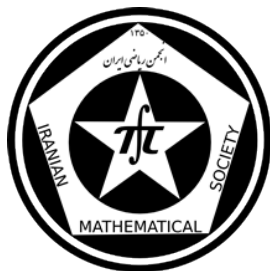


الحمد لله
الرحمن الرحيم

به نام خدا



اولین همایش ملی ریاضیات زیرستی

کتابچه مقالات

۲۱ و ۲۲ اسفند ماه ۱۳۹۷

دانشگاه نیشابور

پیش‌گفتار

با یاری خداوند متعال توفیق اجرای اولین همایش ملی ریاضیات زیستی را با همت ستودنی اساتید و دانشجویان و حمایت گروه ریاضی و آمار و همچنین مسئولین محترم دانشگاه نیشابور و سازمان‌ها و دستگاه‌های مختلف کشور در روزهای ۲۱ و ۲۲ اسفند ۱۳۹۷ در دانشگاه نیشابور یافتیم. تا با ارایه آخرین دست‌آوردهای علمی، بحث و تبادل نظر بین دانشگاهیان، پژوهش‌گران و کارشناسان در این زمینه، زمینه‌ساز ارتباط قوی‌تر بین دانشگاهیان وزارت های علوم، تحقیقات و فناوری و بهداشت، درمان و آموزش پزشکی شده تا باعث کیفیت مطلوب تر زندگی مردم گردیم.

در این همایش، ۳ سخنرانی کلیدی با عناوین نحوه بهینه تجویز داروی تاموکسیفن در مهار متاستاز استخوان، مدل سازی پیامد بیماران آلوده به ویروس *HIV* به کمک سیستم های معادلات دیفرانسیل و آنالیز تک سلولی ها توسط اساتید مدعو در نظر گرفته شد. متقاضیان شرکت در کنفرانس تعداد ۶۲ مقاله به دبیرخانه ارسال نمودند که کمیته علمی کنفرانس به همراه داوران محترم تعداد ۲۴ مقاله را به‌عنوان سخنرانی، ۲۷ مقاله را برای ارایه به صورت پوستر و ۱ کارگاه آموزشی را پذیرفت.

همکاری صمیمانه اساتید دانشگاه‌های سراسر کشور بویژه اعضای کمیته های اجرایی و علمی را قدر می نهیم. همکاری این عزیزان در ارایه پیشنهادهای سازنده، کمک شایانی به هر چه بهتر اجرا شدن همایش و متون علمی این کتابچه بوده است. در پایان، کمیته علمی کنفرانس از تلاش‌های بی‌دریغ و زحمات صادقانه همه مسوولین، همکاران هیات علمی و دانشجویان گروه ریاضی و آمار دانشکده علوم پایه، و کارکنان بخش‌های مختلف دانشگاه نیشابور، تشکر و قدردانی می‌کند.

کمیته برگزاری اولین همایش ملی ریاضیات زیستی

دانشگاه نیشابور

۲۱ و ۲۲ اسفند ماه ۱۳۹۷

اعضای کمیته اجرایی

- | | |
|-----------------|---------------------------------------|
| دانشگاه نیشابور | ۱. دکتر رضا محمدی (دبیر کمیته اجرایی) |
| دانشگاه نیشابور | ۲. دکتر مژگان افخمی گلی |
| دانشگاه نیشابور | ۳. دکتر احسان انجیدنی |
| دانشگاه نیشابور | ۴. دکتر بهاره خطیب آستانه |
| دانشگاه نیشابور | ۵. دکتر محمد حسین رحمانی دوست |
| دانشگاه نیشابور | ۶. دکتر محمد شیرازیان |
| دانشگاه نیشابور | ۷. دکتر سید محسن صالح |
| دانشگاه نیشابور | ۸. دکتر سید مهدی صالحی |
| دانشگاه نیشابور | ۹. دکتر داوود عمارتی مقدم |

اعضای کمیته علمی

۱. دکتر سید محسن صالح (دبیر کمیته علمی) دانشگاه نیشابور
۲. دکتر حبیب الله اسماعیلی دانشگاه علوم پزشکی مشهد
۳. دکتر جواد جمالزاده دانشگاه سیستان و بلوچستان
۴. دکتر محمود حصارکی دانشگاه صنعتی شریف
۵. دکتر حسین خیری دانشگاه تبریز
۶. دکتر امید ربیعی مقدم دانشگاه بیرجند
۷. دکتر محمد حسین رحمانی دوست دانشگاه نیشابور
۸. دکتر علیرضا سپاسیان دانشگاه فسا
۹. دکتر محمد تقی شاکری دانشگاه علوم پزشکی مشهد
۱۰. دکتر موسی الرضا شمسیه زاهدی دانشگاه پیام نور
۱۱. دکتر محمد شیرازیان دانشگاه نیشابور
۱۲. دکتر قاسم صادقی بجستانی دانشگاه امام رضا (ع)
۱۳. دکتر روزبه عابدینی نسب دانشگاه نیشابور
۱۴. دکتر علیرضا فخارزاده جهرمی دانشگاه صنعتی شیراز
۱۵. دکتر عباس فخاری قوچانی دانشگاه شهید بهشتی تهران
۱۶. دکتر آتنا قاسم آبادی مرکز آموزش عالی اسفراین
۱۷. دکتر مهدیه قاسمی دانشگاه نیشابور
۱۸. دکتر فاطمه هلن قانع استاد قاسمی دانشگاه فردوسی مشهد
۱۹. دکتر محمدرضا محمودی دانشگاه فسا
۲۰. دکتر محمد رضا مولایی دانشگاه شهید باهنر کرمان
۲۱. دکتر علی وحیدیان کامیاد دانشگاه فردوسی مشهد

فهرست مقالات

صفحه	عنوان
۱	مدلسازی ریاضی زیست‌حسگر گلوکز و حل عددی با استفاده از روش هم‌مکانی توابع پایه شعاعی مریم ابجدیان و آمنه طالعی
۷	کاربرد نمونه‌گیری مجموعه رتبه‌دار در مسائل پزشکی و زیستی حسن اسفندیاری‌فر و سید مهدی صالحی
۱۲	کاربرد ضریب همبستگی درون گروهی (ICC) در تعیین حجم نمونه برای داده‌های خوشه‌ای و تاثیر آن بر توان مطالعه در پژوهش‌های علوم پزشکی فرزانه برومند، حبیب الله اسماعیلی، محمد تقی شاکری و نوشین اکبری شارک
۱۸	معادلات دیفرانسیل کسری سوسن امیری و محمد حسین رحمانی دوست
۲۳	مقایسه‌های چندگانه در کارآزمایی‌های بالینی ناپست‌تری سعیده حاجبی خانیکی، دکتر محمدتقی شاکری، احسان برادران سیرجانی و نگار سنگ سفیدی
۲۹	معادلات ساختاری بی‌زی نجمه شکیبایی و عقیل علایی
۳۶	تجزیه و تحلیل پایداری مدل شکار-شکارچی با واکسیناسیون و مهاجرت در شکار مرضیه شمس آبادی، محمد حسین رحمانی دوست و محمد شیرازیان
۴۳	مطالعه عددی یک مدل ریاضی بیماری ناشی از آلودگی آب یونس طالعی و حجت اله عبادیزاده
۴۹	بررسی مدل ریاضی شیوع یک عامل بیماری حجت‌الله عبادی‌زاده، علی عزیزی مایوان و سید مهدی کاظمی تربقان
۵۳	یک مدل ریاضی برای هپاتیت B و بررسی اثر واکسیناسیون در رشد بیماری حجت‌الله عبادی‌زاده، علی عزیزی مایوان و سید مهدی کاظمی تربقان
۵۹	کاربرد رگرسیون لجستیک جریمه شده با لاسوی تطبیقی در تحلیل داده‌های با بعد بالا نجمه عسگری فرسنگی، راضیه یوسفی و محمد تقی شاکری
۶۵	تحلیلی بر مدل بیماری سرطان مبتنی بر نظریه آشوب میلاد فهیمی، کاظم نوری و لیلا ترک زاده

- ۷۲ **تحلیلی بر مدل تصادفی اپیدمی با نرخ شیوع غیرخطی بیماری**
میلاذ فهیمی، کاظم نوری و لیلا ترک زاده
- ۷۹ **نتایجی از C^3 -کدهای خودمکمل ماکسیمال**
فریبا فیاضی، سید علیرضا اشرفی، احمد غلامی
- ۸۴ **مدل‌های ریاضی ایمن درمانی سرطان**
فائزه قربانی و سید محسن صالح
- ۸۹ **بررسی نقاط تعادل و پایداری یک مدل شکار - شکارچی به همراه یک منبع زیستی**
فرزانه لطفی و محمد شیرازیان
- ۹۵ **مقایسه فواصل اطمینان برای نسبت شانس**
اعظم مسلمی، محمد رفیعی، راشد پورحمیدی، کورش رضایی، سردار جهانی و مهشید عسکری
- ۱۰۴ **برآورد طول دوره کمون در بیماران مبتلا به HIV در استان خراسان رضوی**
نجمه نخعی راد و وحید فکور
- ۱۱۲ **خواص مجانبی برآوردگرهای کاپلان-مه-یر و نلسون-آلن برای داده های سانسور شده ی وابسته**
نجمه نخعی راد و وحید فکور
- ۱۲۰ **کاربرد تکنیک مدلسازی معادلات ساختاری در شناسایی عوامل خطر بیماری فشار خون**
راضیه یوسفی و آزاده ساکی
- ۱۲۹ **مدلسازی معادلات ساختاری بیزی و کاربرد آن در مدلسازی عوامل خطر بیماری ها**
راضیه یوسفی و آزاده ساکی

مدلسازی ریاضی زیست‌حسگر گلوکز و حل عددی با استفاده از روش هم‌مکانی توابع پایه شعاعی آمنه طالعی^۱، دانشکده ریاضی، دانشگاه صنعتی شیراز، شیراز، ایران

چکیده: زیست‌حسگر گلوکز از اولین و پرکاربردترین زیست‌حسگرهای موجود در بازار است، که مدل ریاضی آن را می‌توان به صورت معادلات دیفرانسیل پاره‌ای بیان کرد. در این مقاله قصد داریم به بررسی حل عددی معادله نفوذ-واکنش زیست‌حسگر بپردازیم. برای حل عددی، از روش هم‌مکانی با استفاده از توابع پایه شعاعی برای گسسته‌سازی متغیر متناظر با غلظت گلوکز مایعاتی مانند خون و غلظت محصولات حاصل از واکنش استفاده می‌کنیم. همچنین متغیر زمان را به روش تفاضل‌متناهی پسر و گسسته‌سازی می‌کنیم.

واژه‌های کلیدی: زیست‌حسگر، روش توابع پایه شعاعی، معادله نفوذ-واکنش.

۱ مقدمه

زیست‌حسگر نخستین بار توسط پروفیسور لی‌لاند. سی. کلارک در سال ۱۹۵۶ معرفی شد و پس از آن در سال ۱۹۷۵ اولین زیست‌حسگر تجاری که برای اندازه‌گیری گلوکز به روش آمپرومتری بود، معرفی شد (۴). گلوکز با آنزیم گلوکز اکسیداز در فرایند اکسیداسیون کاتالستی به اکسولاکتون و آب اکسیژنه تبدیل می‌شود. از طریق فرایند تجزیه‌ی آب اکسیژنه حاصل از این واکنش بر روی سطح الکتروود و انتقال الکترونی حاصل از آن، سیگنالی در مدار ایجاد می‌شود. زیست‌حسگر آمپرومتری وسیله‌ای است که این سیگنالها را اندازه‌گیری می‌کند و در نتیجه میزان گلوکز مایعاتی مانند خون مشخص می‌شود. این عملکرد را می‌توان به صورت واکنش زیر نشان داد:



که E ، S ، ES و P به ترتیب متناظر با آنزیم گلوکز اکسیداز، گلوکز خون (سوبسترا)، کمپلکس آنزیم-سوبسترا و محصول حاصل از واکنش است. با افزایش جمعیت و درگیر شدن افراد با بیماری‌های مختلف، اهمیت استفاده از زیست‌حسگرها در تشخیص بیماری به دلیل تشخیص بیولوژیکی بسیار خاص و پایدار و قابل اعتماد آنها، مشخص می‌شود. شبیه‌سازی عددی فرصت مطالعه‌ی ویژگی‌های زیست‌حسگر گلوکز بدون نیاز به انجام آزمایش‌های پرهزینه و زمان‌بر را فراهم می‌کند.

مدل ریاضی زیست‌حسگر اغلب به صورت معادلات دیفرانسیل غیرخطی است که به ندرت جواب دقیق برای آن وجود دارد. بنابراین روش‌های عددی برای حل اینگونه معادلات بسیار موثر خواهد بود. در این مقاله به حل عددی مدل ریاضی یک بعدی برای زیست‌حسگر آمپرومتری گلوکز بر پایه‌ی آنزیم گلوکز اکسیداز می‌پردازیم. روش هم‌مکانی با استفاده از توابع پایه‌شعاعی برای گسسته‌سازی غلظت گلوکز مایعاتی مانند خون و غلظت محصولات حاصل از واکنش و روش تفاضل‌متناهی پسر و گسسته‌سازی متغیر زمان پیشنهاد می‌شود.

^۱آمنه طالعی: a.taleei@sutech.ac.ir

۲ مدل ریاضی زیست‌حسگر

مدل ریاضی زیست‌حسگر شامل معادلات نفوذ-واکنش با جمله‌ی غیرخطی طرح میکائیلیس-منتن است (۳). مطابق قانون فیک، اتصال واکنش کاتالیزشده با آنزیم در لایه‌ی آنزیم گلوکز اکسیداز با نفوذ در فضای یک بعدی تشریح شده و به معادلات ذیل منجر می‌شود:

$$\frac{\partial \hat{S}}{\partial \hat{t}} = D_{\hat{S}} \frac{\partial^2 \hat{S}}{\partial \hat{x}^2} - \frac{V_{max} \hat{S}}{K_M + \hat{S}}, \quad (1.2)$$

$$\frac{\partial \hat{P}}{\partial \hat{t}} = D_{\hat{P}} \frac{\partial^2 \hat{P}}{\partial \hat{x}^2} - \frac{V_{max} \hat{S}}{K_M + \hat{S}}, \quad (2.2)$$

که $\hat{x}$ و $\hat{t}$ به ترتیب فاصله و زمان را نشان می‌دهند. $D_{\hat{S}}$ و $D_{\hat{P}}$ ضرایب انتشار هستند. V_{max} حداکثر سرعت آنزیم و K_M ثابت میکائیلیس است. ثابت میکائیلیس غلظت سوبسترای است که در آن نصف حداکثر چگالی واکنش کاتالیز شده با آنزیم به دست آمده است. شرایط اولیه و مرزی را می‌توان به فرم زیر در نظر گرفت:

$$\hat{S}(\hat{x}, 0) = 0, \quad \hat{x} \in [0, d], \quad \hat{S}(d, 0) = \hat{S}_0, \quad (3.2)$$

$$\hat{P}(\hat{x}, 0) = 0, \quad \hat{x} \in [0, d], \quad \hat{P}(d, 0) = \hat{P}_0, \quad (4.2)$$

$$\frac{\partial \hat{S}}{\partial \hat{x}}(0, \hat{t}) = 0, \quad \hat{P}(0, \hat{t}) = 0, \quad \hat{t} > 0, \quad (5.2)$$

$$\hat{S}(d, \hat{t}) = \hat{S}_0, \quad \hat{P}(d, \hat{t}) = \hat{P}_0, \quad \hat{t} > 0. \quad (6.2)$$

جریان سنجش شده به عنوان پاسخ زیست‌حسگر که به شارش ماده‌ی الکترواکتیو در سطح الکتروود وابسته است به طور صریح از قوانین فیک و فارادی برحسب اکسایش و یاکاهش $\hat{P}$ به دست می‌آید (۱):

$$\hat{i} = n_e F D_{\hat{P}} \frac{\partial \hat{P}}{\partial \hat{x}}(0, \hat{t}), \quad (7.2)$$

که n_e تعداد الکترون‌های وارده در انتقال شار و F ثابت فارادی است. برای تعیین پارامترهای اصلی حاکم بر مدل ریاضی (۱)-(۷)، متغیرهای بعد دار $\hat{x}$ و $\hat{t}$ و غلظت‌های مجهول $\hat{S}$ و $\hat{P}$ را با پارامترهای بدون بعد زیر جایگزین می‌کنیم:

$$S = \frac{\hat{S}}{\hat{S}_0}, \quad P = \frac{\hat{P}}{\hat{P}_0}, \quad x = \frac{\hat{x}}{d}, \quad t = \frac{D_{\hat{S}} \hat{t}}{d^2}, \quad r = \frac{D_{\hat{P}}}{D_{\hat{S}}}, \quad \mu = \frac{v_m d^2}{D_{\hat{S}} \hat{S}_0}, \quad k = \frac{k_m}{\hat{S}_0}. \quad (8.2)$$

معادلات (۱) و (۲) به صورت زیر نوشته می‌شوند:

$$\frac{\partial S}{\partial t} = \frac{\partial^2 S}{\partial x^2} - \frac{\mu S}{k + S}, \quad (9.2)$$

$$\frac{\partial P}{\partial t} = r \frac{\partial^2 P}{\partial x^2} + \frac{\mu S}{k + S}. \quad (10.2)$$

شرایط اولیه و مرزی مدل بی بعد به صورت زیر است:

$$S(x, 0) = 0, \quad P(x, 0) = 0, \quad x \in [0, 1], \quad (11.2)$$

$$\frac{\partial S}{\partial x}(0, t) = 0, \quad P(0, t) = 0, \quad t > 0, \quad (12.2)$$

$$S(1, t) = 1, \quad P(1, t) = 0, \quad t \geq 0, \quad (13.2)$$

همچنین چگالی جریان بی بعد به صورت زیر می باشد:

$$\frac{\hat{id}}{n_e F D_{\hat{P}} \hat{S}} = \frac{\partial P}{\partial x}(0, t). \quad (14.2)$$

در صورتی که غلظت گلوکز در لایه ی آنزیمی خیلی کمتر از ثابت میکائلیس شود ($S \ll k$)، تابع غیرخطی میکائلیس - منتن به صورت مرتبه اول ساده می شود و معادلات (۹) و (۱۰) به صورت زیر بدست می آید:

$$\frac{\partial S}{\partial t} = \frac{\partial^2 S}{\partial x^2} - \frac{\mu S}{k}, \quad (15.2)$$

$$\frac{\partial P}{\partial t} = r \frac{\partial^2 P}{\partial x^2} + \frac{\mu S}{k}. \quad (16.2)$$

۳ اعمال روش پیشنهادی

در این مقاله، گسسته سازی متغیر زمان به روش تفاضلات متناهی پسرو و گسسته سازی متغیر مکان به روش هم مکانی بر پایه ی توابع پایه شعاعی در نظر گرفته می شود. با انتخاب N گره در بازه $[0, 1]$ ، توابع مجهول S و P و مشتقات آن ها را در هر گام زمانی درونیابی می کنیم. انواع متعددی از توابع شعاعی وجود دارد. در این مقاله تابع چندمربعی تحت پارامتر شکل c ، را به عنوان تابع شعاعی انتخاب می کنیم (۲). با قرار دادن تقریب تابع و مشتقات آن در معادلات دیفرانسیل مفروض و با استفاده از روش هم مکانی، به یک دستگاه معادلات جبری دست می یابیم.

با فرض Δt به عنوان گام زمانی، مشتق متغیر زمان را در معادلات (۱۵) و (۱۶) به صورت $\frac{S^{n+1} - S^n}{\Delta t}$ تقریب می زنیم و با به کاربردن تقریب $S^{n+1}(x) = \sum_{j=1}^N \lambda_j^{n+1} \phi(r_j)$ و $P^{n+1}(x) = \sum_{j=1}^N \alpha_j^{n+1} \phi(r_j)$ داریم:

$$\sum_{j=1}^N \lambda_j^{n+1} [-\Delta t \phi_{xx}(r_{ij}) + (1 + \Delta t \frac{\mu}{k}) \phi(r_{ij})] = \sum_{j=1}^N \lambda_j^n \phi(r_{ij}), \quad (1.3)$$

$$\sum_{j=1}^N \alpha_j^{n+1} [-r \Delta t \phi_{xx}(r_{ij}) + \phi(r_{ij})] - \Delta t \frac{\mu}{k} \sum_{j=1}^N \lambda_j^{n+1} \phi(r_{ij}) = \sum_{j=1}^N \alpha_j^n \phi(r_{ij}). \quad (2.3)$$

با اعمال شرایط مرزی به دستگاه $L_1 Y^{n+1} = L_2 Y^n + C$ می رسیم به طوریکه:

$$L_1 = \begin{bmatrix} A & O \\ M & D \end{bmatrix}; \quad L_2 = \begin{bmatrix} B & O \\ O & B \end{bmatrix}; \quad Y^k = [\lambda; \alpha]^T, \quad k = n, n+1,$$

$$C(1 : N-1) = 0; \quad C(N) = 1; \quad C(N+1 : 2N) = 0,$$

$$A = \begin{bmatrix} \phi_x(r_{11}) & \dots & \phi_x(r_{1N}) \\ -\Delta t \phi_{xx}(r_{11}) + (1 + a\Delta t)\phi(r_{11}) & \dots & -\Delta t \phi_{xx}(r_{1N}) + (1 + a\Delta t)\phi(r_{1N}) \\ \vdots & \vdots & \vdots \\ -\Delta t \phi_{xx}(r_{(N-1)1}) + (1 + a\Delta t)\phi(r_{(N-1)1}) & \dots & -\Delta t \phi_{xx}(r_{(N-1)N}) + (1 + a\Delta t)\phi(r_{(N-1)N}) \\ \phi(r_{N1}) & \dots & \phi(r_{NN}) \end{bmatrix},$$

$$D = \begin{bmatrix} \phi_x(r_{11}) & \dots & \phi_x(r_{1N}) \\ -r\Delta t \phi_{xx}(r_{11}) + \phi(r_{11}) & \dots & -r\Delta t \phi_{xx}(r_{1N}) + \phi(r_{1N}) \\ \vdots & \vdots & \vdots \\ -r\Delta t \phi_{xx}(r_{(N-1)1}) + \phi(r_{(N-1)1}) & \dots & -r\Delta t \phi_{xx}(r_{(N-1)N}) + \phi(r_{(N-1)N}) \\ \phi_x(r_{N1}) & \dots & \phi_x(r_{NN}) \end{bmatrix},$$

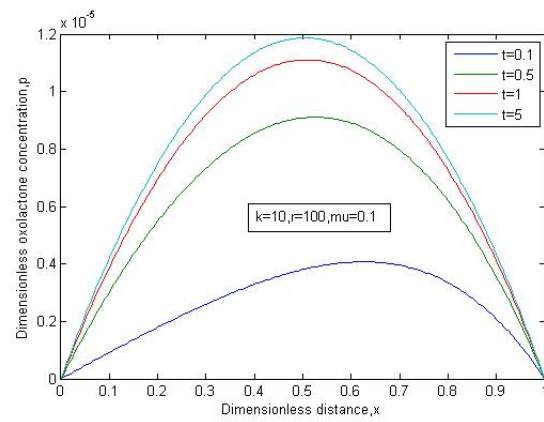
$$M = \begin{bmatrix} \cdot & \dots & \cdot \\ -a\Delta t \phi_{xx}(r_{11}) + \phi(r_{11}) & \dots & -a\Delta t \phi_{xx}(r_{1N}) + \phi(r_{1N}) \\ \vdots & \vdots & \vdots \\ -a\Delta t \phi_{xx}(r_{(N-1)1}) + \phi(r_{(N-1)1}) & \dots & -a\Delta t \phi_{xx}(r_{(N-1)N}) + \phi(r_{(N-1)N}) \\ \cdot & \dots & \cdot \end{bmatrix},$$

$$B = \begin{bmatrix} \cdot & \dots & \cdot \\ \phi(r_{11}) & \dots & \phi(r_{1N}) \\ \vdots & \vdots & \vdots \\ \phi(r_{(N-1)1}) & \dots & \phi(r_{(N-1)N}) \\ \cdot & \dots & \cdot \end{bmatrix}.$$

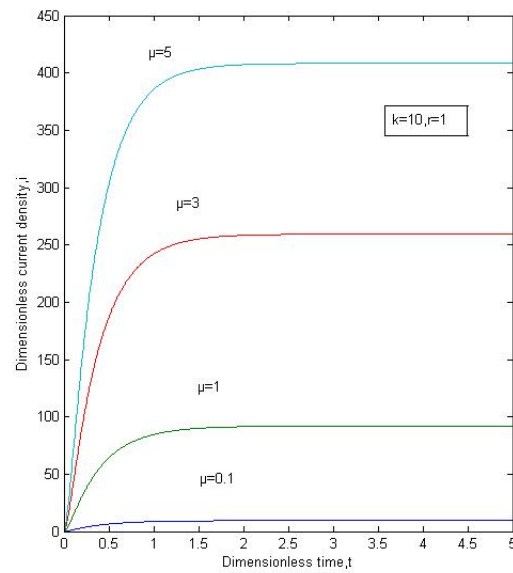
۴ نتایج عددی

غلظت گلوکز، محصول اکسولاکتون و آب‌اکسیژنه حاصل از واکنش و چگالی جریان زیست‌حسگر به پارامتر بی‌بعد $a = \mu/k$ ، بستگی دارند. μ ، سرعت واکنش آنزیم را با سرعت انتقال توده در سراسر لایه‌ی آنزیمی می‌سنجد و پارامتر مهمی در تعیین پارامترهای بهینه مانند ضخامت لایه‌ی آنزیمی است.

برای محاسبه پاسخ زیست‌حسگر، تغییر غلظت اکسولاکتون و آب‌اکسیژنه در سطح مبدل ارزیابی می‌شود. همانطور که در شکل (۱) مشاهده می‌شود غلظت محصول در $t = 5$ ، به حالت پایا می‌رسد. شکل (۲) نشان می‌دهد که با افزایش مدول μ ، چگالی جریان زیست‌حسگر افزایش می‌یابد. با توجه به تاثیر مستقیم ضخامت غشای آنزیمی بر مدول μ ، می‌توان نتیجه گرفت که واکنش حالت پایا همیشه در غشای ضخیم‌تر به نسبت غشای نازک‌تر، بالاتر است.



شکل ۱: غلظت اکسولاکتون در مقادیر زمانی مختلف



شکل ۲: چگالی جریان با مقادیر مختلف مدول μ

۵ نتیجه گیری

با استفاده از روش هم مکانی توابع پایه شعاعی به بررسی مدل ریاضی زیست حسگر آمپرومتری پرداختیم. چگالی جریان زیست حسگر به واکنش گلوکز با آنزیم گلوکز اکسیداز بررسی شد و نشان دادیم که پاسخ به مدول μ وضخامت غشای آنزیمی وابسته است.

مراجع

- [1] R. Baronas, , F. Ivanauskas, , J. Kulys, *Mathematical Modeling of Biosensors: An Introduction for Chemists and Mathematicians*, Springer series on chemical sensors and biosensors, 2010.
- [2] G. E. Fasshauer, *Meshfree Approximation Methods with Matlab*, Interdisciplinary Mathematical Sciences - Vol. 6, World Scientific Publishers, Singapore, 2007.
- [3] O.M. Kirthiga, L. Rajendran, Approximate analytical solution for non-linear reaction diffusion equations in a mono-enzymatic biosensor involving Michaelis- Menten kinetics, *Journal of Electroanalytical Chemistry*, 751, (2015), 119-127.
- [4] J. Wang, Glucose biosensors: 40 years of advances and challenges, *Electroanalysis*, 13 (2001), 983-988.
- [5] M. Ylilammi, L. Lehtinen, Numerical analysis of a theoretical one-dimensional amperometric enzyme sensor, *Medical & Biological Engineering & Computing*, 26 (1988), 81-87.

کاربرد نمونه‌گیری مجموعه رتبه‌دار در مسائل پزشکی و زیستی

حسن اسفندیاری فر، گروه ریاضی، دانشگاه امام علی (ع)، تهران، ایران و سید مهدی صالحی^۱، گروه ریاضی، دانشگاه نیشابور، نیشابور، ایران

چکیده: دونه‌گیری مجموعه‌ی رتبه‌دار برای اولین بار توسط مک‌اینتایر (۱۹۵۲) و با انگیزه‌ی یافتن روشی جدید در کنار نمونه‌گیری تصادفی ساده، معمول‌ترین روش در علم نمونه‌گیری، معرفی شد. یکی از ضعف‌های روش سنتی نمونه‌گیری تصادفی ساده این است که در صورت کم بودن اندازه‌ی نمونه، به‌سختی می‌توان یک نمونه‌ی مناسب معرف جامعه بر حسب آن یافت. درحالی‌که نمونه‌گیری مجموعه‌ی رتبه‌دار، که بر پایه‌ی رتبه‌بندی قضاوتی یا چشمی می‌باشد، بر این مشکل غلبه کرده است. بنابراین در صورت برقراری شرایط اعمال این روش رتبه‌دار در مواردی که اندازه‌گیری دقیق واحدهای جامعه هزینه‌بر یا زمان‌بر است، بهتر است جایگزین روش تصادفی ساده شود. در این مقاله به برخی از کاربردهای طرح نمونه‌گیری مجموعه‌ی رتبه‌دار در مسائل پزشکی و زیستی اشاره می‌کنیم.

واژه‌های کلیدی: نمونه‌گیری مجموعه‌ی رتبه‌دار، نمونه‌گیری تصادفی ساده، آزمایش‌های کلینیکی، آلودگی‌های زیستی

۱ مقدمه

در مبحث نمونه‌گیری، پژوهش‌گر همواره با مسأله‌ی انتخاب روش مناسب برای جمع‌آوری داده‌ها مواجه است، به‌طوری‌که وی به‌دنبال روش‌هایی است که دقت را برای برآورد ویژگی‌های یک جامعه‌ی آماری افزایش داده و هزینه را کمینه کند. شرایط زیر را در نظر بگیرید: اندازه‌گیری دقیق واحدهای جامعه‌ی آماری سخت، پرهزینه و یا زمان‌بر است، اما می‌توان آن‌ها را به‌سادگی و با کمترین هزینه رتبه‌بندی کرد. این رتبه‌بندی را می‌توان به‌صورت چشمی یا استفاده از یک نظر کارشناسانه روی متغیر مورد نظر یا حتی یک متغیر دیگر مرتبط با آن انجام داد. در چنین شرایطی، روش نمونه‌گیری مجموعه‌ی رتبه‌دار می‌تواند جایگزین مناسبی برای روش سنتی تصادفی ساده باشد. دلیل این امر این است که در اغلب موارد، برآوردگرهای حاصل از روش نمونه‌گیری مجموعه‌ی رتبه‌دار عملکرد بهتری نسبت به برآوردگرهای حاصل از روش نمونه‌گیری تصادفی دارند، بنابراین با تعداد نمونه‌ی کمتر در روش اول، می‌توان به دقت یکسان در روش دوم رسید. کاربرد عمده‌ی این روش نمونه‌گیری در مطالعات زیست محیطی، کشاورزی و آزمایشگاهی است. به‌طور مثال فرض کنید مایل به تخمین میانگین طول عمر درختان یک جنگل هستیم. این آزمایش به‌دلیل قطع کردن درختان هزینه‌بر خواهد بود. اما درختان را بر اساس ارتفاع آن‌ها (که با طول عمر آن‌ها وابستگی دارد) می‌توان به‌صورت چشمی مرتب کرد و در نتیجه با به‌کارگیری روش نمونه‌گیری

^۱ سید مهدی صالحی: salehi2sms@gmail.com

مجموعه‌ی رتبه‌دار متحمل هزینه‌ی کمتر شد. ارزیابی میزان خطرپذیری برخی اماکن همانند مراکز تبدیل پسماند، کارخانه‌های پتروشیمی و سایت‌های هسته‌ای معمولاً بسیار هزینه‌بر است. اما در اکثر موارد، اطلاعاتی از قبیل رکوردهای ثبت شده، عکس‌ها و برخی شناسه‌های فیزیکی وجود دارد که می‌توان بر اساس آن‌ها به هدف اصلی دست یافت. در برخی موارد خاص، سطح آلودگی این قبیل اماکن را می‌توان بر اساس سرنخ‌های چشمی یا حتی کاغذهای شیمیایی-واکنشی و مطالعات الکترومغناطیسی به‌دست آورد. در این راستا یک مثال ملموس در تاریخ وجود دارد: چند سال پس از ۱۹۷۰، پژوهشی برای برآورد میزان پلوتونیم (یک عنصر رادیواکتیو که در بسیاری از فرایندهای هسته‌ای مورد استفاده قرار می‌گیرد) در سطح خاک ناحیه‌ای مشخص و محصور شده در مجاورت سایت نوادا (واقع در آمریکا) انجام گرفت. نمونه‌های خاک از نقاط متفاوتی از آن ناحیه‌ی محصور جمع‌آوری شدند تا با به‌کار بردن روش‌های رادیوشیمیایی میزان تجمع عنصر پلوتونیم را در هر متر مربع سطح خاک معین کنند. اما این آزمایش بسیار گران است. در حالی که یک متغیر کمکی به نام *FIDLER* وجود دارد که با استفاده از آن می‌توان نمونه‌های خاک را به آسانی رتبه‌بندی کرد که از انجام آزمایش رادیوشیمیایی بسیار ارزاتر می‌باشد. محققانی روش نمونه‌گیری مجموعه‌ی رتبه‌دار با متغیر کمکی ذکر شده را روی داده‌های تحقیق بالا اعمال کردند و نتیجه گرفتند که روش نمونه‌گیری مجموعه‌ی رتبه‌دار کارایی بسیار بالاتری نسبت به نمونه‌گیری تصادفی دارد (برای مشاهده‌ی مثال‌های بیشتر به چن و همکاران، ۲۰۰۴ مراجعه شود).

در این مقاله ابتدا در بخش دوم به مقایسه اندازه‌ی نمونه‌ی مورد نیاز برای برآورد میانگین جامعه‌ی نرمال با استفاده از روش‌های نمونه‌گیری مجموعه‌ی رتبه‌دار (که ازین به بعد با *RSS* می‌شناسیم) و نمونه‌گیری تصادفی ساده (*SRS*) می‌پردازیم. سپس در بخش سوم چند مثال و پیشنهاد برای کاربرد روش *RSS* در مسائل پزشکی از قبیل آزمایش‌های کلینیکی و همچنین مسائل مربوط به آلودگی‌های زیست محیطی ارائه می‌دهیم.

۲ مقایسه‌ی اندازه‌ی نمونه‌ی مورد نیاز *RSS* و *SRS* در برآورد میانگین جامعه‌ی نرمال

تعریف: فرض کنید $(X_1, X_2, \dots, X_m)$ یک نمونه‌ی تصادفی باشد. این نمونه را به‌صورت چشمی یا قضاوتی مرتب کرده، سپس کوچک‌ترین آن‌ها یعنی $X_{1:m}$ را انتخاب می‌کنیم و با $X_{(1)1}$ نشان می‌دهیم و بقیه را کنار می‌گذاریم. نمونه‌ی تصادفی دیگری به اندازه‌ی m گرفته و پس از مرتب کردن آن به‌صورت چشمی یا قضاوتی، دومین کوچک‌ترین واحد آن را انتخاب می‌کنیم و با $X_{(2)2}$ نشان می‌دهیم و سپس بقیه را کنار می‌گذاریم. این کار را m بار انجام می‌دهیم و $X_{(m)m}$ را از نمونه‌ی m ام انتخاب می‌کنیم و با $X_{m,m}$ نشان می‌دهیم. در نهایت به $(X_{(1)1}, X_{(2)2}, \dots, X_{(m)m})$ یک نمونه‌ی مجموعه‌ی رتبه‌دار (*RSS*) به اندازه‌ی m گفته می‌شود. در واقع این تعریف روش نمونه‌گیری مجموعه‌ی رتبه‌دار یک حلقه‌ای را نشان می‌دهد. اگر این رویه را r بار تکرار کنیم، یک نمونه‌ی مجموعه‌ی رتبه‌دار r -حلقه‌ای به اندازه‌ی نمونه‌ی rm به‌صورت

$$\mathbf{X}_{RSS} = \{X_{(i)i,j}, i = 1, \dots, m, j = 1, \dots, r\} \quad (1.2)$$

حاصل خواهد شد. لازم به‌ذکر است که اندیس j در این‌جا بیانگر شمارنده‌ی حلقه می‌باشد.

حال یک جامعه‌ی آماری با میانگین مجهول μ و واریانس σ^2 را در نظر بگیرید. فرض کنید مایلیم μ را برآورد کنیم. در این صورت بدیهی است که

$$\bar{X}_{SRS} = \frac{1}{n} \sum_{i=1}^n X_i,$$

میانگین یک نمونه‌ی تصادفی با اندازه‌ی $n = rm$ برای μ ناریب است. میانگین یک نمونه‌ی مجموعه‌ی رتبه‌دار r -حلقه‌ای با اندازه‌ی زیرنمونه‌ی m ، یعنی

$$\bar{X}_{RSS} = \frac{1}{rm} \sum_{j=1}^r \sum_{i=1}^m X_{(i)i,j}, \quad (2.2)$$

نیز یک برآوردگر نارایب برای μ می باشد. در این صورت می توان نشان داد (تاکاهاشی و واکیموتو، ۱۹۶۸)

$$Var(\bar{X}_{RSS}) = \frac{\sigma^2}{n} - \frac{1}{nm} \sum_{i=1}^m (\mu_{i:m} - \mu)^2, \quad (3.2)$$

به طوری که $\mu_{i:m}$ بیانگر امید ریاضی i امین آماره ی ترتیبی از یک نمونه ی تصادفی ساده به اندازه ی m می باشد. بنابراین بدون در نظر گرفتن توزیع جامعه همواره داریم

$$Var(\bar{X}_{RSS}) < Var(\bar{X}_{SRS}). \quad (4.2)$$

پس $\bar{X}_{RSS}$ کارایی بیشتری نسبت به $\bar{X}_{SRS}$ دارد. در ادامه می خواهیم شدت برتری برآوردگر اول را نسبت به دومی بیابیم زمانی که توزیع جامعه ی آماری نرمال استاندارد باشد. برای این منظور، اندازه ی SRS را N در نظر بگیریم، در نتیجه $Var(\bar{X}_{SRS}) = \frac{1}{N}$. همچنین اندازه ی RSS را n در نظر بگیریم (یک حلقه ای)، در این صورت خواهیم داشت

$$Var(\bar{X}_{RSS}) = \frac{1}{n} - \frac{1}{n^2} \sum_{i=1}^n \mu_{i:n}^2. \quad (5.2)$$

آرنولد و همکاران (۲۰۰۸) مقادیر عددی $\mu_{i:n}$ را برای توزیع نرمال ارائه داده اند. بر این اساس، می توان واریانس برآوردگرهای دو طرح مد نظر را در شکل ۱ مقایسه نمود. همان طور که مشاهده می شود استفاده از RSS بسیار به صرفه تر از SRS می باشد، چرا که به عنوان مثال یک RSS به اندازه ی ۱۵ کارایی مشابهی با یک SRS با اندازه ی ۱۰۰ دارد.

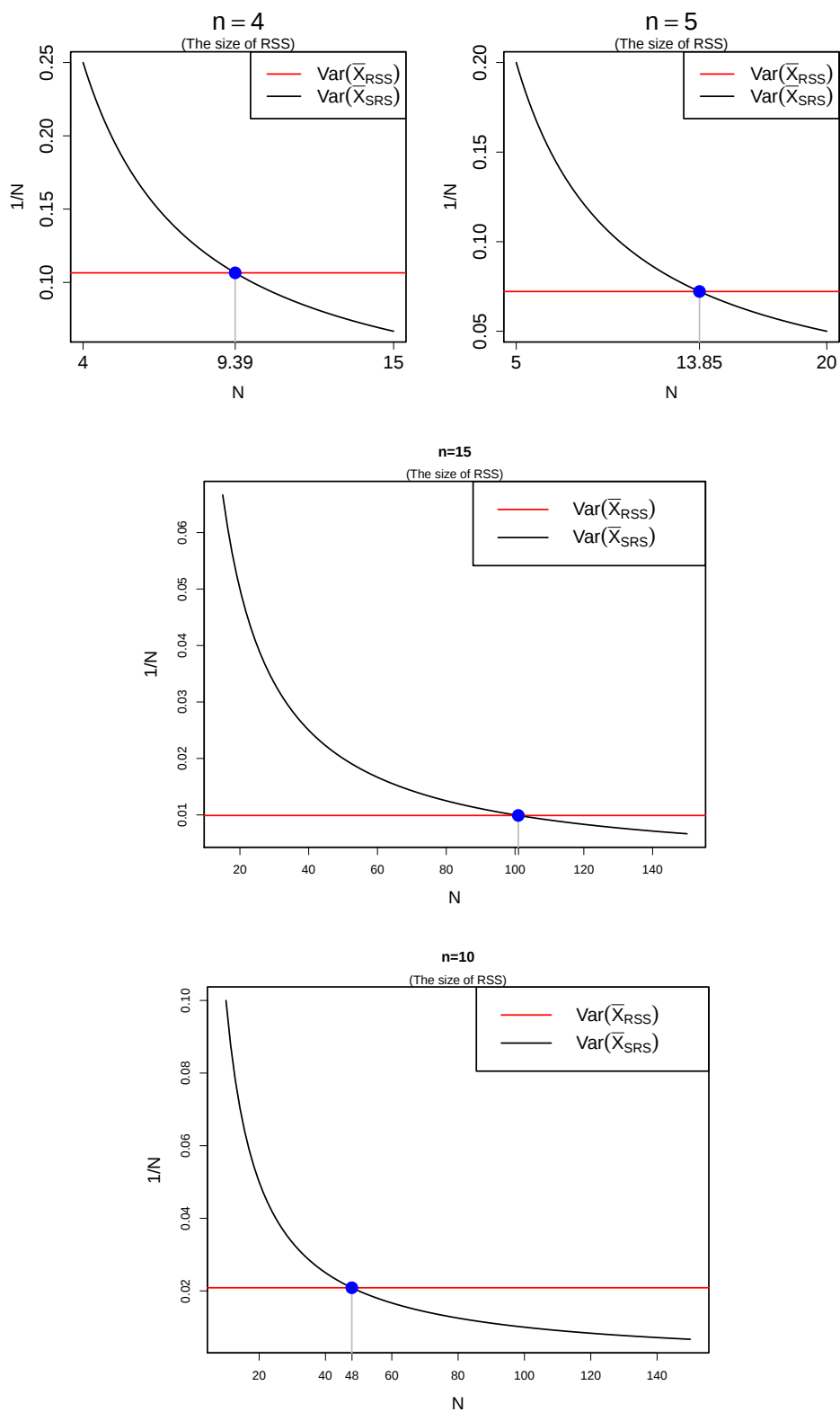
۳ کاربرد RSS در مسائل پزشکی و زیست محیطی

در این بخش چند مثال در رابطه با مسائل بالینی و همچنین زیست محیطی ارائه می شود که در همگی آن ها امکان اجرای طرح RSS وجود دارد.

