارد. این وزنها عبارتند از:

$$w_i = \begin{cases} (\pi_i^A + \pi_i^B)^{-1} & i \in ab \\ (\pi_i^A)^{-1} & i \in a \\ (\pi_i^B)^{-1} & i \in b \end{cases}$$

آن گاه برآوردگر زیر را پیشنهاد کردند:

$$Y_{SF} = \sum_{i \in s_A} w_i y_i + \sum_{i \in s_B} w_i y_i \quad (11)$$

همانطور که ملاحظه می شود این برآوردگر از هیچ گونه اطلاعات کمکی در ارتباط با N_A و N_B استفاده نمی کند. از آنجایی که استفاده از اطلاعات کمکی باعث کاهش واریانس و افزایش دقت برآوردگر می شود (طبق نتایجی که بانکیر [۶] در سال (۱۹۸۶) به دست آورد) اسکینر و راثو [۸] در سال (۱۹۹۶) دو روش برای تصحیح این برآوردگر ارائه دادند که عبارتند از:

(۱) روش تعدیل نسبتی^۱

(۲) روش رگرسیونی

اسکینر و راثو پیرو اثبات قضیه (۱) اسکینر [۷] در سال (۱۹۹۱) نشان دادند که فرآیند تعدیل به برآوردگر تجربی زیر همگراست:

$$\hat{Y}_{SFrake} = \frac{N_A - \hat{N}_{ab}^{rake}}{\hat{N}_a} \hat{Y}_a + \frac{N_B - \hat{N}_{ab}^{rake}}{\hat{N}_b} \hat{Y}_b + \frac{\hat{N}_{ab}^{rake}}{\hat{N}_{abS}} \hat{Y}_{abS} \quad (12)$$

که در آن:

$$\hat{Y}_{abS} = \sum_{s_A} w_i \delta_i^B y_i + \sum_{s_B} w_i \delta_i^A y_i$$

$$\hat{N}_{abS} = \sum_{s_A} w_i \delta_i^B + \sum_{s_B} w_i \delta_i^A$$

و $\hat{N}_{ab}^{rake}$ کوچکترین ریشه معادله درجه دو بفرم زیر می باشد:

$$\hat{N}_{abS} x^2 - [\hat{N}_{abS} (N_A + N_B) + \hat{N}_{aS}^A \hat{N}_{bS}^B] x + \hat{N}_{abS} N_A N_B = 0 \quad (13)$$

و $\hat{N}_{ab,srs}$ کوچکترین ریشه معادله درجه دو بفرم زیر است:

$$\begin{aligned} & (\tilde{n}_A + \tilde{n}_B) x^2 - (\tilde{n}_A N_B N_A + \tilde{n}_{ab}^A N_A + \tilde{n}_{ab}^B N_A + \tilde{n}_{ab}^B N_B) x \\ & + (\tilde{n}_{ab}^A + \tilde{n}_{ab}^B) N_A N_B = 0 \end{aligned} \quad (6)$$

فولر و بورمیستر [۳] در سال (۱۹۷۲) نشان دادند که $\hat{N}_{ab,srs}$ بعنوان یک برآوردگر حداکثر درستنمایی از N_{ab} بوده و نیز اسکینر [۷] در سال (۱۹۹۱) از $\hat{Y}_{FB,srs}$ تفسیری به صورت یک برآوردگر حداکثر درستنمایی داشت. بنابر این خواص مجانبی بهینه داشته و چون از وزنهای یکسانی برای هر متغیر پاسخ استفاده می کند به $\hat{Y}_{FB}$ ارجحیت دارد. اما این برآوردگر مستقیماً با طرحهای نمونه ای پیچیده بکار نمی رود. چون عموماً به صورتی ناسازگار با این نوع نمونه گیرها طراحی می شود. به همین دلیل آنها تصمیم گرفتند این برآوردگر را طوری اصلاح کنند که به برآوردگرهای سازگار با طرحهای پیچیده نائل آیند. آنها برآورد زیر را به دست آوردند:

$$\begin{aligned} \tilde{Y}_{PML}(\theta) = & \frac{N_A - \hat{N}_{ab}^{PML}(\theta)}{\hat{N}_a^A} \hat{Y}_a^A + \hat{N}_{ab}^{PML}(\theta) \frac{\hat{Y}_{ab}(\theta)}{\hat{N}_{ab}(\theta)} \\ & + \frac{N_A - \hat{N}_{ab}^{PML}(\theta)}{\hat{N}_a^B} \tilde{Y}_b^B \end{aligned} \quad (7)$$