تعیین سطح بیلروبین خون در نوزادان: روش استاندارد برای تعیین شدت زردی، اندازه گیری بیلی روبین سرم می باشد که مستلزم نمونه گیری از ورید محیطی است که این کار باعث تحمیل استرس و درد به نوزاد و هم باعث صرف هزینه و وقت میگردد (گلستان و همکاران، ۱۳۸۶). این در حالی است که می توان به سادگی سطح بیلروبین را بر اساس میزان زردرنگی پوست بدن نوزاد رتبه بندی کرد. در واقع بر اساس قانون کرامر هر چه میزان زردی پوست از سر به سمت نواحی پایین بدن سرایت کند، بیلروبین خون بیشتر است. آزمایش های کلینیکی: آزمایش های کلینیکی به طور معمول از هزینه ی بالایی برخوردار هستند. بنابراین برای این که نتیجه ی تحقیق در این زمینه دقیق تر و ارزان تر باشد، می توان از طرح های رتبه دار استفاده کرد. برای این منظور، شکل ظاهری (میزان چاقی)، سابقه ی فامیلی و یا سن افراد می تواند به عنوان یک متغیر کمکی برای رتبه بندی افراد و در نتیجه اعمال نمونه گیری مجموعه ی رتبه دار مورد استفاده قرار گیرد.

برآورد میزان آلودگی هسته ای در یک منطقه: وسیله ای که اندازه گیری دقیق را انجام می دهد، دزتر نام دارد. این آزمایش هزینه بر است. اما می توان نمونه های مشکوک را ابتدا بر اساس روش های آشکارساز از قبیل $IM - ۱۴۷$ ، $X50B$ ، $ANPDR27J$ (اشعه گاما)، $ND - RAD - 302$ رتبه بندی کرد، سپس با استفاده از دزتر میزان دقیق اثر را محاسبه نمود.

برآورد میزان آلودگی شیمیایی در یک منطقه: برای آگاهی از میزان آلودگی شیمیایی روش های فیزیکی گوناگونی وجود دارد که غالباً پرهزینه هستند. این روش ها عبارتند از روش های طیفسنجی تداخل امواج نور، تشخیص تغییرات pH ، قانون نفوذ گراهام و فعالیت های کاتالیکی. در حالی که لوله های مخصوصی وجود دارند که دارای نوارهای رنگی برای مشخص کردن نوع و شدت عوامل می باشند. بنابراین می توان نمونه ها را ابتدا با کمک آن ها رتبه بندی کرد سپس اندازه گیری دقیق فیزیکی را انجام داد.



شکل ۱: مقایسه‌ی اندازه‌ی نمونه‌ی مورد نیاز RSS و SRS در برآورد میانگین جامعه‌ی نرمال

برآورد میزان آلودگی میکروبی و بیولوژیکی: تشخیص آلودگی‌های میکروبی نیز نیازمند صرف هزینه‌های بسیار می‌باشد. یکی از روش‌های تشخیص میزان این نوع آلودگی‌ها، استفاده از لوله‌های GK می‌باشد که محیط کشت را برای آن میکروب فراهم می‌آورد. این در حالی است که می‌توان با استفاده از ابزار $HS - 3$ نمونه‌ها را رتبه‌بندی کرد.

۴ نتیجه‌گیری

در بخش دوم این مقاله مشاهده شد که در صورت نرمال بودن جامعه، استفاده از RSS بسیار به‌صرفه‌تر از SRS در برآورد میانگین این جامعه‌ی آماری می‌باشد. زیرا با یک میزان دقت یکسان داده‌شده، نیاز به واحدهای بیشتری از SRS نسبت به RSS داریم. به‌طور دقیق‌تر اگر (N, n) به‌ترتیب بیانگر اندازه‌های نمونه‌ی SRS و RSS باشد، آن‌گاه میانگین‌های حاصل از این دو نمونه با اندازه‌های $(104, 145)$ ، $(104, 145)$ و $(104, 145)$ و ... کارایی یکسانی خواهند داشت. در نتیجه استفاده از RSS در آزمایش‌های پزشکی و زیست‌محیطی مطرح شده در بخش سوم، که همگی هزینه‌بر هستند، به‌صرفه خواهد بود.

مراجع

- [6] Chen, Z., Bai, Z. and Sinha, B. K. (2004), Ranked Set Sampling: Theory and Applications. Springer-Verlag, New York.
- [7] McIntyre, G.A. (1952), A method for unbiased selective sampling, using ranked sets. *Australian Journal of Agricultural Research*, **3**, 385–390.
- [8] Takahasi, K. and Wakimoto, K. (1968), On unbiased estimates of the population mean based on the sample stratified by means of ordering. *Annals of the Institute of Statistical Mathematics*, **20**, 1–31.
- [9] Wolfe DA. *Ranked set sampling: its relevance and impact on statistical inference*. ISRN Probab. Stat. 2012.

کاربرد ضریب همبستگی درون گروهی (ICC) در تعیین حجم نمونه برای داده‌های خوشه‌ای و تاثیر آن بر توان مطالعه در پژوهش‌های علوم پزشکی

فرزانه برومند^۱، حبیب الله اسماعیلی^۲، محمد تقی شاکری^۳، نوشین اکبری شارک^۳، دکترای تخصصی آمار زیستی، استاد مرکز تحقیقات عوامل اجتماعی موثر بر سلامت، دانشگاه علوم پزشکی مشهد، مشهد، ایران، ^۳دانشجوی دکترای آمار زیستی، کمیته تحقیقات دانشجویی، مرکز تحقیقات عوامل اجتماعی موثر بر سلامت، دانشکده بهداشت، دانشگاه علوم پزشکی مشهد، مشهد، ایران.

چکیده: در پژوهش‌های بالینی، برای دستیابی به برآورد دقیق و قابل اعتماد حجم نمونه، لازم است در چارچوب طراحی اجرای تحقیق، آزمون آماری مناسبی برای فرضیات مورد نظر به کار گرفته شود. در عمل به موارد زیادی برمیخوریم که میان فرضیات مورد آزمون، طراحی مطالعه، روش تحلیل آماری و محاسبه تعداد نمونه هماهنگی لازم وجود ندارد. قطعاً این ناهماهنگی‌ها می‌تواند موجب مخدوش شدن اعتبار و جامعیت مورد نظر در پژوهش‌های بالینی شود. در بسیاری از مطالعات پزشکی، طراحی مطالعه منجر به وجود آمدن ساختار خاصی در داده‌ها می‌شود. این ساختار با نام ساختار خوشه‌ای شناخته می‌شود. در این شرایط اگر پژوهشگر بدون توجه به آنچه گفته شد از روش‌های معمول برای تعیین حجم نمونه استفاده کند، حجم نمونه را بیشتر از حد نیاز برآورد خواهد کرد. ضریب همبستگی درون گروهی یا ICC یک شاخص آماری است که به واسطه آن می‌توان درجه توافق و همبستگی اعضای درون یک گروه (خوشه) را تعیین کرد. مقدار ICC بر تعیین حجم نمونه برای داده‌هایی با ساختار خوشه‌ای تاثیرگذار است. حجم نمونه موثر اصطلاحی است که برای توصیف رابطه حجم نمونه درون هر خوشه در مقایسه با حجم نمونه کل استفاده می‌شود. با توجه به اندازه ICC و اثر طرح (DE) ممکن است حجم نمونه مورد نیاز برای بدست آمدن معنی داری آماری کاهش یابد. حتی مقدار کوچک ICC می‌تواند با افزایش حجم درون هر خوشه، منجر به بزرگ شدن DE و به تبع آن کاهش ESS شود. بنابراین در تعیین حجم نمونه در داده‌های خوشه‌ای از کنار ICC کوچک هم نباید به راحتی گذشت. افزایش تعداد خوشه‌ها، بیشتر از افزایش حجم درون خوشه‌ها می‌تواند به بالا بردن توان مطالعه منجر شود. بنابراین یکی از راه‌های افزایش توان مطالعه در داده‌هایی با ساختار خوشه‌ای افزایش تعداد خوشه‌هاست. توجه داریم فرمول‌های استاندارد برای تعیین حجم نمونه باید بکار گرفته شود تنها در صورت وجود ساختار خوشه‌ای در داده‌ها برای رسیدن به حجم نمونه موثر می‌توان تعدیل‌سازی لازم را انجام داد.

واژه‌های کلیدی: ضریب همبستگی درون گروهی، حجم نمونه موثر، داده‌های خوشه‌ای، توان مطالعه، کارآزمایی بالینی.

۱ مقدمه

در پژوهش‌های بالینی، برای دستیابی به برآورد دقیق و قابل اعتماد حجم نمونه، لازم است در چارچوب طراحی اجرای تحقیق، آزمون آماری مناسبی برای فرضیات مورد نظر به کار گرفته شود. فرضیه مورد آزمون باید بیانگر اهداف مورد مطالعه باشد. در عمل به موارد زیادی برمیخوریم که میان فرضیات مورد آزمون، طراحی مطالعه، روش تحلیل آماری و محاسبه تعداد نمونه هماهنگی لازم وجود ندارد. قطعا این ناهماهنگی‌ها می‌تواند موجب مخدوش شدن اعتبار و جامعیت مورد نظر در پژوهش‌های بالینی شود. (چاو و همکاران (۲۰۰۷)) در بسیاری از مطالعات پزشکی، طراحی مطالعه منجر به وجود آمدن ساختار خاصی در داده‌ها می‌شود. در این ساختار که با نام ساختار خوشه‌ای شناخته می‌شود، مجموعه‌ی داده‌ها در گروه‌ها یا در اصطلاح خوشه‌هایی قرار می‌گیرد به طوری که مشاهدات درون هر خوشه عمدتاً دارای همبستگی مثبت هستند. تحلیل‌هایی که این همبستگی را به حساب نمی‌آورند غالباً منجر به استنباط‌هایی گمراه‌کننده می‌شوند. (دیگل و همکاران (۱۹۹۴) و فیتزموریس و همکاران (۲۰۰۸)) در این شرایط اگر پژوهشگر بدون توجه به آنچه گفته شد از روش‌های معمول برای تعیین حجم نمونه استفاده کند، حجم نمونه را بیشتر از حد نیاز برآورد خواهد کرد. زیرا شباهت ایجاد شده به دلیل ساختار طراحی مطالعه بین اعضای درون یک خوشه، می‌تواند تغییر پذیری متغیر پاسخ بدست آمده از نمونه‌های مربوط به یک خوشه را کاهش می‌دهد. این تغییرپذیری باعث کاهش حجم نمونه مورد نیاز برای رسیدن به یک توان معین است. این موضوع برای پژوهش‌هایی با طرح کارآزمایی بالینی و پژوهش‌های حوزه سلامت بسیار اهمیت دارد، زیرا طراحی مطالعه در بسیاری از این پژوهش‌ها، باعث ایجاد خوشه در داده‌ها می‌شود. (دائر و کلار (۲۰۰۰)) در آمار ضریب همبستگی درون گروهی^۱ یا ICC یک شاخص آماری است که به واسطه آن می‌توان درجه توافق و همبستگی اعضای درون یک گروه (خوشه) را تعیین کرد. (دائر و کلار (۲۰۰۰)) این شاخص اولین بار توسط فیشر^۲ در سال ۱۹۵۴ به عنوان یک اصلاح برای ضریب همبستگی پیرسن معرفی شد. در حالی که ICC مدرن با استفاده از میانگین مربعات آنالیز واریانس محاسبه می‌شود. (کو و لی (۲۰۱۶)) ICC را میتوان با استفاده از نرم‌افزارهای مختلف آماری به آسانی محاسبه کرد. هدف این مقاله اهمیت استفاده از ICC در تعیین حجم نمونه برای داده‌های خوشه‌ای و تاثیر آن بر توان آماری مطالعه است.

۲ روش

همانطور که پیش‌تر در مقدمه بیان شد ICC ، اندازه‌ای از همبستگی درون گروهی یا درون خوشه‌ای را نشان می‌دهد. محاسبه این شاخص با بررسی واریانس بین خوشه و درون خوشه محاسبه می‌شود. با تقسیم واریانس بین خوشه‌ها به مجموع واریانس بین و درون خوشه‌ها مقدار ICC بدست می‌آید. (کلیپ و همکاران (۲۰۰۴))

$$ICC = \frac{S_b^2}{S_b^2 + S_w^2} \quad (۱.۲)$$

که در آن S_b^2 واریانس بین خوشه‌ای و S_w^2 واریانس درون خوشه‌ای است. مقدار ICC براساس این فرمول مقداری بین صفر و یک است. کوچک بودن S_w^2 باعث می‌شود ICC به یک نزدیک شود که نشان دهنده ی بزرگ بودن همبستگی درون خوشه‌ای است. همینطور بزرگ بودن S_w^2 باعث می‌شود ICC به صفر نزدیک شود که نشان دهنده کوچک بودن همبستگی درون خوشه‌ای است. (دائر و کلار (۲۰۰۰)، کلیپ و همکاران (۲۰۰۴))

اینکه مقدار ICC چگونه بر حجم نمونه در داده‌های خوشه‌ای تاثیرگذار است را با بیان مفهومی به نام حجم نمونه موثر^۳ توضیح خواهیم داد. حجم نمونه موثر اصطلاحی است که برای توصیف رابطه حجم نمونه درون هر خوشه در مقایسه با حجم نمونه کل استفاده

^۱Interclass correlation coefficient

^۲Fisher

^۳Effective Sample Size

^۴Design Effect

می‌شود. با توجه به اندازه ICC و اثر طرح DE ^۴ ممکن است حجم نمونه مورد نیاز برای بدست آمدن معنی داری آماری کاهش یابد. حجم نمونه موثر (ESS) از رابطه زیر بدست می‌آید.

$$ESS = \frac{m.k}{DE} \quad (2.2)$$

که در آن m حجم خوشه‌ها، k تعداد خوشه‌ها، $m.k$ حجم نمونه کل و DE اثر طرح است. اثر طرح از رابطه زیر بدست می‌آید:

$$DE = 1 + ICC(m - 1) \quad (3.2)$$

اگر $ICC = 0$ در این صورت $DE = 1$ و حجم نمونه موثر با حجم نمونه کل برابر خواهد بود. ($ESS = m.k$)
اگر $ICC > 0$ (حتی اگر ICC خیلی کوچک باشد) در این صورت مقدار DE با بزرگ بودن حجم درون خوشه (m) بزرگ می‌شود در این حالت ESS یا حجم نمونه موثر کاهش می‌یابد بنابراین به حجم نمونه کمتری برای معنی داری نیاز است. (کلیپ و همکاران (۲۰۰۴)، مورای و همکاران (۱۹۹۴)، مورای و شورت (۱۹۹۶)).
اگر $ICC = 1$ در نتیجه $DE = m$ و به تبع آن $ESS = K$ به عبارت دیگر در این شرایط تنها کافی است از هر خوشه یک نمونه اختیار شود. پیش تر هم ذکر شد در صورتی که $ICC = 1$ به این معناست که تمامی اعضای درون خوشه باهم برابرند. استفاده از این روابط در مرحله قبل آزمایش برای تعیین حجم نمونه مناسب در پژوهش‌های بالینی کار برد فراوان دارد. (کاجران (۲۰۰۷))

۳ مثال کاربردی

برای بررسی تاثیر ماساژ بر سطح بیلی روبین نوزادان نارس یک مطالعه کارآزمایی طراحی شد. در این مطالعه نوزادان نارس که سن حاملگی مادرانشان بین ۳۰ الی ۳۶ هفته بود و در بخش $NICU$ بستری بودند به طور تصادفی به دو گروه ماساژ درمانی با روغن کنجد و ماساژ درمانی با روغن متداول (کنترل) تخصیص داده شدند. چهار پرستار، ماساژ درمانی را بر عهده داشتند. تخصیص این چهار پرستار به گروه ماساژ درمانی با روغن کنجد و کنترل به صورت تصادفی انجام شد. (۲ نفر گروه کنترل ۲ نفر گروه ماساژ درمانی با روغن کنجد) تعیین حجم نمونه برای این مطالعه به طور معمول بر اساس فرمول مقایسه میانگین‌ها در دو گروه مستقل، در سطح معنی داری ۰/۰۵ و توان ۸۰ درصد و اطلاعات بدست آمده از مقالات مشابه $S_1 = 1/9$ و $S_2 = 2/34$ و $\bar{X}_1 = 11/01$ و $\bar{X}_2 = 9/5$ برای هر گروه ۳۲ بدست آمد. ۳۲ نوزاد گروه ماساژ درمانی ۳۲ نوزاد گروه کنترل که در مجموع حجم نمونه مورد نیاز ۶۴ می‌باشد.

$$n = \frac{(Z_{1-\alpha/2} + Z_{1-\beta})^2 (S_1^2 + S_2^2)}{(\bar{X}_2 - \bar{X}_1)^2} \quad (1.3)$$

تعیین حجم نمونه بدون در نظر گرفتن ساختار خوشه‌ای داده‌ها برای رسیدن به توان ۸۰ درصد بدست آمده است. حال آنکه بین نوزادانی که توسط پرستار خاص ماساژ می‌بینند همبستگی وجود دارد و این همبستگی باعث بوجود آمدن خوشه می‌شود. طراحی این مطالعه باعث بوجود آمدن چهار خوشه (به تعداد پرستارانی که نوزادان را ماساژ می‌دهند) شده است. پس از محاسبه ICC برای این مطالعه (۰/۰۲۵) و برآورد حجم نمونه موثر بر اساس فرمول دو، مشخص شد توان آزمون در این شرایط ۰/۶۹ درصد است که این یعنی با در نظر نگرفتن ساختار خوشه‌ای داده‌ها در تعیین حجم نمونه، مقدار قابل توجهی از توان آزمون از دست رفته است.

۴ تاثیر ICC و DE بر محاسبه توان و حجم نمونه

برای واضح شدن تاثیر ICC و طراحی مطالعه بر محاسبه توان و حجم نمونه، سه استراتژی متفاوت را در مورد m ، k و mk بکار می‌بریم تا تاثیر تغییرات هریک را بر توان آزمون بررسی کنیم.

۱- ثابت بودن حجم نمونه کل، mk

در این قسمت توان آزمون برای مثال کاربردی را به ازای مقادیر مختلف m و k به طوری که mk ثابت باشد، شبیه سازی می کنیم. همانطور که در جدول ۱ مشاهده می کنید مقدار کوچک ICC هم می تواند تاثیر بزرگی بر توان آزمون داشته باشد. مقدار mk برابر با عدد ۶۴، بدست آمده براساس فرمول چهار اختیار شده است. همانطور که در سطر اول جدول شماره یک مشاهده می کنید در این شرایط توان آزمون برابر ۶۹ درصد است. اما با افزایش تعداد خوشه ها و کاهش حجم درون هر خوشه می توان توان آزمون را به مقدار مطلوب ۸۰ درصد رساند. پاسخ این سوال که چرا با افزایش تعداد خوشه ها و کاهش حجم درون خوشه ها توان افزایش می یابد، پیش تر در مقدمه گفته شد، تغییر پذیری درون خوشه ها کم است بنابراین با افزایش تعداد خوشه می توان تغییر پذیری را افزایش داد از طرفی چون در این استراتژی mk مقدار ثابتی است. با افزایش تعداد خوشه ها حجم درون خوشه کاهش می یابد. در این مثال اگر بخواهیم با ۶۴ نمونه به توان مطلوب ۸۰ درصد دست یابیم در طراحی مطالعه باید هر نوزاد توسط یک پرستار ماساژ داده شود که یعنی در این مطالعه به ۶۴ پرستار نیاز است. ممکن است این روش برای افزایش توان مطالعه شدنی نباشد.

جدول ۱: حجم نمونه موثر و توان در صورت ثابت بودن mk

$tTest^* POWER$	$ICC = 0.25$		تعداد کل mk	تعداد نوزادان m	تعداد ماساژ دهندگان k
	ESS	DE			
۰/۶۹	۴۷	۱/۳۸	۶۴	۱۶	۴
۰/۷۵	۵۵	۱/۱۸	۶۴	۸	۸
۰/۷۸	۶۰	۱/۰۸	۶۴	۴	۱۶
۰/۷۹	۶۳	۱/۰۳	۶۴	۲	۳۲
۰/۸۰	۶۴	۱/۰۰	۶۴	۱	۶۴

در جدول فوق مقادیر محاسبه شده برای ESS به بالا گرد شده است و به دلیل وجود دو گروه، به ESS های فرد یک عدد اضافه شده است.

*توان محاسبه شده برای آزمون تی مستقل براساس حجم نمونه مورد نیاز (ESS) محاسبه شده است.

۲- ثابت بودن تعداد خوشه ها، k

در این قسمت توان آزمون برای مثال کاربردی را به ازای مقادیر مختلف m و mk به طوری که k ثابت باشد، شبیه سازی می کنیم. با توجه به جدول ۲ در صورت ثابت بودن تعداد خوشه ها با افزایش حجم درون هر خوشه توان مطالعه افزایش می یابد. قابل ذکر است حتی با وجود ICC کوچک در این مطالعه تاثیر قابل توجه آن بر توان آزمون قابل مشاهده است. در جدول فوق مقادیر محاسبه شده برای ESS به بالا گرد شده است و به دلیل وجود دو گروه، به ESS های فرد یک عدد اضافه شده است.

*توان محاسبه شده برای آزمون تی مستقل براساس حجم نمونه مورد نیاز (ESS) محاسبه شده است.

ثابت بودن حجم درون خوشه ها، m

در این قسمت توان آزمون برای مثال کاربردی را به ازای مقادیر مختلف k و mk به طوری که m ثابت باشد، شبیه سازی می کنیم.

کاربرد ضریب همبستگی درون گروهی (ICC) در تعیین حجم نمونه برای داده‌های خوشه‌ای و تاثیر آن بر توان مطالعه در پژوهش‌های علوم پزشکی

جدول ۲: حجم نمونه موثر و توان در صورت ثابت بودن k

$Test^* tPOWER$	$ICC = 0.25$		تعداد کل mk	تعداد نوزادان m	تعداد ماساژ دهندگان k
	ESS	DE			
۰/۵۴	۳۳	۱/۲۳	۴۰	۱۰	۴
۰/۷۵	۵۵	۱/۴۸	۸۰	۲۰	۴
۰/۸۴	۷۰	۱/۷۳	۱۲۰	۳۰	۴
۰/۸۹	۸۲	۱/۹۸	۱۶۰	۴۰	۴

همانطور که در جدول ۳ مشاهده می‌کنید افزایش تعداد خوشه منجر به افزایش توان مطالعه خواهد شد.

جدول ۳: حجم نمونه موثر و توان در صورت ثابت بودن m

$Test^* t POWER$	$ICC = 0.25$		تعداد کل mk	تعداد نوزادان m	تعداد ماساژ دهندگان k
	ESS	DE			
۰/۵۴	۳۳	۱/۲۳	۴۰	۱۰	۴
۰/۷۰	۴۹	۱/۲۳	۶۰	۱۰	۶
۰/۸۲	۶۶	۱/۲۳	۸۰	۱۰	۸
۰/۸۹	۸۲	۱/۲۳	۱۰۰	۱۰	۱۰
۰/۹۳	۹۸	۱/۲۳	۱۲۰	۱۰	۱۲

در جدول فوق مقادیر محاسبه شده برای ESS به بالا گرد شده است و به دلیل وجود دو گروه، به ESS های فرد یک عدد اضافه شده است.

*توان محاسبه شده برای آزمون تی مستقل براساس حجم نمونه مورد نیاز (ESS) محاسبه شده است.

بحث و نتیجه گیری

به طور خلاصه شاخص ICC در آمار نشان دهنده درجه همبستگی و توافق است. ICC اندازه‌ای از همبستگی مقادیر درون یک خوشه را نشان می‌دهد. حجم نمونه در پژوهش‌های بالینی معمولاً کم است با این حال اثر طرح (DE) با افزایش حجم درون هر خوشه، به مقدار قابل توجهی افزایش می‌یابد. افزایش DE حجم نمونه موثر برای رسیدن به معنی داری آماری را کاهش می‌دهد. همچنین مقدار ICC در محاسبه ESS در مطالعاتی که داده‌های آن ساختار خوشه‌ای دارند، موثر است. حتی مقدار کوچک ICC می‌تواند با افزایش m ، منجر به بزرگ شدن DE و به تبع آن کاهش ESS شود. بنابراین در تعیین حجم نمونه در داده‌های خوشه‌ای از کنار ICC کوچک هم نباید به راحتی گذشت. زیاد بودن تعداد خوشه‌ها و کم بودن اعضای درون خوشه می‌تواند به کوچک شدن DE بیانجامد. افزایش تعداد خوشه‌ها، بیشتر از افزایش حجم درون خوشه‌ها می‌تواند به بالا بردن توان مطالعه منجر شود. بنابراین یکی از راه‌های افزایش توان مطالعه در داده‌هایی با ساختار خوشه‌ای افزایش تعداد خوشه‌هاست. توجه داریم فرمول‌های استاندارد برای تعیین حجم نمونه باید بکار گرفته شود تنها در صورت وجود ساختار خوشه‌ای در داده‌ها برای رسیدن به حجم نمونه موثر می‌توان این تعدیل سازی را انجام داد.

- [10] CHOW, S.-C., WANG, H. & SHAO, J. 2007. *Sample size calculations in clinical research*, CRC press.
- [11] COCHRAN, W. G. 2007. *Sampling techniques*, John Wiley & Sons.
- [12] DIGGLE, P., LIANG, K.-Y. & ZEGER, S. L. 1994. Longitudinal data analysis. *New York: Oxford University Press*, 5, 13.
- [13] DONNER, A. & KLAR, N. 2000. Design and analysis of cluster randomization trials in health research.
- [14] FITZMAURICE, G., DAVIDIAN, M., VERBEKE, G. & MOLENBERGHS, G. 2008. *Longitudinal data analysis*, CRC Press.
- [15] KILLIP, S., MAHFOUD, Z. & PEARCE, K. 2004. What is an intraclass correlation coefficient? Crucial concepts for primary care researchers. *The Annals of Family Medicine*, 2, 204-208.
- [16] KOO, T. K. & LI, M. Y. 2016. A guideline of selecting and reporting intraclass correlation coefficients for reliability research. *Journal of chiropractic medicine*, 15, 155-163.
- [17] MURRAY, D. M., ROONEY, B. L., HANNAN, P. J., PETERSON, A. V., ARY, D. V., BIGLAN, A., BOTVIN, G. J., EVANS, R. I., FLAY, B. R. & FUTTERMAN, R. 1994. Intraclass correlation among common measures of adolescent smoking: estimates, correlates, and applications in smoking prevention studies. *American Journal of Epidemiology*, 140, 1038-1050.
- [18] MURRAY, D. M. & SHORT, B. 1996. Intraclass correlation among measures related to alcohol use by school aged adolescents: estimates, correlates and applications in intervention studies. *Journal of Drug Education*, 26, 207-230.

معادلات دیفرانسیل کسری

سوسن امیری^۱، محمد حسین رحمانی دوست گروه ریاضی، دانشگاه نیشابور، نیشابور

چکیده: در سال ۱۹۶۵ میلادی هوپیتال نامه‌ای به لایپ‌نیتز نوشت. و این سؤال را مطرح کرد که «در نمادگذاری $\frac{d^n y}{dx^n}$ ، هنگامی که $n = \frac{1}{p}$ باشد چه رخ می‌دهد؟» لایپ‌نیتز پاسخ داد: «این امر در حال حاضر پارادوکسی آشکار است که روزی به عنوان یکی از مفیدترین مباحث علمی مطرح خواهد گشت.» محاسبات دیفرانسیل و انتگرال کسری به طور گسترده در حوزه‌هایی چون ریاضیات زیستی، اکولوژی، فیزیک، شیمی و غیره استفاده می‌شود. مدل‌هایی با مشتقات مرتبه کسری، بر حسب هدف و نوع سیستم با یکدیگر متفاوت هستند. در این میان، معادلات دیفرانسیل کسری^۱ به طور گسترده در مدل‌های بیولوژی استفاده می‌شود. این دو معادله از جنبه‌های مختلفی همچون ارائه راه‌حل‌های معادلات دیفرانسیل کسری و نیز ویژگی‌های رفتار مختلف راه‌حل‌ها، بررسی شده‌اند. یکی از مهم‌ترین کاربردهای این معادلات در فیزیک در شبیه‌سازی حرکت و شبیه‌سازی سیستم‌های بیولوژیکی می‌باشد.

واژه‌های کلیدی: مشتق کسری گرانولد-لنتیکوف، مشتق کسری ریمان-لیوویل، مشتق کسری کاپتوی

۱ مقدمه

محاسبات دیفرانسیل و انتگرال کسری، شاخه‌ای از محاسبات کلاسیک بوده که مرتبه مشتقات و انتگرال را به یک مرتبه غیر عددی صحیح تعمیم می‌دهد. این نوع محاسبات در ابتدا بیش از ۳۰۰ سال پیش از سوی ال هوپیتال و لایپ‌نیتز به صورت مشتقاتی با مرتبه $\alpha = \frac{1}{p}$ ارائه شدند. در سال‌های اخیر معادلات دیفرانسیل کسری به دلیل پیدایش مکرر در مکانیسم‌های جریان، بیولوژی، الکتروشیمی و مسائل فنی و فیزیکی دیگر، محققین بسیاری را بر آن داشته است تا در این زمینه کارهای زیادی انجام دهند. بسیاری از پژوهشگران از معادلات دیفرانسیل کسری برای توسعه مدل‌های خود در بیولوژی از جمله تعاملات بین شکار و شکارچی، فارماکوکینتیک (بخشی از داروسازی که با جذب و پخش دارو در بدن سر و کار دارد)، عفونت *HIV* و غیره استفاده کرده‌اند. امروزه روش‌های زیادی برای حل این نوع مسائل ابداع شده‌اند و افراد زیادی این روش‌ها را به کار گرفته‌اند تا بتوانند جواب واقعی این مسائل را بیابند.

^۱Fractional differential equation (FDE)

۱.۱ مشتق کسری گرانولد-لتنیکوف

تابع پیوسته $f(t) \in C[a, b]$ را به صورت زیر در نظر می‌گیریم. مشتق مرتبه اول تابع $f(t)$ چنین است:

$$f'(t) = \frac{df}{dt} = \lim_{h \rightarrow 0} \frac{f(t) - f(t-h)}{h}, \quad (۱.۱)$$

$$f''(t) = \frac{d^2 f}{dt^2} = \lim_{h \rightarrow 0} \frac{f(t) - 2f(t-h) + f(t-2h)}{h^2}, \quad (۲.۱)$$

$$f'''(t) = \frac{d^3 f}{dt^3} = \lim_{h \rightarrow 0} \frac{f(t) - 3f(t-h) + 3f(t-2h) - f(t-3h)}{h^3}, \quad (۳.۱)$$

$$f^n(t) = \frac{d^n f}{dt^n} = \lim_{h \rightarrow 0} \frac{1}{h^n} \sum_{r=0}^n (-1)^{n-r} \binom{n}{r} f(t - (n-r)h) \quad (۴.۱)$$

رابطه (۴.۱) را می‌توان به صورت زیر بازنویسی کرد:

$$f^n(t) = \lim_{h \rightarrow 0} \frac{1}{h^n} \sum_{r=0}^n (-1)^r \binom{n}{r} f(t - rh). \quad (۵.۱)$$

اکنون تعمیم روابط (۱.۱) و (۵.۱) به صورت زیر خواهد بود:

$$f_h^p(t) = \frac{1}{h^p} \sum_{r=0}^n (-1)^r \binom{p}{r} f(t - rh),$$

$${}_a^C D_t^p f(t) = \sum_{k=0}^m \frac{f^{(k)}(a)(t-a)^{-p+k}}{\Gamma(-p+k+1)} + \frac{1}{\Gamma(-p+m+1)} \int_a^x (t-\tau)^{m-p} \times f^{(m+1)}(\tau) d\tau.$$

۲.۱ مشتق کسری ریمان-لیوویل

مشتق کسری ریمان-لیوویل به صورت زیر تعریف می‌شود:

$${}_a D_t^p f(t) = \left(\frac{d}{dt}\right)^{m+1} \frac{1}{\Gamma(m-p+1)} \int_a^t (t-\tau)^{m-p} f(\tau) d\tau, \quad m \leq p < m+1.$$

۳.۱ مشتق کسری کاپتوی

فرض کنید $f \in A^k[t, T]$ و $k \in \mathbb{N}$ باشد که $A^k[t, T]$ مجموعه‌ای از تمام توابع $f : [t, T] \rightarrow \mathbb{R}$ بوده که $f^{(k-1)}$ پیوسته است و $t \in [t, T]$. در این صورت برای هر $k-1 < \alpha < k$ مشتق مرتبه کسری کاپتوی به صورت زیر تعریف می‌شود:

$${}_t^C D_t^\alpha f(t) = \frac{1}{\Gamma(k-\alpha)} \int_t^t (t-\tau)^{k-\alpha-1} f^{(k)}(\tau) d\tau,$$

که $\Gamma(\cdot)$ اشاره به تابع گاما دارد.

۲ نتایج اصلی

لم ۱.۲. اگر $f = (f_1, \dots, f_n) \in C^1$ ، آنگاه مساله مقدار اولیه

$$D^{\alpha_i} y_i(t) = f_i(t, y_1, \dots, y_n), \quad y_i^{(k)} = C_k^i, \quad 1 \leq i \leq n, \quad 1 \leq k \leq m_i$$

معادل معادله انتگرال ولترای زیر است:

$$y_i(t) = \sum_{k=1}^{m_i} C_k^i \frac{t^k}{k!} + I^{\alpha_i} f_i(t, y_1, \dots, y_n), \quad 1 \leq i \leq n. \quad (1.2)$$

□

اثبات. به (۱۹) مراجعه شود.

قضیه ۲.۲. فرض کنید

$$f = (f_1, \dots, f_n) : w \rightarrow \mathbb{R}^n \in C^1$$

که در آن

$$w = [\cdot, \chi^*] \times \prod_{j=1}^n [y_j(\cdot) - l_j, y_j(\cdot) + l_j], \quad \chi^* > \cdot, \quad l_j > \cdot, \quad \forall j$$

در این صورت دستگاه معادلات

$$D^{\alpha_i} y_i(t) = f_i(t, y_1, \dots, y_n), \quad y_i^{(k)} = C_k^i, \quad 1 \leq i \leq n, \quad 1 \leq k \leq m_i \quad (2.2)$$

که $1 < \alpha_i < m_i + 1$ و D^{α_i} یک جواب یکتا به صورت $\bar{y} : [\cdot, \chi] \rightarrow \mathbb{R}$ دارد که

$$\chi = \min \left\{ \chi^*, \left(\frac{l \Gamma(\alpha_i + 1)}{[1 + \alpha] \|f\|} \right)^{\frac{1}{\alpha_i}}, \left(\frac{lk!}{[1 + \alpha] |C_k^i|} \right)^{\frac{1}{k}} \right\}, \quad i = 1, 2, \dots, n$$

$$k = 1, 2, \dots, m_i, \quad l = \min\{l_1, l_2, \dots, l_n\}, \quad \alpha = \max\{\alpha_1, \alpha_2, \dots, \alpha_n\}.$$

لازم به ذکر است که $[1 + \alpha]$ نشان دهنده جزء صحیح $1 + \alpha$ می باشد. اثبات. با توجه به لم (۱.۲) معادله (۲.۲) معادل است با

$$y - i(t) = \sum_{k=1}^{m_i} C_k^i \frac{t^k}{k!} + I^{\alpha_i} f_i(t, \bar{y}), \quad 1 \leq i \leq n.$$

قرار می دهیم $\bar{A} = (A_1(\bar{y}), \dots, A_n(\bar{y}))$ که

$$A_i[\bar{y}(t)] = \sum_{k=1}^{m_i} C_k^i \frac{t^k}{k!} + \frac{1}{\Gamma(\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} f(s, \bar{y}(s)) ds. \quad (3.2)$$

ناگفته نماند که یک نقطه ثابت از عملگر A یک جواب معادله (۲.۲) است. چون $f \in C^1$ ، به طور موضعی با یک همسایگی به شعاع l و با ثابت لیبشیتز است.

فرض کنید

$$U = \{\bar{y} \in B : |y_i - y_i(\cdot)| \leq l, \quad 1 \leq i \leq n, \quad \chi < l\}.$$

به وضوح U یک زیر مجموعه بسته و محدب و غیر تهی از $C([\bullet, \chi]^n)$ که f لپشیتز است. چون f_i ($i = 1, 2, \dots$) روی مجموعه فشرده w پیوسته بوده، بنابراین روی w پیوسته یکنواخت خواهد بود. تعریف پیوستگی یکنواخت را می نویسیم:

$$\forall \varepsilon > \bullet, \exists \delta > \bullet, |\bar{y} - \bar{z}| < \delta, |f_i(t, \bar{y}) - f_i(t, \bar{z})| < \varepsilon^*$$

$$\varepsilon^* = \min\left\{\frac{\varepsilon \Gamma(\alpha_1 + 1)}{\chi^{\alpha_1}}, \dots, \frac{\varepsilon \Gamma(\alpha_n + 1)}{\chi^{\alpha_n}}\right\}.$$

اکنون فرض می کنیم $\bar{y}, \bar{z} \in U$ به طوری که $|\bar{y} - \bar{z}| < \delta$. بنابراین با توجه به نامساوی

$$|f_i(t, \bar{y}) - f_i(t, \bar{z})| < \varepsilon^* \quad (4.2)$$

داریم

$$\begin{aligned} |A_i \bar{y}(t) - z_i \bar{y}(t)| &= \left| \frac{1}{\Gamma(\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} [f_i(s, \bar{y}(s)) - f_i(s, \bar{z}(s))] ds \right| \\ &< \varepsilon^* \left| \frac{1}{\Gamma(\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} ds \right| \leq \varepsilon^* \frac{\chi^{\alpha_i}}{\Gamma(\alpha_i + 1)} \leq \varepsilon, \quad \forall i \end{aligned}$$

با توجه به پیوستگی A_i ها، A نیز پیوسته خواهد بود.

به علاوه برای $\bar{y} \in U$ و $t \in [\bullet, \chi]$ داریم

$$\begin{aligned} |A_i \bar{y}(t) - C_i^i| &= \left| \sum_{k=1}^{m_i} C_k^i \frac{t^k}{k!} + \frac{1}{\Gamma(\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} f(s, \bar{y}(s)) ds \right| \\ &\leq \left| \sum_{k=1}^{m_i} C_k^i \frac{t^k}{k!} \right| + \frac{\|f\| \chi^{\alpha_i}}{\Gamma(\alpha_i + 1)} \leq l, \quad \forall i \end{aligned}$$

بنابراین نشان داده شد که A نگاشتی از مجموعه U به خودش است. اکنون با استقرا ثابت می کنیم که برای هر $\bar{y}, \bar{z} \in U$

$$|A_i^n \bar{y}(s) - A_i^n \bar{z}(s)| \leq \frac{(Ls^{\alpha_i})^n}{\Gamma(1 + n\alpha_i)} |\bar{y} - \bar{z}|, \quad n = \bullet, 1, 2, \dots$$

اگر $n = \bullet$ بدیهی است که تساوی بالا برقرار است. فرض کنیم تساوی در گام $n - 1$ برقرار باشد. ثابت می شود که در گام n نیز برقرار است.

$$\begin{aligned} |A_i^n \bar{y}(t) - A_i^n \bar{z}(t)| &= |A_i(A_i^{n-1} \bar{y}(t)) - A_i(A_i^{n-1} \bar{z}(t))| \\ &= \frac{1}{\Gamma(\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} [f_i(s, A_i^{n-1} \bar{y}(s)) - f_i(s, A_i^{n-1} \bar{z}(s))] ds. \end{aligned}$$

با توجه به این که f روی U لپشیتز است، از استقرا به دست می آید:

$$\begin{aligned} |A_i^n \bar{y}(t) - A_i^n \bar{z}(t)| &\leq \frac{L}{\Gamma(\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} \|A_i^{n-1} \bar{y}(s) - A_i^{n-1} \bar{z}(s)\| ds \\ &\leq \frac{L^n}{\Gamma(\alpha_i) \Gamma(1 + (n-1)\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} s^{\alpha_i(n-1)} |\bar{y} - \bar{z}| ds \\ &\leq \frac{L^n |\bar{y} - \bar{z}|}{\Gamma(\alpha_i) \Gamma(1 + (n-1)\alpha_i)} \int_0^t (t-s)^{\alpha_i-1} s^{\alpha_i(n-1)} ds \\ &= \frac{L^n |\bar{y} - \bar{z}| I^{\alpha_i} t^{\alpha_i(n-1)}}{\Gamma(1 + (n-1)\alpha_i)} = \frac{(Lt^{\alpha_i})^n |\bar{y} - \bar{z}|}{\Gamma(1 + n\alpha_i)} \leq \frac{L^n |\bar{y} - \bar{z}|}{\Gamma(1 + n\alpha_i)} \chi^{n\alpha_i}. \end{aligned}$$

در نتیجه

$$|A_i^n \bar{y}(t) - A_i^n \bar{z}(t)| \leq \frac{(L\chi^{\alpha_i})^n}{\Gamma(1 + n\alpha_i)} |\bar{y} - \bar{z}|.$$

با در نظر گرفتن

$$\beta_n = \frac{(L\chi^{\alpha_i})^n}{\Gamma(1 + n\alpha_i)}$$

و به منظور به کار بردن قضیه نقطه ثابت باناخ تعمیم یافته نیاز به همگرایی داریم:

$$\sum_{n=0}^{\infty} \frac{(L\chi^{\alpha_i})^n}{\Gamma(1 + n\alpha_i)} = E_{\alpha_i}(L\chi^{\alpha_i})$$

که یک تابع میتاگ-لفلر از مرتبه α_i روی $L\chi^{\alpha_i}$ است. بنابراین با توجه به قضیه نقطه ثابت باناخ تعمیم یافته، A_i به یک نقطه ثابت همگرا بوده و یکتاست. $\square$

۳ نتیجه گیری

در این مقاله، وجود و یکتایی جواب‌های دستگاه معادلات دیفرانسیل کسری غیر خطی

$$D^{\alpha_i} y_i(t) = f_i(t, y_1, \dots, y_n), \quad y_i^{(k)} = C_k^i, \quad 1 \leq i \leq n, \quad 1 \leq k \leq m_i$$

که D^{α_i} ، $m_i < \alpha_i < m_i + 1$ بیان کننده مشتق کسری کاپتو می‌باشد، بررسی شده است.

مراجع

- [۱۹] ح. جعفری، مقدمه‌ای بر معادلات دیفرانسیل کسری، چاپ اول، دانشگاه مازندران، بابل، ۱۳۹۲.
- [20] A. D. Obembe, H. Y. Al-Yousef, M. E. Hossain, S. A. Abu-Khamsin, Fractional derivatives and their applications in reservoir engineering problems, *a review. J. Pet. Sci. Eng.*, **157**, 312-327 (2017).
- [21] F. Mainardi, *Fractional Calculus and Waves in Linear Viscoelasticity: An Introduction to Mathematical Models*, Imperial College Press, London (2010).
- [22] R. L. Magin, *Fractional Calculus in Bioengineering*, Begell House, West Redding (2006).

مقایسه‌های چندگانه در کارآزمایی‌های بالینی ناپست‌تری

سعیده حاجبی خانیکی^۱، دانشجوی دکتری آمار زیستی، مرکز تحقیقات عوامل اجتماعی موثر بر سلامت، دانشگاه علوم پزشکی مشهد، مشهد، ایران، دکتر محمدتقی شاکری استاد آمار زیستی، مرکز تحقیقات عوامل اجتماعی موثر بر سلامت، دانشگاه علوم پزشکی مشهد، مشهد، ایران، احسان برادران سیرجانی کارشناس ارشد آمار ریاضی، دانشگاه آزاد اسلامی واحد مشهد، مشهد، ایران، نگار سنگ سفیدی دانشجوی کارشناسی ارشد آمار زیستی، مرکز تحقیقات عوامل اجتماعی موثر بر سلامت، دانشگاه علوم پزشکی مشهد، مشهد، ایران

چکیده: در سال‌های اخیر به علت پیشرفت‌های پزشکی و دارویی نیاز به استفاده از کارآزمایی‌های بالینی ناپست‌تری برای مقایسه‌ی درمان‌های جدید کم هزینه تر با درمان استاندارد قبلی بیشتر احساس می‌شود. در این مقاله روش‌های مقایسه چندگانه درمان استاندارد با درمان‌های جدید برای ارزیابی اثربخشی آنها مورد بررسی قرار گرفته است. اگر ارزیابی درمان‌ها براساس یک یا چند متغیر پاسخ باشد، روش‌های پیشنهادی متفاوت است. در این مقاله آزمون‌های اجتماع-اشتراک، اشتراک-اجتماع، رویه هلم، رویه هوچبرگ، روش آزمون بسته و اصلاحات آنها به عنوان روش‌هایی که خطای نوع اول کلی خانواده‌ی مقایسه‌ها را در کارآزمایی بالینی ناپست‌تری کنترل می‌کنند، پیشنهاد شده‌اند.

واژه‌های کلیدی: کارآزمایی بالینی ناپست‌تری، مقایسه‌های چندگانه، رویه هلم، رویه هوچبرگ، روش آزمون بسته.

۱ مقدمه

کارآزمایی‌های بالینی با گروه کنترل فعال (درمان استاندارد) که در آن هدف این است تا نشان دهیم درمان جدید به طور غیرقابل قبولی بدتر از درمان استاندارد نیست را کارآزمایی‌های "ناپست‌تری"^۱ می‌نامند (۲۳). امروزه نیاز به استفاده از کارآزمایی‌های بالینی ناپست‌تری در تحقیقات پزشکی و داروسازی به دلیل ماهیت طرح و امکان جایگزینی درمان‌های کم هزینه تر با اثرات جانبی کمتر بیشتر از گذشته احساس می‌شود.

اساس کار کارآزمایی‌های بالینی ناپست‌تری به این صورت است که اگر حدپایین فاصله اطمینان برای تفاوت متغیر پاسخ بین

^۱Non-inferiority

^۱سعیده حاجبی خانیکی: hajebis971@mums.ac.ir

گروه‌های مورد مقایسه به طور معنی داری کمتر از یک مقدار از پیش تعیین شده (δ) نباشد، آنگاه با خطای نوع اول α می‌توان اظهار داشت که درمان جدید به طور غیر قابل قبولی بدتر از درمان استاندارد نیست یا بعبارت دیگر دلیلی بر رد فرض ناپست‌تری نخواهیم داشت. مقدار δ یا میزان تفاوت قابل قبول بین درمان استاندارد و درمان جدید نیز از مطالعات مشابه یا متا آنالیزهای انجام شده بدست می‌آید (۲۴).

مقایسه‌های چندگانه در هر کارآزمایی بالینی با مشکلاتی مواجه است. این مشکلات جنبه‌های زیادی دارد، از جمله وجود چندین درمان، چندین متغیر پاسخ برای اندازه‌گیری اثربخشی درمان و چندین بازه زمانی. اما همه منجر به یک مشکل مشابه می‌شوند: عدم به کارگیری روش درست برای مقایسه چندگانه، احتمال اینکه به اشتباه کارآیی یک دارو یا درمان را بپذیریم، را افزایش می‌دهد. در کارآزمایی‌های ناپست‌تری نقش فرضیه‌های صفر و جایگزین به نوعی معکوس می‌شود. در آزمون‌های برتری،^۲ که هدف تایید برتری یک درمان جدید است، خطای نوع اول احتمال نتیجه برتری یک درمان زمانیکه اثرات درمان یکسان است، می‌باشد. در این مطالعات، خطای نوع اول احتمال نتیجه‌گیری ناپست‌تری است وقتی که کنترل فعال به طور قابل توجهی بهتر از درمان جدید است. کنترل میزان خطای نوع I می‌تواند به روش‌های مختلفی تعریف شود. معروف ترین روش برای کارآزمایی‌های بالینی، کنترل خطای خانوادگی^۳ (FWE) است که احتمال رد حداقل یک فرضیه درست صفر است. در مورد آزمون چندین متغیر پاسخ در کارآزمایی‌های ناپست‌تری، FWE به معنای احتمال نتیجه‌گیری حداقل یک ناپست‌تری در یک متغیر پاسخ است زمانی که ناپست‌تری واقعاً درست نباشد. در بخش ۲ به بررسی روش‌های مختلف برای مقایسه درمان استاندارد با چندین درمان در حالتی که فقط یک متغیر پاسخ داریم و در بخش ۳ در حالتی که چندین متغیر پاسخ داریم، پرداخته می‌شود.

۲ مقایسه چندین درمان با درمان استاندارد

تیمارهای آزمایشی چندگانه را می‌توان با یک درمان استاندارد با کارایی شناخته شده مقایسه کرد. مثلاً هنگامی که دوزهای متعدد از یک دارو در نظر گرفته می‌شود یا زمانی که درمان‌های متعدد تجربی با یک درمان استاندارد برای مقایسه کارآیی آن‌ها طرح ریزی می‌شود. (۲۵) در هر مورد، خانواده‌ی فرضیاتی که برای آن نرخ خطای نوع I باید کنترل شود، بایستی در نظر گرفته شده و برای کنترل FWE تعدیل مناسب انجام شود. دو ساختار کلی موجود به صورت زیر است:

۱. پیدا کردن زیرمجموعه‌ای از تیمارها که ناپست‌تر از کنترل نیستند. در این حالت فرضیه‌ها به این صورت نوشته می‌شوند:

$$\begin{aligned} H_0: \mu_C - \mu_{E,i} &> \delta_i && \text{برای تمام } i\text{ها} \\ H_a: \mu_C - \mu_{E,i} &\leq \delta_i && \text{برای حداقل یک } i \end{aligned}$$

که در آن μ_C میانگین متغیر پاسخ در گروه کنترل و $\mu_{E,i}$ میانگین متغیر پاسخ در گروه درمانی i ام و δ_i تفاوت قابل قبول بین آن‌هاست.

۲. نشان دادن اینکه تمام تیمارها ناپست‌تر از گروه کنترل نیستند. در این حالت فرضیه‌ها به این صورت نوشته می‌شوند:

$$\begin{aligned} H_0: \mu_C - \mu_{E,i} &> \delta_i && \text{برای حداقل یک } i \\ H_a: \mu_C - \mu_{E,i} &\leq \delta_i && \text{برای تمام } i\text{ها} \end{aligned}$$

در مورد اول آزمون اجتماع-اشتراک (intersection-union) و رویه‌های مقایسه‌های چندگانه استاندارد مورد نظر قرار می‌گیرند.

²Superiority

³Familywise error

در مورد دوم آزمون اشتراک-اجتماع (union-intersection) نتیجه شده و غالباً هیچ تعدیلی برای خطای نوع I اسمی نیاز نیست. (۲۶))

۱.۲ تیمارهای غیرترتیبی : انتخاب زیرمجموعه

زمانی که هدف مقایسه‌ی چندگانه‌ی تیمارهای آزمایشی (درمان های مختلف) باشد که ترتیب در تیمارها در آزمون اجتماع-اشتراک اهمیت نداشته باشد، راههای زیادی برای کنترل FWE وجود دارد.

۱. یک راه این است که مقایسه‌های مختلف تیمارها با درمان استاندارد را مانند مقایسه‌های غیرمرتبط در نظر گرفته و هیچ تصحیحی انجام ندهیم. اما وقتی که یک گروه کنترل برای مقایسه با چندین تیمار آزمایشی در کارآزمایی تاییدی^۴ استفاده می‌شود این روش نامناسب است زیرا فقط CWE ^۵ را کنترل کرده و FWE کنترل نخواهد شد. یعنی تصحیحی برای مقایسه‌ها به صورت چندگانه اعمال نمیکند.

۲. روش دیگر استفاده از تصحیح بن-فرونی می‌باشد که در آن هر یک از I تیمار آزمایشی با درمان استاندارد در سطح $\frac{\alpha}{I}$ مقایسه می‌شود. در اینجا فرضیه ناپست‌تری زمانی نتیجه می‌شود که حد بالایی فاصله اطمینان در سطح یاد شده کمتر از δ_i باشد. این روش خطای نوع اول را کنترل کرده و فواصل اطمینان همزمان تولید می‌کند.

۳. رویه‌های قوی تری مانند رویه هلم^۶ هم وجود دارد. این رویه نیز FWE را کنترل کرده و بعلاوه بطور یکنواختی قدرتمند تر از رویه بن-فرونی است. کارکرد این رویه به این صورت است که پی-مقدارهای مقایسه‌های مختلف را بدست آورده و آنها را مرتب می‌کنیم. سپس کوچکترین پی-مقدار $p(1)$ با $\frac{\alpha}{I}$ مقایسه می‌شود. اگر $p(1) < \frac{\alpha}{I}$ فرضیه مربوط به تیمار و مقایسه‌ی نظیر آن رد شده و کوچکترین پی-مقدار بعدی با $\frac{\alpha}{I-1}$ مقایسه شده و این روند ادامه پیدا می‌کند. اگر در یکی از مراحل رابطه $(p(i) < \frac{\alpha}{I-i+1})$ برقرار نباشد، آنگاه رویه متوقف شده و نتیجه‌ای در مورد سایر مقایسه‌های آزمون نشده نمی‌توان گرفت. از آنجایی که در کارآزمایی‌های ناپست‌تری نتیجه‌گیری بر اساس فواصل اطمینان صورت می‌گیرد، باید با روشی مشابه کار کرد. به این صورت که برپایه‌ی حفظ خطای نوع اول یکطرفه‌ی $\frac{\alpha}{I}$ اگر ناپست‌تری بر اساس حداقل یکی از فواصل اطمینان دو طرفه‌ی $[1 - \alpha/I] \times 100$ نتیجه گرفته شود، رویه به مرحله دوم می‌رود. در مرحله دوم اگر ناپست‌تری بر اساس حداقل دوتا از فواصل اطمینان دو طرفه‌ی $[(1 - \alpha/(I-1)) \times 100]$ نتیجه گرفته شود، رویه به مرحله‌ی بعدی می‌رود و به همین ترتیب مراحل طی می‌شود. هر زمان این رویه متوقف شود (بعبارت بهتر ناپست‌تری برای حداقل j تا از I مقایسه برقرار نباشد) آنگاه فقط برای $1 - j$ مقایسه از مراحل قبل نتیجه‌ی نهایی گرفته می‌شود (۲۷).

۴. یکی از روش‌های معمول دیگر هم برای مواقعی که هدف مقایسه‌ی بیش از یک گروه آزمایشی با کنترل فعال باشد، استفاده از آزمون دانت^۷ است. این آزمون فاصله اطمینانی را به فرم $\bar{x}_i - \bar{x}_c \pm d_\alpha(I-1, f) \times se$ تولید می‌کند که در آن f درجه آزادی و برابر با تعداد کل مشاهدات منهای تعداد گروههای آزمایشی و I تعداد گروههای آزمایشی و کنترل است. (۲۵)

۲.۲ تیمارهای ترتیبی : انتخاب زیرمجموعه

زمانیکه ترتیب خاصی بین تیمارها وجود داشته باشد، باید جهت کارایی بیشتر، ساختارهای بیشتری به مقایسه‌ها اضافه شود. مثلاً دوزهای چندگانه‌ی یک تیمار آزمایشی کارآیی متفاوتی دارند به این صورت که دوزهای بیشتر کارآیی بالاتری دارد. اگر هدف مقایسه‌ی دوزهای

⁴Confirmatory trial

⁵Comparison-wise error

⁶Holm procedure

⁷Dunnett's test

مختلف با درمان استاندارد جهت تعیین ناپست‌تر بودن تیمارها نسبت به آن باشد، روش دنباله‌ای ثابت^۸ را می‌توان استفاده کرد. به این صورت که بالاترین دوز (بالاترین کارآیی) با درمان استاندارد مقایسه می‌شود، اگر ناپست‌تری نتیجه گرفته شود، دوز بالایی بعدی مقایسه می‌شود و به همین ترتیب ادامه پیدا می‌کند تا زمانی که یک دوز ناپست‌تر از کنترل نباشد. در این روش تا زمانی که مراحل ادامه پیدا میکند می‌توان مقایسه‌ها را در همان سطح α انجام داد. یکی از اصلاحاتی که بر روش دنباله‌ای ثابت صورت گرفته، رویه‌ی فال-بک^۹ می‌باشد. در این روش زمانیکه مقایسه‌های اولیه منجر به رد فرض صفر نشود، مقداری از نرخ خطای نوع I برای مقایسه‌های بعدی نگه داشته می‌شود. (۲۸)

۳ ناپست‌تری روی چندین متغیر پاسخ

مقایسه تیمار آزمایشی و درمان استاندارد برپایه‌ی چندین متغیر پاسخ معمولاً در مطالعات کارآزمایی بالینی بسیار مورد توجه است. در واقع برای افزایش کارآیی، بسیاری از اوقات چندین متغیر جهت تصمیم‌گیری نهایی جمع‌آوری شده و در استنباطها مورد استفاده قرار می‌گیرد. اما مسأله‌ی مهمی که باید در نظر گرفته شود حفظ سطح کل FWE می‌باشد. گام اول در تحلیل، ایجاد اولویت بین مقایسه‌ها است. مهمترین متغیرها که باید تعدادشان نیز بسیار کم باشد را مقایسه‌های اولیه، گروه بعدی مقایسه‌های ثانویه و به همین ترتیب تمام گروههای ممکن از متغیرها برای مقایسه را تشکیل می‌دهیم.

۱.۳ چندین برآمد در یک خانواده^{۱۰}

ابتدا با حالتی شروع می‌کنیم که فقط یک خانواده از متغیرهای پاسخ مهم یعنی متغیرهای پاسخ گروه اولیه را مبنای تصمیم‌گیری قرار دهیم.

۱. اگر بین پاسخ‌ها یک ترتیب معنادار از لحاظ اهمیت متغیرها وجود داشته باشد، آنگاه همان روش رویه دنباله ثابت را می‌توان مورد استفاده قرار داد. با این رویه، ابتدا ناپست‌تری را برای مهمترین متغیر بررسی کرده و متغیرهای دیگر لحاظ نمی‌شوند مگر اینکه ناپست‌تری برای متغیر قبلی ثابت شود.

۲. اصلاح روش دنباله ثابت همانطور که اشاره شد روش فال-بک است. با این روش حتی اگر یک یا چند متغیر پاسخ در نتایج استفاده نشود، تمام مقایسه‌ها را می‌توان لحاظ و محاسبه کرد. (۲۹)

۳. روش هلم

۴. هوچبرگ^{۱۱} نیز روشی را برپایه روش هلم ارائه کرده که در آن بجای رویکرد رو به عقب از روشی رو به جلو استفاده می‌شود. به عبارت ساده تر ناپست‌تری برای متغیر پاسخ نظیر بزرگترین پی-مقدار که در رابطه‌ی $p(j) \leq \frac{\alpha}{(I-j+1)}$ صدق کند و تمام متغیرهای پاسخ با پی-مقدارهای کوچکتر از آن پذیرفته می‌شود. این تعریف شامل تمام متغیرهایی خواهد بود که در روش هلم ناپست‌تری را نتیجه می‌دهند. (۳۰)

۵. روش دیگری که FWE را کنترل می‌کند روش آزمون بسته^{۱۲} نام دارد. روش آزمون بسته تمام زیر مجموعه‌های ممکن از

⁸Fixed sequence approach

⁹Fall back procedure

¹⁰Single family

¹¹Hochberg

¹²Closed testing procedure

¹³global null hypothesis

فرضیه‌ها را در نظر می‌گیرد. در اینجا با J متغیر پاسخ، $1 - 2^j = \sum_{k=1}^j \binom{j}{k}$ زیرمجموعه خواهیم داشت که درون هر زیر مجموعه، یک فرضیه صفر فراموضعی^{۱۳} به این صورت در نظر گرفته می‌شود:

فرضیه صفر: ناپست‌تری در هیچ کدام از متغیرهای پاسخ زیر مجموعه‌ی مورد بررسی وجود ندارد.

فرضیه جایگزین: ناپست‌تری در حداقل یکی از متغیرهای پاسخ زیر مجموعه‌ی مورد بررسی وجود دارد.

و نهایتاً ناپست‌تری در مورد یک برآمد خاص در صورتی نتیجه گرفته می‌شود که آن متغیر در تمام زیرمجموعه‌هایی که در بردارنده‌ی آن هستند، نتیجه گرفته شود (۳۱)

توجه ۱.۳. در استفاده از روش آزمون بسته مطلوب است که آزمون‌ها α -فرسا^{۱۴} باشد. یعنی در آزمون فرضیه صفر در هر زیرمجموعه، خطای نوع اول دقیقاً مقدار α بوده و کمتر از آن نباشد. بعنوان مثال در روش بن-فرونی این مساله لحاظ نمی‌شود. زیرا خطای نوع اول در یک زیرمجموعه‌ی مورد نظر دارای کران بالای $\sum_{k=1}^I \alpha_k I(k)$ می‌باشد که $I(k)$ فقط زمانی مقدار ۱ را می‌پذیرد که متغیر پاسخ I در زیر مجموعه مورد بررسی باشد. در این حالت فقط در مورد زیرمجموعه‌ای که شامل تمام متغیرهای پاسخ است این مجموع برابر α خواهد بود.

۲.۳ چندین متغیر پاسخ در چندین خانواده

اگر چند خانواده از متغیرهای پاسخ مهم بوده (برآمدهای گروه اولیه، ثانویه و ...). و چند خانواده از برآمدها مبنای تصمیم‌گیری قرار گیرد، برای این موضوع ۲ استراتژی را در نظر می‌گیریم که عبارتند از استراتژی کنترل پیایی^{۱۵} و کنترل موازی^{۱۶}.

در کنترل پیایی باید ابتدا تمام برآمدها در یک خانواده معنی دار شوند و سپس خانواده بعدی مورد بررسی قرار گیرد (مراحل کار این روش مانند روش دنباله ثابت است). آزمون‌های درون هر خانواده می‌تواند از هر یک از روش‌های بن-فرونی، هلم، فال-بک و .. که FWE را در سطح α حفظ می‌کند، انجام شود.

در روش کنترل موازی، اگر حداقل یک متغیر پاسخ در یک خانواده معنی دار شود می‌توان خانواده‌ی بعدی را نیز مورد بررسی قرار داد. با این وجود اگر ناپست‌تری برای بعضی از متغیرهای پاسخ در یک خانواده فقط نتیجه گرفته شود، متغیرها در خانواده‌ی بعدی در سطح خطای نوع اول کوچکتری آزمون می‌شوند. همچنین آزمون درون یک خانواده از متغیرها در این روش توان کمتری نسبت به روش کنترل پیایی دارد. (۳۲)

۴ نتیجه‌گیری

در این مقاله با توجه به اهمیت روز افزون کارآزمایی‌های ناپست‌تری، به بررسی روش‌های مقایسه‌های چندگانه در این نوع کارآزمایی‌ها پرداخته شد. در این نوع کارآزمایی‌ها که مبنای تصمیم‌گیری بر اساس دو یا چند متغیر پاسخ مهم می‌باشد روش‌های مختلفی ارائه شد که تمام آنها بر پایه‌ی حفظ سطح خطای نوع اول می‌باشند.

مراجع

[23] FOOD, U. & ADMINISTRATION, D. (2016). Non-inferiority clinical trials to establish effectiveness. Guidance for industry.

¹⁴ α -exhaustive

¹⁵ serial gatekeeping

¹⁶ parallel gatekeeping

- [24] SCHUMI, J. & WITTES, J. T. (2011). Through the looking glass: understanding non-inferiority. *Trials*, 12, 106.
- [25] DUNNETT, C. W. (1964). New tables for multiple comparisons with a control. *Biometrics*, 20, 482-491.
- [26] BERGER, R. L. & HSU, J. C. (1996) Bioequivalence trials, intersection-union tests and equivalence confidence sets. *Statistical Science*, 11, 283-319.
- [27] HOLM, S. 1979. A simple sequentially rejective multiple test procedure. *Scandinavian journal of statistics*, 65-70.
- [28] WIENS, B. L. (2003). A fixed sequence Bonferroni procedure for testing multiple endpoints. *Pharmaceutical Statistics*, 2, 211-215.
- [29] WIENS, B. L. & DMITRIENKO, A. (2005). The fallback procedure for evaluating a single family of hypotheses. *Journal of Biopharmaceutical Statistics*, 15, 929-942.
- [30] HOCHBERG, Y. (1988). A sharper Bonferroni procedure for multiple tests of significance. *Biometrika*, 75, 800-802.
- [31] MARCUS, R., ERIC, P. & GABRIEL, K. R. (1976). On closed testing procedures with special reference to ordered analysis of variance. *Biometrika*, 63, 655-660.
- [32] DMITRIENKO, A., WIENS, B. L., TAMHANE, A. C. & WANG, X. (2007). Tree-structured gatekeeping tests in clinical trials with hierarchically ordered multiple objectives. *Statistics in medicine*, 26, 2465-2478.

معادلات ساختاری بیزی

نجمه شکیبایی^۱، کمیته تحقیقات دانشجویی، گروه آمار زیستی، دانشکده علوم پزشکی دانشگاه اصفهان، اصفهان، عقیل علایی
گروه آمار، دانشکده علوم ریاضی، دانشگاه صنعتی اصفهان

چکیده: مدل سازی معادلات ساختاری بیزی^۱ نسل دوم تجزیه و تحلیل آماری چند متغیره بسیار نیرومند از خانواده رگرسیون چند متغیری است که با اعتبار و قابلیت اعتماد به نمرات مشاهده شده مورد استفاده قرار می گیرد. به عبارتی الگوهای فرضی از ارتباطات مستقیم و غیرمستقیم در میان یک مجموعه از متغیرهای مشاهده شده و پنهان بررسی می شود که این الگوهای فرضی در مطالعات بالینی بسیار به چشم می خورد. این مقاله، روش و کاربرد مدل سازی معادلات ساختاری بیزی همراه با یک مثال در پژوهش های بالینی را معرفی نموده است.

واژه های کلیدی: معادلات ساختاری بیزی، تجزیه و تحلیل آماری، بیز

۱ مقدمه

مدل سازی معادله ساختاری دیدگاهی است که در آن الگوهای فرضی از ارتباطات مستقیم و غیرمستقیم در میان یک مجموعه از متغیرهای مشاهده شده و پنهان بررسی می شود. هم چنین یک متغیر می تواند همزمان هم نقش پیش بینی کننده و هم نتیجه را ایفا کند. کاربرد اصلی آن در موضوعات چند متغیره ای است که نمی توان آنها را به شیوه دومتغیری با در نظر گرفتن هربار یک متغیر مستقل با یک متغیر وابسته انجام داد. از دیگر کاربردهای مدل سازی معادلات ساختاری می توان به موارد تجزیه و تحلیل مدل با سازه های پنهان، اجرای تحلیل عاملی تاییدی، تجزیه رگرسیون با مشکلات چندهم خطی، تجزیه و تحلیل مسیر با چند متغیر وابسته، آزمون همبستگی و کوواریانس در یک مدل، مدل سازی روابط داخلی بین متغیرها به طور همزمان در یک مدل، مدل سازی و آزمون متغیرهای میانجی در یک مدل و تجزیه و تحلیل متغیرهای تعدیلگر در یک مدل زیر اشاره نمود. (۳۳)

اهمیت این تکنیک در پژوهش های بالینی از آنجاست که غالباً در این حوزه از مطالعات، پژوهشگران به بررسی روابط بین متغیرهای مختلف در قالب مدل یا شبکه ای از روابط می پردازند؛ بنابراین آنان، مبتنی بر فرضیه های خود در مورد روابط بین متغیرها، شمای کلی از این روابط را در قالب مدلی از پیش ساخته طراحی می نمایند (۳۴). در چنین وضعیتی پژوهشگر با کمک مدل سازی معادلات ساختاری

¹BSEM

^۱نجمه شکیبایی: najmeh.shakibaei@gmail.com

بیزی و با کمک نرم افزارهای موجود سوالات پیش رو را پاسخ می‌دهد.

از اولین مطالعات بیزی معادلات ساختاری می‌توان به شاینز و همکاران (۳۵)؛ انصاری و جدی (۳۶)، دانسون (۳۷)، لی (۳۸)، اشاره کرد. همچنین با گسترش نرم افزارها از مطالعات کاربرد مدلسازی معادلات ساختاری در حوزه بالینی نیز می‌توان به مدلسازی معادلات ساختاری بیزی: کاربرد سلامت و شاخص سلامت اشاره نمود (۳۹؛ ۴۰). با توجه به مفید بودن این روش آماری در حوزه‌های بالینی بر آن شدیم تا مرور مختصری با استفاده از منابع مطالعه شده در این زمینه به معرفی روش معادلات ساختاری با استفاده از نرم افزار وین باگ^۲ بپردازیم.

۲ متن اصلی

اطلاعات این مطالعه مروری با استفاده از مرور و خلاصه سازی مقالات و کتب مرتبط و از طریق جستجوی هدفمند کتابخانه‌ای جهت دستیابی به کتب و منابع الکترونیکی شامل موتور جستجوگر گوگل و بانک‌های اطلاعاتی جهت دستیابی به کتب و مقالات به دست آمد. نتایج پژوهشگران عرصه علوم انسانی و آموزش علوم سلامت قادرند تا با استفاده از مدلسازی معادلات ساختاری به مناسب بودن مدل‌های مفهومی یا کاربرد آن در جامعه مورد مطالعه پی‌ببرند. یافته‌های این مطالعه مروری در قالب معرفی مفهوم و گام‌های کلی که پژوهشگران لازم است مدنظر قرار دهند ارائه شده است.

۱.۲ تعریف مدلسازی معادلات ساختاری

مدل معادله ساختاری اساساً ترکیب مدل‌های مسیر و مدل‌های تحلیل عاملی تاییدی است. تکنیک‌های آماری مدل سازی معادلات ساختاری از طریق محاسبه واریانس مشترک یا عمومی بین متغیرهای آشکار این کار را انجام می‌شود. مدل سازی معادلات ساختاری همانند تحلیل عاملی نقش تفکیک واریانس مربوط به هر شاخص را به عهده دارد (۳۴). اطلاعات مورد نیاز برای تحلیل در قالب پارامترهای مدل در اختیار نرم افزار قرار داده می‌شوند متغیرها و پارامترهای موجود در تحقیق شامل:

تعریف ۱.۲ (متغیر پنهان یا سازه‌ها یا عامل‌ها). متغیرهایی که نمی‌توان آنها را مستقیماً مشاهده یا مورد سنجش قرار داد. همچون عشق، نفرت، رضایت، کیفیت و ...

تعریف ۲.۲ (متغیر آشکار یا مشاهده شده). متغیرهایی که مستقیماً اندازه‌گیری می‌شوند. مانند قد، وزن، فشارخون و ... این متغیرها رابه منظور تعریف یا استنباط در مورد متغیر پنهان به کار می‌بریم.

تعریف ۳.۲ متغیرهای وابسته و مستقل متغیرها چه آشکار و چه پنهان، همچنین می‌توانند به عنوان متغیرهای مستقل و وابسته تعریف شوند.

تعریف ۴.۲ (متغیر مستقل یا برون‌زا). متغیرهایی هستند که تحت تاثیر متغیرهای موجود درمدل نیستند. این متغیرها حداقل یک مسیر به متغیردیگروارد می‌کنند.

تعریف ۵.۲ (متغیر وابسته یا درون‌زا). متغیرهایی هستند که مقادیر آنها توسط مدل برآورد می‌شود. این متغیرها حداقل یک مسیر را ازمتغیردیگردریافت می‌کنند.

ترکیب اصلی معادلات ساختاری کلی را می‌توان به دو خرده مدل تقسیم بندی کرد:

۱. مدل اندازه گیری

۲. مدل ساختاری.

مدل اندازه‌گیری ارتباط بین متغیرهای مشاهده شده و مشاهده نشده را تعریف می‌کند. به عبارت دیگر، بر مبنای اندازه‌گیری، پژوهش‌گر تعریف می‌کند که کدام متغیرهای مشاهده‌پذیر می‌توانند جهت اندازه‌گیری متغیرهای پنهان مورد استفاده قرار بگیرند. بنابراین مدل اندازه‌گیری، مدل تحلیل عاملی تاییدی را مشخص می‌کند که به وسیله آن هر اندازه روی عامل خاصی بارگذاری می‌شود. درمقایسه، مدل ساختاری، ارتباط بین متغیرهای مشاهده نشده را تعریف می‌کند (۳۳). یکی از رایج‌ترین روش‌های برآورد، روش بیشینه درستنمایی^۳ است. در اکثر برنامه‌های کامپیوتری تحلیل مدل‌های معادلات ساختاری است و نیز اکثر مطالعات منتشر شده استفاده از این روش را گزارش نموده‌اند. از آنجایی که این روش برآورد، بر مبنای پیش فرض‌هایی همچون مطالعه روی حجم نمونه‌های بزرگ، پیوسته بودن متغیرها استوار است پژوهشگران لازم است از برقرار بودن این شرایط اطمینان حاصل نمایند (۳۴). اما در مقابل روش بیزی وابستگی کمتر به این فرضیات را دارد و جوابگو برای نمونه‌های کوچک می‌باشد (۴۲).

۳ تشریح مراحل انجام معادلات ساختاری بیزی

۱.۳ تدوین مدل

فرض کرده $Y = (Y_1, Y_2, \dots, Y_n)$ ماتریس داده‌ها باشند. معادله اندازه‌گیری به صورت $Y_i = \lambda \omega_i + \varepsilon_i$ در نظر گرفته می‌شود. که Y_i معرف بردار $1 \times p$ از متغیرهای مشاهده شده برای توصیف بردار تصادفی $1 \times q$ متغیرهای پنهان ω_i ، λ ماتریس $p \times q$ از ضرایب رگرسیون Y_i روی ω_i و ε_i بردار تصادفی $1 \times p$ از خطای اندازه‌گیری است که متغیر پنهان ω_i به صورت زیر در نظر گرفته می‌شود.

$$\omega_i = \begin{pmatrix} \eta_i \\ \xi_i \end{pmatrix} \quad (۱.۳)$$

η_i معرف بردار تصادفی $1 \times q_1$ از متغیرهای پنهان وابسته و ξ_i معرف بردار تصادفی $1 \times (q_2 = q - q_1)$ از متغیرهای پنهان مستقل است. بنابراین معادله ساختاری به صورت زیر در نظر گرفته می‌شود. که β ماتریس رابطه متغیرهای پنهان درونی، Γ ماتریس $q_1 \times q_2$ رابطه میان متغیرهای پنهان درونی و بیرونی و σ بردار تصادفی $1 \times q_1$ از باقیمانده‌ها است.

$$\eta_i = \beta \eta_i + \Gamma \xi_i + \sigma_i \quad (۲.۳)$$

۲.۳ تشخیص مدل

در مساله تشخیص این سوال مطرح می‌شود که آیا بر اساس داده‌های نمونه‌ای (شامل شده در ماتریس کواریانس نمونه‌ای S) و مدل نظری (تعریف شده بوسیله ماتریس کواریانس جامعه)، می‌توان مجموعه منحصر به فردی از برآورد پارامترها یافت؟

۳.۳ برآورد بیزین پارامترها

فرض کرده M یک SEM دلخواه از پارامترهای ناشناخته θ و داده‌های مشاهده شده با حجم نمونه n باشد. در رویکرد بیزی θ تصادفی با توزیع پیشین $p(\theta)$ است. برآورد بیزین بر اساس داده‌های مشاهده شده Y و توزیع پیشین θ صورت می‌گیرد. یک مسئله انتخاب تابع پیشین مناسب است. مراجعه به (۴۲)

برآوردهای بیزی از θ را معمولاً به عنوان میانگین یا مد تابع پسین $\theta|y$ در نظر می‌گیریم. اگر بتوانیم به اندازه کافی تعدادی از مشاهدات $\theta|y$ را شبیه سازی کنیم می‌توانیم میانگین و آماره‌های مفید دیگر را از طریق مشاهدات شبیه سازی به دست آوریم. برای حل کردن مشکل، کافی است روش‌های قابل اعتماد و کارآمد برای استخراج مشاهدات از توزیع پسین به دست آوریم. به جای کار

³Maximum likelihood

کردن با تابع چگالی $p(\theta|y)$ بر تابع چگالی $p(\theta, \eta|y)$ متمرکز می‌شویم (η مجموعه متغیرهای پنهان). تابع چگالی $\theta, \eta|y$ فرم بسته ای ندارد و دسترسی به آن سخت است، براساس مجموعه داده کامل (y, η) توزیع شرطی $(\theta|\eta, y)$ و توزیع $(\eta|\theta, y)$ به آسانی استخراج می‌شود. بنابراین بعضی روش‌های $MCMC$ اجرا می‌کنیم تا مشاهدات را از طریق مشاهدات تکراری چگالی شرطی آن‌ها $(\theta|\eta, y)$ و $(\eta|\theta, y)$ به دست آوریم. در نهایت توزیع توام از (θ^t, η^t) بعد از تعداد تکرار کافی به توزیع پسین $(\theta|\eta, y)$ همگرا می‌شود. استنتاج‌های آماری از مدل بر اساس نمونه شبیه سازی شده از مشاهدات $p(\theta|\eta, y)$ صورت می‌گیرد. به عبارت دیگر $\{\theta^t, \eta^t; t = 1, 2, \dots, T\}$ برآورد θ و انحراف استاندارد آن

$$\hat{\theta} = T^{-1} \sum_{i=1}^T \theta^i \quad (3.3)$$

$$Var(\theta|Y) = (T - 1)^{-1} \sum_{i=1}^T (\theta^i - \hat{\theta})(\theta^i - \hat{\theta})^{-1} \quad (4.3)$$

زمانی که T به سمت ∞ می‌رود $\hat{\theta}$ به $E(\theta|y)$ میل می‌کند به عبارت دیگر: برای هر Y_i و W_i نشان دهنده بردار از متغیرهای پنهان باشد، $E(W_i|Y_i)$ میانگین پسین است. برآوردگر بیز W_i از طریق: $E(W_i|Y_i) = T^{-1} \sum_{i=1}^T W_i$ که W_i i مین ستون از η^t است. این برآورد W_i می‌تواند برای آنالیز داده‌های پرت و باقی مانده‌ها و ارزیابی نیکویی برازش مدل مورد استفاده قرار گیرد (۴۲).

۴.۳ آزمون مدل

معیارهای تشخیص همگرایی

معیارهای تشخیصی اولیه شامل خطای مونت کارلو، نمودار تریس، نمودار تکرارها در مقابل مقادیر تولید شده را بدست می‌دهد. نوساناتی حول یک محور بدون روند یا تناوب، گویای وضعیت خوب زنجیره است. همچنین می‌توان از *densitybutton* که یک برآورد تقریبی کرنل از تابع چگالی تابع احتمال پسین پارامترها فراهم می‌کند و نمودار اتوکرولیشن^۴ استفاده نمود. برای دستیابی به معیارهای دقیق‌تر همگرایی می‌توان از آزمون گلن رابین و آزمون رفتی و لوئیس استفاده نمود (۴۳).

شاخص‌های نیکویی برازش

انواع شاخص‌های نیکویی برازش مورد استفاده عبارتند از معیار اطلاع دویانس و ارزیابی پیش بینی کننده پسین. معیار اطلاع دویانس توسط اشپیگل هالتر به عنوان اندازه ای برای مقایسه مدلها و بررسی کفایت مدل معرفی شد. مقادیر کمتر DIC اشاره به برازش بهتر دارد.

ارزیابی پیش بینی کننده پسین

نمونه‌های تولید شده از توزیع پیش بینی کننده پسین با داده‌های واقعی مقایسه می‌گردد. در این تکنیک اگر مدل برازش خوبی به داده‌ها داشته باشد داده‌های شبیه‌سازی شده تحت مدل باید نزدیک به داده‌های مشاهده شده باشد. مقادیر نزدیک ۵.۰ ارزیابی کننده پسین نشان دهنده نزدیک بودن توزیع داده‌های واقعی و داده‌های شبیه سازی شده به یکدیگر است. در حالی که مقادیر نزدیک به صفر یا یک این احتمال بیانگر وجود اختلاف بین آنهاست.