بطوریکه $\hat{N}_{ab}^{PML}(\theta)$ تابعی از $\hat{N}_{ab}^A$ و $\hat{N}_{ab}^B$ و θ بوده و کوچکترین ریشه معادله درجه دو به فرم زیر است:

$$\begin{aligned} & \left[\frac{\theta}{N_B} + \frac{(1-\theta)}{N_A} \right] x^2 - \left[1 + \theta \frac{\hat{N}_{ab}^A}{N_B} + (1-\theta) \frac{\hat{N}_{ab}^B}{N_A} \right] x \\ & + \hat{N}_{ab}(\theta) = 0 \end{aligned} \quad (8)$$

اسکینر و راثو [۸] در سال (۱۹۹۶) انتخاب $\theta = \theta_P$ را برای می نیمم کردن واریانس مجانبی $\hat{N}_{ab}^{PML}(\theta)$ با

$$\theta_P = \frac{N_a N_B V(\hat{N}_{ab}^B)}{N_a N_B V(\hat{N}_{ab}^B) + N_b N_A V(\hat{N}_{ab}^A)} \quad (9)$$

پیشنهاد کردند.

همچنین آنها با خطی کردن ریشه های معادله درجه دو به صورت $\hat{N}_{ab}^{PML}(\theta_P) \approx \hat{N}_{ab}(\phi)$ راهی ساده تر برای محاسبه برآوردگر شبه حداکثر درستنمایی ارائه دادند. مقدار بهینه ϕ در این حالت عبارتست از:

$$\phi = \frac{V(\hat{N}_{ab}^B)}{V(\hat{N}_{ab}^A) + V(\hat{N}_{ab}^B)} \quad (10)$$

در عمل N_a و N_b و واریانس ها در روابط (۹) و (۱۰) مجهول بوده و باید از روی داده ها برآورد شوند در نتیجه برآوردگر حاصل عبارتست از:

$$\hat{Y}_{PML}(\hat{\theta}_P)$$

بیان کرد. سپس توکی^[۱۲] پیشنهاد کرد که این روش علاوه بر کاهش اریبی، می‌تواند برای برآورد واریانس نیز استفاده شود. به نظر می‌رسد که استفاده از جک نایف در برآورد واریانس جوامع متناهی اولین بار توسط دوربین^[۱۳] که استفاده از آن را در برآورد نسبتی مورد مطالعه قرار داد، بررسی شد. کتاب جامعی توسط میلر^[۱۴] در سال (۱۹۷۴) و در رساله‌ای از افرون^[۱۵] در سال (۱۹۸۲) داده شده است.

فرض کنید چارچوب‌های A و B به ترتیب دارای H و L طبقه باشند بطوریکه طبقات h و l هر یک از آنها به ترتیب شامل N_h^A و N_l^B واحد بوده و در مجموعه $\hat{N}_h^A$ و $\hat{N}_l^B$ واحد اولیه نمونه‌گیری (psu) دارند که به ترتیب $\tilde{n}_h^A$ و $\tilde{n}_l^B$ واحد از آنها نمونه‌گیری می‌شود. بطوریکه $\tilde{n}_A = \sum_{h=1}^H \tilde{n}_h^A$ و $\tilde{n}_B = \sum_{l=1}^L \tilde{n}_l^B$ کل نمونه‌ای است که به ترتیب از چارچوبهای A و B انتخاب می‌شود.

حال فرض می‌کنیم که نمونه‌هایی به اندازه (≥ 2) $\tilde{n}_l^B$ و $\tilde{n}_h^A$ به ترتیب با احتمالات شمول متناسب با اندازه، $\tilde{\pi}_{hi}^A = \tilde{n}_h^A p_{hi}^A$ و $\tilde{\pi}_{lj}^B = \tilde{n}_l^B p_{lj}^B$ از چارچوبهای A و B انتخاب می‌شود. بطوریکه p_{lj}^B و p_{hi}^A به ترتیب برابر با احتمال انتخاب واحد i ام و j ام در طبقات h و l متناسب با اندازه psu ها بوده و $\sum_i p_{hi}^A = 1$ و $\sum_j p_{lj}^B = 1$ ($p_{hi}^A = u_i / \sum_A u_i$ و $p_{lj}^B = u_j / \sum_B u_j$) u_i اندازه تقریبی I امین psu می‌باشد)

با در نظر گرفتن دو بردار A و B به صورت :

$$A = (Y_a, Y_{ab}, N_{ab}, N_A)^T \quad \text{و} \quad B = (Y_b, Y_{ab}, N_{ab}, N_B)^T$$

مجموع جامعه عبارت خواهد بود از:

$$Y = N\bar{Y} = N_g(\bar{A}, \bar{B})$$

$$= N_\tau$$

بطوریکه $\bar{A} = A/N_A$ و $\bar{B} = B/N_B$ آن گاه $\bar{Y}$ به صورت زیر قابل بیان است :

$$g(\bar{A}, \bar{B}) = \bar{Y} = (N_A/N)(\bar{A}_1 + \theta \bar{A}_\tau) + (N_B/N)(\bar{B}_1 + (1-\theta)\bar{B}_\tau)$$

که θ مقدار ثابت می‌باشد. میانگین‌های $\bar{A}$ و $\bar{B}$ به ترتیب با $\hat{\bar{A}}$ و $\hat{\bar{B}}$ برآورد شده و $\hat{\tau} = g(\hat{\bar{A}}, \hat{\bar{B}})$ نیز برآوردی از τ می‌باشد.

تصحیح با N_A و N_B معلوم از طریق روش رگرسیون برآورگر را نتیجه می‌دهد:

$$\hat{Y}_{SFreg} = \hat{Y}_S + \hat{\beta}^T [N_A - N_S^A, N_B - N_S^B] \quad (14)$$

بطوریکه مقدار بهینه $\hat{\beta}$ عبارتست از:

$$\hat{\beta}_S^T = Cov\{[\hat{N}_{AS}/V(\hat{N}_{AS}), \hat{N}_{BS}], \hat{Y}_S\}$$

تحت طرحهای نمونه‌گیری خودوزن، در حالتی که

$\theta_S = w_i^B / (w_i^A + w_i^B)$ برای همه واحدها مقدار ثابتی است آن گاه $\hat{Y}_{abS} = \hat{Y}_{ab}(\theta_S)$ و $\hat{N}_{abS} = \hat{N}_{ab}(\theta_S)$ است و

$N_A - \hat{N}_{AS} = (1 - \theta_S)\Delta N$, $N_B - \hat{N}_{BS} = -\theta_S\Delta N$ آن گاه برآوردرگر SF و $SFreg$ در این حالت به صورت زیر خواهند بود:

$$\hat{Y}_{SF} = \hat{Y}(\circ) + \theta_S \Delta Y = \hat{Y}(\theta_S)$$

$$\hat{Y}_{SFreg} = \hat{Y}_S + \beta_{S\tau} \Delta Y$$

که مقدار بهینه $\beta_{S\tau}$ عبارتست از:

$$\beta_{S\tau} = -(c_\tau + \theta_S D_{\tau\tau}) / D_{\tau\tau}$$

در این حالت نیز با خطی کردن ریشه معادله (۱۳) به صورت $\hat{N}_{ab}^{rake} \approx \hat{N}_{ab}(\psi)$ می‌توان برآوردرگر نسبتی تعدیل شده را به فرم ساده تری محاسبه کرد. مقدار بهینه ψ در این حالت عبارت از $\psi = N_b(\theta_S N_a + N_{ab}) / (N_A N_A - N_{ab}^2)$ است و در صورتی که N_a/N و N_b/N هر دو نسبتاً کوچک باشند (حدود ۰/۱ یا ۰/۲) آن گاه $\psi \approx N_b / (N_a + N_b)$

۳. برآورد واریانس

اسکینر و راثو [۸] در سال (۱۹۹۶) با استفاده از بسط سری تیلور برآوردرگری خطی برای برآورد واریانس برآوردرگرهای چارچوب دوگانه ارائه دادند در این بخش ابتدا برآورد واریانس به روش جک نایف برای برآوردرگرهای چارچوب دوگانه را ارائه داده و سپس نشان می‌دهیم که با برآوردرگر خطی واریانس بطور مجانبی هم‌ارز می‌باشد.

نخستین بار کوئولی [۱۱] جک نایف را به عنوان روشی بر اساس حذف هر بار یک مشاهده از مجموعه داده‌های اولیه و محاسبه مجدد برآوردرگر با استفاده از بقیه داده‌ها برای کاهش اریبی برآوردرگر ضریب همبستگی پیاپی^۱ معرفی نمود. در مقاله‌ای در سال (۱۹۵۶) این روش را تعمیم داده و خواص کلی کاهش اریبی آن را در یک جامعه متناهی

Tukey - ۲
Durbin - ۳
Miller - ۴
Efron - ۵

۱- Serial Correlation Coefficient

قضیه فوق بیان می دارد که برآوردگر جک نایف واریانس برای شکل بهینه هر برآوردگر سازگار است.

۴. نتایج شبیه سازی

در این بخش با استفاده از شبیه سازی به مطالعه خواص تجربی برآوردگرهای واریانس، می پردازیم.