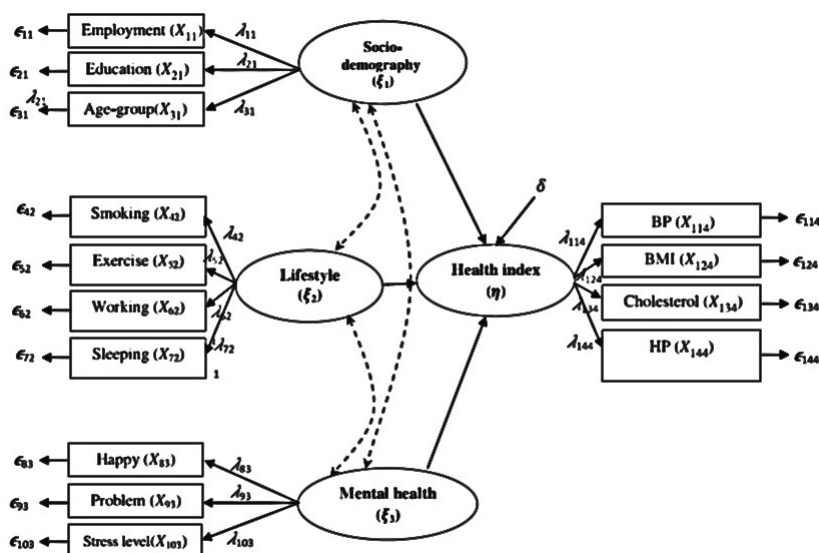
⁴ Autocorrelation

۵.۳ اصلاح مدل

اگر برازش یک مدل نظری به قوتی نبود که انتظار داشتیم، آنگاه گام بعدی، اصلاح مدل و ارزیابی مدل جایگزین و اصلاح شده است. و در نهایت به تفسیر نتایج می‌پردازیم.

۴ بررسی یک مطالعه

هدف این مطالعه مدل سازی شاخص سلامت بر اساس معادله ساختاری بیزی است. مطالعات مختلف نشان می‌دهند که بسیاری از عوامل بر سلامت افراد از جمله ویژگی‌های جمعیت‌شناسی، سبک زندگی سلامت روان و... تاثیر می‌گذارند. داده‌های مطالعه از ۵۰۳۵ نفر افراد بزرگتر از ۱۸ سال در *huluLangat* (منطقه‌ای در مالزی) جمع آوری شده است. متغیرهای پنهان این مطالعه ویژگی‌های جمعیت‌شناختی، سبک زندگی و سلامت ذهن است. مدل سازی معادلات این مطالعه برای $i = 1, 2, \dots, n$ به صورت زیر است.



شکل ۱:

که در آن λ ماتریس ضرایب عاملی به دست آمده از رگرسیون Y_i بر ω_i است.

$$y_i = \lambda \omega_i + \varepsilon_i$$

بنابراین متغیرهای پنهان به صورت $\omega_i = (\eta, \xi_1, \xi_2, \xi_3)^{-1}$ و معادله ساختاری برای توضیح روابط میان متغیرهای پنهان به صورت $\eta_i = \Gamma \xi_i + \sigma_i$ است.

نتایج این مطالعه نشان می‌دهد که میان شاخص سلامت و ویژگی‌های جمعیت‌شناختی رابطه معنادار وجود دارد هم‌چنین میان شاخص سلامت و سبک زندگی اما میان شاخص سلامت و سلامت روان رابطه معناداری یافت نشد که ممکن است تحت تاثیر متغیرهای مهم دیگر چون رنج و مشقت، جنون جوانی، افسردگی، اضطراب و اختلال اجتماعی و... باشد و باید آن‌ها نیز در مدل گنجانده شوند (۴۴).

۵ نتیجه گیری

در مجموع، مدل یابی معادلات ساختاری به دلیل دارا بودن قابلیت های متعدد و نیز غلبه بر محدودیت های روش های کلاسیک، کاربردهای بسیاری در مطالعات بالینی دارد. از آنجا که پدیده های انسانی غالباً دارای علت های چندگانه بوده و مطالعه آنها مستلزم لحاظ کردن متغیرها و سازه های متعدد و در برخی کوریت ها به دلیل نحوه جواب دادن افراد به سوالات ممکن است با تعداد زیادی داده گمشده مواجه شویم بنابراین در این موارد رویکرد بیزی پاسخ بهتری می دهد. همچنین با توجه به استفاده از داده ها برای تازه کردن اطلاعات پیشین موجود از پارامتر و قرار گرفتن اطلاعات تازه شده در قالب توزیع پسین پارامتر، وابستگی کمتر به فرضیات، و هم چنین زمانی که تعداد عوامل زیاد می شود، کشف ارتباط بین متغیرها پیچیده تر می شود. در این شرایط معمولاً مدل بندی بیزی در تفکیک واریانس ها موفق عمل می کند. (۴۵؛ ۳۴) امید است پژوهشگران فعال در حوزه بالینی بر اصول علمی مربوطه احاطه داشته و بتوانند از قابلیت های این تکنیک ها به طور مناسبی در مطالعات خود بهره مند گردند.

مراجع

- [33] Lomax RG, Schumacker RE. A beginner's guide to structural equation modeling. psychology press; 2004 Jun 24.
- [34] Alavi,M, Structural Equation Modeling (SEM) in Health Sciences Education Researches: An Overview of the Method and Its Application , Iranian Journal of Medical Education 2013: 13(6) / 530
- [35] R. Scheines, H. Hoijtink, and A. Boomsma, Bayesian estimation and testing of structural equation models, Psychometrika 64 (1999), pp. 37–52.
- [36] Ansari, K. Jedidi, and S. Jagpal, A hierarchical Bayesian methodology for treating heterogeneity in structural equation models, Market. Sci. 19 (2000), pp. 328–347.
- [37] S.Y. Lee and J.Q. Shi, Bayesian analysis of structural equation model with fixed covariates, Struct. Equ. Model. 7 (2000), pp. 411–430
- [38] Dunson DB, Palomo J, Bollen K. Bayesian structural equation modeling. SAMSI TR2005-5. 2005 Jul 27
- [39] S.Y. Lee, X.Y. Song, S. Skevton, and Y.T. Hao, Application of structural equation models to quality of life, Struct. Equ. Model. 12 (2005), pp. 435–453
- [40] Stojanovski, E. and K. Mengersen, Bayesian Structural Equation Models: A Health Application,
- [41] Cook, J., 2012. Conjugate prior relationships.
- [42] Xin-Yuan Song and Sik-Yum Lee, Basic and Advanced Bayesian 10. Structural Equation Modeling With Applications in the Medical and Behavioral Sciences

- [43] Ntzoufras I. Bayesian modeling using WinBUGS. John Wiley, Sons; 2011 Sep 20.
- [44] Ferra Yanuar a , Kamarulzaman Ibrahim b, Abdul Aziz Jemain, Bayesian structural equation modeling for the health index, , Journal of Applied Statistics, 2014
- [45] Fon Sim Ong, Kok Wei Khong, Ken Kyid Yeoh, Osman Syuhaily, Othman Mohd. Nor, A comparison between Structural Equation Modelling (SEM) and Bayesian SEM approaches on in-store behaviour, ,

تجزیه و تحلیل پایداری مدل شکار-شکارچی با واکسیناسیون و مهاجرت در شکار

محمد شیرازیان، محمد حسین رحمانی دوست و مرضیه شمس آبادی^۱، دانشگاه نیشابور، نیشابور، ایران

چکیده: در این مقاله به بررسی یک مدل سه بعدی که از شکار مستعد، شکار واکسینه شده و شکارچی تشکیل شده و فرمول بندی می شود می پردازیم. پاسخ عملکردی از نوع لوتکا-ولترا می باشد. همچنین پویایی دستگاه از قبیل جواب های کراندار، شرایط وجود تعادل و پایداری نقاط تعادلی به صورت ریاضی وار تجزیه و تحلیل خواهد شد.

واژه های کلیدی: لوتکا-ولترا، مدل شکار-شکارچی، پایداری.

۱ مقدمه

برای مطالعه ی رفتار پویایی یک مدل، مدل سازی ریاضی به عنوان یک ابزار مؤثر برای توصیف و تحلیل، مورد استفاده قرار می گیرد. در سال ۱۷۹۸، اقتصاددان بریتانیایی، مالتوس یک مدل نوع گونه ای راشکل داد و سپس توسط ورهاست اصلاح شد. در ابتدا لوتکا و ولترا مدل شکار-شکارچی را پیشنهاد دادند. بعد از آن مدل شکار-شکارچی به یک بخش تحقیقاتی مهم در ریاضیات کاربردی تبدیل شد. بخاطر مدل کرماک-مکندیک روی دستگاه های $SIRS$ ریاضی اپیدمیولوژیکی به یک موضوع جالب تحقیقاتی تبدیل شد. نویسندگان بسیاری در مورد مدل شکار شکارچی مطالعه کرده اند و مقاله هایی در مطبوعات منتشر نموده اند. از جمله جهری که یک مدل شکار شکارچی از نوع لوتکا-ولترا را با وجود بیماری در گونه ی شکار مطرح و پایداری موضعی و سراسری آن را بررسی نمودند (۴۸). در این مقاله، یک مدل سه بعدی با وجود واکسیناسیون و مهاجرت در شکار بررسی می شود.

۲ بررسی مدل ریاضی

مدل شامل دو جمعیت است. شکار که تراکم جمعیت آن با $N(t)$ مشخص می شود و شکارچی که تراکم جمعیت آن با $P(t)$ نشان داده می شود. اینجا t متغیر زمان است. فرضیات زیر برای ساختن مدل در نظر گرفته می شود.

(۱) جمعیت شکار به طور منطقی با نرخ رشد $r(r > 0)$ رشد می یابد و ظرفیت حمل $K(K > 0)$ در غیاب بیماری،

^۱مرضیه شمس آبادی: m.shamsabadi.96@gmail.com

واکسیناسیون و شکارگری است. لذا داریم:

$$\frac{dS}{dt} = rS\left(1 - \frac{S}{K}\right)$$

(۲) شکار واکسینه شده ی $V(t)$ گروه جداگانه ای است و ϕ میزان واکسیناسیون شکار مستعد است θ میزان بی اثر شدن واکسن می باشد.

(۳) جمعیت شکار با وجود واکسیناسیون به دو دسته تقسیم می شود. که شامل شکار مستعد و شکار واکسینه شده است. بنابراین کل جمعیت شکار در لحظه t عبارتند از: $N(t) = S(t) + V(t)$.

(۴) شکارچی با شکار کردن شکار مستعد و واکسینه شده آشکارا به نتیجه ی مشابهی می رسد با تفاوت بازده ی جستجو که به ترتیب با P_1, P_2 نشان داده می شود.

(۵) پاسخ عملکردی شکارچی از نوع لوتکا-ولترا است.

(۶) ضرایب تبدیل شکار مستعد و واکسینه شده به شکارچی به ترتیب با q_1, q_2 نشان داده می شوند که متفاوت هستند.

(۷) جمعیت شکار دارای نرخ مهاجرت m_1, m_2 می باشند که به ترتیب مربوط به نرخ مهاجرت شکار مستعد و واکسینه شده است.

(۸) نرخ مرگ و میر طبیعی سه گونه ی شکار مستعد، شکار واکسینه و شکارچی متفاوت است و به ترتیب با d_1, d_2, d_3 نشان داده می شوند.

از مدل ریاضی با فرضیات فوق معادلات دیفرانسیل زیر نتیجه می شود:

$$\begin{aligned}\frac{dS}{dt} &= rS\left(1 - \frac{S}{k}\right) - \phi S + \theta V - p_1 PS - m_1 S - d_1 S, \\ \frac{dV}{dt} &= \phi S - \theta V - p_2 PV - m_2 V - d_2 V, \\ \frac{dP}{dt} &= q_1 p_1 PS + q_2 p_2 PV - d_3 P.\end{aligned}\quad (1.2)$$

با شرایط اولیه:

$$P(0) = P_0 \geq 0, V(0) = V_0 \geq 0, S(0) = S_0 \geq 0,$$

توجه ۱.۲. اگر $\phi = 0$ آنگاه واکسیناسیون وجود ندارد. بنابراین داریم: $\lim_{t \rightarrow \infty} V(t) = 0$.

توجه ۲.۲. در این مدل بین مرگ و میر با مهاجرت تفاوت قائل می شویم ولی جالب است توجه داشته باشیم که شرایط مهاجرت در مدل فوق مشابه شرایط مرگ و میر است.

کرانداری جواب ها

قضیه ۳.۲. همه ی جواب های دستگاه (۱.۲) همواره در ناحیه ی مثبت R_+^3 هستند.

اثبات.

چون تمام پارامترها نامنفی هستند. سمت راست (۱.۲) یک تابع هموار از متغیرهای (S, V, P) در $\frac{1}{\lambda}$ مثبت فضای $\Omega = \{(S, V, P) : S \geq 0, V \geq 0, P \geq 0\}$ است. چون دستگاه فوق همگن است لذا $(S = 0, V = 0, P = 0)$ یک جواب دستگاه است. طبق قضیه وجود و یکتایی جواب؛ منحنی جواب مورد نظرا از اولین ربع شروع می شود و در آن باقی می ماند. و این یعنی هیچ منحنی فضای مختصاتی را قطع نمی کند. توجه داریم که Ω یک مجموعه ی ثابت است. $\square$

قضیه ۴.۲. همه ی جوابهای دستگاه (۱.۲) به طور یکنواخت کراندار هستند.

اثبات. فرض می کنیم که $\chi = S + V + P$ بنابراین مشتق χ نسبت به زمان عبارتند از

$$\frac{d\chi}{dt} = \frac{dS}{dt} + \frac{dV}{dt} + \frac{dP}{dt}$$

حال برای هر $\mu > 0$ داریم:

$$\begin{aligned} \frac{d\chi}{dt} + \mu\chi &= rS\left(1 - \frac{S}{K}\right) - p_1PS - m_1S - d_1S - p_2PV - m_2V - d_2V \\ &\quad - q_1p_1PS + q_2p_2PV - d_3P + \mu\chi \end{aligned}$$

بنابراین

$$\frac{d\chi}{dt} + \mu\chi \leq k \frac{(r + \mu)^r}{r} - (m_2 + d_2 - \mu)V - (d_3 - \mu)P$$

اگر $\mu < \min(m_2 + d_2, d_3)$ آنگاه داریم:

$$\frac{d\chi}{dt} + \mu\chi \leq k \frac{(r + \mu)^r}{r} = \eta$$

بنابراین با استفاده از قضیه نابرابری دیفرانسیل داریم:

$$0 < \chi(S, V, P) < \left(\frac{\eta}{\mu}\right)(1 - \exp(-\mu t)) + \chi(S_0, V_0, P_0) \exp(-\mu t)$$

بنابراین زمانی که $t \rightarrow \infty$ نتیجه می شود $0 < \chi(S, V, P) < \frac{\eta}{\mu}$ لذا همه ی جواب های دستگاه (۱.۲) به فضای زیر محدود می شوند.

$$\Lambda = \{(S, V, P) \in R_+^r : \chi = \frac{\eta}{\mu} + \varepsilon, \forall \varepsilon > 0\}$$

□

۱.۲ نقاط تعادلی مدل

برای دستگاه (۱.۲) نقاط تعادلی زیر وجود دارد.

(۱) تعادل بدیهی $E_0 = (0, 0, 0)$ همواره وجود دارد.

(۲) تعادل بدون شکارچی $E_1 = (S_1, V_1, 0)$ وجود دارد اگر شرط زیر برقرار باشد.

$$(r - \phi - d_1 - m_1 + \frac{\theta\phi}{\theta + d_2 + m_2}) > 0. \quad (2.2)$$

$$S_1 = \frac{K}{r}(r - \phi - m_1 - d_1 + \frac{\theta\phi}{\theta + m_2 + d_2})$$

$$V_1 = \frac{\phi K}{r(\theta + m_2 + d_2)}(r - \phi - m_1 - d_1 + \frac{\theta\phi}{\theta + m_2 + d_2})$$

(۳) تعادل درونی $E_2 = (S_2, V_2, P_2)$ وجود دارد اگر شرایط زیر برقرار باشد.

$$(d_2 - p_2 q_2 V_2) > 0$$

و P_2 ریشه ی مثبت معادله ی زیر باشد.

$$\left(-\frac{rd_2(\theta + d_2 + m_2 + p_2 P_2)^2}{K}\right) + [r\theta + (r - \phi)(d_2 + m_2) - (\theta + d_2 + m_2)(d_1 + m_1 + p_1 P_2) - P_2(-r + \phi + d_1 + m_1 + p_1 P_2)p_2][p_1(\theta + d_2 + m_2 + p_2 P_2)q_1 + \phi p_2 q_2] = 0. \quad (۳.۲)$$

در این نقطه داریم:

$$S_2 = \frac{d_2 - p_2 q_2 V_2}{p_1 q_1},$$

$$V_2 = \frac{\phi d_2}{(\theta p_1 q_1 + d_2 p_1 q_1 + m_2 p_1 q_1 + \phi p_2 q_2 + p_1 p_2 q_1 P_2)}.$$

۲.۲ بررسی پایداری نقاط تعادل

(۱) نقطه تعادل بدیهی $E_0 = (0, 0, 0)$ همواره وجود دارد.

می دانیم برای تجزیه و تحلیل پایداری مدل از ماتریس ژاکوبین استفاده می کنیم لذا داریم:

$$J = \begin{bmatrix} r - \frac{rS}{K} - \phi - p_1 P - m_1 - d_1 & \theta & -p_1 S \\ \phi & -\theta - p_2 P - m_2 - d_2 & -p_2 V \\ q_1 p_1 P & q_2 p_2 P & q_1 p_1 S + q_2 p_2 V - d_2 \end{bmatrix}_{3 \times 3}$$

بنابراین ماتریس ژاکوبین ارزیابی شده در E_0 به صورت زیر می باشد.

$$J(E_0) = \begin{bmatrix} r - \phi - m_1 - d_1 & \theta & 0 \\ \phi & -\theta - m_2 - d_2 & 0 \\ 0 & 0 & -d_2 \end{bmatrix}$$

از اینرو چند جمله ای مشخصه ماتریس فوق عبارتند از:

$$(-\lambda - d_3)(-\theta\phi - (r - \phi - d_1 - m_1 - \lambda)(\lambda + \theta + d_2 + m_2)) = 0 \quad (۴.۲)$$

لذا یکی از مقادیر ویژه ی $\hat{J}(E_0)$ برابر $-d_3$ می باشد و دو ریشه ی باقیمانده ی (۴.۲) از معادله درجه دوم زیر به دست می آید.

$$(\lambda - r + \phi + d_1 + m_1)(\lambda + \theta + d_2 + m_2) - \theta\phi = 0$$

اکنون با استفاده از محک روت -هرویتز اگر شرایط زیر برقرار باشد E_0 پایدار موضعی است .

$$\begin{cases} (-r + \theta + \phi + d_1 + d_2 + m_1 + m_2) > 0, \\ (-r + \phi + d_1 + m_1)(\theta + d_2 + m_2) - \theta\phi > 0. \end{cases} \quad (۵.۲)$$

(۲) در نقطه تعادلی بدون شکارچی $E_1 = (S_1, V_1, 0)$ ماتریس ژاکوبین به صورت زیر می باشد.

$$\hat{J}(E_1) = \begin{bmatrix} r - \frac{rS_1}{K} - \phi - m_1 - d_1 & \theta & -p_1S_1 \\ \phi & -\theta - m_2 - d_2 & -p_2V_1 \\ 0 & 0 & q_1p_1S_1 + q_2p_2V_1 - d_3 \end{bmatrix}$$

لذا معادله مشخصه ی $\hat{J}(E_1)$ به صورت زیر می باشد.

$$\frac{1}{K}(-\lambda - d_3 + p_1q_1S_1 + p_2q_2V_1)[(rS_1 + K(-r + \lambda + \phi + d_1 + m_1))(\lambda + \theta + d_2 + m_2) - K\theta\phi] = 0. \quad (۶.۲)$$

بنابراین یکی از مقادیر ویژه برابر است با $\lambda_1 = -d_3 + p_1q_1S_1 + p_2q_2V_1$

و دو ریشه ی باقیمانده ی چندجمله ای مشخصه ی (۵.۲) از معادله ی درجه دوم زیر به دست می آید.

$$(\lambda + \frac{rS_1}{K} - r + \phi + d_1 + m_1)(\lambda + \theta + d_2 + m_2) - \theta\phi = 0. \quad (۷.۲)$$

بنا به محک روت - هرویتز E_1 پایدار است اگر

$$\begin{cases} (-r + \frac{rS_1}{K} + \theta + \phi + d_1 + d_2 + m_1 + m_2) > 0, \\ (\frac{rS_1}{K} - r + \phi + d_1 + m_1)(\theta + d_2 + m_2) - \theta\phi > 0, \\ \lambda_1 = (-d_3 + p_1q_1S_1 + p_2q_2V_1) < 0 \end{cases} \quad (۸.۲)$$

که اینجا $S_1 > 0$ ، $V_1 > 0$.

(۳) در نقطه تعادل درونی $E_2 = (S_2, V_2, P_2)$ ماتریس ژاکوبین برابر است با

$$J(E_2) = \begin{bmatrix} \alpha_{11} & \alpha_{12} & \alpha_{13} \\ \alpha_{21} & \alpha_{22} & \alpha_{23} \\ \alpha_{31} & \alpha_{32} & \alpha_{33} \end{bmatrix}$$

که اینجا

$$\begin{aligned} \alpha_{11} &= r - \frac{rS_2}{K} - \phi - p_1P_2 - m_1 - d_1, \alpha_{12} = \theta, \alpha_{13} = -p_1S_2, \\ \alpha_{21} &= \phi, \alpha_{22} = -\theta - p_2P_2 - m_2 - d_2, \alpha_{23} = -p_2V_2, \\ \alpha_{31} &= q_1p_1P_2, \alpha_{32} = q_2p_2P_2, \alpha_{33} = q_1p_1S_2 + q_2p_2V_2 - d_3 \end{aligned}$$

لذا معادله مشخصه ماتریس فوق به صورت زیر می باشد.

$$\lambda^3 + B_1\lambda^2 + B_2\lambda + B_3 = 0$$

$$B_1 = -\text{tr}A = -(\alpha_{11} + \alpha_{22} + \alpha_{33}),$$

مجموع مینورهای مرتبه ی دوم برابر است با B_2

$$B_2 = (\alpha_{11}\alpha_{22} - \alpha_{12}\alpha_{21}) + (\alpha_{11}\alpha_{33} - \alpha_{13}\alpha_{31}) + (\alpha_{22}\alpha_{33} - \alpha_{23}\alpha_{32})$$

$$B_3 = -\det A = -[(\alpha_{11}(\alpha_{22}\alpha_{33} - \alpha_{23}\alpha_{32}) + \alpha_{12}(\alpha_{31}\alpha_{23} - \alpha_{33}\alpha_{21})) + \alpha_{13}(\alpha_{21}\alpha_{32} - \alpha_{31}\alpha_{22})]$$

با توجه به محک روت -هرویتز E_2 پایدار موضعی است اگر داشته باشیم:

$$\begin{cases} B_1B_2 > B_3, \\ B_i > 0 \end{cases}$$

برای $i = 1, 2, 3$.

۳ نتیجه گیری

در مدل ارائه شده نقطه تعادلی بدیهی E همواره وجود دارد. و اگر در شرط (۵.۲) صدق کند آنگاه این نقطه پایدار است. از لحاظ ریاضی این حالت ناپایدار است زیرا اگر این حالت پایدار باشد به این معنی است که گونه منقرض شده است. نقطه تعادلی E_1 اگر پایدار باشد به این معنی است که جمعیت شکار بدون شکارچی مدت طولانی زنده می ماند. و مهمترین نقطه تعادلی؛ تعادل E_2 است زیرا این نقطه همزیستی همه ی گونه ها در اکوسیستم را نشان می دهد. این نقطه برای تعادل زیستی بسیار ضروری است.

تشکر و قدردانی

بدینوسیله نویسنده تشکر و سپاسگزاری می نماید از مسئولین دانشگاه نیشابور، بویژه اساتید محترم گروه ریاضی که از محضرشان بهرمند شدیم.

مراجع

- [46] H. W. Hethcote, W. D. Wang, L. T. Han and M. Zhien, *A predator- prey model with infected prey*, Theor. Popul. Biol. 66 (2004), 259- 268.
- [47] S. Jana , T. K. Kar, *Modeling and analysis of a prey- predator system with disease in the prey*, Chaos Solitons Fractals. 47 (2013), 42- 53.
- [48] A. Johri, et al. *Study of a prey- predator model with diseased prey*, Int.J.Contemp. Math. Sci. 7 (2012), no. 9- 12, 489- 498.
- [49] S. Kumar, H. Kharbanda, *Stability analysis of prey- predator model with infection, migration and vaccination in prey*, arXiv preprint arXiv:1709.10319 (2017).

مطالعه عددی یک مدل ریاضی بیماری ناشی از آلودگی آب

یونس طالعی دانشگاه تبریز، تبریز، ایران، حجت اله عبادیزاده^۱

دانشگاه امام علی (ع)، ایران

چکیده: بیماری تب تیفوئید یکی از شایع ترین بیماری ناشی از آلودگی مواد غذایی و منابع آب و یکی از مشکلات زیست محیطی در حال حاضر جهان صنعتی می باشد. این بیماری عفونی مانند سایر بیماری ها از جمله ایدز، هپاتیت و ... در قالب یک دستگاه معادلات دیفرانسیل غیرخطی مدل بندی می شود. در این مقاله، یک روش عددی طیفی ماتریسی برای مطالعه عددی مدل ریاضی تب تیفوئید معرفی می شود. در اینجا، هدف پیدا کردن جواب تقریبی با دقت مناسب برای تابع میزان افراد محافظت شده، تحت درمان و احیا شده از طریق آموزش برای جلوگیری و درمان بیماری تیفوئید می باشد. نتایج حاصل از اجرای روش، دقت و کارایی روش را نشان می دهد.

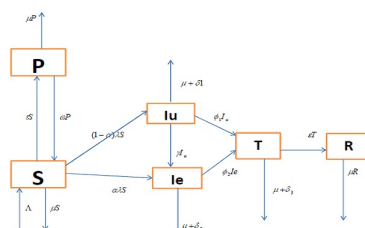
واژه های کلیدی: بیماری تب تیفوئید، عوامل آلودگی آب، مدل ریاضی، روش گالرکین ماتریسی

۱ مقدمه

بیماری تب تیفوئید یک بیماری عفونی است، که در اثر باکتری سالمونلا تیفی ایجاد می شود، این بیماری واگیردار بوده و از طریق آب و غذای آلوده گسترش می یابد و معمولاً با تب بالا، اسهال، بی اشتها و سردرد بروز می کند. باکتری تیفوئید در آب های گل آلود تا یک ماه و در یخ تا سه ماه زنده می ماند، در اثر گرمای ۶۰ تا ۱۰۰ درجه از میان می رود و بخصوص نور آفتاب به سرعت سبب انهدام باکتری می شود. در برابر خشکی هم تا دو ماه مقاومت دارد. این بیماری عفونی ممکن است به شکل تک گیر یا همه گیر درآید. در فصل پائیز و تابستان بیشتر به صورت همه گیر در می آید، این همه گیری بخصوص در اجتماعات مانند سربازخانه ها و مدارس دیده می شود. این بیماری در ایران در تمام فصول فراوان است ولی در تابستان و پائیز بیشتر است. در زمان جنگ که تمام این شرایط موجود است شیوع این بیماری چشمگیر است. افراد تازه وارد که به مناطق آلوده وارد می شوند بیشتر از افراد بومی مبتلا می گردند، کم شدن ترشح اسید معده هم در آلودگی مؤثر است زیرا اسید معده خاصیت از بین بردن این باکتری را دارد. مدل های ریاضی در تحلیلی رفتار دینامیکی این بیماری مانند بیماری های ایدز، مالاریا، هپاتیت و ... مورد استفاده قرار می گیرد. این مقاله، به بررسی جواب های عددی یک مدل ریاضی بیماری تیفوئید که یک دستگاه معادلات دیفرانسیل غیر خطی می باشد، می پردازیم. روش گالرکین ماتریسی، یک روش طیفی با دقت بالا برای تقریب مسائل فیزیکی و مدل های ریاضی می باشد که در آن، جواب های مساله با پایه های استاندارد از بعد متناهی تقریب زده

^۱حجت اله عبادیزاده: ebadizadeh.h@gmail.com

می شود.



شکل ۱: مدل اپیدمی تب تیفوئید

۲ مدل ریاضی

در این بخش به معرفی مدل ریاضی رفتار دینامیکی بیماری تب تیفوئید (شکل ۱) می پردازیم. این مدل ریاضی، توسط پیرلسن^۱ در سال ۱۹۸۹ بصورت یک دستگاه معادلات دیفرانسیل غیر خطی بصورت زیر معرفی شد:

$$\left\{ \begin{array}{l} \frac{dP}{dt} = \tau S - (\mu + \omega)P - \mu S \\ \frac{dS}{dt} = \Lambda - S(\beta_1 I_u + \beta_2 I_e + \beta_3 T) + \omega P - (\tau + \mu)S \\ \frac{dI_u}{dt} = (1 - \alpha)S(\beta_1 I_u + \beta_2 I_e + \beta_3 T) - (\mu + \delta_1 + \phi_1 + \gamma)I_u \\ \frac{dI_e}{dt} = \alpha S(\beta_1 I_u + \beta_2 I_e + \beta_3 T) - (\mu + \delta_2 + \phi_2)I_e + \gamma I_u \\ \frac{dT}{dt} = \phi_1 I_u + \phi_2 I_e - (\mu + \delta_3 + \varepsilon)T \\ \frac{dR}{dt} = \varepsilon T - \mu R \end{array} \right. \quad (1.2)$$

در اینجا $P(t)$ میزان افراد محافظت شده از عفونت در لحظه t ، $S(t)$ میزان افراد حساس در لحظه t ، $I_u(t)$ میزان افراد بیمار آموزش ندیده در لحظه t ، $I_e(t)$ میزان افراد بیمار آموزش دیده در لحظه t ، $T(t)$ میزان افراد تحت درمان در لحظه t و R میزان افراد احیا شده در لحظه t می باشند. در این مقاله، هدف اصلی ارایه یک روش عددی ماتریسی بر پایه استاندارد $t^1, t^2, \dots, t^N$ (روش گالرکین ماتریسی) برای مدل ریاضی عفونت بیان شده در (۱.۲) می باشد. قرار می دهیم

$$\left\{ \begin{array}{l} A_1 = \tau S - (\mu + \omega)P - \mu S \\ A_2 = \Lambda - S(\beta_1 I_u + \beta_2 I_e + \beta_3 T) + \omega P - (\tau + \mu)S \\ A_3 = (1 - \alpha)S(\beta_1 I_u + \beta_2 I_e + \beta_3 T) - (\mu + \delta_1 + \phi_1 + \gamma)I_u \\ A_4 = \alpha S(\beta_1 I_u + \beta_2 I_e + \beta_3 T) - (\mu + \delta_2 + \phi_2)I_e + \gamma I_u \\ A_5 = \phi_1 I_u + \phi_2 I_e - (\mu + \delta_3 + \varepsilon)T \\ A_6 = \varepsilon T - \mu R \end{array} \right. \quad (2.2)$$

در اینصورت خواهیم داشت:

$$\left| \frac{\partial A_1}{\partial P} \right| = \mu + \omega < \infty, \quad \left| \frac{\partial A_1}{\partial S} \right| = \tau, \quad \left| \frac{\partial A_1}{\partial I_e} \right| = 0, \quad \left| \frac{\partial A_1}{\partial I_u} \right| = 0, \quad \left| \frac{\partial A_1}{\partial T} \right| = 0, \quad \left| \frac{\partial A_1}{\partial R} \right| = 0,$$

بطور مشابه برای $i = 2, \dots, 6$ ، A_i حاصل مشتق های نسبی کراندار خواهد بود، لذا مدل دینامیکی ۱.۲ بنا به قضیه زیر دارای جواب یکتا و پیوسته می باشد.

قضیه ۱.۲. (وجود و یکتایی جواب) برای معادله دیفرانسیل مرتبه اول غیر خطی

$$x'(t) = f(t, x(t)), \quad x(0) = x. \quad (۳.۲)$$

اگر تابع $f(t, x(\cdot))$ در شرط لیپشیتز صدق کند، عبارتی دیگر $|f(t, x(\cdot)) - f(t, y(\cdot))| < L|x - y|$ یا اینکه تغییرات f نسبت به مولفه دوم کراندار باشد، یعنی داشته باشیم $|\frac{\partial f(t, x(\cdot))}{\partial x}| < \infty$. آنگاه مساله مقدار اولیه (۳.۲) دارای جواب پیوسته یکتا خواهد بود. (ر. ک. معادلات دیفرانسیل عادی، محمود حصارکی (۱۳۹۴).

۳ نتایج اصلی

در این قسمت به پیاده سازی روش گالرگین ماتریسی برای حل عددی مساله پرداخته و سپس روش پیشنهادی در قالب یک مثال عددی بررسی می شود.

۱.۳ روش گالرگین ماتریسی

اگر فرض می کنیم $T_N(t) = \sum_{k=0}^N e_k t^k$, $I_{e,N}(t) = \sum_{k=0}^N e_k t^k$, $I_{u,N}(t) = \sum_{k=0}^N u_k t^k$, $S_N(t) = \sum_{k=0}^N s_k t^k$, $P_N(t) = \sum_{k=0}^N p_k t^k$, $R_N(t) = \sum_{k=0}^N r_k t^k$, $T_N(t) = \sum_{k=0}^N T_k t^k$ تقریب هایی از جواب های مساله در فضای توابع با بعد متناهی باشند، آنگاه دارای نمایش ماتریسی بصورت زیر خواهند بود:

$$P_N(t) = \mathbf{E} \mathbf{X}_N(t) \mathbf{P}, \quad S_N(t) = \mathbf{E} \mathbf{X}_N(t) \mathbf{S}, \quad I_{u,N}(t) = \mathbf{E} \mathbf{X}_N(t) \mathbf{I} \mathbf{u}, \quad I_{e,N}(t) = \mathbf{E} \mathbf{X}_N(t) \mathbf{I} \mathbf{e},$$

$$T_N(t) = \mathbf{E} \mathbf{X}_N(t) \mathbf{T}, \quad R_N(t) = \mathbf{E} \mathbf{X}_N(t) \mathbf{R}.$$

بطوریکه

$$\mathbf{P} = [p_0, p_1, \dots, p_N]^T, \quad \mathbf{S} = [s_0, s_1, \dots, s_N]^T, \quad \mathbf{I} \mathbf{u} = [u_0, u_1, \dots, u_N]^T,$$

$$\mathbf{I} \mathbf{e} = [e_0, e_1, \dots, e_N], \quad \mathbf{T} = [T_0, T_1, \dots, T_N]^T, \quad \mathbf{R} = [r_0, r_1, \dots, r_N]^T, \quad \mathbf{E} \mathbf{X}_N(t) = [1, t, \dots, t^N].$$

قضیه ۱.۳. با فرض اینکه ماتریس M و $\hat{\mathbf{S}}$ بصورت

$$\mathbf{M} = \begin{pmatrix} \cdot & \cdot & \cdot & \dots & \cdot \\ \cdot & \cdot & \cdot & \dots & \cdot \\ \cdot & \cdot & \cdot & \dots & \cdot \\ \cdot & \cdot & \cdot & \dots & \cdot \\ \cdot & \cdot & \cdot & \dots & \cdot \end{pmatrix}, \quad \hat{\mathbf{S}} = \begin{pmatrix} s_0 & \cdot & \cdot & \dots & \cdot \\ s_1 & s_0 & \cdot & \dots & \cdot \\ s_2 & s_1 & s_0 & \dots & \cdot \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ s_N & s_{N-1} & s_{N-2} & \dots & s_0 \\ \cdot & s_N & s_{N-1} & \dots & s_1 \\ \cdot & \cdot & s_N & \dots & s_2 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \cdot & \cdot & \cdot & \dots & s_N \end{pmatrix}.$$

باشند، آنگاه خواهیم داشت:

۱.

$$\frac{dP_N}{dt} = \mathbf{P} \mathbf{M} \mathbf{E} \mathbf{X}_N(t), \quad \frac{dS_N}{dt} = \mathbf{S} \mathbf{M} \mathbf{E} \mathbf{X}_N(t), \quad \frac{dI_{u,N}}{dt} = \mathbf{I} \mathbf{u} \mathbf{M} \mathbf{E} \mathbf{X}_N(t),$$

$$\frac{dI_{e,N}}{dt} = \mathbf{I} \mathbf{e} \mathbf{M} \mathbf{E} \mathbf{X}_N(t), \quad \frac{dT_N}{dt} = \mathbf{T} \mathbf{M} \mathbf{E} \mathbf{X}_N(t), \quad \frac{dR_N}{dt} = \mathbf{R} \mathbf{M} \mathbf{E} \mathbf{X}_N(t).$$

۲.

$$S_N I_{u,N} = \mathbf{E} \mathbf{X}_N(t) \hat{\mathbf{S}} \mathbf{I} \mathbf{u}, \quad S_N I_{e,N} = \mathbf{E} \mathbf{X}_N(t) \hat{\mathbf{S}} \mathbf{I} \mathbf{e}, \quad S_N T_N = \mathbf{E} \mathbf{X}_N(t) \hat{\mathbf{S}} \mathbf{T}.$$

۲.۳ بیان روش

با جایگذاری نمایش های ماتریسی بیان شده در قضیه ۱.۳ در جملات سیستم غیر خطی ۱.۲ خواهیم داشت:

$$\left\{ \begin{array}{l} G_1 := \text{PMEX}_N(t) - \tau \text{EX}_N(t)S + (\mu + \omega) \text{EX}_N(t)P + \mu \text{EX}_N(t)S \\ G_2 := \text{SMEX}_N(t) - \Lambda + \beta_1 \text{EX}_{1N}(t)\hat{S}\text{Iu} + \beta_2 \text{EX}_{1N}(t)\hat{S}\text{Ie} + \beta_3 \text{EX}_{1N}(t)\hat{S}\text{T} - \omega \text{EX}_N(t)P - (\tau + \mu) \text{EX}_N(t)S \\ G_3 := \text{IuMEX}_N(t) - (1 - \alpha)(\beta_1 \text{EX}_{1N}(t)\hat{S}\text{Iu} + \beta_2 \text{EX}_{1N}(t)\hat{S}\text{Ie} + \beta_3 \text{EX}_{1N}(t)\hat{S}\text{T}) + (\mu + \delta_1 + \phi_1 + \gamma) \text{EX}_N(t)\text{Iu} \\ G_4 := \text{IeMEX}_N(t) - \alpha(\beta_1 \text{EX}_{1N}(t)\hat{S}\text{Iu} + \beta_2 \text{EX}_{1N}(t)\hat{S}\text{Ie} + \beta_3 \text{EX}_{1N}(t)\hat{S}\text{T}) + (\mu + \delta_1 + \phi_1) \text{EX}_N(t)\text{Ie} - \gamma \text{EX}_N(t)\text{Iu} \\ G_5 := \text{TMEX}_N(t) - \phi_1 \text{EX}_N(t)\text{Iu} - \phi_2 \text{EX}_N(t)\text{Ie} + (\mu + \delta_2 + \varepsilon) \text{EX}_N(t)\text{T} \\ G_6 := \text{RMEX}_N(t) - \varepsilon \text{EX}_N(t)\text{T} + \mu \text{EX}_N(t)\text{R} \end{array} \right. \quad (۱.۳)$$

بنا به روش طیفی گالرکین به ازای $k = 0, 1, \dots, N-1$ خواهیم داشت:

$$\langle G_1, t^k \rangle = \langle G_2, t^k \rangle = \langle G_3, t^k \rangle = \langle G_4, t^k \rangle = \langle G_5, t^k \rangle = \langle G_6, t^k \rangle = 0. \quad (۲.۳)$$

و از طرفی برای شرط های اولیه خواهیم داشت:

$$P_N(0) = p, \quad S_N(0) = s, \quad I_{u,N}(0) = u, \quad I_{e,N}(0) = e, \quad T_N(0) = T, \quad R_N(0) = r.$$

در نتیجه با محاسبه ضرب های داخلی (۲.۳) و استفاده از شرط های اولیه یک دستگاه معادلات جبری غیرخطی حاصل می شود، در نتیجه با حل این دستگاه مقادیر بردارهای مجهول P, S, Iu, Ie, T, R حاصل می شوند.