فرض می شود که جامعه نامتناهی بوده و طرح نمونه گیری برای هر چارچوب یک طبقه دارد. یک نمونه خوشه ای دو مرحله ای با $\tilde{n}_A$ خوشه و m عنصر از هر خوشه به عنوان نمونه حاصل از چارچوب A ($n_A = \tilde{n}_A \cdot m$) و یک نمونه تصادفی ساده با n_B مشاهده به عنوان نمونه حاصل از چارچوب B تولید می شود. همانطور که گفته شد فرض بر این است که جامعه نامتناهی بوده و $\frac{N_a}{N}$ و $\frac{N_b}{N}$ را به ترتیب با γ_a و γ_b (معلومند) جایگزین می کنیم. نمونه حاصل از چارچوب A شامل مقادیر

$$\{(y_{ij}, m_{ai}), i=1, \dots, n_A, j=1, \dots, m\}$$

می باشد که مقادیر m_{ai} عبارت از تعداد عناصر نمونه گیری شده از I امین خوشه که متعلق به حوزه a است و y_{ij} نیز مقدار مرتبط با j امین عنصر در i امین خوشه می باشد. نمونه حاصل از چارچوب B شامل

$$\{(y_j, n_b), j=1, \dots, n_B\}$$

است که n_b عبارت از تعداد نمونه ای متعلق به حوزه b است و y_j نیز مقدار مرتبط با j امین عنصر است. برای به دست آوردن نمونه حاصل از چارچوب A ، ابتدا m_{ai} را از یک توزیع بتا - دو جمله ای بصورت زیر تولید می کنیم:

(۱) $\tilde{\gamma}_{ai}$ نسبت مورد انتظار مشاهداتی از خوشه ای i ام که در حوزه a قرار دارد، از توزیع بتا با پارامترهای a_1 و a_2 قرار داده می شود بطوریکه:

$$E(\tilde{\gamma}_{ai}) = \frac{\gamma_a}{1 - \gamma_b}$$

$$V(\tilde{\gamma}_{ai}) = \frac{\gamma_a^2 (1 - \gamma_a - \gamma_b)}{\{[a_1 (1 - \gamma_b) + \gamma_a]\} (1 - \gamma_b)^2}$$

بنابر این مقدار کوچکتر a_1 تغییر پذیری $\tilde{\gamma}_{ai}$ را از یک خوشه به خوشه دیگر افزایش می دهد.

(۲) به ازای مقداری از $\tilde{\gamma}_{ai}$ ، m_{ai} را از یک توزیع دو جمله ای با پارامترهای m و $\tilde{\gamma}_{ai}$ تولید می کنیم.

(۳) مراحل (۱) و (۲) را مستقلاً تکرار کرده تا $\{m_{ai}, i=1, \dots, \tilde{n}_A\}$ به دست آید.

برآوردگرهای چارچوب دوگانه در این حالت می توانند به صورت تابعی از $\hat{\tau}$ بیان شوند. به عنوان مثال $\tilde{Y}_H(\theta_H) = N g(\hat{A}, \hat{B})$ برآوردگر PML نیز به ازای θ ثابت، می توانند به صورتی مشابه به عنوان تابعی از $\hat{A}$ و $\hat{B}$ و با توجه به اینکه $\hat{N}_{ab}^{PML}(\theta)$ تابعی از $\hat{A}_\tau$ و $\hat{B}_\tau$ است، و سپس با استفاده از رابطه (۷) بیان شود.

با فرض اینکه $g(\hat{A}, \hat{B})$ تابعی مشتق پذیر و دارای مشتقات اول و دوم پیوسته و کراندار در یک همسایگی از $(\bar{A}, \bar{B})$ باشد، برآوردگر خطی واریانس $\hat{\tau}$ عبارتست از:

$$v_L(\hat{\tau}) = g_A^T(\hat{A}, \hat{B}) S^A g_A(\hat{A}, \hat{B}) + g_B^T(\hat{A}, \hat{B}) S^B g_B(\hat{A}, \hat{B}) \quad (15)$$

که $g_A(\bar{A}, \bar{B})$ و $g_B(\bar{A}, \bar{B})$ به ترتیب بردارهای ۴ بعدی از مشتق های اول $g(\circ)$ نسبت به مؤلفه های $\hat{A}$ و $\hat{B}$ بوده و S^A و S^B به ترتیب برابر با ماتریس واریانس - کواریانس نمونه ای بردارهای $\hat{A}$ و $\hat{B}$ می باشند. برای محاسبه برآورد جک نایف واریانس $\hat{\tau}$ فرض می کنیم $\hat{\tau}_{hi}^A$ برآوردگری به فرم $\hat{\tau}$ باشد که بعد از حذف مشاهدات i امین psu نمونه طبقه h چارچوب A به دست آمده و $\hat{\tau}_{(hi)}^A = g(\hat{A}, \hat{B}_{(lj)})$ که $\hat{B}_{(lj)}$ برآورد $\bar{B}$ محاسبه شده بعد از حذف j امین psu نمونه طبقه l چارچوب B است. برآورد جک نایف واریانس $\hat{\tau}$ با توجه به استقلال نمونه ها عبارتست از:

$$v_J(\hat{\tau}) = \sum_{h=1}^H \frac{\tilde{n}_h^A - 1}{\tilde{n}_h^A} \sum_{i=1}^{\tilde{n}_h^A} (\hat{\tau}_{(hi)}^A - \hat{\tau})^2 + \sum_{l=1}^L \frac{\tilde{n}_l^B - 1}{\tilde{n}_l^B} \sum_{j=1}^{\tilde{n}_l^B} (\hat{\tau}_{(lj)}^B - \hat{\tau})^2 \quad (16)$$

رائو و لهر [۱۶] طی قضیه ای ثابت کردند که $v_L(\hat{\tau})$ و $v_J(\hat{\tau})$ بطور مجانبی هم ارزند بعبارت دیگر:

$$v_L(\hat{\tau}) = V(\hat{\tau}) + o_p(\tilde{n}^{-1}) \quad (1)$$

$$v_J(\hat{\tau}) = v_L(\hat{\tau}) + o_p(\tilde{n}^{-1}) \quad (2)$$

جک نایف با توابع هموار قابل کاربرد است. در بین برآوردگرهای دو چارچوبی، برآوردگرهای SF و $SFrake$ بصورت توابع همواری از میانگین جوامع قابل بیان هستند. دیگر برآوردگرها نیز در حالتی که پارامترهای θ_p و θ_H و β_{FB} و $\beta_{S\tau}$ ثابت بوده و از روی داده ها برآورد نشوند، نیز توابع همواری از میانگین جوامع خواهند بود. بنابر این

جامعه Y/N و بجای برآورد مجموع $\hat{Y}$ با برآورد میانگین جامعه $\hat{Y}/N$ کار می کنیم. برای مثال:

$$\frac{Y}{N} = \gamma_a \mu_a + \gamma_b \mu_b + \gamma_{ab} \mu_{ab} \quad (\gamma_{ab} = 1 - \gamma_a - \gamma_b)$$

و برآوردگر میانگین جامعه در حوزه ها عبارتند از:

$$\frac{\hat{Y}_a}{N} = \frac{1 - \gamma_b}{n_A} \sum_{s_a} y_{ij} \quad \frac{\hat{Y}_{ab}^A}{N} = \frac{1 - \gamma_b}{n_A} \sum_{s_{ab}^A} y_{ij}$$

$$\frac{\hat{Y}_b}{N} = \frac{1 - \gamma_a}{n_B} \sum_{s_b} y_j \quad \frac{\hat{Y}_{ab}^B}{N} = \frac{1 - \gamma_a}{n_B} \sum_{s_{ab}^B} y_j$$

(لازم به ذکر است که برآوردگرهای فوق تحت طرح نمونه گیری خودوزن تعریف شده اند)

$EMSE$ خود منجر به خطای مونت کارلو شده و واریانس مونت کارلو آن بصورت زیر برآورد می شود:

$$s^2(EMSE) = \frac{1}{R-1} = \sum_{r=1}^R (\hat{Z}_r - \bar{\hat{Z}})^2$$

که $\hat{Z} = (\hat{Y}_r - Y)^2$ و $\bar{\hat{Z}} = \sum_r \hat{Z}_r / R$ می باشد با گرفتن جذر از رابطه فوق، خطای معیار مونت کارلو $s(EMSE)$ را می توان محاسبه کرد. نتایج شبیه سازی برای میانگین توان دوم خطای تجربی $EMSE$ و خطای استاندارد مونت کارلو نظیرش $s(EMSE)$ تحت طرحهای نمونه گیری فوق به ازای $\tilde{n}_A = 10$ ، $m = 30$ ، $n_B = 100$ و $\tilde{n}_A = 20$ آورده شده است. در مقایسه با برآوردگرهای دو چارچوبی، برآوردگر چارچوب منفرد افزایش قابل ملاحظه ای را در $EMSE$ وقتی $\gamma_a \leq \gamma_b$ یا n_A خیلی بزرگتر از n_B باشد، نشان می دهد. برای هر شش برآوردگر چارچوب دو گانه، دو برآوردگر واریانس خطی و جک نایف، را محاسبه می کنیم. برای محاسبه برآورد جک نایف واریانس، در صورت وابسته بودن پارامترهای $\hat{\theta}_H$ ، $\hat{\theta}_{FB}$ ، $\hat{\theta}_P$ و $\hat{\beta}_{S^2}$ به ماتریسهای واریانس - کواریانس نمونه، $S_{(hi)}^A$ و $S_{(hi)}^B$ را با استفاده از یک جک نایف جداگانه درون هر تکرار جک نایف محاسبه کرده و سپس برآورد جک نایف واریانس برآوردگرها را با استفاده از رابطه (۱۷) به دست می آوریم. نایچ شبیه سازی نشان می دهد که اربیبی نسبی برآوردگر خطی واریانس در صورت افزایش حجم نمونه میل به کاهش دارد. به هر حال، در نمونه های کوچکتر، برآوردگر خطی واریانس، واریانس را کم برآورد می کند در حالی که جک نایف آن را بیش برآورد می کند. همچنین نتایج حاصل از خطای معیار نسبی حاکی از آن است که برآوردگر جک نایف واریانس از پایداری کمتری نسبت به برآورد خطی برخوردار است. بعنوان مثال وقتی $\tilde{n}_A = 10$ خطای معیار نسبی