۴ کران خطا

در این بخش، یک کران بالا برای خطای مطلق برای جواب تقریبی بدست آمده از روش پیشنهادی بر اساس درجه همواری جواب مساله ارائه می دهیم. ابتدا یک کران بالا برای خطای مطلق $P_N(t)$ تعیین می کنیم. تابع خطا را بصورت

$$e_{N,P}(t) = P(t) - P_N(t)$$

تعریف می کنیم. از آنجاییکه جواب تقریبی $P_N(t)$ بعنوان یک بهترین تقریب برای $P(t)$ از یک زیر فضای با بعد متناهی $\square$ می باشد، لذا طبق قضیه تیلور با باقیمانده لاگرانژ می توان نوشت:

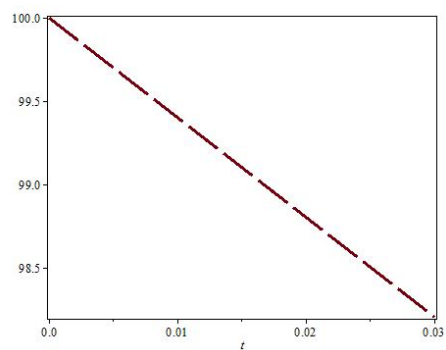
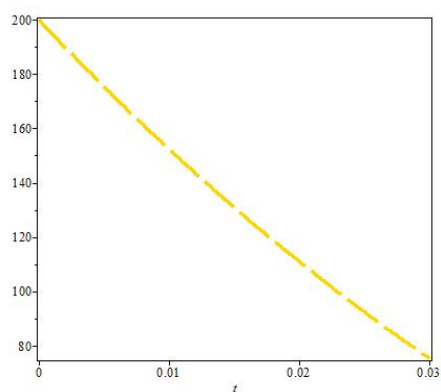
$$|e_{N,P}(t)| = |P(t) - P_N(t)| \leq \frac{t^{N+1}}{(N+1)!} P^{(N+1)}(\zeta) \leq \frac{M_p}{(N+1)!}$$

که در آن $M_p = \max_{t \in [0,1]} |P^{(N+1)}(t)|$ می باشد. بنابراین، مشابه با روند بالا، خواهیم داشت:

$$\begin{aligned} |e_{N,S}(t)| &\leq \frac{M_S}{(N+1)!}, \quad |e_{N,Iu}(t)| \leq \frac{M_{Iu}}{(N+1)!}, \quad |e_{N,Ie}(t)| \leq \frac{M_{Ie}}{(N+1)!}, \\ |e_{N,T}(t)| &\leq \frac{M_T}{(N+1)!}, \quad |e_{N,R}(t)| \leq \frac{M_R}{(N+1)!}. \end{aligned}$$

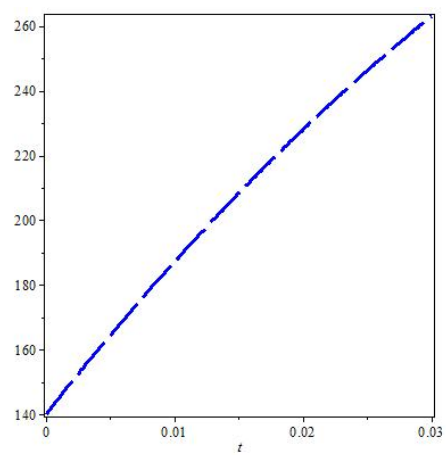
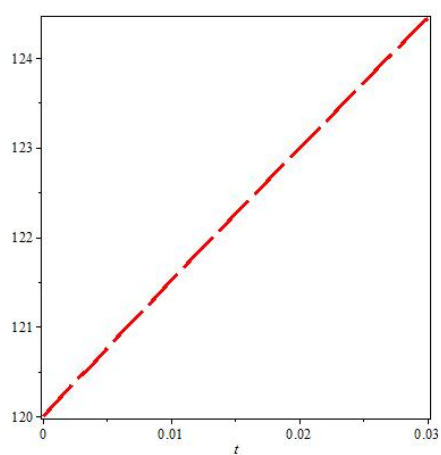
۵ نتایج عددی

مثال ۱.۵. به عنوان مثال شرایط $P(0) = 100, S(0) = 200, Iu(0) = 140, Ie(0) = 120, T(0) = 80, R(0) = 60$ را در نظر می گیریم که در آن مقادیر پارامترها در جدول ۱ داده شده است. نتایج عددی بدست آمده در شکل ها، دقت و کارایی روش پیشنهادی را به ازای $N = 5$ نشان می دهد.



شکل ۲: رفتار تقریبی تابع $P(t)$

شکل ۳: رفتار تقریبی تابع $S(t)$

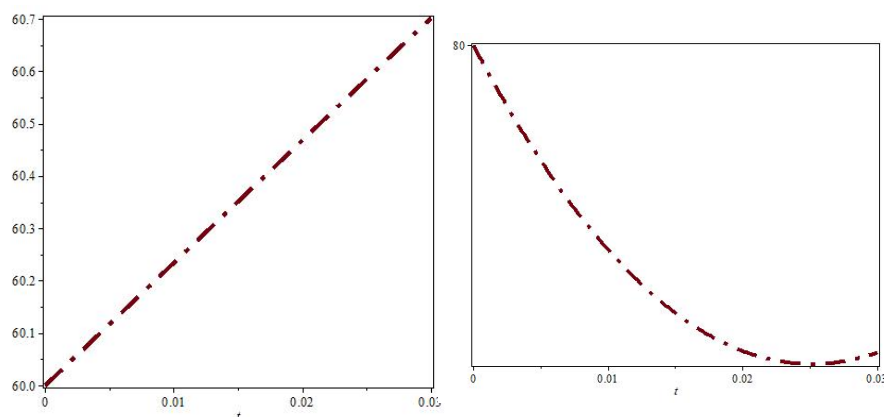


شکل ۴: رفتار تقریبی تابع $Iu(t)$

شکل ۵: رفتار تقریبی تابع $Ie(t)$

جدول ۱: مقادیر پارامترها

$\delta_1, \delta_2, \delta_3$	α	τ	ω	β_3	β_2	β_1	μ	γ	ε	ϕ_2	ϕ_1	Λ
۰۰۷۵.۰	۰۰۷۲.۰	۱.۰	۵.۰	۰۵.۰	۰۵.۰	۵.۱	۱۴۲.۰	۶.۰	۴.۰	۳.۰	۰۴.۰	۲۰۰

شکل ۶: رفتار تقریبی تابع $T(t)$ شکل ۷: رفتار تقریبی تابع $R(t)$

۶ نتیجه گیری

مدل‌های ریاضی مسائل و بررسی رفتار آن‌ها در قالب الگوریتم‌های عددی یکی از ابزارهای مناسب و قابل اعتماد پژوهش‌های علمی-ریاضی و بین رشته‌ای می‌باشد. قابلیت تفسیر رفتارهای دینامیکی سیستم‌های خطی و غیر خطی با زبان الگوریتم ریاضی می‌تواند شامل روش‌های عددی مختلف باشد. در این مقاله، یک روش عددی به زبان ساده ماتریسی برای بررسی مدل بیماری تب تیفوئید معرفی شد که می‌تواند برای سیستم‌های مشابه نیز قابل تعمیم باشد.

مراجع

- [۵۰] غ. قربانی، ع. توانا و ح. پرهیزگار، بررسی اپیدمی تیفوئید در یک واحد نظامی، مجله طب نظامی، ۵ (۱۳۸۲)، ۳۹-۴۲.
- [۵۱] م. ملکوتیان و همکاران، بررسی میزان شیوع و مرگ و میر ناشی از بیماری‌های منتقله از آب آشامیدنی و غذا مجله ارتقای ایمنی و پیشگیری از مصدومیت‌ها، ۳ (۱۳۹۴)، ۱۵-۲۴.
- [۵۲] م. حصارکی، و. رومی، معادلات دیفرانسیل عادی، (۱۳۹۴).

بررسی مدل ریاضی شیوع یک عامل بیماری

حجت‌الله عبادی‌زاده^۱، دانشکده علوم پایه، گروه ریاضی، دانشگاه افسری امام علی (ع)، علی عزیزی مایوان، دانشکده علوم پایه، گروه ریاضی، دانشگاه افسری امام علی (ع)، سید مهدی کاظمی تربقان، دانشکده علوم پایه، گروه ریاضی، دانشگاه بجنورد

چکیده: در این مقاله به ارائه و بررسی یک مدل ریاضی برای شیوع یک عامل بیماری در جامعه می‌پردازیم. در این مدل افراد جامعه را به چهار دسته‌ی: افراد مستعد دریافت عامل بیماری، افراد آلوده به عامل بیماری، افراد بیماری که تحت درمان قرار گرفته‌اند و افرادی که بهبود یافته‌اند، تقسیم می‌کنیم. سپس دستگاه معادلات دیفرانسیل مدل را ارائه داده و با استفاده از ماتریس نسل دوم عدد شیوع را محاسبه می‌نماییم. در پایان، در کنار تحلیل دینامیکی مدل، به بررسی شرایط پایداری حول نقطه تعادل دستگاه معادلات می‌پردازیم.

واژه‌های کلیدی: عدد شیوع، ماتریس نسل دوم، مدل ریاضی بیماری. **رده‌بندی ریاضی (۲۰۱۰):** $92D30$, $34D23$

۱ مقدمه

ویروس یک عامل عفونی ذره‌بینی است که تنها قادر است درون سلول‌های زنده یک جاندار رشد کرده و تکثیر شود. ویروس‌ها تقریباً در تمام اکوسیستم‌های روی کره زمین یافت می‌شوند و به‌عنوان یکی از عوامل مهم ایجاد بیماری در یک جاندار و انتقال آن به سایر جانداران شناخته می‌شود. از جمله بیماری‌هایی که دارای منشأ ویروسی هستند می‌توان به سرماخوردگی و آنفولانزا اشاره کرد که بیماری‌هایی شایع و تقریباً بی‌خطر هستند، اما بیماری‌های خطرناکی چون ایدز و انواع هپاتیت نیز دارای منشأ ویروسی هستند. شیوع ویروس یک بیماری در جامعه، در صورت فراهم بودن شرایط، می‌تواند به یک همه‌گیری (اپیدمی) منجر شود. اما برای کنترل و جلوگیری از گسترش یک عامل ویروسی در جامعه، فهمیدن مکانیسم و فرایند دینامیکی انتقال ویروس یک امر بسیار مهم است.

در سال ۱۷۶۰، دانیل برنولی اولین مدل ریاضی در زمینه اپیدمیولوژی را ارائه داد که نشان می‌داد واکسیناسیون در برابر بیماری آبله سودمند است [۷]. مدل برنولی زمینه‌ساز ورود مدل‌های ریاضی به مطالعات زیستی و بیولوژیک شد به‌طوری‌که در اوایل قرن بیستم مطالعات گسترده و جامعی بر روی شیوع بیماری‌های عفونی، با استفاده از مدل‌های ریاضی، صورت پذیرفت [۳، ۴]. به‌عنوان مثال، در سال ۱۹۰۶ هامر یک مدل برای بررسی گسترش بیماری سرخک ارائه داد، در سال ۱۹۱۱ یک فیزیكدان به نام راس، با استفاده از معادلات دیفرانسیل، به طراحی یک مدل برای شیوع بیماری مالاریا از طریق پشه پرداخت و در سال ۱۹۲۶ این کرماک و مک کندریک بودند که مدل معروف SIR را برای اپیدمی بیماری طاعون در لندن در سال ۱۶۶۶ ارائه نمودند [۱، ۶، ۷].

در اپیدمیولوژی، مدل‌های ریاضی فهم باارزشی از روند گسترش یک بیماری در میان انسان‌ها ارائه دادند و فرضیه‌هایی برای بهبود

^۱حجت‌الله عبادی‌زاده: ebadizadeh.h@gmail.com

استراتژی کنترل بیماری‌ها مطرح نمودند. امروزه مدل‌های ریاضی در تحلیل دینامیک بیماری‌های زیادی مانند ایدز، مالاریا، هپاتیت و غیره مورد استفاده قرار می‌گیرند. این مدل‌ها نقش مهمی در درک بهتر از الگوهای اپیدمی جهت کنترل بیماری دارند.

در این مقاله یک مدل ریاضی برای شیوع یک عامل بیماری (ویروس) در جامعه ارائه می‌دهیم که در این مدل جمعیت را به چهار دسته تقسیم کرده‌ایم: افرادی که مستعد دریافت عامل بیماری هستند، افرادی که به عامل بیماری آلوده شده‌اند، افراد بیماری که تحت درمان قرار می‌گیرند و کسانی که بهبود یافته‌اند. با توجه به اینکه با گذشت زمان تعداد جمعیت مورد نظر تغییر می‌کند، مدل ما یک مدل وابسته به زمان خواهد بود.

مباحث ارائه شده در این مقاله به این صورت است که: در بخش دوم به مباحث مربوط به مدل‌سازی ریاضی که شامل تعیین متغیرها و پارامترها، فرمول‌بندی معادلات و تشکیل دستگاه معادلات متناظر با آن که در اینجا یک دستگاه معادلات معمولی غیرخطی است پرداخته می‌شود. در بخش سوم نیز نقطه تعادل عاری از بیماری و نقطه تعادل شیوع بیماری ارائه شده و با استفاده از ماتریس نسل دوم مدل ریاضی، فرمول عدد شیوع محاسبه می‌گردد. علاوه بر این، پایداری دستگاه معادلات مدل موردنظر نیز مورد بحث قرار داده می‌شود.

۲ مدل‌سازی ریاضی

با توجه به مطالب گفته شده در بخش قبل، کل جمعیت جامعه مورد مطالعه را در زمان t با $N(t)$ نمایش می‌دهیم. حال اگر $S(t)$ مجموعه افرادی از جامعه باشد که مستعد دریافت عامل بیماری هستند، $I(t)$ مجموعه افرادی که آلوده به عامل بیماری باشند، $T(t)$ مجموعه همه بیمارانی باشد که تحت درمان قرار می‌گیرند و $R(t)$ مجموعه همه کسانی باشد که بهبود یافته‌اند، آنگاه داریم:

$$N(t) = S(t) + I(t) + T(t) + R(t).$$

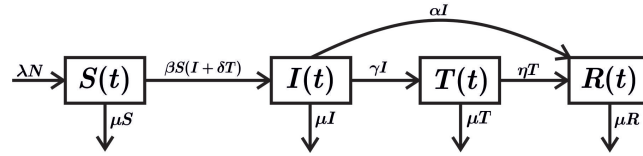
این مدل را به اختصار $SITR$ می‌نامیم. در این مدل، نرخ انتقال بیماری را β و نرخ مرگ و میر طبیعی و زاد و ولد را به ترتیب μ و λ در نظر می‌گیریم. همچنین نرخ انتخاب افراد بیمار برای درمان را γ ، نرخ بهبود بیماران تحت درمان را η و نرخ بهبود بیماران بدون استفاده از خدمات درمانی را α در نظر می‌گیریم. به این ترتیب اگر کسر δ بیانگر تعداد افراد بیمار تحت درمانی باشد که همچنان می‌توانند بیماری را انتقال دهند، آنگاه تعداد افراد جدیدی که در زمان t آلوده می‌شوند برابر خواهد بود با $\beta S(t)(I(t) + \delta T(t))$. حال با توجه به متغیرها و پارامترهای تعریف شده در بالا، به فرمول‌بندی مدل ریاضی $SITR$ می‌پردازیم که منجر به دستگاه معادلات دیفرانسیل غیرخطی زیر می‌شود:

$$\begin{aligned}\dot{S} &= -\beta S(I + \delta T) - \mu S + \lambda N, \\ \dot{I} &= \beta S(I + \delta T) - (\alpha + \gamma + \mu)I, \\ \dot{T} &= \gamma I - (\eta + \mu)T, \\ \dot{R} &= \eta T + \alpha I - \mu R.\end{aligned}\tag{۱.۲}$$

توجه کنید که تمام پارامترها در معادلات بالا نامنفی هستند. در شکل ۱ نمودار گردش این مدل ارائه شده است.

اگر N را برابر با یک در نظر بگیریم، در این صورت مجموعه‌های S ، I ، T و R کسری از N بوده و داریم:

$$S + I + T + R = ۱,$$



شکل ۱: نمودار گردش مدل SITR

بنابراین $R = 1 - S - I - T$. به این ترتیب می‌توان دستگاه معادلات (۱) را به صورت زیر در نظر گرفت:

$$\begin{aligned}\dot{S} &= -\beta S(I + \delta T) - \mu S + \lambda, \\ \dot{I} &= \beta S(I + \delta T) - (\alpha + \gamma + \mu)I, \\ \dot{T} &= \gamma I - (\eta + \mu)T.\end{aligned}\quad (2.2)$$

۳ نقاط تعادل و عدد شیوع

جواب‌های دستگاه معادلات (۲) بصورت $Q = (S, I, T)$ هستند که قضیه زیر مثبت بودن این جواب‌ها را نشان می‌دهد [۵].

قضیه ۱: اگر مقادیر $S(0)$ ، $I(0)$ و $T(0)$ نامنفی باشند، آنگاه برای $t > 0$ مقادیر $S(t)$ ، $I(t)$ و $T(t)$ نامنفی هستند.

اگر فرض کنیم که در جمعیت N هیچ بیماری وجود نداشته باشد، در این صورت نقطه تعادل عاری از بیماری دستگاه (۲) به صورت $Q_0 = (S_0, 0, 0)$ خواهد بود که $S_0 = \frac{\lambda}{\mu}$. سمت راست معادلات دستگاه (۲) را برابر با صفر قرار داده و نقطه تعادل دیگر دستگاه (۲) که به صورت (S^*, I^*, T^*) است را محاسبه می‌کنیم. پس داریم

$$\begin{aligned}S^* &= \frac{-(\alpha + \gamma + \mu)I^* - \lambda}{\mu}, \\ T^* &= \frac{\gamma I^*}{\eta + \mu}.\end{aligned}$$

عدد شیوع بیماری، تعداد افرادی است که یک فرد بیمار، در طول دوره بیماری خود، آنها را مبتلا می‌کند. با توجه به [۲]، عدد شیوع R_0 در واقع همان شعاع طیفی ماتریس نسل دوم FV^{-1} است. ماتریس‌های F و V از دستگاه (۲) به صورت زیر بدست می‌آوریم:

$$F = \begin{bmatrix} \frac{\beta\lambda}{\mu} & \frac{\beta\lambda\delta}{\mu} \\ \cdot & \cdot \end{bmatrix},$$

و

$$V = \begin{bmatrix} \alpha + \gamma + \mu & \cdot \\ -\gamma & \eta + \mu \end{bmatrix}.$$

پس داریم:

$$V^{-1} = \frac{1}{(\alpha + \gamma + \mu)(\eta + \mu)} \begin{bmatrix} \eta + \mu & \gamma \\ \cdot & \alpha + \gamma + \mu \end{bmatrix},$$

که در نهایت

$$R_* = \frac{\beta\lambda}{\mu(\alpha + \gamma + \mu)}.$$

در ادامه، با استفاده از ژاکوبین دستگاه معادلات (۲)، پایداری سیستم را حول نقطه تعادل عاری از بیماری، یعنی $(\frac{\lambda}{\mu}, 0, 0)$ ، بررسی می‌کنیم.

قضیه ۲: نقطه تعادل عاری از بیماری بطور مجانبی پایدار است اگر $R_* < 1$ و ناپایدار است اگر $R_* > 1$. اثبات. ماتریس ژاکوبی دستگاه (۲) را به صورت زیر محاسبه می‌کنیم:

$$J_* = \begin{bmatrix} -\mu & \frac{-\beta\lambda}{\mu} & \frac{-\beta\lambda\delta}{\mu} \\ \cdot & \alpha & \frac{\beta\lambda\delta}{\mu} \\ \cdot & \gamma & -(\eta + \mu) \end{bmatrix}.$$

پس چندجمله‌ای مشخصه J_* به صورت $(x + \mu)(x^2 + bx - c)$ خواهد بود که

$$b = \eta + \alpha + \gamma + 2\mu - \frac{\beta\lambda}{\mu}$$

$$c = \left(\frac{\beta\lambda}{\mu} - (\alpha + \gamma + \mu)\right)(\eta + \mu) - \frac{\beta\lambda\gamma\delta}{\mu}.$$

با توجه به [۲]، در صورتی که جواب‌های چندجمله‌ای مشخصه دارای بخش حقیقی منفی باشند، آنگاه دستگاه پایدار مجانبی خواهد بود و این در صورتی اتفاق می‌افتد که $R_* < 1$. همچنین اگر $R_* > 1$ آنگاه دستگاه ناپایدار خواهد بود. $\square$

مراجع

- [53] F. Brauer, P. Van den, J. Wu, *Mathematical Epidemiology*, Springer-Verlag Berlin, 2008.
- [54] P.V.D. Driessche and J. Watmough, *Reproduction numbers and subthreshold endemic equilibria for compartmental models of disease transmission*, Math Biosci. 180 : 29 – 48, 2002
- [55] M. Farkas, *Dynamical Models in Biology*, Academic Press, 2001.
- [56] Z. Ma, J. Li, *Dynamical Modeling and Analysis of Epidemics*, World Scientific Publishing Co. Pte. Ltd., 2009.
- [57] I. Shaban, E. S. Massaw, O. Daniel Makinde, *Modelling the effect of screening on the spread of HIV infection in a population with variable inflow of infective immigrants*, Scientific Research and Essays, 6 : 4397 – 4405, 2011.
- [58] E. Vynnycky, R. G. White, *An Introduction to Infectious Disease Modelling*, Oxford: Oxford University Press, 2010.
- [59] X. Q. Zhao, *Dynamical Systems in Population Biology*, Springer-Verlag New York, 2003.

یک مدل ریاضی برای هپاتیت B و بررسی اثر واکسیناسیون در رشد بیماری

حجت‌الله عبادی‌زاده^۱، دانشکده علوم پایه، گروه ریاضی، دانشگاه افسری امام علی (ع)، علی عزیزی مایوان، دانشکده علوم پایه، گروه ریاضی، دانشگاه افسری امام علی (ع)، سید مهدی کاظمی تربقان، دانشکده علوم پایه، گروه ریاضی، دانشگاه بجنورد

چکیده: این مقاله یک مدل ریاضی برای بیماری هپاتیت B ارائه می‌شود. در این مدل جامعه مورد مطالعه بر اساس نوع فعالیت افراد که ممکن است منجر به ابتلای آنها به هپاتیت B گردد، به دو دسته تقسیم می‌شود که نرخ واکسیناسیون در این دو دسته می‌تواند متفاوت باشد. در پایان، فرمول عدد شیوع بر اساس پارامترهای مدل محاسبه می‌شود.

واژه‌های کلیدی: بیماری هپاتیت B ، عدد شیوع، ماتریس نسل دوم.

۱ مقدمه

در سال ۱۷۶۰، دانیل برنولی اولین مدل ریاضی در زمینه اپیدمیولوژی را ارائه داد که نشان می‌داد واکسیناسیون در برابر بیماری آبله سودمند است [۱۲]. مدل برنولی زمینه‌ساز ورود مدل‌های ریاضی به مطالعات زیستی و بیولوژیک شد. در اپیدمیولوژی، مدل‌های ریاضی فهم بالارزشی از روند گسترش یک بیماری در میان انسان‌ها ارائه دادند و فرضیه‌هایی برای بهبود استراتژی کنترل بیماری‌ها مطرح نمودند. امروزه مدل‌های ریاضی در تحلیل دینامیک بیماری‌ها زیادی مانند ایدز، مالاریا، هپاتیت و غیره مورد استفاده قرار می‌گیرند. این مدل‌ها نقش مهمی در درک بهتر از الگوهای اپیدمی جهت کنترل بیماری دارند.

هپاتیت B یک عفونت کبدی خطرناک است که ناشی از ویروس HBV از خانواده هپادنا ویریده می‌باشد. عفونت B ممکن است کوتاه مدت باشد (هپاتیت B حاد) یا مدت زمان زیادی طول بکشد (هپاتیت B مزمن). در برخی بیماران که عفونت هپاتیت B مزمن شده، ممکن است منجر به نارسایی کبدی، سرطان کبد یا سیروز شود (سیروز کبدی بیماری است که باعث زخم و سخت شدن پایدار بافت کبدی می‌شود). راه‌های انتقال این بیماری همانند بیماری ایدز است با این تفاوت که شیوع هپاتیت B در دنیا در حدود ۵۰ تا ۶۰ برابر بیشتر از شیوع بیماری ایدز است. آمارهای جهانی نشان می‌دهد که نزدیک به ۲ میلیارد انسان در سرتاسر دنیا آلوده به ویروس هپاتیت B هستند و سالانه ۵۰ میلیون نفر به این تعداد اضافه می‌شوند. طبق تخمین‌هایی که زده می‌شود، سالانه ۶۰۰۰۰۰ نفر نیز از عوارض ناشی از این بیماری جان خود را از دست می‌دهند [۱، ۹].

اصلی‌ترین راه‌های انتقال HBV روابط جنسی، انجام تزریقات با سرنگ و سوزن آلوده و انتقال از مادر به جنین است. به این ترتیب، معنادارترین تزریقی که از سرنگ و سوزن مشترک استفاده می‌کنند، کارکنان بخش‌های مراقبت‌های بهداشتی، کادر درمان و کسانی که

^۱حجت‌الله عبادی‌زاده: ebadizadeh.h@gmail.com

با خون انسان در تماس هستند و افرادی که روابط جنسی محافظت نشده دارند، در معرض خطر ابتلا به هپاتیت B قرار دارند. البته فرد می‌تواند با یک واکسن از بروز بیماری پیشگیری کند. اگر فردی مبتلا به هپاتیت B باشد، باید اقدامات لازم را به کار گیرد تا از انتقال این بیماری به دیگران پیشگیری کند. آمارهای گردآوری شده در چند سال اخیر نشان می‌دهد که در ایران شیوع هپاتیت B کمتر از ۳ درصد است [۴]. برنامه‌هایی که در ایران برای کنترل و کاهش شیوع بیماری هپاتیت B اجرا شدند، از جمله برنامه آگاهی‌بخشی عمومی نسبت به هپاتیت B و راه‌های انتقال آن و همچنین طرح واکسیناسیون که از سال ۱۹۹۳ برای نوزادان، کارکنان مراکز بهداشتی و سایر افراد در معرض خطر صورت پذیرفت، باعث شدند که درصد شیوع هپاتیت B در ایران از ۵/۳ درصد در سال ۱۹۹۰ به ۱۴/۲ درصد در سال ۲۰۰۰ برسد [۲].

تحقیقات و مقالات بسیاری در زمینه اپیدمیولوژی بیماری هپاتیت B وجود دارد که در آنها از مدل‌های ریاضی استفاده شده است. بعنوان نمونه می‌توان به [۳] اشاره کرد که در آن از یک مدل ساده ریاضی برای بررسی اثر حامل‌های ویروس در انتقال ویروس HBV استفاده شده است. در [۵] از مدل‌های ریاضی برای بهبود استراتژی ریشه‌کنی هپاتیت B در نیوزلند کمک گرفته شده است. استفاده از مدل‌های ریاضی در بررسی دینامیک انتقال و شیوع هپاتیت B در چین [۱۱] و در بررسی اثر واکسیناسیون در کنترل هپاتیت B [۶] نیز از جمله کاربردهای مدل‌های ریاضی در تحقیقات مربوط به اپیدمیولوژی هپاتیت B است. چند محقق ایرانی نیز اخیراً در [۱۰] به ارائه یک مدل ریاضی برای بررسی دینامیک شیوع هپاتیت B در بخشی از استان کرمان پرداخته‌اند.

در این مقاله یک مدل ریاضی برای بیماری هپاتیت B ارائه می‌دهیم که در این مدل جمعیت را به شش دسته تقسیم کرده‌ایم: افرادی که مستعد بیماری هستند، افرادی که در دوره نهفتگی بیماری قرار دارند، افرادی که تحت واکسیناسیون قرار می‌گیرند، افرادی که دچار هپاتیت B حاد هستند، افرادی که دچار هپاتیت B مزمن هستند و کسانی که بهبود یافته‌اند. همچنین کل جمعیت را به دو دسته‌ی کسانی که معناد تزریقی هستند و یا به طریقی با تزریقات و خون انسان سروکار دارند و کسانی که این‌گونه نیستند، تقسیم می‌کنیم. با توجه به اینکه با گذشت زمان تعداد جمعیت مورد نظر تغییر می‌کند، مدل ما یک مدل وابسته به زمان خواهد بود.

مباحث ارائه شده در این مقاله به این صورت است که: مطالب مربوط به مدل‌سازی ریاضی شامل تعیین متغیرها و پارامترها، فرمول‌بندی معادلات و تشکیل دستگاه معادلات متناظر با آن که در اینجا یک دستگاه معادلات معمولی غیرخطی است، در بخش دوم ارائه می‌شود. در بخش سوم، ماتریس نسل دوم مدل ریاضی در نقطه تعادل عاری از بیماری ارائه شده و فرمول عدد شیوع محاسبه می‌گردد.

۲ مدل‌سازی ریاضی

در این بخش به ارائه مدل ریاضی هپاتیت B برای دو گروه متمایز از افراد می‌پردازیم. گروه اول معتادان تزریقی، کارکنان بخش بهداشت و کسانی هستند که به طریقی با تزریقات و خون انسان سروکار دارند که این گروه را N_b می‌نامیم. گروه دوم نیز سایر افراد جامعه هستند، یعنی کسانی که متعلق به N_b نباشند، که گروه دوم را N_n می‌نامیم. کل جمعیت مورد نظر را با فرض مهاجرت ثابت، در زمان t ، $N(t)$ در نظر می‌گیریم و داریم $N(t) = N_b(t) + N_n(t)$.

گروه N_b را به شش دسته تقسیم می‌کنیم: $S_b(t)$: افرادی که مستعد بیماری هستند، $E_b(t)$: افراد آلوده به ویروس HBV که بیماری آنها در مرحله نهفتگی بیماری قرار دارد، $I_b(t)$: افرادی که دچار بیماری هپاتیت B حاد گشته‌اند، $C_b(t)$: افرادی که بیماری هپاتیت B آنها مزمن شده است، $V_b(t)$: افرادی که تحت واکسیناسیون قرار گرفته‌اند، $R_b(t)$: افرادی که بهبود یافته‌اند. به همین ترتیب گروه N_n را نیز به شش دسته $S_n(t)$ ، $E_n(t)$ ، $I_n(t)$ ، $C_n(t)$ ، $V_n(t)$ و $R_n(t)$ تقسیم می‌کنیم. به عبارت دیگر:

$$N_b(t) = S_b(t) + E_b(t) + I_b(t) + C_b(t) + V_b(t) + R_b(t),$$

$$N_n(t) = S_n(t) + E_n(t) + I_n(t) + C_n(t) + V_n(t) + R_n(t).$$

این دسته‌ها متغیرهای مدل هستند. این نکته قابل ذکر است که طرح واکسیناسیون هپاتیت B برای نوزادان بطور کامل انجام می‌شود و این‌که واکسیناسیون به معنی مصونیت کامل از هپاتیت B نیست. در ادامه پارامترهای مؤثر در مدل را برای گروه N_b بصورت زیر ارائه می‌دهیم:

β : نرخ انتقال بیماری،

p_b : نرخ واکسیناسیون افراد مستعد بیماری،

α_b : نرخ واکسیناسیون افرادی که در مرحله نهفتگی بیماری قرار دارند،

t : نرخ جابجایی افراد از مرحله E_b به مرحله I_b ،

γ : نرخ خارج شدن از مرحله I_b ،

r : نرخ جابجایی از مرحله C_b به R_b ،

δ : احتمال بیمار شدن افرادی که تحت واکسیناسیون قرار گرفته‌اند،

η : نرخ جابجایی از V_b به R_b ،

μ : نرخ مرگ و میر طبیعی،

λ : نرخ زاد و ولد.

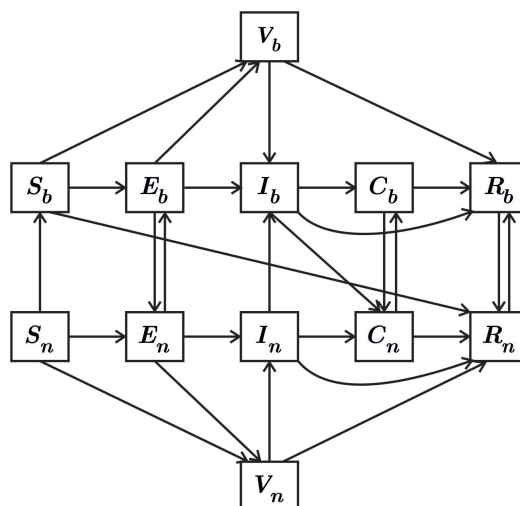
تمام پارامترهای بالا برای گروه N_n نیز برقرارند بجز نرخ‌های واکسیناسیون که برای گروه N_n این نرخ‌ها p_n و α_n خواهند بود. اگر $I = I_b + I_n$ و $C = C_b + C_n$ باشند، θ را احتمال جابجایی افراد از I به C در نظر می‌گیریم و همچنین فرض می‌کنیم که u نرخ جابجایی افراد از گروه N_n به گروه N_b باشد.

حال با توجه به متغیرها و پارامترهای تعریف شده در بالا، به فرمول‌بندی مدل ریاضی بیماری هپاتیت B می‌پردازیم که منجر به دستگاه معادلات دیفرانسیل غیرخطی زیر می‌شود:

$$\begin{aligned}
 \dot{S}_b &= -\beta(1 - p_b)S_b(I + C) - (\mu + p_b)S_b + uS_n \\
 \dot{E}_b &= \beta(1 - p_b)S_b(I + C) - (\alpha_b + (1 - \alpha_b)t + \mu)E_b + uE_n \\
 \dot{I}_b &= (1 - \alpha_b)tE_b - (\mu + \gamma(1 - \theta) + \theta)I_b + \beta\delta(I + C)V_b + uI_n \\
 \dot{C}_b &= \theta I_b - (\mu + r)C_b + uC_n \\
 \dot{V}_b &= \alpha_b E_b + p_b S_b - (\beta\delta(I + C) + \eta)V_b + uV_n \\
 \dot{R}_b &= \gamma(1 - \theta)I_b + \eta V_b + rC_b - \mu R_b + uR_n \\
 \dot{S}_n &= -\beta(1 - p_n)S_n(I + C) - (u + \mu + p_n)S_n \\
 \dot{E}_n &= \beta(1 - p_n)S_n(I + C) - (\alpha_n + (1 - \alpha_n)t + \mu + u)E_n + \lambda(E_b + E_n) \\
 \dot{I}_n &= (1 - \alpha_n)tE_n - (\mu + \gamma(1 - \theta) + \theta + u)I_n + \beta\delta(I + C)V_n \\
 \dot{C}_n &= \theta I_n - (\mu + r + u)C_n + \lambda(I + C) \\
 \dot{V}_n &= \alpha_n E_n + p_n S_n - (\beta\delta(I + C) + \eta + u)V_n \\
 \dot{R}_n &= \gamma(1 - \theta)I_n + \eta V_n + rC_n - (\mu + u)R_n + \lambda(R_b + R_n + S_b + S_n)
 \end{aligned} \tag{۱.۲}$$

در شکل ۱ نمودار گردشی این مدل ارائه شده است.

توجه داشته باشید، در این مدل فرض بر این است که هر نوزاد متولد شده از مادر مبتلا به هپاتیت B حاد و یا مزمن، دچار هپاتیت B مزمن می‌شود. همچنین نرخ جابجایی از N_n به N_b را می‌توان همان نرخ گرایش افراد گروه N_n به اعتیاد تزریقی در نظر گرفت.



شکل ۱: نمودار گردش مدل ریاضی بیماری هپاتیت B

۳ نقطه تعادل رهایی از بیماری و عدد شیوع

جواب‌های دستگاه معادلات (۱) بصورت $Q = (S_b, E_b, I_b, C_b, V_b, S_n, E_n, I_n, C_n, V_n)$ هستند که قضیه زیر مثبت بودن این جواب‌ها را نشان می‌دهد [۷].

قضیه ۱: اگر مقادیر $S_b(\cdot), E_b(\cdot), I_b(\cdot), C_b(\cdot), V_b(\cdot), S_n(\cdot), E_n(\cdot), I_n(\cdot), C_n(\cdot), V_n(\cdot)$ نامنفی باشند، آنگاه برای $t > 0$ مقادیر $S_b(t), E_b(t), I_b(t), C_b(t), V_b(t), S_n(t), E_n(t), I_n(t), C_n(t), V_n(t)$ نامنفی هستند.

اگر فرض کنیم که در جمعیت N هیچ بیماری وجود نداشته باشد، در این صورت نقطه تعادل عاری از بیماری دستگاه (۱) به صورت $Q_0 = (S_{b,0}, 0, 0, 0, V_{b,0}, S_{n,0}, 0, 0, 0, V_{n,0})$ خواهد بود.

عدد شیوع بیماری، تعداد افرادی است که یک فرد بیمار، در طول دوره بیماری خود، آنها را مبتلا می‌کند. با توجه به [۷]، عدد شیوع R_0 درواقع همان شعاع طیفی ماتریس نسل دوم FV^{-1} است. ماتریس‌های F و V را از دستگاه (۱) به صورت زیر بدست می‌آوریم:

$$F = \begin{bmatrix} \cdot & \beta(1-p_b)S_{b,0} & \beta(1-p_b)S_{b,0} & \cdot & \beta(1-p_b)S_{b,0} & \beta(1-p_b)S_{b,0} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \beta(1-p_n)S_{n,0} & \beta(1-p_n)S_{n,0} & \cdot & \beta(1-p_n)S_{n,0} & \beta(1-p_n)S_{n,0} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{bmatrix},$$

$$V = \begin{bmatrix} \alpha_b + (1 - \alpha_b)t + \mu & \cdot & \cdot & -u & \cdot & \cdot \\ - (1 - \alpha_b)t & m & -\beta\delta V_b & \cdot & -\beta\delta V_b - u & -\beta\delta V_b \\ \cdot & -\theta & \mu + r & \cdot & \cdot & -u \\ -\lambda & \cdot & \cdot & \alpha_n + (1 - \alpha_n)t + \mu - \lambda & \cdot & \cdot \\ \cdot & -\beta\delta V_n & -\beta\delta V_n & - (1 - \alpha_n)t & n & -\beta\delta V_n \\ \cdot & -\lambda & -\lambda & \cdot & -\theta - \lambda & \mu + r + u - \lambda \end{bmatrix},$$

که $m = \mu + \gamma(1 - \theta) + \theta - \beta\delta V_b$ و $n = \mu + \gamma(1 - \theta) + \theta + u - \beta\delta V_n$ است. شعاع طیفی ماتریس

FV^{-1} به صورت زیر خواهد بود:

$$R_1 = \beta(1 - p_n)S_n \left[\left(\frac{1 - \alpha_n t}{bx} \right) \left(1 - z + \frac{\lambda + \theta}{f} \right) + \left(\frac{1 - n_1 - n_2 - n_3}{bn_1} \right) \left(\frac{-y(1 - \alpha_n t)}{x} + u(1 - \alpha_b t) \right) \right] \\ + \beta(1 - p_b)S_b \left[\left(\frac{1 - \alpha_n t}{bx} \right) \left(\frac{\lambda}{(1 - t)\alpha_b + t + \mu} \right) \left(1 - z + \frac{\lambda + \theta}{f} \right) \right. \\ \left. + \left(\frac{1 - n_1 - n_2 - n_3}{n_1} \right) \left(\frac{\lambda u(1 - \alpha_b t)}{b((1 - t)\alpha_b + t + \mu)} - \frac{\lambda y(1 - \alpha_n t)}{b((1 - t)\alpha_b + t + \mu)} + g \right) \right].$$

$$d = c = \beta\delta V_n \left(\frac{\mu + r - \theta}{\theta} \right), b = (1 - t)\alpha_n + t + \mu - \lambda(1 + u), a = \gamma(1 - \gamma)\theta + \mu - \beta\delta V_b, \\ g = \frac{u\lambda}{\theta} + \mu + r + u - \lambda, f = \frac{-\lambda(\mu + r + \theta)}{\theta}, e = \left(\frac{\mu + r - \theta}{\theta} \right) (\gamma(1 - \gamma)\theta + \mu - \beta\delta V_b), \gamma(1 - \gamma)\theta + u - \beta\delta V_n, \\ n_1 = \frac{h}{x} - \frac{cg}{fx}, z = \left(-\frac{(\mu + r)}{\theta} \right) \frac{(\theta + \lambda)}{f}, y = \frac{e(\theta + \lambda)}{f} \beta\delta V_b - u, x = d + \frac{c(\theta + \lambda)}{f}, h = \beta\delta V_n \left(\frac{u - \theta}{\theta} \right) \\ n_2 = \frac{g}{f} \left(\frac{\mu + r}{\theta} \right) + \frac{u}{\theta} + zn_1, n_3 = \frac{g}{f} + \left(\frac{\lambda + \theta}{f} \right) n_1, n_4 = -\beta\delta V_b - \frac{eg}{f} + yn_1$$

مراجع

- [60] R. Akbari, A. Vahidian, A. A. Heydari, A. Heydari, *The Analysis of a Disease-free Equilibrium of Hepatitis B model*, Sahand communication in mathematical analysis, 3(2) : 1 – 11, 2016.
- [61] S. Alavian, B. Hajarizadeh, M. Ahmadzad-Asl, A. Kabir, K. Bagheri-Lankarani, *Hepatitis B Virus Infection in Iran: A Systematic Review*, Hepatitis Monthly, 8 : 281 – 294, 2008.
- [62] R.M. Anderson and R.M. May, *Infectious Disease of Humans: Dynamics and Control*, Oxford University Press, Oxford, 1991.
- [63] H. Keyvani, M. Sohrabi, F. Zamani, H. Poustchi, H. Ashrafi, F. Saeedian, M. Mooadi, N. Motamed, H. Ajdarkosh, M. Khonsari, G. Hemmasi, M. Ameli, A. Kabir, M. Khodadost, *A Population Based Study on Hepatitis B Virus in Northern Iran*, Amol, Hepat Mon, 14 : 1 – 8, 2014.
- [64] G.F. Medley and N.A. Lindop, *Hepatitis-B virus endemicity: heterogeneity, cat-astrophic dynamics and control*, Nature Medicine, 7(5) : 619 – 624, 2001.
- [65] J. Pang, J.A. Cui, and X. Zhou, *Dynamical behavior of a hepatitis B virus transmission model with vaccination*, Journal of Theoretical Biology, 265 : 572 – 578, 2010.