به ازای مقدار از m_{ai} مقادیر y_{ij} از مدل خطای لانه ای^۱ زیر تولید می شوند:

$$y_{ij} = \mu_a + a_{ai} + \varepsilon_{ij} \quad j = 1, \dots, m_{ai}$$

$$y_{ij} = \mu_{ab} + a_{abi} + \varepsilon_{ij} \quad j = m_{ai} + 1, \dots, m$$

از مقادیر $\mu_a = 9$ و $\mu_{ab} = 10$ و $\mu_b = 11$ برای میانگین جامعه در حوزه های a و ab و b استفاده می کنیم. ε_{ij} ها متغیرهای تصادفی مستقل و هم توزیع (iid) با $\varepsilon_{ij} \sim N[0, \sigma^2(1 - \rho)]$ به ازای $\sigma^2 = 1$ و $\rho = 0.1$ و نیز مستقل از $(\alpha_{ai}, \alpha_{abi}, m_{ai})$ هستند. $(\alpha_{ai}, \alpha_{abi})$ نیز از یک توزیع نرمال دو متغیره با میانگین 0 و ماتریس کواریانس 2×2 با عناصر غیر قطری $\rho\sigma^2$ و عناصر غیرقطری $\rho\delta\sigma^2$ به ازای $\delta = 0.5$ به دست می آیند. این مدل، یک همبستگی یکسان ρ را درون حوزه های a و ab و یک همبستگی یکسان $\rho\delta$ را بین حوزه ها تضمین می کند.

برای به دست آوردن نمونه حاصل از چارچوب B ابتدا n_b را از یک توزیع دو جمله ای با پارامترهای n_B و $\gamma = \frac{\gamma_b}{1 - \gamma_a}$ تولید می کنیم. به ازای مقداری از n_B مقادیر نمونه ای $\{y_j, j = 1, \dots, n_B\}$ از چارچوب B را از مدل

$$y_j = \mu_b + \delta_j \quad j = 1, \dots, n_b$$

$$y_j = \mu_{ab} + \delta_j \quad j = n_b + 1, \dots, n_B$$

تولید می کنیم بطوریکه δ_j ها متغیرهای تصادفی مستقل و هم توزیع با $\delta_j \sim N(0, \sigma^2)$ می باشند.

به ازای مقادیر مختلف $\tilde{n}_A$ ، n_B ، a_1 ، γ_a و γ_b $R = 10000$ بار این شبیه سازی را تکرار می کنیم. در هر تکرار برآوردگرهای $\hat{Y}_H(\hat{\theta}_H)$ ، $\hat{Y}_{FB}(\hat{\beta}_{FB})$ و $\hat{Y}_{PML}(\hat{\theta}_P)$ و برآوردگرهای تک چارچوبی $\hat{Y}_{SF}$ ، $\hat{Y}_{SFrake}$ و $\hat{Y}_{SFreg}(\hat{\beta}_{S^2})$ را با استفاده از مقادیر بهینه $\hat{\theta}_H$ و $\hat{\beta}_{FB}$ و $\hat{\theta}_P$ و $\hat{\beta}_{S^2}$ محاسبه می کنیم. سپس میانگین توان دوم خطای تجربی هر برآوردگر $(EMSE)$ را به صورت زیر محاسبه می کنیم:

$$EMSE = \frac{1}{R} \sum_{r=1}^R (\hat{Y}_r - Y)^2$$

که $\hat{Y}_r$ عبارت از مقدار $\hat{Y}$ برای r امین تکرار شبیه سازی و R تعداد تکرارهای شبیه سازی می باشد. لازم به ذکر است که تمام محاسبات با استفاده از نرم افزار $SPLUS/2000$ صورت گرفته است. چون جامعه را نامتناهی فرض کردیم و $N \rightarrow \infty$ ، بجای مجموع Y با میانگین

است خصوصاً وقتی از صدک t استفاده می شود احتمال پوشانی جک نایف خیلی نزدیک ۰/۹۵ است.