- [66] J. Shaban, S. S. Massaw, O. Daniel Makinde, *Modelling the effect of screening on the spread of HIV infection in a population with variable inflow of infective immigrants*, Scientific Research and Essays, 6 : 4397 – 4405, 2011.
- [67] P. Van den Driessche and J. Watmough, *Reproduction numbers and sub-threshold endemic equilibria for compartmental models of disease transmission*, Math. Biosci., 180(1) : 29 – 48, 2002.
- [68] K. Wang, A. Fan, A. Torres, *Global Properties of an Improved Hepatitis B Virus Model*, Non-linear analysis: Real World applications, 11(4) : 3131 – 3138, 2010.
- [69] T. Waezizadeh, M. Mohammad Rezaei, *Dynamical Behavior of HBV in a Population*, Cankaya University Journal of Science and Engineering, 14 : 29 – 41, 2017.
- [70] S.J. Zhao, Z.Y. Xu, and Y. Lu, *A mathematical model of hepatitis B virus transmission and its application for vaccination strategy in China*, Int.J.Epidemiol., 29 : 744 – 752, 2000.
- [71] X. Q. Zhao, *Dynamical Systems in Population Biology*, Springer-Verlag New York, 2003.

کاربرد رگرسیون لجستیک جریمه شده با لاسوی تطبیقی در تحلیل داده‌های با بعد بالا
نجمه عسگری فرسنگی^۱، کارشناس ارشد آمار زیستی، گروه اپیدمیولوژی و آمار زیستی، دانشگاه علوم پزشکی مشهد، مشهد، ایران
راضیه یوسفی کارشناس ارشد آمار زیستی، گروه اپیدمیولوژی و آمار زیستی، دانشگاه علوم پزشکی مشهد، مشهد، ایران
محمد تقی شاکری استاد آمار زیستی، گروه اپیدمیولوژی و آمار زیستی، دانشگاه علوم پزشکی مشهد، مشهد، ایران

چکیده: با تغییر سبک زندگی و افزایش شیوع سرطان، تشخیص اولیه بیماری از اهمیت زیادی برخوردار می‌باشد. در راستای شناسایی ژن‌های افتراق دهنده‌ی بیماران از افراد سالم در مطالعات سرطان‌شناسی، الگوهای بیان ژن در دو بافت سالم و سرطانی با یکدیگر مقایسه می‌شوند که به دلیل بالا بودن بعد داده‌های بیان ژنی استفاده از روش‌های معمول آماری از جمله رگرسیون لجستیک مناسب نیست. به این منظور از روش‌های رگرسیونی اصلاح شده استفاده می‌شود. رگرسیون لجستیک جریمه شده با لاسو یکی از محبوب‌ترین روش‌ها است که انتخاب متغیر و برآورد احتمالات را به‌طور همزمان انجام می‌دهد. لاسو تمامی متغیرها را به یک اندازه جریمه می‌کند در صورتی که دارای درجه اهمیت یکسانی نیستند. در این راستا رگرسیون لجستیک جریمه شده با لاسوی تطبیقی در جهت اصلاح جریمه‌ی لاسو معرفی گردید که وزن‌های متفاوتی برای ضرایب در عبارت جریمه در نظر می‌گیرد و در نتیجه نتایج لاسو را بهبود می‌بخشد.

واژه‌های کلیدی: رگرسیون لجستیک جریمه شده، لاسوی تطبیقی، داده‌های با بعد بالا

۱ مقدمه

با پیشرفت تکنیک‌هایی با فناوری بالا و تولید حجم عظیم داده در بیولوژی، به کار بردن روش‌هایی برای تحلیل این داده‌ها در جهت پی بردن به سیستم پیچیده پویای حیات مورد توجه است. در بین فناوری‌های بسیار پیشرفته‌ی اخیر، ریزآرایه‌ی^۱ دی‌ان‌ای (DNA)^۲ از مشهورترین‌ها است که بیان هزاران ژن را به‌طور همزمان در بافت موردنظر اندازه‌گیری می‌کند. در مطالعات سرطان‌شناسی، الگوهای بیان ژن در دو بافت سالم و سرطانی با یکدیگر مقایسه می‌شوند. غالباً به دلیل هزینه نسبتاً بالا، آزمایش‌های ریزآرایه‌ای معمولاً در تعداد اندکی نمونه انجام می‌شوند.

انتخاب متغیر در شرایطی که تعداد متغیرها (p) بسیار بیشتر از تعداد نمونه‌ها (n) است (داده‌های با بعد بالا) در روش‌های آماری جدید

^۱Microarray

^۲Deoxyribonucleic acid

و یادگیری ماشین از اهمیت ویژه‌ای برخوردار بوده و استفاده از روش‌های آماری در طبقه‌بندی داده‌های با بعد بالا (شرایط $p > n$) کاری چالش برانگیز است (۷۲).

یکی از محبوب‌ترین مدل‌های خطی تعمیم یافته، برای تحلیل رابطه‌ی یک یا چند متغیر توضیحی بر متغیر پاسخ دو یا چندتایی رگرسیون لجستیک است که برای دستیابی به نتایج پایا^۳ و قابل قبول به حجم داده‌هایی بیش از آنچه که در رگرسیون خطی معمولی مواجه هستیم نیاز می‌باشد (۷۳). از دیدگاه آماری نیز وجود تعداد زیادی متغیر، منجر به بیش‌برازشی شده و عملکرد طبقه‌بندی را کاهش می‌دهد (۷۴)، (۷۵).

مدل‌های رگرسیونی جریمه شده در این شرایط کارا هستند که عبارات جریمه مختلفی در مقالات مورد بررسی قرار گرفته‌اند (۷۶) - (۷۸). یکی از محبوب‌ترین عبارات جریمه، عملگر انتخاب و انقباض کمترین قدرمطلق (LASSO) ^۴ است که توسط تیشیرانی در سال ۱۹۹۶ برای رگرسیون خطی معرفی شد (۷۷). لاسو با انتخاب متغیرهای کمتر در داده‌های با بعد بالا مدل‌های ساده، تفسیرپذیر و دارای صحت پیش‌بینی بالا ارائه می‌دهد.

جریمه کردن تمام متغیرها به یک اندازه در صورتی که دارای درجه‌ی اهمیت یکسانی نیستند یکی از ضعف‌های لاسو است. به منظور بهبود نتایج لاسو، لاسوی تطبیقی^۵ را پیشنهاد کرد که وزن‌های متفاوتی برای ضرایب در نظر می‌گیرد و نشان داد اگر وزن‌ها وابسته به داده‌ها بوده، لاسوی وزن‌دار شده دارای خاصیت الهام بخشی شده، ارببی در انتخاب^۶ کاهش یافته و مدل پایدارتری حاصل می‌شود (۷۹). در این مقاله رگرسیون لجستیک جریمه شده با لاسوی تطبیقی معرفی شده و کارایی روش از نظر تعداد متغیرهای انتخابی و صحت پیش‌بینی مدل در مجموعه داده‌ی ریزآرایه بررسی گردید.

۲ رگرسیون لجستیک جریمه شده با لاسوی تطبیقی

رگرسیون لجستیک یکی از روش‌های محبوب آماری در پزشکی است که در تحلیل‌ها به‌طور گسترده استفاده می‌شود ولی شواهدی وجود دارد که در داده‌های باحجم نمونه کم و تعداد متغیر زیاد نتایج پایایی نداشته و کند عمل می‌کند. به منظور استفاده از رگرسیون لجستیک در داده‌های با بعد بالا، با اضافه کردن عبارتی نامنفی تحت عنوان جریمه به منفی لگاریتم تابع درست‌نمایی نمونه، مدل قابل استفاده می‌شود.

عبارات جریمه متفاوتی در متون مورد بررسی قرار گرفته‌اند که یکی از محبوب‌ترین عبارات جریمه، جریمه l_1 معروف به لاسو است که این عبارت جریمه قابل تعمیم به مدل‌های خطی تعمیم یافته، از جمله رگرسیون لجستیک نیز است. علی‌رغم مزیت لاسو برای استفاده در داده‌های با بعد بالا، یکی از ضعف‌های آن در نظر گرفتن جریمه‌ی یکسان برای تمام متغیرها می‌باشد (۸۰). به منظور بهبود نتایج لاسو، لاسوی تطبیقی را پیشنهاد کرد (۷۹). لاسوی تطبیقی وزن‌های متفاوتی برای ضرایب مختلف در نظر گرفته و برآورد ضرایب مدل از رابطه (۱.۲) به‌دست می‌آید.

$$\hat{\beta}_{APLR} = \underset{\beta}{argmin} \left[- \sum_{i=1}^n \{y_i \ln(\pi_i) + (1 - y_i) \ln(1 - \pi_i)\} + \lambda \sum_{j=1}^p w_j |\beta_j| \right] \quad (1.2)$$

در رابطه (۱.۲) $(|\beta_j|)^{-\gamma}$ بردار وزن‌های معلوم بوده و γ یک مقدار مثبت ثابت فرض می‌شود. زو نشان داد، اگر وزن‌ها وابسته به داده‌ها باشند لاسوی وزن‌دار شده دارای خاصیت الهام بخشی می‌شود یعنی احتمال انتخاب مجموعه‌ای از متغیرها به‌درستی، همگرا به یک است که در نتیجه ارببی در انتخاب کاهش یافته و مدل پایدارتری به‌دست می‌آید (۸۱). زو از معکوس برآوردهای حداکثر

³ Stable

⁴Least Absolute Shrinkage and Selection Operator

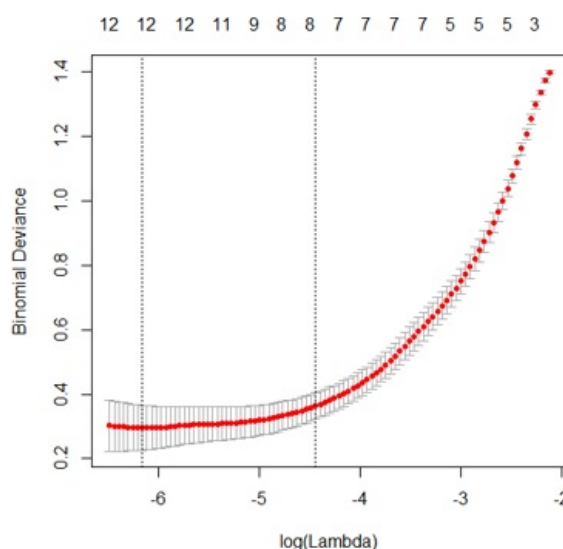
⁵Adaptive LASSO

⁶Selection bias

درست‌نمایی به‌عنوان وزن اولیه استفاده کرد، که در داده‌های با بعد بالا معتبر نیستند (۷۹). بوهلمن و ون‌دگیر به‌منظور به‌کارگیری لاسوی تطبیقی در مجموعه داده‌های با بعد بالا از برآوردهای لاسو برای وزن اولیه لاسوی تطبیقی استفاده کردند (۸۲).

۳ مثال کاربردی در داده‌های واقعی

به منظور بررسی کارکرد لاسوی تطبیقی در داده‌های با بعد بالا، از داده‌های ریزآرایه سرطان روده بزرگ شامل ۲۲۲۷۷ متغیر (ژن) و حجم نمونه ۱۱۱ با کد GSE44861 در پایگاه عمومی GEO^۷ استفاده شد. داده‌ها در محیط نرم‌افزاری R3.5.1 فراخوانی و با استفاده از بسته‌ی آماری glmnet، مدل لاسوی تطبیقی به داده‌ها برازش داده‌شد. به این منظور ابتدا مدل رگرسیون لجستیک جریمه شده با لاسو به داده‌ها برازش داده شد سپس از معکوس ضرایب مدل لاسو به عنوان وزن اولیه در لاسوی تطبیقی استفاده گردید. پارامتر لاندای مناسب برای برازش مدل، مقداری است که به‌ازای آن مدل برازش شده دارای کمترین خطا باشد. با استفاده از اعتبارسنجی متقاطع پارامتر لاندای بهینه برای برازش مدل تعیین و مدل نهایی برازش داده شد و تنها ۱۲ متغیر توسط مدل انتخاب گردید (شکل ۱).



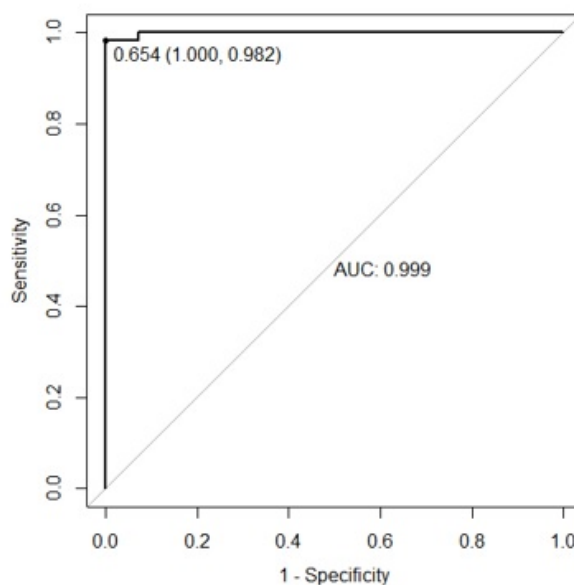
شکل ۱: اعتبارسنجی متقاطع مدل لاسوی تطبیقی

صحت مدل لاسوی تطبیقی با استفاده از سطح زیر منحنی مشخصه عملکرد (ROC) (۰/۹۹۹، ۰/۹۹۶)، حساسیت و ویژگی مدل نیز به ترتیب ۰/۹۸ و یک بود (شکل ۲).

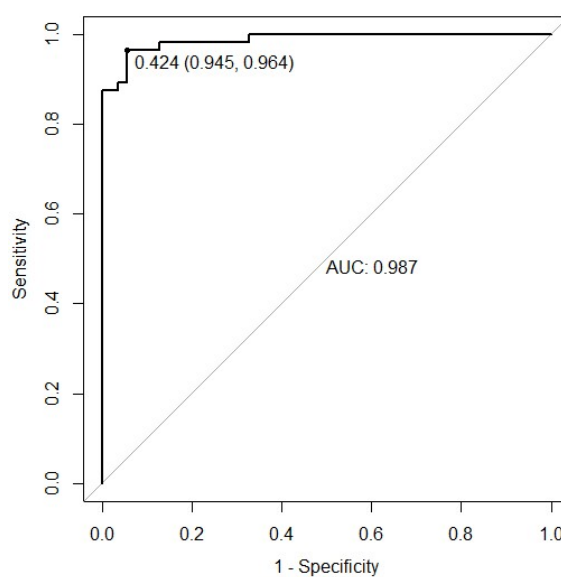
۴ بحث

کاربرد مهم داده‌های ریزآرایه شناسایی ژن‌های مرتبط با سرطان است که به دلیل بالا بودن بعد این داده‌ها، روش‌های رگرسیونی معمول کارایی لازم را ندارند. به منظور استفاده از رگرسیون لجستیک برای طبقه بندی افراد به دو دسته بیمار و سالم در داده‌های با بعد بالا از روش‌های رگرسیونی جریمه شده استفاده می‌شود. عبارت جریمه‌ی لاسو به دلیل برآورد ضرایب و انتخاب متغیر به طور همزمان از محبوب‌ترین‌ها است. لاسو دارای معایبی می‌باشد که زو در راستای بهبود بخشیدن به نتایج آن لاسوی تطبیقی را معرفی کرد.

⁷Gene Expression Omnibus (<http://www.ncbi.nlm.nih.gov/geo>)



شکل ۲: منحنی مشخصه عملکرد مدل لاسوی تطبیقی



شکل ۳: منحنی مشخصه عملکرد مدل لاسو

زو نشان داد که لاسوی تطبیقی در شرایط $p < n$ ضرایب غیر صفر را با احتمال یک به درستی انتخاب می‌کند (۷۹). هوآنگ و همکاران در مطالعه‌ای عملکرد لاسوی تطبیقی در مقایسه با لاسو را در شرایط $n > p$ با استفاده از شبیه‌سازی و در داده‌های واقعی بررسی کردند و نشان دادند در داده‌های با بعد بالا، لاسو و لاسوی تطبیقی از نظر عملکرد یکسان بوده هرچند لاسوی تطبیقی از لاسو از نظر تعداد متغیرهای انتخابی و صحت پیش‌بینی کمی بهتر عمل می‌کند (۸۳). در مطالعه‌ای الجمال و لی رگرسیون لجستیک جریمه شده با لاسوی تطبیقی با وزن اولیه منطبق بر همبستگی را معرفی کرده و نشان دادند از لاسوی تطبیقی با وزن اولیه برآوردگر لاسو از نظر تعداد متغیرهای انتخابی و صحت پیش‌بینی، مدل بهتر عمل می‌کند، که البته در صورت وجود تنها یک همبستگی کامل بین متغیرها نمی‌توان از این روش استفاده کرد (۸۴). در پژوهشی رئیسی و همکاران به منظور شناسایی ژن‌هایی که نقش مهمی در سرطان مثانه دارند

از دو روش جریمه شده‌ی لاسو و لاسوی تطبیقی استفاده کردند. در این مطالعه‌ی مورد شاهدی ۲۵ بیمار و ۲۳ فرد سالم شرکت داشتند که میزان بیان ۲۲ ژن در آن‌ها اندازه‌گیری شد. ۱۲ ژن با صحت پیش‌بینی ۰/۷۱ در روش لاسو و ۴ ژن با صحت پیش‌بینی ۰/۸۹ با استفاده از روش لاسوی تطبیقی شناسایی شد که صحت پیش‌بینی دو مدل اختلاف معنی داری داشتند ($p\text{-value}=0/009$) (۸۵). نتایج مطالعه حاضر نیز نشان داد که لاسوی تطبیقی در داده‌های با بعد بالا قابلیت انتخاب تعداد کمی متغیر با صحت پیش‌بینی بالا را داشته و مدلی ساده و تفسیرپذیر نتیجه می‌دهد.

مراجع

- [72] Piao Y, Piao M, Park K, Ryu K-H. *An ensemble correlation-based gene selection algorithm for cancer classification with gene expression data*. Bioinformatics (Oxford, England). 2012;28:3306–15.
- [73] Agresti A. *An introduction to categorical data analysis*. 2nd ed. New Jersey: John Wiley & Sons 2007.
- [74] Chen K-H, Wang, K.-J., Wang, K.-M., & Angelia, M.-A. *Applying particle swarm optimization-based decision tree classifier for cancer classification on gene expression data* Appl Soft Comput. 2014;24:773-80.
- [75] Peng H, Fu Y, Liu J, Fang X, Jiang C. *Optimal gene subset selection using the modified SFFS algorithm for tumor classification*. Neural Comput Appl. 2013;23:1531–8.
- [76] Liang Y, Liu C, Luan X-Z, Leung K-S, Chan T-M, Xu Z-B, et al. *Sparse logistic regression with a L1/2 penalty for gene selection in cancer classification*. BMC Bioinformatics. 2013;14:198-210.
- [77] Tibshirani R. *Regression shrinkage and selection via the lasso*. J Royal Statist Soc B. 1996;58(1):267-88.
- [78] Zhenqiu L, Feng J, Guoliang T, Suna W, Fumiaki S, Ming T. *Sparse logistic regression with Lp penalty for biomarker identification*. Stat Appl Genet Mol Biol. 2007;6:1-22.
- [79] Zou H. *The adaptive lasso and its oracle properties*. J Am Stat Assoc. 2006;101(476):1418-29.
- [80] Zou H, Hastie T. *Regularization and variable selection via the elastic net*. J R Stat Soc Series B Stat Methodol. 2005;67(2):301-20.
- [81] Fan J, Li R. *Variable Selection via Nonconcave Penalized Likelihood and its Oracle Properties*. J Am Stat Assoc. 2001;96(456):1348-60.
- [82] Bühlmann P, VanDeGeer S. *Statistics for high-dimensional data: Methods, theory and applications [Internet]*. Switzerland: Springer, Berlin, Heidelberg; 2011. Available from: <https://link.springer.com/book/10.1007/978-3-642-20192-9>.

- [83] Huang J, Ma S, Zhang C-H. *Adaptive lasso for sparse high-dimensional regression models*. Stat Sin. 2008;18:1603-18.
- [84] Algama Z, Lee M. *Regularized logistic regression with adjusted adaptive elastic net for gene selection in high dimensional cancer classification*. Comput Biol Med. 2015;(6 (136-45.
- [85] Raeisi H, Jaberipoor M, Zare N, Hosseini A. *The Role of 22 Genes Expression in Bladder Cancer by Adaptive LASSO*. Iran J Cancer Prev. 2016;9(6).

تحلیلی بر مدل بیماری سرطان مبتنی بر نظریه آشوب

میلاد فهیمی^۱، کارشناسی ارشد ریاضیات کاربردی دانشگاه سمنان کاظم نوری عضو هیات علمی دانشکده ریاضی و علوم کامپیوتر دانشگاه سمنان لیلا ترک زاده عضو هیات علمی دانشکده ریاضی و علوم کامپیوتر دانشگاه سمنان

چکیده: در این مقاله، یک مدل ریاضیاتی سه بعدی مبتنی بر نظریه آشوب-بی نظمی برای بیماری سرطان مورد بررسی و تحلیل قرار گرفته است؛ بدین منظور در راستای واقع گرایانه شدن نتایج حاصل از تحلیل مقدار پارامترها، در تشکیل فرم دیفرانسیلی مدل، از مشتقات مرتبه کسری کاپوتو-لیوویل استفاده شده است و سپس با استفاده از خواص موجود در حسابان کسری، شرایط لازم برای وجود جوابی برای مدل مطرح شده مورد بررسی قرار گرفته است.

واژه‌های کلیدی: مدل سرطان، نظریه آشوب، حسابان کسری، مشتق کسری

۱ مقدمه

مدل های ریاضیاتی در ساختار های مختلف به مطالعه و تحلیل رفتار بیماری ها از قبیل آنفولانزا، ایدز و سرطان پرداخته اند و هر یک جنبه های خاصی از آن بیماری از جمله نحوه کنترل، انتشار و انقراض آن را مورد بررسی قرار داده اند. در همه این مدل ها هدف اصلی و اساسی مطالعات، رسیدن به فهم واقع بینانه از مکانیسم رشد یا کنترل بیماری و پیش بینی رفتار آن است. امروزه با توسعه علوم ریاضی و ظهور مفاهیم جدید در این علم، همواره سعی شده مدل هایی ارائه شود که تطابق بیشتری با واقعیت دارند [۳، ۴]. برای نمونه با استفاده از مفاهیم معادلات دیفرانسیل کسری می توان مدل های واقع گرایانه تری ارائه داد. مبداء حساب و دیفرانسیل کسری به اواخر قرن هفدهم بازمی گردد که نیوتن و لایب نیتز از پیشگامان این حوزه بودند. استفاده از مشتقات کسری در مسائل اپیدمی، رویکردی سودمند و واقع گرایانه برای مطالعه رفتار بیماری ها است که در همین راستا در سال های اخیر پژوهش هایی مبتنی بر نظریه آشوب-بی نظمی برای مطالعه رفتار بیماری ها انجام شده است [۱، ۲، ۳].

^۱میلاد فهیمی: fahimilad@gmail.com

۱.۱ ساختار ریاضیاتی مدل بیماری سرطان:

در این مقاله مدل مطالعاتی رفتار بیماری سرطان به شرح زیر نظر گرفته شده است :

$$\begin{aligned}\dot{x}_1 &= x_1(t)(1 - x_1(t)) - Ax_1(t)x_2(t) - B(x_1(t))x_3(t), \\ \dot{x}_2(t) &= Cx_2(t)(1 - x_2(t)) - Dx_1(t)x_2(t), \\ \dot{x}_3(t) &= E\frac{x_1(t)x_2(t)}{x_1(t) + F} - Gx_1(t)x_3(t) - Hx_3(t),\end{aligned}$$

که در آن

$$A, B, C, D, E, F, G, H$$

پارامترهای مورد بررسی در معادله هستند [۲]. در این مدل سه دسته سلول های درگیر بیماری مورد مطالعه قرار می گیرد که در آن $x_1(t)$ تعداد سلولهای سرطانی در لحظه t و $x_2(t)$ تعداد سلولهای سالم در لحظه t و $x_3(t)$ تعداد سلولهای دارای مصونیت نسبت به بیماری در لحظه t را مشخص می کند . معادله اول میزان تغییر در تعداد سلول های سرطانی در لحظه t را نشان می دهد و معادله دوم میزان تغییر در تعداد سلول های سالم که دارای روند رشد طبیعی هستند را نشان می دهد و معادله سوم میزان آمادگی سیستم ایمنی سلول ها در برابر بیماری را نشان می دهد. در این مقاله به منظور واقع گرایانه شدن بررسی ها و نتایج از عملگرهای مشتق کسری استفاده شده است ؛ بدین منظور به مقدماتی از حسابان کسری می پردازیم.

تعریف ۱.۱. عملگر مشتق کسری کاپوتو - لیوویل به صورت زیر قابل تعریف است [۴] :

$$CFCD_t^\gamma \{f(t)\} = \frac{M(\gamma)}{n - \gamma} \int_0^t f^n(\theta) \exp\left[-\frac{\gamma}{n - \gamma}(t - \theta)\right] d\theta, \quad n - 1 < \gamma \leq n,$$

که در آن $M(\gamma)$ تابع نرمال است به گونه ای که : $M(0) = M(1) = 1$

۲ نتایج اصلی

حال در این بخش مدل مورد مطالعه برای بیماری سرطان را با استفاده از مشتقات کسری از نوع کاپوتو بررسی می کنیم و لذا به اعمال شرایط آن می پردازیم.

۱.۲ ساختار تغییر یافته مدل بیماری سرطان با اعمال عملگرها در حیطه حسابان کسری :

بنابراین شکل تغییر یافته مدل با در نظر گرفتن کرنل های نمایی به صورت زیر قابل ارائه خواهد بود: مدل (۲)

$$\begin{aligned}CFCD_t^\gamma x_1(t) &= x_1(t)(1 - x_1(t)) - Ax_1(t)x_2(t) - Bx_1(t)x_3(t), \\CFCD_t^\gamma x_2(t) &= Cx_2(t)(1 - x_2(t)) - Dx_1(t)x_2(t), \\CFCD_t^\gamma &= E \frac{x_1(t)x_2(t)}{x_1(t) + F} - Gx_1(t)x_2(t) - Hx_2(t),\end{aligned}$$

با شرایط اولیه به صورت :

$$x_1(\cdot)(t) = x_1(\cdot); \quad x_2(\cdot)(t) = x_2(\cdot); \quad x_3(\cdot)(t) = x_3(\cdot).$$

که در آن $CFCD_t$ نشان دهنده مشتق کسری از نوع کاپوتو می باشد. حال معادلات دیفرانسیل در مدل (۲) را به فرم معادلات انتگرال تغییر می دهیم و لذا داریم : مدل (۳)

$$\begin{aligned}x_1(t) - x_1(\cdot) &= CFI_t^\gamma [x_1(t)(1 - x_1(t)) - Ax_1(t)x_2(t) - Bx_1(t)x_3(t)], \\x_2(t) - x_2(\cdot) &= CFI_t^\gamma [Cx_2(t)(1 - x_2(t)) - Dx_1(t)x_2(t)], \\x_3(t) - x_3(\cdot) &= CFI_t^\gamma [E \frac{x_1(t)x_2(t)}{x_1(t) + F} - Gx_1(t)x_2(t) - Hx_2(t)],\end{aligned}$$

پس با اعمال مشتق کسری طبق تعریف (۱) خواهیم داشت:

$$\begin{aligned}\Rightarrow x_1(t) &= x_1(\cdot) + \frac{\gamma(1 - \gamma)}{(\gamma - \gamma)M(\gamma)} [x_1(t)(1 - x_1(t)) - Ax_1(t)x_2(t) - Bx_1(t)x_3(t)] \\&\quad + \frac{\gamma\gamma}{(\gamma - \gamma)M(\gamma)} \int_0^t [x_1(s)(1 - x_1(s)) - Ax_1(s)x_2(s) - Bx_1(s)x_3(s)] ds, \\ \Rightarrow x_2(t) &= x_2(\cdot) + \frac{\gamma(1 - \gamma)}{(\gamma - \gamma)M(\gamma)} [Cx_2(t)(1 - x_2(t)) - Dx_1(t)x_2(t)] \\&\quad + \frac{\gamma\gamma}{(\gamma - \gamma)M(\gamma)} \int_0^t [Cx_2(s)(1 - x_2(s)) - Dx_1(s)x_2(s)] ds, \\ \Rightarrow x_3(t) &= x_3(\cdot) + \frac{\gamma(1 - \gamma)}{(\gamma - \gamma)M(\gamma)} [E \frac{x_1(t)x_2(t)}{x_1(t) + F} - Gx_1(t)x_2(t) - Hx_2(t)] \\&\quad + \frac{\gamma\gamma}{(\gamma - \gamma)M(\gamma)} \int_0^t [E \frac{x_1(s)x_2(s)}{x_1(s) + F} - Gx_1(s)x_2(s) - Hx_2(s)] ds.\end{aligned}$$

بنابراین کرنل ها به صورت زیر قابل تعریف خواهند بود :

$$\begin{aligned}\tau(t, x_1(t)) &= x_1(1 - x_1(t)) - Ax_1(t)x_2(t) - Bx_1(t)x_3(t), \\v(t, x_2(t)) &= Cx_2(t)(1 - x_2(t)) - Dx_1(t)x_2(t), \\\phi(t, x_2(t)) &= E \frac{x_1(t)x_2(t)}{x_1(t) + F} - Gx_1(t)x_2(t) - Hx_2(t).\end{aligned}$$

قضیه ۱.۲. شرایط لیب شیتز برای هر یک کرنل های سه گانه برقرار است .

اثبات. برای کرنل های سه گانه مورد نظر داریم :

$$\begin{aligned}\|\tau(t, x_1(t)) - \tau(t, X_1(t))\| &= \|(x_1(t) - X_1(t))[1 - (x_1(t) - X_1(t))] \\ &\quad - A(x_1(t) - X_1(t))x_2(t) - B(x_1(t) - X_1(t))x_2(t)\|, \\ \|v(t, x_2(t)) - v(t, X_2(t))\| &= C\|(x_2(t) - X_2(t))[1 - (x_2(t) - X_2(t))] \\ &\quad - Dx_1(t)(x_2(t) - X_2(t))\|,\end{aligned}$$

حال با استفاده از نامساوی کوشی داریم :

$$\begin{aligned}\|\phi(t, x_2(t)) - \phi(t, X_2(t))\| &= E \left\| \frac{x_1(t)(x_2(t) - X_2(t))}{x_1(t) + F} \right. \\ &\quad \left. - Gx_1(t)(x_2(t) - X_2(t)) - H(x_2(t) - X_2(t)) \right\|.\end{aligned}$$

$$\begin{aligned}\|\tau(t, x_1(t)) - \tau(t, X_1(t))\| &\leq \|(x_1(t) - X_1(t))[1 - (x_1(t) - X_1(t))] \\ &\quad - A(x_1(t) - X_1(t))x_2(t) - B(x_1(t) - X_1(t))x_2(t)\|, \\ \|v(t, x_2(t)) - v(t, X_2(t))\| &= C\|(x_2(t) - X_2(t))[1 - (x_2(t) - X_2(t))] \\ &\quad - Dx_1(t)(x_2(t) - X_2(t))\|,\end{aligned}$$

$$\begin{aligned}\|\phi(t, x_2(t)) - \phi(t, X_2(t))\| &\leq E \left\| \frac{x_1(t)(x_2(t) - X_2(t))}{x_1(t) + F} \right. \\ &\quad \left. - Gx_1(t)(x_2(t) - X_2(t)) - H(x_2(t) - X_2(t)) \right\|;\end{aligned}$$

همچنین داریم :

$$\begin{aligned}x_1(t) &= \frac{\gamma(1-\gamma)}{(\gamma-\gamma)M(\gamma)}\tau(t, x_{1(n-1)}) + \frac{\gamma\gamma}{(\gamma-\gamma)M(\gamma)}\int_0^t \tau(s, x_{1(n-1)})ds, \\ x_2(t) &= \frac{\gamma(1-\gamma)}{(\gamma-\gamma)M(\gamma)}\tau(t, x_{2(n-1)}) + \frac{\gamma\gamma}{(\gamma-\gamma)M(\gamma)}\int_0^t \tau(s, x_{2(n-1)})ds, \\ x_3(t) &= \frac{\gamma(1-\gamma)}{(\gamma-\gamma)M(\gamma)}\phi(t, x_{3(n-1)}) + \frac{\gamma\gamma}{(\gamma-\gamma)M(\gamma)}\int_0^t \phi(s, x_{3(n-1)})ds.\end{aligned}$$

حال با در نظر گرفتن نرم هر یک از پارامترها و با استفاده از نامساوی مثلثی داریم:

$$\begin{aligned}\|Y_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \|\tau(t, x_{\mathfrak{I}(n-\mathfrak{I})}(t)) - \tau(t, X_{\mathfrak{I}(n-\mathfrak{I})}(t))\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \left\| \int_{\cdot}^t [\tau(s, x_{\mathfrak{I}(n-\mathfrak{I})}(s)) - \tau(s, X_{\mathfrak{I}(n-\mathfrak{I})}(s))] \right\| ds,\end{aligned}$$

$$\begin{aligned}\|\Phi_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \|v(t, x_{\mathfrak{I}(n-\mathfrak{I})}(t)) - v(t, X_{\mathfrak{I}(n-\mathfrak{I})}(t))\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \left\| \int_{\cdot}^t [v(s, x_{\mathfrak{I}(n-\mathfrak{I})}(s)) - v(s, X_{\mathfrak{I}(n-\mathfrak{I})}(s))] \right\| ds,\end{aligned}$$

$$\begin{aligned}\|\Psi_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \|\phi(t, x_{\mathfrak{I}(n-\mathfrak{I})}(t)) - \phi(t, X_{\mathfrak{I}(n-\mathfrak{I})}(t))\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \left\| \int_{\cdot}^t [\phi(s, x_{\mathfrak{I}(n-\mathfrak{I})}(s)) - \phi(s, X_{\mathfrak{I}(n-\mathfrak{I})}(s))] \right\| ds,\end{aligned}$$

که در آن :

$$x_{\mathfrak{I}(n)}(t) = \sum_{m=\cdot}^{\infty} Y_m(t); \quad x_{\mathfrak{I}(n)}(t) = \sum_{m=\cdot}^{\infty} \Phi_m(t); \quad x_{\mathfrak{I}(n)}(t) = \sum_{m=\cdot}^{\infty} \Psi_m(t).$$

از طرفی چون کرنل ها دارای خاصیت لیپ شیتز هستند لذا داریم :

$$\begin{aligned}\|Y_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \|\Delta \|(x_{\mathfrak{I}(n-\mathfrak{I})}(t)) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \Delta_{\mathfrak{I}} \int_{\cdot}^t \|x_{\mathfrak{I}(n-\mathfrak{I})}(s) - X_{\mathfrak{I}(n-\mathfrak{I})}(s)\| ds,\end{aligned}$$

$$\begin{aligned}\|\Phi_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \Delta_{\mathfrak{I}} \|x_{\mathfrak{I}(n-\mathfrak{I})}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \Delta_{\mathfrak{I}} \int_{\cdot}^t \|x_{\mathfrak{I}(n-\mathfrak{I})}(s) - X_{\mathfrak{I}(n-\mathfrak{I})}(s)\| ds,\end{aligned}$$

$$\begin{aligned}\|\Psi_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \Delta_{\mathfrak{I}} \|x_{\mathfrak{I}(n-\mathfrak{I})}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \Delta_{\mathfrak{I}} \int_{\cdot}^t \|x_{\mathfrak{I}(n-\mathfrak{I})}(s) - X_{\mathfrak{I}(n-\mathfrak{I})}(s)\| ds,\end{aligned}$$

$$\begin{aligned}\|Y_n(t)\| &= \|x_{\mathfrak{I}(n)}(t) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \leq \frac{\mathfrak{I}(1-\gamma)}{(\mathfrak{I}-\gamma)M(\gamma)} \|\Delta \|(x_{\mathfrak{I}(n-\mathfrak{I})}(t)) - X_{\mathfrak{I}(n-\mathfrak{I})}(t)\| \\ &\quad + \frac{\mathfrak{I}\gamma}{(\mathfrak{I}-\gamma)M(\gamma)} \Delta_{\mathfrak{I}} \int_{\cdot}^t \|x_{\mathfrak{I}(n-\mathfrak{I})}(s) - X_{\mathfrak{I}(n-\mathfrak{I})}(s)\| ds,\end{aligned}$$

$$\begin{aligned}\|\Phi_n(t)\| &= \|x_{\mathfrak{r}(n)}(t) - X_{\mathfrak{r}(n-1)}(t)\| \leq \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{r}} \|x_{\mathfrak{r}(n-1)}(t) - X_{\mathfrak{r}(n-1)}(t)\| \\ &\quad + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{r}} \int_0^t \|x_{\mathfrak{r}(n-1)}(s) - X_{\mathfrak{r}(n-1)}(s)\| ds,\end{aligned}$$

$$\begin{aligned}\|\Psi_n(t)\| &= \|x_{\mathfrak{r}(n)}(t) - X_{\mathfrak{r}(n-1)}(t)\| \leq \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{d}} \|x_{\mathfrak{r}(n-1)}(t) - X_{\mathfrak{r}(n-1)}(t)\| \\ &\quad + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{e}} \int_0^t \|x_{\mathfrak{r}(n-1)}(s) - X_{\mathfrak{r}(n-1)}(s)\| ds,\end{aligned}$$

□

و لذا اثبات کامل است .

قضیه ۲.۲. بررسی وجود جواب های مدل (۲): مدل شماره (۲) دارای جواب است .

اثبات. حال با فرض این که :

$$x_{\mathfrak{v}}(t) = x_{\mathfrak{v}(n)}(t) - \Theta_{\mathfrak{v}(n)}(t); \quad x_{\mathfrak{r}}(t) = x_{\mathfrak{r}(n)}(t) - \Theta_{\mathfrak{r}(n)}(t); \quad x_{\mathfrak{v}}(t) = x_{\mathfrak{r}(n)}(t) - \Theta_{\mathfrak{r}(n)}(t)$$

$$\begin{aligned}\|x_{\mathfrak{v}}(t) - \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \tau(t, x_{\mathfrak{v}}) - x_{\mathfrak{v}}(\cdot) - \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \int_0^t \tau(s, x_{\mathfrak{v}}(s)) ds\| \\ \leq \|\Theta_{\mathfrak{v}(n)}(t)\| + \left\{ \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{v}} + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{r}t} \right\} \|\Theta_{\mathfrak{v}(n)}(t)\|, \\ \|x_{\mathfrak{r}}(t) - \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} v(t, x_{\mathfrak{r}}) - x_{\mathfrak{r}}(\cdot) - \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \int_0^t v(s, x_{\mathfrak{r}}(s)) ds\| \\ \leq \|\Theta_{\mathfrak{r}(n)}(t)\| + \left\{ \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{r}} + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{r}t} \right\} \|\Theta_{\mathfrak{r}(n)}(t)\|, \\ \|x_{\mathfrak{r}}(t) - \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \phi(t, x_{\mathfrak{r}}) - x_{\mathfrak{r}}(\cdot) - \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \int_0^t \phi(s, x_{\mathfrak{r}}(s)) ds\| \\ \leq \|\Theta_{\mathfrak{r}(n)}(t)\| + \left\{ \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{d}} + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \Delta_{\mathfrak{e}t} \right\} \|\Theta_{\mathfrak{r}(n)}(t)\|,\end{aligned}$$

اگر از دو طرف معادلات فوق حد بگیریم به گونه ای که : $n \rightarrow \infty$

$$\begin{aligned}x_{\mathfrak{v}}(t) &= \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \tau(t, x_{\mathfrak{v}}(t)) + x_{\mathfrak{v}}(\cdot) + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \int_0^t \tau(s, x_{\mathfrak{v}}(s)) ds, \\ x_{\mathfrak{r}}(t) &= \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} v(t, x_{\mathfrak{r}}(t)) + x_{\mathfrak{r}}(\cdot) + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \int_0^t v(s, x_{\mathfrak{r}}(s)) ds, \\ x_{\mathfrak{r}}(t) &= \frac{\mathfrak{r}(1-\gamma)}{(\mathfrak{r}-\gamma)M(\gamma)} \phi(t, x_{\mathfrak{r}}(t)) + x_{\mathfrak{r}}(\cdot) + \frac{\mathfrak{r}\gamma}{(\mathfrak{r}-\gamma)M(\gamma)} \int_0^t \phi(s, x_{\mathfrak{r}}(s)) ds,\end{aligned}$$

□

پس ، مدل (۲) دارای جواب است و لذا اثبات کامل است .

۳ نتیجه گیری

در این مقاله با توجه به اعمال عملگرهای مشتقات کسری بر مدل سه بعدی مطالعه بیماری سرطان و بهره گیری از تعاریف موجود در حیطه حسابان کسری، می توان این گونه استنتاج کرد که مدل سه بعدی بیماری سرطان، با توجه به شرایط تعریف شده برای کرنل های آن، قطعاً دارای جواب است. همچنین بهره گیری از مشتقات مرتبه کسری در این مدل، امکان بررسی دقیق تر تغییر هر یک از پارامترهای اساسی را فراهم می آورد.

مراجع

- [86] S. P. Banks, Analysis and structure of a new chaotic system, *Int. J. Bifurc. Chaos* 9, (1999): 1571–1583.
- [87] K. Itik ; S. P. Banks, Chaos in a three-dimensional cancer model , *Int. J. Bifurc. Chaos* 20, (2010): 71–79 .
- [88] J. Francisco Gómez-Aguilar, M. Guadalupe López , Chaos in a Cancer Model via Fractional Derivatives with Exponential Decay and Mittag-Leffler Law , *Entropy. J.*, (2017) :3-19
- [89] N. Sheikh; Ali. Saqib, M. Khan, A comparative study of Atangana-Baleanu and Caputo-Fabrizio fractional derivatives to the convective flow of a generalized Casson fluid *Eur. Phys. J.* 132 (2017) 10-39

تحلیلی بر مدل تصادفی اپیدمی با نرخ شیوع غیرخطی بیماری

میلاد فهیمی^۱، کارشناسی ارشد ریاضیات کاربردی دانشگاه سمنان، کاظم نوری عضو هیات علمی دانشکده ریاضی و علوم کامپیوتر دانشگاه سمنان، لیلا ترک زاده عضو هیات علمی دانشکده ریاضی و علوم کامپیوتر دانشگاه سمنان ایمیل

چکیده: در این مقاله معادله دیفرانسیل معمولی برای مدل اپیدمی بیماری به صورت (SIQ) مورد بررسی قرار می گیرد و با اعمال نرخ یکسان سازی شیوع و ایجاد شرایط مناسب لم ایتو دو بلین، یک معادله دیفرانسیل تصادفی تحت عنوان مدل تصادفی $(SIQS)$ اپیدمی بیماری تشکیل می شود و سپس وجود جواب و یکتایی آن برای معادله دیفرانسیل تصادفی ارائه شده مورد بررسی قرار می گیرد. در پایان با استفاده از روش شبیه سازی عددی اوپلر ماریاما و برنامه نویسی متلب نمودار نوسانات پارامترها مورد بررسی قرار می گیرد.

واژه های کلیدی: مدل اپیدمی، نرخ شیوع غیرخطی، مدل تصادفی، فرمول ایتو، معادله دیفرانسیل تصادفی

۱ مقدمه

مدل سازی ریاضی در علوم مختلف از جمله مهندسی و شیمی و زیست به دفعات استفاده شده است. در سال های اخیر استفاده از مدل های ریاضی برای مطالعه مسائل گوناگون علوم زیستی مورد توجه پژوهشگران قرار گرفته است. در مسائل اپیدمی، مدل های ریاضی برای تحلیل رفتار بیماری هایی از قبیل آنفولانزا و ایدز و سرطان استفاده شده اند تا جنبه هایی از نحوه انتشار و انقراض بیماری ها را ارائه دهند. یکی از راه های رایج برای مدل سازی مسائل اپیدمی، استفاده از معادلات دیفرانسیل معمولی می باشد. این دسته از مدل ها نخستین بار توسط $Kermak$ و $McKendrick$ در سال ۱۹۷۲ میلادی معرفی شد که در آن افراد مورد مطالعه در ۳ گروه مستعد پذیرش بیماری و بهبود یافته تقسیم می شوند. [۱] در همین راستا مدل هایی از قبیل SIR و $SIRV$ و $SEIR$ و $SEIRS$ به عنوان مدل هایی در حیطه مدلسازی قطعی مطرح و توسعه یافتند [۵-۷]. از آنجا که در مدل سازی قطعی همواره با محدودیت هایی روبرو می شویم که امکان تطابق مدل ها با مسائل واقعی در آن وجود ندارد؛ لذا باید از مدل هایی مبتنی بر معادلات دیفرانسیل تصادفی استفاده شود تا مدل سازی به صورت واقع گرایانه تری انجام شود. در این مقاله ابتدا با تعریف نویز سفید یا اختلالات تصادفی و ایجاد زمینه ظهور حرکت براونی، مدلی تصادفی برای مدل سازی رفتار بیماری ها تعریف می شود.

^۱میلاد فهیمی: fahimimilad@gmail.com

۱.۱ مبانی نظری پژوهش

اگر فرض کنیم که $\mathbb{R}_+^d = \{x \in \mathbb{R}^d : x_i > 0, 1 \leq i \leq d\}$ ، در این صورت معادله دیفرانسیل تصادفی d -بعدی به صورت زیر قابل تعریف خواهد بود:

$$dX(t) = f(t, X(t))dt + g(t, \alpha(t))dB_t, \quad t \geq t_0.$$

با مقدار اولیه $X(t_0) = x_0$ و حرکت براونی m بعدی B_t .

بنابراین عملگر دیفرانسیلی L را به صورت زیر تعریف می کنیم:

$$L = \frac{\partial}{\partial t} + \sum_{i=1}^d f_i(x, t) \frac{\partial}{\partial x_i} + \frac{1}{2} \sum_{i,j=1}^d [g^T(x, t)g(x, t)]_{i,j} \frac{\partial^2}{\partial x_i \partial x_j}$$

حال اگر L بر روی تابع $V \in C^{2,1}(\mathbb{R}^d \times [t_0, \infty); \mathbb{R}_+)$ اعمال شود به گونه ای که:

$$LV(x, t) = V_t(x, t) + V_x(x, t)f(x, t) + \frac{1}{2} \text{trace}[g^T(x, t)V_{xx}(x, t)g(x, t)]$$

که در آن:

$$V_t = \frac{\partial V}{\partial t}, \quad V_x = \left(\frac{\partial V}{\partial x_1}, \dots, \frac{\partial V}{\partial x_d} \right)^T, \quad V_{xx} = \left(\frac{\partial^2 V}{\partial x_i \partial x_j} \right)$$

بنابراین لم ایتو - دوبلین به صورت زیر قابل ارائه خواهد بود:

$$dV(x, t) = LV(x, t)dt + V_x(x, t)g(x, t)dB_t$$

۲.۱ ساختار ریاضیاتی مدل تصادفی اپیدمی

حال معادله دیفرانسیل تصادفی (SIQS) با اعمال نرخ واقع گرایانه شیوع اپیدمی (incidence rate)

$$g(I)S = \frac{SI}{1 + \alpha I}$$

به فرم زیر قابل ارائه خواهد بود: (مدل شماره (۱))

$$\begin{cases} dS(t) = \left[A - \frac{SI}{1 + \alpha I} - dS + \gamma I + PQ \right] dt + \sigma_1(s - s^*)dB_1(t) \\ dI(t) = \left[\frac{SI}{1 + \alpha I} - (\gamma + \delta + d)I \right] dt + \sigma_2(I - I^*)dB_2(t) \\ dQ(t) = \left[\delta I - (P + d)Q \right] dt + \sigma_3(Q - Q^*)dB_3(t) \end{cases}$$

که در آن: A نرخ همبستگی شاخص S به نرخ تولد و مهاجرت و d نرخ مرگ طبیعی و P و γ نشان دهنده متوسط تعداد افراد بهبود یافته هستند و σ نشان دهنده متوسط (نرخ) اشخاص مبتلا به بیماری عفونی است و همچنین B_i به ازای $(i = 1, 2, 3)$ حرکت های براونی استاندارد مستقل هستند که در آن σ_i ها به ازای $(i = 1, 2, 3)$ نویزها (white noise) هستند حال اگر در مدل بیماری فرض کنیم که $\sigma_i = 0$ ($i = 1, 2, 3$) لذا مدل فوق به صورت زیر خواهد بود:

$$\begin{cases} S(t) = A - \frac{SI}{1 + \alpha I} - dS + \gamma I + PQ \\ I(t) = \frac{SI}{1 + \alpha I} - (\gamma + \delta + d)I \\ Q(t) = \delta I - (P + d)Q \end{cases}$$

حال با توجه به روش استخراج عدد مولد عفونت [۳،۲]، در این مدل تصادفی عدد مولد عفونت (Reproduction number) به صورت $R_* = \frac{A}{d(\gamma + \delta + d)}$ خواهد بود.

تعریف ۱.۱. اگر $X(t) = (S(t), I(t), Q(t))$ جواب معادلات مدل (۱) باشد، گوییم این جواب برای مدل به طور تصادفی کراندار است هرگاه به ازای هر $\varepsilon \in (0, 1)$ ، مقدار ثابت $\delta = \delta(\varepsilon) > 0$ وجود داشته باشند به گونه ای که به ازای هر مقدار اولیه $(S(0), I(0), Q(0))$ ، جواب $X(t)$ دارای خاصیت زیر باشد: $\limsup P\{|X(t)| > \delta\} < \varepsilon$.

۲ نتایج اصلی

در این بخش، وجود، یکتایی و کرانداری جواب معادله دیفرانسیل ذکر شده در مدل و نیز شبیه سازی عددی مدل مورد بررسی قرار می گیرد.

۱.۲ بررسی وجود، یکتایی و کرانداری جواب

قضیه ۱.۲. معادله (۱) دارای جواب منحصر به فرد $(S(t), I(t), Q(t))$ می باشد و این جواب با احتمال یک در $\mathbb{R}_+^3$ بسته خواهد بود.

اثبات. تابع V را به صورت $V: \mathbb{R}_+^3 \rightarrow \mathbb{R}_+$ تعریف می کنیم به گونه ای که:

$$V(S, I, Q) = \ln(S(t), I(t), Q(t)) = \ln(S(t)) + \ln(I(t)) + \ln(Q(t))$$

حال با اعمال لم ایتو-دوبلین به صورت زیر داریم:

$$\begin{aligned} \Rightarrow dV(S, I, Q) &= V_t dt + V_S dS + V_I dI + V_Q dQ + \\ &+ \frac{1}{2} V_{SS} dS dS + \frac{1}{2} V_{II} dI dI + \frac{1}{2} V_{QQ} dQ dQ + \\ &+ \frac{1}{2} V_{SQ} dS dQ + \frac{1}{2} V_{SI} dS dI + \frac{1}{2} V_{IQ} dI dQ \end{aligned}$$

لذا با اعمال آن داریم:

$$\begin{aligned} \Rightarrow V_S dS + V_I dI + V_Q dQ &= \frac{S'(t)}{S(t)} + \frac{I'(t)}{I(t)} + \frac{Q'(t)}{Q(t)} = \\ &= \left[\frac{A}{S} - \frac{I}{1 + \alpha I} - \frac{dS}{S} + \frac{\gamma I}{S} + \frac{PQ}{S} \right] dt - \frac{\sigma_1(S - S^*)}{S} dB_1(t) \\ &+ \left[\frac{S}{1 + \alpha I} - (\gamma + \delta + d) \right] dt + \frac{\sigma_2(I - I^*)}{I} dB_2(t) \\ &+ \left[\frac{\delta I}{Q} - (P + d) \right] dt + \frac{\sigma_3(Q - Q^*)}{Q} dB_3(t) \\ &= \left[\frac{A}{S} - \frac{I}{1 + \alpha I} - \frac{dS}{S} + \frac{\gamma I}{S} + \frac{PQ}{S} \frac{S}{1 + \alpha I} - (\gamma + \delta + d) + \frac{\delta I}{Q} \right. \\ &\left. - (P + d) \right] dt - \sigma_1 \left(1 - \frac{S^*}{S} \right) dB_1(t) + \sigma_2 \left(1 - \frac{Q^*}{Q} \right) dB_2(t) \\ &+ \sigma_3 \left(1 - \frac{Q^*}{S} \right) dB_3(t) \end{aligned}$$

همچنین داریم :

$$\begin{aligned} & \Rightarrow \frac{1}{\gamma} V_{SS} dS dS + \frac{1}{\gamma} V_{II} dI dI + \frac{1}{\gamma} V_{QQ} dQ dQ \\ & = \frac{1}{\gamma} \left(\sigma_1 (S - S^*) + \sigma_2 (I - I^*) + \sigma_3 (Q - Q^*) \right) d_t \\ & \Rightarrow dV(S, I, Q) \leq C \end{aligned}$$

که در آن C یک عدد مثبت ثابت می باشد. حال اگر از دو طرف معادله فوق انتگرال نسبت به t بگیریم، لذا با توجه به [۷] نتیجه زیر حاصل می شود:

$$+\infty > V(S(\cdot), I(\cdot), Q(\cdot)) + KT \geq +\infty$$

□

و لذا با وجود تناقض فوق، حکم ثابت است.

قضیه ۲.۲ (بررسی کرانداري). به ازای هر مقدار اولیه $(S(\cdot), I(\cdot), Q(\cdot)) \in \mathbb{R}_+^3$ ، جواب معادله های مدل (۱) به طور تصادفی کراندار است.

اثبات. با فرض این که $V(S, I, Q) = e^t (S^\theta, I^\theta, Q^\theta)$ باشد و $\theta \in (0, 1)$ ، در این صورت با اعمال شرط ایتو-دوبلین، فرم دیفرانسیلی زیر حاصل می گردد:

$$\begin{aligned} dV(S, I, Q) &= e^t \left\{ (S^\theta + I^\theta + Q^\theta) + HS^{\theta-1} \left(A - \frac{\beta SI}{1 + \alpha I} - dS \right. \right. \\ &\quad \left. \left. + \gamma I + PQ \right) + HI^{\theta-1} \left[-\frac{\beta SI}{1 + \alpha I} - (\gamma + \delta + d)I \right] \right. \\ &\quad \left. + \theta Q^{\theta-1} [\delta I - (P + d)Q] + \frac{\theta(\theta-1)}{\gamma} [\sigma_1^\gamma S^\theta (1 - \frac{S^*}{S})^\gamma \right. \\ &\quad \left. + \sigma_2^\gamma I^\theta (1 - \frac{I^*}{I})^\gamma + \sigma_3^\gamma Q^\theta (1 - \frac{Q^*}{Q})^\gamma] \right\} dt \\ &\quad + e^t \theta [\sigma_1 S^\theta (1 - \frac{S^*}{S}) dB_1(t) + \sigma_2 I^\theta (1 - \frac{I^*}{I}) dB_2(t) + \sigma_3 Q^\theta (1 - \frac{Q^*}{Q}) dB_3(t)] \\ &\leq e^t \left\{ \gamma + A - d(S + I + Q) + \frac{\theta(\theta-1)}{\gamma} [\sigma_1^\gamma S^\theta (1 - \frac{S^*}{S})^\gamma \right. \\ &\quad \left. + \sigma_2^\gamma I^\theta (1 - \frac{I^*}{I})^\gamma + \sigma_3^\gamma Q^\theta (1 - \frac{Q^*}{Q})^\gamma] \right\} dt \\ &\quad + e^t \theta [\sigma_1 S^\theta (1 - \frac{S^*}{S}) dB_1(t) + \sigma_2 I^\theta (1 - \frac{I^*}{I}) dB_2(t) + \sigma_3 Q^\theta (1 - \frac{Q^*}{Q}) dB_3(t)] \\ &\leq (\gamma + A) e^t dt + e^t \theta [\sigma_1 S^\theta (1 - \frac{S^*}{S}) dB_1(t) + \sigma_2 I^\theta (1 - \frac{I^*}{I}) dB_2(t) \\ &\quad + \sigma_3 Q^\theta (1 - \frac{Q^*}{Q}) dB_3(t)]. \end{aligned}$$

و لذا با انتگرال گرفتن از دو طرف عبارت فوق به نامساوی زیر دست می یابیم:

$$\begin{aligned} E(V(S, I, Q)) &\leq V(S(\cdot), I(\cdot), Q(\cdot)) + (\gamma + A)(e^t - 1) \\ &\Rightarrow e^{-t} E(V(S, I, Q)) \leq e^{-t} V(S(\cdot), I(\cdot), Q(\cdot)) + \gamma + A \end{aligned} \quad (۱.۲)$$

لذا با تعریف نرم به صورت $|X|$ داریم:

$$|X(t)| = (S^\gamma(t) + I^\gamma(t) + Q^\gamma(t))^{\frac{1}{\gamma}}$$

و لذا داریم:

$$\begin{aligned} |X(t)|^\theta &= (S^\gamma(t) + I^\gamma(t) + Q^\gamma(t))^{\frac{\theta}{\gamma}} \leq \gamma^{\frac{\theta}{\gamma}} \max\{S^\theta, I^\theta, Q^\theta\} \\ &\leq \gamma^{\frac{\theta}{\gamma}} (S^\theta, I^\theta, Q^\theta) \end{aligned}$$

بنابراین با توجه به رابطه (۲):

$$\limsup E|X(t)|^\theta \leq \gamma^{\frac{\theta}{\gamma}} (\gamma + A) < +\infty$$

حال با در نظر گرفتن $\theta = \frac{1}{\gamma}$ ، ثابت $\delta_1 > 0$ وجود دارد به گونه ای که:

$$\limsup E|X(t)|^{\frac{1}{\gamma}} \leq \delta_1$$

بنابراین به ازای هر $\varepsilon > 0$ و قرار دادن $\delta = \frac{\delta_1^\gamma}{\varepsilon}$ در اعمال نامساوی چبیشف، نتیجه می گیریم که:

$$P\{|X(t)| > \delta\} \leq \frac{E|X(t)|^{\frac{1}{\gamma}}}{\delta^{\frac{1}{\gamma}}} = \frac{\delta_1}{\delta^{\frac{1}{\gamma}}} = \varepsilon$$

□

و لذا با توجه به تعریف (۱): جواب $X(t)$ برای مدل (۱) به طور تصادفی کراندار خواهد بود.

۲.۲ شبیه سازی عددی

با استفاده از طرح گسسته سازی اوایلر ماریاما [۴] داریم:

$$S_{k+1} = S_k + [A - \frac{S_k I_k}{1 + \alpha I_k} - dS_k + \gamma I_k + PQ_k] \Delta t + \sigma_1 (S_k - S^*)$$

$$I_{k+1} = I_k + [\frac{S_k I_k}{1 + \alpha I_k} - (\gamma + \delta + d)I_k] \Delta t + \sigma_2 (I_k - I^*) \Delta t$$

$$Q_{k+1} = Q_k + [\delta I_k - (P + d)Q_k] \Delta t + \sigma_3 (Q_k - Q^*) \Delta t$$

که در آن $\Delta w_t \sim \sqrt{\Delta t} Z$ و $Z \sim N(0, 1)$ به ازای $t = 1, 2, \dots, n$ به عبارت دیگر Z متغیر تصادفی گاوسی می باشد.

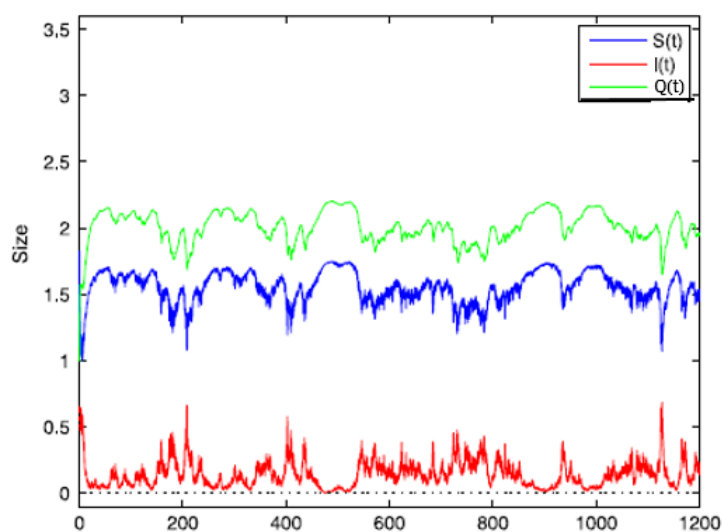
حال با قراردادن مقادیر اولیه به صورت $(S(0), I(0), Q(0)) = (1, 1, 0.5)$ و در نظر گرفتن:

$$A = 1, \quad \alpha = 0.07, \quad d = 0.05, \quad \delta = 0.05, \quad P = 0.03, \quad \gamma = 0.01$$

لذا با توجه به شکل (۱) می توان نتیجه گرفت که به ازای $\sigma_1 = \sigma_2 = \sigma_3 = 0$ ، مدل (۱) به طور تصادفی پایدار بوده و این امر با

تغییر مقادیر پارامترها تغییری نخواهد کرد.

شکل (۱) - نمودار نوسانات پارامترهای اساسی بر حسب زمان



۳ نتیجه گیری

در این مقاله با توجه به مدل تعیینی-قطعی اپیدمی، مدل تصادفی اپیدمی تشکیل شده است. در فرآیند تشکیل این معادله از راهبرد اعمال نرخ شیوع غیرخطی اپیدمی برای واقعی تر شدن پارامترها و نیز با در نظر گرفتن نویزها (نوسانات) و حرکت براونی برای دستیابی به معادلات دیفرانسیل تصادفی استفاده شده است. با توجه به قضایای مطرح شده، وجود، یکتایی و کرانداری تصادفی جواب معادلات موجود در مدل تصادفی استنتاج می شود. همچنین با شبیه سازی عددی مبتنی بر طرح گسسته سازی اوایلر - ماریاما در نرم افزار متلب و تحلیلی نتایج عددی حاصل از ترسیم نمودار نوسانات پارامترهای اساسی مدل، می توان پایداری مدل به ازای تغییرات پارامترها را نتیجه گرفت.

مراجع

- [90] V. Capasso, A generalisation of the Kermack–McKendrick deterministic epidemic model, *Math. Biosci.* 42, (1978), 43-61 .
- [91] O. Diekmann , J. A. P. Heesterbeek , J. A. J. Metz , On the definition and the computation of the basic reproduction ratio R_0 in models for infectious diseases in heterogeneous populations , *Journal of Mathematical Biology* , Volume 28 ,(1990) : 365–382 .
- [92] J. Stoer and R. Bulirsch, *Introduction to Numerical Analysis*, 2nd ed., Springer-Verlag, Berlin: 1993.
- [93] D.J. Higham , An algorithmic introduction to numerical simulation of stochastic differential equations, *SIAM Rev.* 43 , (2001) : 525–546 .
- [94] Q.liu , Q.chen , Analysis of the deterministic and stochastic SIRS epidemic models with nonlinear incidence , *physica A* (2015) : 1-24

- [95] Q.liu , Q.chen , Dynamics of a stochastic SIR epidemic model with saturated incidence , *App.math.comput* 282, (2016) : 155-166
- [96] F.Wei , F.Fangxiang chen , Stochastic permanence of an SIQS epidemic model with saturated incidence and independent random perturbations , *physica* , 4453, (2016):99 - 107 .

نتایجی از C^3 - کدهای خودمکمل ماکسیمال

فریبا فیاضی^۱، دانشکده علوم ریاضی، دانشگاه قم، قم، ایران، سید علیرضا اشرفی دانشکده علوم ریاضی، دانشگاه کاشان، کاشان، ایران، احمد غلامی دانشکده علوم ریاضی، دانشگاه قم، قم، ایران

چکیده: در سال ۲۰۱۴، النافیل و همکارانش، C^3 - کد خودمکمل ماکسیمال را به ۲۷ کلاس هم ارزی شامل ۸ کد افزاز کردند و زیرگروهی از گروه متقارن روی پایه های نوکلئوتید (A ، آدنین، T ، تیمین، C ، سیتوزین و G ، گوانین) ساختند. در این مقاله، ما نمایش جایگشتی و جدول سرشت را برای گروه ساخته شده متناظر با C^3 - کدهای خودمکمل ماکسیمال بدست می آوریم و سپس با استفاده از آن نوع ارتعاشات IR و رامن را بین کدون های مورد نظر مشخص می کنیم.

واژه های کلیدی: C^3 - کد خودمکمل ماکسیمال، نمایش جایگشتی، جدول سرشت، ارتعاشات، کدهای ژنتیکی

۱ مقدمه

اسیدهای نوکلئوتید بسپارهایی (پلیمرهایی) با زنجیر طولانی و وزن مولکولی بالا متشکل از نوکلئوتید ها هستند. هر نوکلئوتید از یک مولکول قند ۵ کربنی و یک مولکول باز نیتروژن دار، تشکیل شده است. در هسته هر سلول مولکول هایی قرار دارند که اساسی ترین اطلاعات حیات را در خود ذخیره کرده اند. این مولکول ها، دئوکسی ریبونوکلئوتید اسید (DNA) نام دارند. DNA از چهار نوع مولکول ساخته شده است، که به آنها "نوکلئوتید" می گوئیم. این چهار نوکلئوتید عبارتند از: آدنین (A)، گوانین (G)، سیتوزین (C) و تیمین (T). این مولکول ها خود به دو زیر گروه تقسیم می شوند: آدنین و گوانین در گروه پورین قرار می گیرند و سیتوزین و تیمین در گروه پیرامیدین جای داده می شوند. نوکلئوتیدها از طریق پیوند هیدروژنی به هم متصل می شوند. سه کدهای ژنتیکی با سه حرف از پایه های اسید نوکلئوتید ساخته می شوند. فرض کنیم $\beta = \{T(U), C, A, G\}$ گروه متقارن S_β روی β به صورت زیر تعریف می شود:

$$S_\beta = \{\pi : \beta \rightarrow \beta \mid \pi \text{ است سویی دو} \}$$

(S_β, o) یک گروه است که در آن o عمل ترکیب روی توابع است. حال به معرفی نگاشت π می پردازیم. منظور از π_{AGCT} نگاشت زیر است:

$$\pi : (A, T, C, G) \rightarrow (G, A, T, C)$$

^۱فریبا فیاضی: fayazifariba@yahoo.com

$$\pi(A) = G, \pi(T) = A, \pi(C) = T, \pi(G) = C.$$

نگاشت دو سویی $\pi: \beta \rightarrow \beta$ برای مجموعه کدون ها (یک دنباله از سه نوکلئوتید در دنباله DNA) نگاشت زیر را القا می کند که آن را نیز با π نشان می دهیم:

$$\pi: \beta^3 \rightarrow \beta^3$$

در (۱۰۱) چهار نگاشت از انتقال های اسید نوکلئوتیدها به صورت زیر تعیین و نام گذاری شده اند:

$$i = I(or id) : (A, T, C, G) \rightarrow (A, T, C, G)$$

$$c = SW : (A, T, C, G) \rightarrow (T, A, G, C)$$

$$p = YR : (A, T, C, G) \rightarrow (G, C, T, A)$$

$$r = KM : (A, T, C, G) \rightarrow (C, G, A, T)$$

که در آن ها اولی نگاشت همانی، دومی نگاشت مکمل، سومی نگاشت پیرامیدین- پورین و چهارمی نگاشت آمینو- کیتو است. برای هر زیرگروه X از β^3 ، ۶۴ کدون در β^3 وجود دارد.

حال می خواهیم گروه متقارن (S_3, o) را بررسی کنیم که نمایش آن به صورت زیر است:

$$S_3 = \{\alpha : \{1, 2, 3\} \rightarrow \{1, 2, 3\} | \alpha \text{ است سویی دو}\}$$

به وضوح هر α یک نگاشت روی مجموعه کدون ها (β^3) القا می کند. به عنوان مثال جایگشت (۱۳۲) کدون (b_1, b_2, b_3) را به (b_2, b_1, b_3) می برد. بنابراین برای هر مجموعه دلخواه $X \subseteq \beta^3$ ، $\pi(X)$ یک تبدیل روی نوکلئوتیدها است در حالی که $\alpha(X)$ یک جایگشت از موقعیت کدون است. همانطور که در زیر نشان خواهیم داد، این تفاوت از دیدگاه بیولوژیکی نقش مهمی دارد. در ادامه روی گروه (A_3, o) به عنوان زیرگروهی از (S_3, o) تمرکز می کنیم که یکی از نمایش های آن به صورت زیر است:

$$A_3 = \{\alpha_0 = (1)(2)(3), \alpha_1 = (123), \alpha_2 = (132)\}.$$

همانطور که در بالا ذکر کردیم گروه S_3 روی مجموعه کدون ها عمل می کند و لذا می توان گفت دو کدون $x_1, x_2 \in \beta^3$ به طور دوری هم ارزند هرگاه نگاشت $\alpha \in A_3$ وجود داشته باشد به طوری که:

$$\alpha(x_1) = x_2$$

و در این حالت می نویسیم $x_1 \sim x_2$. برای مثال، برای کدون دلخواه ATG داریم:

$\alpha_2(ATG) = GAT$ و $\alpha_1(ATG) = TGA$. لذا: $ATG \sim GAT \sim TGA$. در این مقاله از هم ارزی کلاس های تزویجی استفاده خواهیم کرد. واضح است که کدون های AAA, TTT, CCC, GGG به یک کلاس هم ارزی تعلق دارند و بقیه ۲۰ کلاس های هم ارزی نیز هر کدام شامل سه عضو هستند. جایگشت (۱۳)(۲) از S_3 را در نظر بگیرید که عضوی از A_3 نیست. وارون این جایگشت خودش است. این وارون را با $\leftarrow$ نشان می دهیم و جایگشت وارون ساز می نامیم. برای کدون دلخواه $x = (b_1, b_2, b_3) \in \beta^3$ داریم: $x^{\leftarrow} = (b_3, b_2, b_1)$. جایگشت های وارون ساز نقش مهمی در بیولوژی دارند زیرا برخی از پروتئین ها (که معمولاً در میتوکندری هستند)، می توانند در جهت وارون کد گذاری شوند.

در زیست شناسی مولکولی دنباله ها را در سه قالب متفاوت در جهت $3 \rightarrow 5$ می خوانند، هر کدام از یک نوکلئوتید متفاوت در یک سه گانه آغاز می شود. دنباله $ATGACGGTACGATTG...$ را در نظر بگیرید:

قالب ۰: خواندن دنباله از اولین کدون موجود $\{ATG, ACG, \dots\}$.

قالب ۱: در قالب ۰ یک نوکلئوتید را رد می کنیم و بقیه کدون ها را می خوانیم $\{TGA, CCG, \dots\}$.

قالب ۲: در قالب ۰ دو نوکلئوتید را رد می کنیم و بقیه کدون ها را می خوانیم $\{GAC, GGT, \dots\}$.

توزیع ۶۴ نوکلئوتید در سه قالب از ژن ها یکسان نیست و لذا می توان سه مجموعه از نوکلئوتیدها را به صورت X_2, X_1, X_0 مشخص نمود (۹۷).

تعریف ۱.۱. نگاشت مکمل $\pi: \beta^3 \rightarrow \beta^3$ با ضابطه ی $\pi(uv) = \pi(u)\pi(v)$ تعریف می شود که در آن $u \in \beta^n$ و $v \in \beta^{n-1}$ تعریف می شود.

تعریف ۲.۱. یک کلمه از مجموعه $S \subseteq S_\beta^n$ یک کد نامیده می شود هرگاه برای هر $x_1, x_2, \dots, x_n$ و $y_1, y_2, \dots, y_m$ در S و به ازای $n, m \geq 1$ داشته باشیم:

$$x_i = y_j, n = m, \forall i = 1, \dots, n, j = 1, \dots, m.$$

توجه کنید که همواره کدها به صورت خطی خوانده می شوند.

جایگشت تری نوکلئوتید $W_0 = l_1 l_2 l_3$ که در آن $l_i \in \beta$ به صورت زیر است:

$$P(W_0) = W_1 = l_2 l_1 l_3, \quad P(P(w_0)) = P(w_1) = W_2 = l_1 l_2 l_3.$$

مجموعه های $X_0, X_1 = P(X_0)$ و $X_2 = P(X_1)$ شامل ۲۰ کد دایره ای تری نوکلئوتید ماکسیمال هستند. همچنین:

$$C(X_0) = X_0, C(X_1) = X_2, C(X_2) = X_1.$$

تعریف ۳.۱. فرض کنیم $X \subseteq \beta^3$. مجموعه کدون های X را یک کد مستدیر تری نوکلئوتید گوئیم اگر هر کلمه از حروف β که به صورت مستدیر نوشته شده است حداکثر یک تجزیه از کلمات X داشته باشد. کد مستدیر تری نوکلئوتید ماکسیمال است هرگاه برای هر $x \in S_\beta$ و $x \notin X$ ، $x \in X \cup \{x\}$ کد مستدیر تری نوکلئوتید نباشد. همچنین برای کلمات ۳ یا ۴ حرفی می توان گفت کد مستدیر تری نوکلئوتید ماکسیمال است هرگاه دقیقاً شامل ۲۰ کدون باشد یعنی $|X| = 20$.

لازم به ذکر است که، اگر X کد مستدیر باشد آنگاه لزوماً $X_1 = P(X_0)$ و $X_2 = P(X_1)$ کدهای مستدیر نیستند. مثال بعدی این موضوع را ثابت می کند.

فرض کنیم $Y_0 = \{GGC, CCG\}$ کد مستدیر باشد.

$Y_1 = P(Y_0) = \{GCG, CGC\}$ کد مستدیر نیست زیرا دارای دو تجزیه زیر است:

$$w = GCG, CGC, \quad w' = CGC, GCG.$$

تعریف ۴.۱. $X \subseteq \beta^3$ یک C^3 -کد از X است، هرگاه X_2, X_1 کدهای مستدیر باشند به طوری که داشته باشیم:

$$\alpha_1(X) = X_1, \quad \alpha_2(X) = X_2 \quad \forall \alpha \in S_3.$$

بنابر تعریف هر C^3 -کد یک کد مستدیر است.

اگر X C^3 -کد باشد، آنگاه $X_1 = P(X_0)$ و $X_2 = P(X_1)$ کدهای مستدیر هستند. مثال بعدی این موضوع را ثابت می کند.

فرض کنیم $Y_1 = \{CGG, GCC\}$ کد مستدیر باشد اما C^3 -کد نیست. $P(Y_1) = Y_1 = \{GGC, CCG\}$ کد مستدیر است اما $P(P(Y_1)) = Y_2 = \{GCG, CGC\}$ کد مستدیر نیست زیرا دارای دو تجزیه زیر است:

$$w = GCG, CGC, w' = CGC, GCG.$$

برای هر کدون $x = N_1 N_2 N_3$ آنتی کدون x را با $N'_1 N'_2 N'_3 = \overleftarrow{c(x)}$ نشان می دهیم که در آن N'_i ها، مکمل نوکلئوتیدهای N_i هستند یعنی $N'_i = c(N_i)$.

تعریف ۵.۱. $X \subseteq \beta^3$ یک خود مکمل است هرگاه برای هر $x \in X$ آنتی کدون $\overleftarrow{c(x)}$ وجود داشته باشد به طور ی که:

$$x \in X \iff \overleftarrow{c(x)} \in X.$$

قضیه ۶.۱. (۱) همانی و جایگشت های وارون ساز در موقعیت پایه های یک کدون تنها جایگشت هایی هستند که مستدیر بودن را در یک کدون مستدیر $X \subseteq \beta^3$ حفظ می کنند.

(۲) فرض کنید $X \subseteq \beta^3$ یک کد دایره ای تری نوکلئوتید باشد. برای هر $\pi \in S_\beta$ ، $\pi(X)$ نیز یک کد دایره ای تری نوکلئوتید است. همچنین اگر X یک C^3 -کد باشد $\pi(X)$ نیز C^3 -کد است.

قضیه ۷.۱. فرض کنید $\pi \in S_\beta$. آنگاه $\pi oc = co\pi$ اگر و فقط اگر $\pi(X)$ یک کد خود مکمل مستدیر تری نوکلئوتید باشد که در آن $X \subseteq \beta^3$ هر کد خود مکمل مستدیر تری نوکلئوتید دلخواه است.

بنابر قضایای فوق می توان گفت ۲۱۶ C^3 -کد خودمکمل ماکسیمال به ۲۷ کلاس هم ارزی که در هر کلاس ۸ کد وجود دارد، تقسیم می شود. لذا می توان زیرگروه L از S_β را بصورت زیر تشکیل داد.

$$L = \{id, c, p, r, \pi_{AT}, \pi_{CG}, \pi_{ACTG}, \pi_{AGTC}\}$$

گروه L زیرگروه نرمال S_β نیست اما زیرگروهی یکریخت با گروه دو وجهی از مرتبه ۸ است. این زیرگروه را با D_8 نشان می دهیم. در ادامه می خواهیم جدول سرشت متناظر گروه L متناظر با C^3 -کد خودمکمل ماکسیمال را بدست آوریم و با کمک آن نوع ارتعاش بین کدون ها را مشخص کنیم.

۲ نتایج اصلی

در ابتدا مفاهیمی از نظریه سرشت را بیان می کنیم (۱۰۲). یک همریختی از یک گروه به توی گروهی از ماتریس ها و جایگشت ها را به ترتیب یک نمایش ماتریسی یا جایگشتی آن گروه می نامند. در حالت اول تابعی از گروه به توی میدان را که هر عضو را به اثر ماتریس متناظر در نمایش تصویر می کند را سرشت متناظر با آن نمایش می نامند. سرشت χ از یک نمایش جایگشتی ρ تعداد اعضایی است که تحت عمل روی گروه ثابت می مانند. $\chi(1)$ را درجه سرشت گوئیم و سرشت ها از درجه یک را سرشت خطی نامیم. بنابر نتیجه ای معروف سرشت ها روی کلاس های تزویج ثابت هستند و اگر g و h مزدوج باشند آنگاه $\chi_\rho(h) = \chi_\rho(g)$. در حالت کلی تابع $\phi: G \rightarrow \mathbb{C}$ را یک تابع کلاسی گویند هرگاه ϕ روی کلاس های تزویج G ثابت باشد. سرشت متناظر با یک نمایش تحویل ناپذیر را یک سرشت تحویل ناپذیر می نامند. بنابر نتیجه ای معروف در نظریه نمایش سرشت χ تحویل ناپذیر است اگر و تنها اگر $\frac{1}{|G|} \sum_{g \in G} \chi(g) \chi(g^{-1}) = 1$. مجموعه همه سرشت های تحویل ناپذیر گروه را با $Irr(G)$ نشان می دهند. بنابر نتیجه ای معروف، هر سرشت یک گروه را می توان به صورت ترکیبی خطی از سرشت های تحویل ناپذیر نوشت که ضرایب آن اعدادی صحیح و نامنفی هستند. به علاوه $|G| = \sum_{i=1}^k \chi_i(1)^2$. جدول سرشت یک گروه، یک جدول r در r است که در آن r تعداد کلاس های تزویج گروه مورد نظر و سطرها جدول در حقیقت مقادیر سرشت های تحویل ناپذیرند و در بالای هر کدام از ستون های آن نماینده یکی از کلاس های تزویج گروه قرار دارد. قضایای زیر را نیز از مرجع (۱۰۲) داریم:

قضیه ۱.۲. (روابط تعامد کلی [قضیه ۲.۱۳، ص. ۱۹]): برای هر عضو $h \in G$ داریم:

$$\frac{1}{|G|} \sum_{g \in G} \chi_i(gh) \chi_j(g^{-1}) = \delta_{ij} \frac{\chi_i(h)}{\chi_i(1)}.$$

لم ۲.۲. ([لم ۲.۱۵، ص. ۲۰]): فرض کنید که G یک گروه متناهی باشد و χ_i سرشت از G و ρ نمایشی که χ_i از آن حاصل شده است. اگر $g \in G$ و $|G| = n$ آنگاه:

۱. $\rho(g)$ با یک ماتریس قطری به شکل $diag(v_1, \dots, v_f)$ متشابه است.

$$v_i^n = 1.$$

$$\chi(g) = \sum v_i, |\chi(g)| \leq 1.$$

$$\chi(g^{-1}) = \overline{\chi(g)}.$$

قضیه ۳.۲. (رابطه تعامد دوم [قضیه ۲.۱۸، ص. ۲۱]): فرض کنید g و h دو عضو دلخواه از گروه G باشند. در این صورت داریم:

حال فرض کنیم گروه L متناظر با C^3 - کد خود مکمل ماکسیمال باشد. یک نمایش جایگشتی از گروه L روی مجموعه متقارن

S_β^3 ، همریختی $\rho: L \rightarrow S_\beta^3$ است. گروه L دارای پنج کلاس تزویجی زیر است:

$$\{e\}, \{\pi_{AGTC}, \pi_{ACTG}\}, \{\pi_{AT}, \pi_{CG}\}, \{c\}, \{r, p\}$$

لذا دارای پنج سرشت تحویل ناپذیر خواهد بود و می توان جدول سرشت گروه مورد نظر را با استفاده از مفاهیم بیان شده، بدست آورد.

۳ نتیجه گیری

در این مقاله جدول سرشت گروه L متناظر با C^3 - کد خود مکمل ماکسیمال را بدست آوردیم و سپس با کمک آن نوع ارتعاشات بین کدون های مورد نظر را مشخص نمودیم.

مراجع

- [97] Arquès DG, Michel CJ, A complementary circular code in the protein coding genes, *J Theor Biol.*(1996), no. 182, 45–58.
- [98] Benard E, Michel CJ, Transition and transversion on the common trinucleotide circular code, *Comput Biol J ID.* (2013), no. 795418:10.
- [99] Christian J. Michel, The maximal C^3 self-complementary trinucleotide circular code X in genes of bacteria, eukaryotes, plasmids and viruses, *Journal of Theoretical Biology.* (2015), no. 380: 156-177.
- [100] Fimmel E, et al, Circular codes, symmetries and transformations, *J. Math. Biol.* DOI 10.1007/s00285-014-0806-7 (2014).
- [101] Gonzalez D. L., Giannerini S., Rosa R. Strong short-range correlations and dichotomic codon classes in coding DNA sequences, *Phys. Rev. E* (2008) 78(5, ID 051918).
- [102] Isaacs, Irving Martin, Character theory of finite groups, *Academic press, INC* (1976).