جک نایف معمولاً با توابع غیر خطی نظیر نسبت مجموع دو جامعه بکار می رود. در این حالت مشتقات جزئی که در محاسبه برآورد خطی چنین مقادیر غیر خطی استفاده می شود، در بررسی های دو چارچوبی پیچیده تر از تک چارچوبی است. در صورتی که با استفاده از جک نایف می توان از این محاسبات پیچیده اجتناب کرد. برآوردهای جک نایف اخیراً در اداره آمار کانادا برای محاسبه برآورد واریانس در آمارگیری وسیع ملی کودکان و جوانان^۱ استفاده می شود. روشهای باز نمونه گیری دیگر برآورد واریانس نظیر تکرار نامتعادل پاسخ^۲ (BRR) و بوتسترپ^۳ نیز می توانند به موازات جک نایف در این گونه طرحها مورد استفاده قرار گیرند. حسن این روش ها این است که برخلاف جک نایف با توابع غیر هموار نظیر میانه نیز قابل کاربرد هستند.

برآوردگر جک نایف واریانس دو برابر خطای معیار نسبی برآورگر خطی است. با این وجود با افزایش حجم نمونه این پایداری بهبود می یابد. با استفاده از فاصله اطمینان $\hat{t} \pm 1/96\sqrt{\hat{V}}(\hat{t})$ احتمال پوشایی تجربی جک نایف بیشتر از خطی است. با استفاده از توزیع t با $1 - \tilde{n}_A$ درجه آزادی بجای ۱/۹۶، جک نایف احتمال پوشایی نزدیک ۰/۹۵ دارد در حالی که روش خطی به وضوح احتمال پوشایی کمتری دارد.

۵. نتیجه گیری

برآوردهای مجموع در بررسی های دو چارچوبی را بطور خلاصه مطرح و مقایسه کردیم. این برآوردها قابل تعمیم به بررسی هایی با بیش از دو چارچوب نیز می باشند. برآوردر جک نایف واریانس را بطور نظری توجیه کرده با استفاده از شبیه سازی نشان دادیم که دارای اریبی کوچکتری از برآوردر خطی واریانس است. همچنین برآوردر جک نایف به وضوح دارای احتمال پوشایی بالاتری از برآوردر خطی

احتمال پوشایی							
t		نرمال		اریبی نسبی		$EMSE$	برآوردگر
J	L	J	L	J	L		
۰/۹۴۳	۰/۹۲۳	۰/۹۱۹	۰/۸۸۷	۱۵/۴۵ (۱/۹۰)	-۱۹/۹۴ (۰/۷۸)	۷/۱۶	H
۰/۹۴۶	۰/۹۳۱	۰/۹۲۱	۰/۹۰۱	۱۰/۴۶ (۱/۶۴)	-۱۲/۷۵ (۰/۷۷)	۶/۸۳	FB
۰/۹۵۰	۰/۹۳۹	۰/۹۲۴	۰/۹۱۲	۱۱/۴۶ (۱/۸۷)	-۹/۹۱ (۰/۸۲)	۶/۵۸	PML
۰/۹۴۴	۰/۹۳۰	۰/۹۱۹	۰/۸۹۹	۱۲/۲۱ (۱/۸۱)	-۱۳/۲۵ (۰/۸۰)	۶/۸۳	$PMLlin$
۰/۹۷۲	۰/۹۷۲	۰/۹۴۸	۰/۹۴۸	۴/۱۷ (۰/۴۷)	۴/۱۷ (۰/۴۷)	۱۱/۲۶	SF
۰/۹۴۳	۰/۹۲۸	۰/۹۱۸	۰/۸۹۷	۱۱/۴۲ (۱/۷۸)	-۱۳/۷۳ (۰/۸۲)	۶/۹۷	$SFreg$
۰/۹۵۰	۰/۹۴۴	۰/۹۲۳	۰/۹۱۸	۴/۸۵ (۱/۱۹)	-۰/۱۱ (۱/۰۸)	۶/۵۵	$SFrake$
۰/۹۴۴	۰/۹۲۳	۰/۹۲۸	۰/۹۰۶	۱/۸۰ (۰/۷۴)	-۱۵/۹۷ (۰/۴۱)	۳/۵۸	H
۰/۹۴۳	۰/۹۳۰	۰/۹۲۸	۰/۹۱۳	-۰/۶۳ (۰/۶۷)	-۱۲/۶۹ (۰/۴۰)	۳/۵۲	FB
۰/۹۴۵	۰/۹۳۵	۰/۹۳۰	۰/۹۱۹	-۱/۱۰ (۰/۷۴)	-۱۱/۲۶ (۰/۴۲)	۳/۴۴	PML
۰/۹۴۳	۰/۹۲۸	۰/۹۲۸	۰/۹۱۱	-۰/۰۵ (۰/۷۲)	-۱۳/۱۰ (۰/۴۲)	۳/۵۲	$PMLlin$
۰/۹۵۸	۰/۹۵۸	۰/۹۴۴	۰/۹۴۴	-۰/۶۹ (۰/۲۳)	-۰/۶۹ (۰/۲۳)	۵/۸۳	SF
۰/۹۴۲	۰/۹۲۸	۰/۹۲۸	۰/۹۱۳	۰/۰۸ (۰/۷۲)	-۱۲/۹۲ (۰/۴۲)	۳/۵۷	$SFreg$
۰/۹۴۴	۰/۹۴۰	۰/۹۲۹	۰/۹۲۵	-۳/۰۷ (۰/۵۵)	-۵/۳۳ (۰/۵۳)	۳/۴۸	$SFrake$

۱- National Longitudinal Survey of Children and Youth

۲- Balanced Repeated Replication

۳- Bootstrap

لازم به توجه است که در این جدول نتایج تنها به ازای $a_1=1$ و $a_1=0.1$ = 0.1 و γ_a و $\gamma_b = 0.2$ ارائه شده است. در بخش بالایی، $\tilde{n}_A = 10$ و $n_B = 100$ و در بخش پایینی $\tilde{n}_A = 20$ و $n_B = 200$ می باشد. $EMSE$ عبارت از میانگین مربعات خطای مونت کارلو برای ۱۰۰۰۰ بار تکرار شبیه سازی است. اریبی نسبی برای روش های خطی (L) و جک نایف (J) بصورت $EMSE(EV-EMSE)$ ۱۰۰ محاسبه شده است، بطوریکه EV عبارت از میانگین ۱۰۰۰۰ برآورد واریانس برای آن روش است. احتمال پوشایی هر روش با استفاده از $[SE(Estimate)]$

$[Estimate + 1/96]$ و همچنین یک صدک t با $1 - \tilde{n}_A$ درجه آزادی محاسبه شده است. خطای معیار نسبی برای هر روش (= انحراف معیار ۱۰۰۰۰ برآورد واریانس $(\sqrt{EMSE})$ در پراتز در زیر اریبی نسبی داده شده است. همچنین شبیه سازی را به ازای مقادیر مختلف a_1 ، (ρ, δ) و $(n_B$ و $\tilde{n}_A)$ بررسی کرده ایم. و نتایج شبیه نتایج بالا است.

مراجع

- [1] Hartley, H. O., 1962. Multiple Frame Surveys, in Proceedings of the Social Statistics Section, *American Statistical Association*. 203-206.
- [2] Lund, R. E., 1968. Estimators in Multiple Frame Surveys, in *Proceedings of the Social Statistics Section, American Statistical Association*. 282-288.
- [3] Fuller, W. A., & Burmeister, L. F., 1972. Estimators for Samples Selected From two Overlapping Frames *Proceedings of the Social Statistics Section, American Statistical Association*. 245-249.
- [4] Vogel, Frederic A., 1975. Survey With Overlapping Frame-Problems in Application, *Proceedings of the Social Statistics Section, American Association*. 694-699.
- [5] Ford, B. L., & Bosecker, R. R., 1979. Multiple Frame Estimation With Stratified Overlap Domain *Proceedings of the Social Statistics Section, American Statistical Association*. 219-224.
- [6] Bankier, M. D., 1986. Estimators Based on Several Stratified Samples with Applications to Multiple Frame Surveys, *Journal of the American Statistical Association*, 81, 1074-1079.
- [7] Skinner, C. J., 1991, On the Efficiency of Raking Ratio Estimation for Multiple Frame Surveys'' *Journal of the American Statistical Association*, 86. 779-784.
- [8] Skinner, C. J., & Rao, J. N. K., 1996. Estimation in Dual Frame Survey with Complex Designs, *Journal of the American Statistical Association*, 91. 349-356.
- [9] Kalton, G., & Anderson, T. W., 1986. ''Sampling Rare Populations'' *Journal of the Royal Statistical Society, Ser.A*, 149. 65-82.
- [10] Sarndal, C-E., Swensson, B., & Wretman, J. 1992., *Model Assisted Survey Sampling*, New York: Springer-Verlag.
- [11] Quenouille, M. H., 1956, Notes on Bias in Estimation, *Biometrika*, 43. 353-360.
- [12] Tukey, J. W., 1958. Bias and Confidence in Note Quite Large Samples, *Annals of Mathematical Statistics*, 29. 614.
- [13] Durbin, J., 1959, *A Note on the Application of Quenouille's Method of Bias Reduction To the Estimation of Ratios*, *Biometrika*, 46. 477-480.
- [14] Miller, R. G. Jr, 1974, The Jackknife-A Review , *Biometrika*, 61, 1-15.
- [15] Efron, B., 1982. The Jackknife, *The Bootstrap, and The other Resampling Plans*, Philadelphia: Society for Industrial and Applied Mathematics.
- [16] Rao, J. N. K., & Lohr, Sharon, 2000. Inference from Dual Frame Survey, *Journal of The American Statistical Association*, 95. 271-280.