مدل‌های ریاضی ایمن درمانی سرطان

فائزه قربانی^۱، سید محسن صالح گروه ریاضی، دانشگاه نیشابور، نیشابور

چکیده: طبق تحقیقات اخیر، امیدوارکننده‌ترین مسیرها در عرصه تحقیقات سرطانی پیش‌گیری از ابتلا به سرطان و ایمن درمانی است. ایمن درمانی در واقع نوع خاصی از شیمی درمانی است. با این تفاوت که در روش شیمی درمانی سیتوتاکسیک، سلول‌های تومور مستقیماً مورد هدف قرار می‌گیرند؛ در حالی که ایمن درمانی سلول‌های تومور را به طور غیرمستقیم و با افزایش کارایی سلول‌های ایمنی بدن از بین می‌برد. هدف از این مقاله، طراحی الگوهای ایمن درمانی و شیمی درمانی سرطان است.

واژه‌های کلیدی: مدل‌سازی ریاضی، سرطان، شیمی درمانی، ایمن درمانی

۱ مقدمه

در این مقاله، برخی از انواع تلاش‌های پژوهشگران در طراحی مدل‌های ریاضی و شبیه‌سازی شیمی درمانی، ایمن درمانی و ترکیبی از آن دو برای سرطان تشریح شده است.

سرطان مهم‌ترین عامل تلفات در جهان است. ارتقای وضعیت واکنش ایمنی با استفاده از تزریق و ارائه مؤثر درمان‌های ایمنی، مسیری امیدوارکننده برای افزایش شانس بقا و همچنین بهبود کیفیت معیشتی بیماران سرطانی است. با توجه به واکنش‌های متفاوت بیماران نسبت به درمان‌های مختلف، ایجاد ابزارهای پیش‌بینی‌کننده برای تعیین بهترین روش درمانی برای هرکدام از بیماران حائز اهمیت است. الگوهای ریاضی می‌توانند به کشف تعاملات بین درمان‌ها و درک واکنش‌ها به دوزدهی کمک کند. هدف نهایی از این مطالعات، ارائه چارچوبی برای عصر جدید در اندازه‌گیری دوزها و همچنین بهینه کردن تعداد دفعات و میزان استفاده از روش‌های درمانی از جمله شیمی درمانی یا پرتو درمانی است.

۲ طراحی الگوهای شیمی درمانی سیتوتاکسیک

در اوایل دهه نود میلادی، مدل ریاضی واکنش لَمفوسیت T سیتوتاکسیک نسبت به رشد تحریک سلول‌های ایمنی در برابر تومور به عنوان مدلی دو جمعیتی مطرح شد. از نتایج بسیار مهم این مدل، ثبت پدیده‌های رکود و فرار تومورها است. پیرو همین الگو، تحقیقاتی

^۱فائزه قربانی: nastaranvafayi74@gmail.com

صورت گرفت که نشان داد روش‌های شیمی درمانی بسیار شدید در کنترل تومور ضعیف عمل می‌کنند و شیمی درمانی بسیار تهاجمی به سیستم ایمنی بدن آسیب وارد می‌کند (۱۰۳).

تحقیقات دیگری نشان می‌دهد که برخی از بیماران واکنش‌های ناهماهنگ و غیرمنتظره‌ای نسبت به درمان‌های سرطانی در مقایسه با روش‌های شیمی درمانی نشان می‌دهند.

برای درک بهتر این اتفاق، الگوی ریاضی تعاملات تومور با سیستم ایمنی و ردیابی سه جمعیت سلولی، سلول‌های تومور، سلول‌های ایمنی و سلول‌های عادی مورد بررسی قرار گرفت. از اکتشافات جالب آنالیز این الگو، آن بود که تنها در هنگام وجود سیستم ایمنی و قدرت کافی آن، تعادل کاملاً عاری از تومور مشاهده می‌شد (۱۰۴؛ ۱۰۵). سپس با افزودن جمعیت چهارم که میزان غلظت دارو در سیستم را نشان می‌دهد، معادله واکنش دارو در بدن معرفی شد که در آن هر معادله جمعیت سلولی با عبارتی برای تشریح آسیب‌های وارده به سلول‌ها توسط شیمی درمانی را داراست.

$$\frac{dE}{dt} = s - dE \frac{rET}{\sigma + T} - c_1 ET - k_1(1 - e^{-u}), \quad (1.2)$$

$$\frac{dT}{dt} = a_1 T(1 - b_1 T) - c_2 ET - c_3 NT - k_2(1 - e^{-u}), \quad (2.2)$$

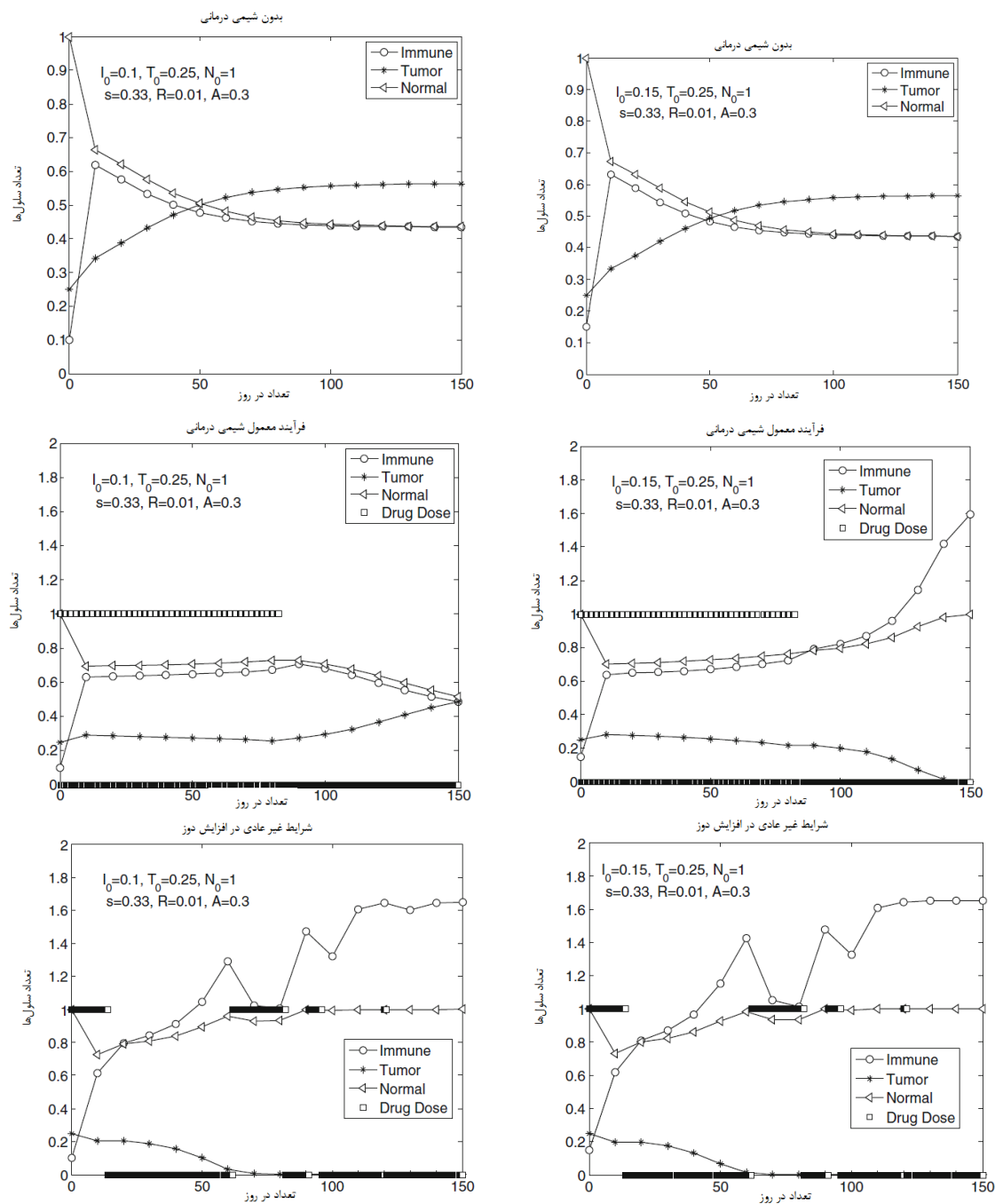
$$\frac{dN}{dt} = a_2 N(1 - b_2 N) - c_4 NT - k_3(1 - e^{-u}), \quad (3.2)$$

$$\frac{du}{dt} = v(t) - d_4 u. \quad (4.2)$$

در این معادلات E بیانگر جمعیت سلول‌های ایمنی تاثیر گذار، T جمعیت سلول‌های تومور، N نشان دهنده جمعیت سلول‌های نرمال و u میزان غلظت داروی شیمی درمانی می‌باشد. نرخ آسیب رسانی به هر جمعیت سلولی با حروف k_1 ، k_2 و k_3 در معادلات (۱.۲)–(۳.۲) نشان داده شده است که با کمترین نرخ آسیب رسانی به سلول‌های نرمال و بیشترین نرخ آسیب رسانی به سلول‌های تومور معرفی شده‌اند. با در نظر داشتن شرایط درمانی، شبیه‌سازی با برنامه‌های دوز دهی فرضی امکان‌پذیر شد. سادگی نسبی محاسبات این الگو کارایی تئوری کنترل بهینه را میسر می‌نماید. به شکل ۱ رجوع شود.

در این مدل، هدف به حداقل رساندن اندازه کل تومور در بازه زمانی از پیش تعیین شده است در حالی که به جمعیت سلول‌های نرمال آسیب چندانی نرسد. در این مورد، هدف دست یافتنی است. اما شبیه‌سازی‌ها نوساناتی را در اندازه‌های جمعیت سلولی به وجود آورد.

خوشبختانه هدف تابع، منحصر به فرد نیست و گزینه‌های ناشی از عوارض جانبی دیگر می‌تواند به نتایج متفاوتی منجر شود. برای مثال ممکن است اندازه تومور را به حداقل رساند در حالی که نوسانات در جمعیت سلول نیز به حداقل رسیده باشد. در سایر موارد، هدف به حداقل رساندن همزمان اندازه تومور و کل میزان دوزدهی شیمی درمانی است. علاوه بر این، می‌خواهیم بدانیم در صورتی که نقطه پایانی زمان درمان تومور نیز بهینه گردد نتایج درمانی چگونه تغییر می‌کنند؟ (۱۰۴؛ ۱۰۵؛ ۱۰۶).



شکل ۱: مقایسه پروتکل‌های شیمی درمانی و ایمن درمانی

۳ طراحی الگوهای ایمن درمانی

هدف از ایمن درمانی، تقویت توان دفاعی خود سیستم در برابر تومور است. از مزایای این رویکرد نسبت به شیمی درمانی سیتوتاکسیک، مسمومیت نسبتاً کمتر و در برخی موارد توانایی مورد هدف قراردادن سلول‌های تومور است که داروهای سیتوتاکسیک قادر به نفوذ در آن‌ها نیستند.

تحقیقات اصلی در این زمینه الگویی از معادلات دیفرانسیل عادی سه جمعیتی درمان IL_2 است. این الگو سلول‌های تومور،

جمعیتی از سلول‌های تاثیرگذار ایمنی فعال شده شامل سیتوتاکسیک T ، ماکروفاژ و سلول‌های کشنده که نقش نهاجمی دارند و تراکمی از سیتوکین $IL2$ که نقش افزایش قدرت سیستم ایمنی را دارد، ردیابی می‌کند.

$IL2$ یکی از سیتوکین‌های ضد التهابی است که در یک سلول میزبان به عنوان فعالیت ایمنی طبیعی ایجاد می‌گردد. تزریق تراکم‌های $IL2$ به عنوان بخشی از روش‌های خارجی درمان سرطان به کار گرفته می‌شود که می‌تواند بخشی از سلول‌های تومور را نابود کند (۱۰۸).

$IL2$ داخلی نیز از طریق تعاملات میان جمعیت تاثیرگذار و سلول‌های سرطانی تحریک می‌شود. در این الگو مشاهده می‌شود که میزان بسیار پایین $IL2$ فعالیت سیستم ایمنی را به میزان کافی بالا نمی‌برد در حالی که میزان بی رویه آن نیز می‌تواند تاثیرات مخربی به وجود آورد. پس می‌توان نتیجه گرفت که $IL2$ به تنهایی در هر دوز نمی‌تواند رشد تومور را کنترل کند، اما میزان صحیح ترکیب با نوعی از داروهای شیمی درمانی می‌تواند منجر به از بین رفتن موفقیت آمیز تومور شود.

۴ طراحی الگوهای ترکیبی شیمی درمانی و ایمن درمانی

ترکیب شیمی درمانی با ایمن درمانی می‌تواند حداقل از دو جنبه حائز اهمیت باشد:

۱-محافظت بیمار از عفونت

۲-مبارزه با خود سرطان.

درمان‌های سیتوتاکسیک سیستم ایمنی بدن را تضعیف می‌کنند و همین امر، بیمار را مستعد عفونت‌های خطرناکی می‌کند. حفظ سیستم ایمنی قوی به ویژه در زمان شیمی درمانی برای کنترل موفقیت آمیز رشد تومور امری ضروری می‌باشد. بنابراین مدلی برپایه ترکیب شیمی درمانی سیتوتاکسیک با ایمن درمانی ایجاد شد. این مدل شش جمعیت را ردیابی می‌کند:

۱-سلول‌های تومور،

۲-سلول‌های NK (لنفوسیت‌هایی از سیستم ایمنی بدن هستند که کشنده طبیعی هستند)،

۳-لنفوسیت $CD8 + T$ (نوعی لنفوسیت T (گلوبول سفید خون) که سلول‌های سرطانی را از بین می‌برد و نقش کشنده دارد)،

۴-جمعیت لنفوسیت وارد شده به گردش خون،

۵-غلظت شیمی درمانی سیتوتاکسیک

۶-غلظت $IL2$ خارجی (۱۰۷).

ایمن درمانی متشکل از تقویت سلول‌های $CD8 + T$ و تزریق $IL2$ هستند و شیمی درمانی با یک سیستم خاص سیتوتاکسیک و چرخه‌های غیرسلولی معرفی شده است. این مدل ثابت کرد که تزریق داروهای ایمن درمانی همزمان با شیمی درمانی موثرترین روش نابودی سرطان است. سوال مهمی که در بررسی روش‌های درمانی ترکیبی مطرح می‌شود، وجود راهی برای بهینه‌سازی برنامه‌های درمانی است.

۵ نتیجه‌گیری

در این مقاله، مدل‌های شیمی درمانی و همچنین ایمن درمانی سرطان مورد بحث قرار گرفت و بهترین و کاملترین نوع درمان ارائه شد. امید داریم مطالعه این مقاله تلاشی الهام‌بخش در زمینه الگوهای ریاضی درمان سرطان برای خوانندگان باشد.

مراجع

- [103] K. Roesch, D. Hasenclever, M. Scholz, Modelling Lymphoma Therapy and Outcome, *Bull Math Biol*, **76**(2) (2014) 401-430.
- [104] L. de Pillis, A. Radunskaya, *A mathematical tumor model with immune resistance and drug therapy*, an optimal control approach. *J Theor Med*, **3** (2001) 79-100.
- [105] L. de Pillis, A. Radunskaya, *The dynamics of an optimally controlled tumor model*, a case study. *Math Comput Model (Special Issue)*, **37** (2003) 1221-1244.
- [106] L. de Pillis, W. Gu, K. Fister, T. Head, K. Maples, A. Murugan et al, *Chemotherapy for tumors*, an analysis of the dynamics and a study of quadratic and linear optimal controls. *Math Biosci*, **209** (2007) 292-315.
- [107] L. de Pillis, A. Radunskaya, C. Wiseman, A validated mathematical model of cell-mediated immune response to tumor growth, *Cancer Res*, **65**(1) (2005) 7950-7958.
- [108] D. Kirschner, J. Panetta, Modeling immunotherapy of the tumor-immune interaction, *J Math Biol*, **37**(3) (1998) 235-252.

بررسی نقاط تعادل و پایداری یک مدل شکار - شکارچی به همراه یک منبع زیستی فرزانه لطفی^۱، محمد شیرازیان گروه ریاضی - دانشگاه نیشابور - نیشابور - ایران

چکیده: این کار به مدل شکار - شکارچی در محیطی که در آن ظرفیت حمل به صورت یک متغیری از زمان فرض شده است می پردازد و یکی از گونه ها از دیگری تغذیه می کند. یک منبع مشترک توسط دو گونه ی رقابتی مصرف می شود و همزمان شکارچی شکار را مصرف می کند. هدف اصلی ما شامل دو حالت می باشد: مدلی با ظرفیت حمل متغیر و دیگری رقابت درون گونه ای. حالت های پایدار را مشخص می کنیم و رفتارهای دینامیکی را مورد بررسی قرار می دهیم. همچنین پایداری سراسری و محلی نقطه تعادل داخلی را توسط معیار روت-هورویتز بررسی می کنیم و سپس یک تابع لیپانوف مناسب متناظر با آن را ارائه می دهیم.

واژه های کلیدی: مدل شکار - شکارچی، ظرفیت حمل، نقطه تعادل، تابع لیپانوف

۱ مقدمه

در اطراف ما بسیاری از مسائل جالب زیستی وجود دارد که نیاز به مدل سازی ریاضی برای درک رفتار بحرانی و ماهیت ذاتی سیستم دارد. هدف اصلی توسعه مدل، محاسبه تغییرات در جمعیت خاصی نیست بلکه مطالعه میزان پیچیدگی موجود در محیط زیست است. در جمعیت زیستی عموماً ظرفیت حمل میزان واقعی جمعیت را بر اساس محیط آن پیش بینی می کند. ظرفیت حمل یک عامل بسیار مهم است زیرا نشان دهنده ی سرعت و بالاترین سطح رشد جمعیت می باشد. در تحقیقات اخیر، ظرفیت حمل مقدار ثابت فرض شده است اما در محیط متغیر پویا ظرفیت حمل به عنوان یک متغیر حالت در نظر گرفته شده است در مدلی که توسط هایل^۱ و بورز^۲ معرفی شده است جمعیت شکار - شکارچی را به طور لجستیک با ظرفیت حمل ثابت به صورت زیر رشد می کند:

$$\begin{aligned}\frac{dx}{dt} &= r_1 x \left(1 - \frac{x}{k_1}\right) - aF(x)y, \\ \frac{dy}{dt} &= r_2 y \left(1 - \frac{y}{k_2}\right) + bF(x)y,\end{aligned}$$

¹Hoyle

²Bowers

که در آن X و Y توده‌ی زیستی جمعیت شکار و شکارچی است که به صورت لجستیک با نرخ‌های رشد r_1 و r_2 ظرفیت‌های حمل k_1 و k_2 به ترتیب رشد می‌کنند. $F(x)$ تابع پاسخ را نشان می‌دهد (۱۱۰). در مدل شکار-شکارچی معرفی شده توسط لسل^۳ و گاور^۴، ظرفیت حمل جمعیت شکار مقدار ثابت k فرض شده است و ظرفیت حمل شکارچی با ظرفیت حمل شکار متناسب است (۱۱۱؛ ۱۱۲). این کار می‌تواند به بیشتر از دو گونه در مدل زنجیره‌غذایی گسترش یابد.

گاهی اوقات مشاهده می‌شود که غنی‌سازی منبع زیستی از چند طریق بر محیط زیست تاثیر می‌گذارد. غنی‌سازی منبع که برای رشد جمعیتی سودمند است ممکن است باعث کاهش تنوع گونه‌ها گردد. درک درستی از برهم‌کنش گونه‌های مختلف در یک محیط زیست برای مدیریت منابع زیستی قابل برداشت خیلی مهم است. هر گونه، جزئی از سیستم‌های زیستی پیچیده است که با دیگر گونه‌ها برهم‌کنش دارد این برهم‌کنش‌ها ممکن است شکل‌های مختلفی از قبیل رقابتی، شکار کردن، همزیستی (انگلی)، همزیستی دو موجود را داشته باشند. شکار درون گروهی ترکیبی از اولی و دومی است یعنی رقابتی و شکار کردن که به معنی کشتن و مصرف کردن گونه‌هایی که سهمی در منبع زیستی محدود و مشترک دارند است. شکارچی درون گروهی غالب، دو سود به دست می‌آورد تغذیه کردن و حذف رقیب بالقوه. شکار درون گروهی از هر دو شکار کلاسیک و شکار رقابتی متمایز است زیرا رقبای بالقوه را کاهش می‌دهد. بنابراین تاثیر قویتری روی دینامیک‌های جمعیت به نسبت شکار و رقابت به همراه دارد. برداشت از گونه‌های مختلف روی سیستم محیط زیست تاثیر زیادی می‌گذارد. در اینجا شکار درون گروهی به همراه دینامیک‌های شکار-شکارچی در سه معادله دیفرانسیل برای شکار، شکارچی و منابع رخ می‌دهد. نتایج غنی‌سازی منبع زیستی رادر سطح کم، متوسط و بالا بررسی می‌کنیم همچنین نقاط تعادل زیستی احتمالی را بررسی می‌کنیم.

۲ فرمول‌بندی مدل

برای مطالعه سیستم شکار - شکارچی مدلی را ارائه می‌دهیم که توسعه‌ی مدل پیشنهاد شده توسط لسل^۳ و بیسنر و راس (۱۰۹) و سافوان و همکاران (۱۱۳) می‌باشد. استراتژی برداشت مستقل را به عنوان تعمیم ممکن برای یافته‌های جدید در نظر می‌گیریم:

$$\begin{aligned}\frac{dX}{dt} &= r_1 X \left(1 - \frac{X}{mZ}\right) - aXY - E_1 X, \\ \frac{dY}{dt} &= r_2 Y \left(1 - \frac{Y}{nZ}\right) + bXY - E_2 Y, \\ \frac{dZ}{dt} &= Z(c - uX - vY)\end{aligned}\quad (۱.۲)$$

X و Y توده‌ی زیستی جمعیت شکار و شکارچی است که به صورت لجستیکی به ترتیب با نرخ رشد r_1 و r_2 رشد می‌کند. Z منبع زیستی را نشان می‌دهد. سومین معادله منبع زیستی که به صورت پیوسته با زمان تغییر می‌کند جمله‌های mZ و nZ به ترتیب ظرفیت‌های حمل زیست‌محیطی برای جمعیت شکار و شکارچی هستند در اینجا، $m + n = 1$ اگر $m < n$ باشد بدین معناست که گونه‌های شکار دسترسی به منبع کمتری در مقایسه با شکارچی دارند. منبع زیستی با نرخ توسعه c رشد می‌کند. uZX و vZY به ترتیب نرخ‌های مصرف منبع توسط گونه‌های شکار و شکارچی هستند که در آن u و v ثابت هستند. E_1 و E_2 به ترتیب برداشت‌های مستقل از گونه‌های شکار و شکارچی هستند.

³Leslie

⁴Gower

۳ تجزیه و تحلیل مدل

با استفاده از تبدیلات $\epsilon = \frac{v}{a}$ ، $\delta = \frac{u}{b}$ ، $\gamma = \frac{c}{r_1}$ ، $\beta = \frac{bp}{aq}$ ، $\alpha = \frac{r_2}{r_1}$ ، $\tau = r_1 t$ ، $z = \frac{bpZ}{r_1}$ ، $y = \frac{aY}{r_1}$ ، $x = \frac{bX}{r_1}$ در معادلات (۱.۲) به شکل بی بعد زیر نشان داده می‌شود:

$$\begin{aligned}\frac{dx}{d\tau} &= x \left(1 - \frac{x}{z} \right) - xy - \mu x, \\ \frac{dy}{d\tau} &= \alpha y \left(1 - \frac{\beta y}{z} \right) + xy - \nu y \\ \frac{dz}{d\tau} &= z(\gamma - \delta x - \epsilon y)\end{aligned}\quad (۱.۳)$$

برای به دست آوردن نقاط تعادل، قرار می‌دهیم $\frac{dx}{d\tau} = 0$ ، $\frac{dy}{d\tau} = 0$ و $\frac{dz}{d\tau} = 0$ بنابراین

$$\begin{aligned}x \left(1 - \frac{x}{z} \right) - xy - \mu x &= 0 \Rightarrow \begin{cases} x = 0 \\ \left(1 - \frac{x}{z} \right) - y - \mu = 0 \end{cases} \\ \alpha y \left(1 - \frac{\beta y}{z} \right) + xy - \nu y &= 0 \Rightarrow \begin{cases} y = 0 \\ \alpha \left(1 - \frac{\beta y}{z} \right) + x - \nu = 0 \end{cases} \\ z(\gamma - \delta x - \epsilon y) &= 0 \Rightarrow \begin{cases} z = 0 \\ \gamma - \delta x - \epsilon y = 0 \end{cases}\end{aligned}$$

دو نقطه تعادل به صورت $\left(\frac{\gamma}{\delta}, 0, \frac{\gamma}{\delta(1-\mu)} \right)$ ، $\left(0, \frac{\gamma}{\epsilon}, \frac{\alpha\beta\gamma}{\epsilon(\alpha-\nu)} \right)$ و نقطه تعادل داخلی توسط $(\hat{x}, \hat{y}, \hat{z})$ داده می‌شود، که در آن

$$\begin{aligned}\hat{y} &= \frac{\gamma - \delta\hat{x}}{\epsilon} \\ \hat{z} &= \frac{\epsilon\hat{x}}{\phi + \delta\hat{x}} \\ \phi &= \epsilon(1 - \mu) - \gamma\end{aligned}\quad (۲.۳)$$

و $\hat{x}$ در معادله‌ی زیر صدق می‌کند.

$$(\alpha\beta\delta\epsilon^2 + \epsilon^2)\hat{x}^2 + ((\alpha - \nu)\epsilon^2 + \alpha\beta\delta(\phi - \gamma))\hat{x} - \alpha\beta\gamma\phi = 0.$$

۱.۳ بررسی پایداری نقاط تعادل

لم ۱.۳. فرض کنید $\delta_1 = r_1 - E_1$ و $\delta_2 = r_2 - E_2$ آنگاه:

□. اگر $\delta_2 > 0$ سیستم مدل (۱.۳) یک نقطه تعادل بدون شکار $\left(0, \frac{\gamma}{\epsilon}, \frac{\alpha\beta\gamma}{\epsilon(\alpha-\nu)} \right)$ دارد که پایدار است هرگاه $\gamma > (1 - \mu)\epsilon$ ، ناپایدار است هرگاه $\gamma < (1 - \mu)\epsilon$ ، در این حالت نقطه تعادل غیربدهی، پایدار است بر حسب میزان برداشت این نقطه پایدار است هرگاه $E_1 > r_1 - \frac{ac}{\nu}$ ، ناپایدار است هرگاه $E_1 < r_1 - \frac{ac}{\nu}$.

□□. اگر $\delta_1 > 0$ سیستم (۱.۳) یک نقطه تعادل بدون شکارچی دارد $\left(\frac{\gamma}{\delta}, 0, \frac{\gamma}{\delta(1-\mu)} \right)$ پایدار است هرگاه $\nu > \alpha + \frac{\gamma}{\delta}$ ، ناپایدار است هرگاه $\nu < \alpha + \frac{\gamma}{\delta}$ ، در این حالت نقطه تعادل غیربدهی، پایدار است. بر حسب میزان برداشت این نقطه پایدار است هرگاه $E_2 > r_2 + \frac{bc}{u}$ ، ناپایدار است هرگاه $E_2 < r_2 + \frac{bc}{u}$.

□□□ اگر $\hat{x} < \frac{\gamma}{\delta}$ و $\epsilon(1 - \mu) > \gamma$ آنگاه سیستم (۱.۳) دارای یک نقطه تعادل داخلی است.

اثبات. ژاکوبین در حالت کلی برای معادله (۱.۳) به شکل زیر است:

$$J(x, y, z) = \begin{pmatrix} 1 - \frac{\gamma x}{z} - y - \mu & -x & \frac{x^2}{z^2} \\ y & \alpha - \frac{\gamma \alpha \beta y}{z} + x - \nu & \frac{\alpha \beta y^2}{z^2} \\ -\delta z & -\epsilon z & \gamma - \delta x - \epsilon y \end{pmatrix} \quad (۳.۳)$$

□ با جایگذاری نقطه تعادل بدون شکار داریم:

$$J\left(\cdot, \frac{\gamma}{\epsilon}, \frac{\alpha \beta \gamma}{\epsilon(\alpha - \nu)}\right) = \begin{pmatrix} 1 - \frac{\gamma}{\epsilon} - \mu & \cdot & \cdot \\ \frac{\gamma}{\epsilon} & \nu - \alpha & \frac{(\nu - \alpha)^2}{\alpha \beta} \\ -\frac{\delta \alpha \beta \gamma}{\epsilon(\nu - \alpha)} & \frac{\alpha \beta \gamma}{(\nu - \alpha)} & \cdot \end{pmatrix} \quad (۴.۳)$$

که مقادیر ویژه آن عبارتند از:

$$1 - \frac{\gamma}{\epsilon} - \mu, \quad \frac{1}{2}(-\alpha + \nu \pm \sqrt{\alpha^2 - 2\alpha\nu + \nu^2 + 4\gamma\nu - 4\gamma\alpha}).$$

سیستم پایدار است هر گاه تمام مقادیر ویژه فوق منفی باشند، یعنی

$$1 - \frac{\gamma}{\epsilon} - \mu < 0 \quad \square \text{ در نتیجه } \epsilon(1 - \mu) < \gamma$$

$$\frac{1}{2}(-\alpha + \nu \pm \sqrt{\alpha^2 - 2\alpha\nu + \nu^2 + 4\gamma\nu - 4\gamma\alpha}) < 0 \quad \square \text{ بنابرین } 4\gamma\nu - 4\gamma\alpha < 0 \text{ یا } \gamma(\nu - \alpha) < 0$$

و در نتیجه با فرض $\delta_2 > 0$ داریم $\nu - \alpha < 0$ یا $\nu < \alpha$

در غیر این صورت سیستم ناپایدار خواهد بود.

□□ با جایگذاری نقطه تعادل بدون شکارچی در ماتریس ژاکوبین (۳.۳) داریم:

$$J\left(\frac{\gamma}{\delta}, \cdot, \frac{\gamma}{\delta(1 - \mu)}\right) = \begin{pmatrix} \mu - 1 & -\frac{\gamma}{\delta} & (\mu - 1)^2 \\ \cdot & \alpha + \frac{\gamma}{\delta} - \nu & \cdot \\ \frac{\gamma}{\mu - 1} & \frac{\gamma\epsilon}{\delta(\mu - 1)} & \cdot \end{pmatrix} \quad (۵.۳)$$

که مقادیر ویژه آن عبارتند از:

$$\alpha - \nu + \frac{\gamma}{\delta}, \quad \frac{1}{2}(-1 + \mu \pm \sqrt{1 - 2\mu + \mu^2 + 4\gamma\mu - 4\gamma})$$

سیستم پایدار است هر گاه تمام مقادیر ویژه فوق منفی باشند یعنی

$$\alpha - \nu + \frac{\gamma}{\delta} < 0 \quad \square \text{ در نتیجه } \alpha + \frac{\gamma}{\delta} < \nu$$

$$\frac{1}{2}(-1 + \mu \pm \sqrt{1 - 2\mu + \mu^2 + 4\gamma\mu - 4\gamma}) < 0 \quad \square \text{ بنابرین } 4\gamma\mu - 4\gamma < 0 \text{ یا } \gamma(\mu - 1) < 0 \text{ و در نتیجه با}$$

فرض $\delta_1 > 0$ داریم $\mu - 1 < 0$ یا $\mu < 1$

در غیر این صورت سیستم ناپایدار است.

□

قضیه ۲.۳. اگر $\hat{z} \leq \min\{\frac{\epsilon}{\delta}, \frac{1}{\epsilon}, 1\}$ آنگاه سیستم (۱.۳) پایدار مجانبی محلی است.

اثبات. معادله مشخصه ماتریس ژاکوبین سیستم مدل (۱.۳) در نقطه تعادل داخلی، یعنی $\det(J(\hat{x}, \hat{y}, \hat{z}) - \lambda I) = 0$ داریم:

$$\lambda^3 + a_1 \lambda^2 + a_2 \lambda + a_3 = 0$$

که در آن

$$\begin{aligned} a_1 &= \frac{\hat{x} + \hat{y}\alpha\beta}{\hat{z}^2} \\ a_2 &= \frac{\hat{x}\hat{y}\hat{z}^2 + \hat{x}\hat{y}\alpha\beta + \hat{x}^2\hat{z}\delta + \hat{y}^2\hat{z}\alpha\beta\epsilon}{\hat{z}^2} \\ a_3 &= \frac{\hat{x}^2\hat{y}\alpha\beta\delta - \hat{x}\hat{y}^2\hat{z}\alpha\beta\delta + \hat{x}^2\hat{y}\hat{z}\epsilon + \hat{x}\hat{y}^2\alpha\beta\epsilon}{\hat{z}^2} \end{aligned}$$

به وضوح $a_1 > 0$ ، $a_2 > 0$ اگر $\hat{z} < \frac{\epsilon}{\delta}$ آنگاه $a_3 > 0$ اکنون اگر $1 - \hat{z}\epsilon > 0$ و $1 - \hat{z} > 0$ آنگاه

$$\begin{aligned} a_1 a_2 - a_3 &= \frac{\hat{x} + \hat{y}\alpha\beta}{\hat{z}^2} \frac{\hat{x}\hat{y}\hat{z}^2 + \hat{x}\hat{y}\alpha\beta + \hat{x}^2\hat{z}\delta + \hat{y}^2\hat{z}\alpha\beta\epsilon}{\hat{z}^2} \\ &\quad - \frac{\hat{x}^2\hat{y}\alpha\beta\delta - \hat{x}\hat{y}^2\hat{z}\alpha\beta\delta + \hat{x}^2\hat{y}\hat{z}\epsilon + \hat{x}\hat{y}^2\alpha\beta\epsilon}{\hat{z}^2} \\ &= [\hat{x}^2\hat{y}\hat{z}^2(1 - \hat{z}\epsilon) + \hat{x}^2\hat{y}\alpha\beta + \hat{x}^2\hat{z}\delta + \hat{x}\hat{y}^2\hat{z}\alpha\beta\epsilon(1 - \hat{z}) + \alpha\beta\hat{x}\hat{y}^2\hat{z}^2 \\ &\quad + \hat{x}\hat{y}^2\alpha^2\beta^2 + \alpha\beta\delta\hat{x}^2\hat{y}\hat{z}(1 - \hat{z}) + \alpha^2\beta^2\epsilon\hat{y}^2\hat{z} + \hat{x}\hat{y}^2\hat{z}^2\alpha\beta\delta]/\hat{z}^4 \\ &> 0, \end{aligned}$$

بنابراین $a_1 > 0$ ، $a_2 > 0$ و $a_1 a_2 - a_3 > 0$ برقرار است هرگاه $\hat{z} \leq \{\min\{\frac{\epsilon}{\delta}, \frac{1}{\epsilon}, 1\}$ با استفاده از معیار روت-هورویتز سیستم، پایدار مجانبی محلی است. $\square$

قضیه ۳.۳. سیستم (۱.۳) در ناحیه‌ای که در آن $z(\epsilon + \delta) > \alpha\beta y$ باشد پایدار مجانبی سراسری است که z چگالی تعادل منابع زیستی را نشان می‌دهد.

اثبات. برای اثبات پایداری سراسری سیستم یک تابع لیاپانوف می‌سازیم.

$$V(x, y, z) = \left[(x - \hat{x}) - \hat{x} \log\left(\frac{x}{\hat{x}}\right) \right] + H_1 \left[(y - \hat{y}) - \hat{y} \log\left(\frac{y}{\hat{y}}\right) \right] + H_2 \left[(z - \hat{z}) - \hat{z} \log\left(\frac{z}{\hat{z}}\right) \right]$$

اکنون داریم :

$$\begin{aligned}
 \frac{dv}{dt} &= \frac{dv}{dx} \frac{dx}{dt} + \frac{dv}{dy} \frac{dy}{dt} + \frac{dv}{dz} \frac{dz}{dt} = \frac{x - \hat{x}}{x} \frac{dx}{dt} + H_1 \frac{y - \hat{y}}{y} \frac{dy}{dt} + H_2 \frac{z - \hat{z}}{z} \frac{dz}{dt} \\
 &= (x - \hat{x}) \left[1 - \frac{x}{z} - y - \mu - 1 + \frac{\hat{x}}{\hat{z}} + \hat{y} + \mu \right] + H_1 (y - \hat{y}) \left[\alpha - \alpha \beta \frac{y}{z} \right. \\
 &\quad \left. + x - \nu - \alpha + \alpha \beta \frac{\hat{y}}{\hat{z}} - \hat{x} + \nu \right] + H_2 (z - \hat{z}) [\gamma - \delta x - \epsilon y - \gamma + \delta \hat{x} + \epsilon \hat{y}] \\
 &= -(x - \hat{x}) \left[\frac{z(x - \hat{x}) - x(z - \hat{z})}{z\hat{z}} + (y - \hat{y}) \right] + H_1 (y - \hat{y}) \left[\alpha \beta \frac{y(z - \hat{z}) - z(y - \hat{y})}{z\hat{z}} + x - \hat{x} \right] \\
 &\quad - H_2 \delta (z - \hat{z})(x - \hat{x}) - H_2 \epsilon (z - \hat{z})(y - \hat{y}) \\
 &= -\frac{(x - \hat{x})^2}{\hat{z}} + \frac{x(x - \hat{x})(z - \hat{z})}{z\hat{z}} - (x - \hat{x})(y - \hat{y}) - H_1 \frac{\alpha \beta (y - \hat{y})^2}{\hat{z}} + H_1 \frac{\alpha \beta (y - \hat{y})(z - \hat{z})y}{z\hat{z}} \\
 &\quad + H_1 (y - \hat{y})(x - \hat{x}) - H_2 \delta (x - \hat{x})(z - \hat{z}) - H_2 \epsilon (y - \hat{y})(z - \hat{z}) \\
 &\leq -\left[\frac{(x - \hat{x})^2}{\hat{z}} + H_1 \alpha \beta \frac{(y - \hat{y})^2}{\hat{z}} \right] + \frac{1}{\nu} \left[H_2 \epsilon + \frac{H_1 \alpha \beta y}{z\hat{z}} \right] [(y - \hat{y})^2 + (z - \hat{z})^2] \\
 &\quad + \frac{1}{\nu} \left[-H_2 \delta + \frac{x}{z\hat{z}} \right] [(x - \hat{x})^2 + (z - \hat{z})^2] \\
 &= \left[\frac{1}{\nu} \left(-H_2 \delta + \frac{x}{z\hat{z}} \right) - \frac{1}{\hat{z}} \right] (x - \hat{x})^2 + \left[\frac{1}{\nu} \left(-H_2 \epsilon + \frac{H_1 \alpha \beta y}{z\hat{z}} \right) - \frac{H_1 \alpha \beta}{\hat{z}} \right] (y - \hat{y})^2 \\
 &\quad + \frac{1}{\nu} \left[-H_2 \epsilon + \frac{H_1 \alpha \beta y}{z\hat{z}} - H_2 \delta + \frac{x}{z\hat{z}} \right] (z - \hat{z})^2 < 0.
 \end{aligned}$$

هرگاه $\frac{H_1 \alpha \beta y}{z\hat{z}} + \frac{x}{z\hat{z}} > H_2 \epsilon + H_2 \delta$ سیستم مدل (۱.۳) در ناحیه‌ای که در $\hat{z}(\epsilon + \delta)z > \alpha \beta y + x$ پایدار مجانبی سراسری است.

□

مراجع

- [109] B. Basener and D. S. Ross, Booming and crashing populations and Easter Island, SIAM J. Appl. Math. 65 (2005) 684
- [110] A. Hoyle and R. Bowers, When is evolutionary branching in predator-prey systems possible with an explicit carrying capacity? Math. Biosci. 210 (2007) 1.
- [111] P. H. Leslie, A stochastic model for studying the properties of certain biological systems by numerical methods, Biometrika 45 (1958) 16.
- [112] P. H. Leslie and J. C. Gower, The properties of a stochastic model for the predator-prey type of interaction between two species, Biometrika 47 (1960) 219.
- [113] H. Safuan, I. Towers, Z. Jovanoski and H. Sidhu, ANZIAMJ 53 (2012) C172.

مقایسه فواصل اطمینان برای نسبت شانس

اعظم مسلمی، استادیار، هیئت علمی گروه آمار زیستی، دانشکده پزشکی، دانشگاه علوم پزشکی اراک، اراک، ایران محمد رفیعی، دانشیار، هیئت علمی گروه آمار زیستی، دانشکده پزشکی، دانشگاه علوم پزشکی اراک، اراک، ایران راشد پورحمیدی، دانشجوی ارشد آمار زیستی، دانشکده پزشکی، دانشگاه علوم پزشکی اراک، اراک، ایران سردار جهانی، دانشجوی ارشد آمار زیستی، دانشکده پزشکی، دانشگاه علوم پزشکی اراک، اراک، ایران کورش رضایی، مربی، هیئت علمی گروه پرستاری، دانشکده پرستاری، دانشگاه علوم پزشکی اراک، اراک، ایران مهشید عسکری^۱، دانشجوی ارشد آمار زیستی، دانشکده پزشکی، دانشگاه علوم پزشکی اراک، اراک، ایران

چکیده: در مسائل تجزیه و تحلیل داده ها بخش کوچکی از جامعه که در قالب نمونه استخراج شده است کاربرد دارد. در استنباط آماری از برآوردهای نقطه ای و فاصله ای (فاصله اطمینان) برای برآورد پارامتر مورد نظر استفاده می کنند. برآورد فاصله ای دقت بالاتری نسبت به برآورد نقطه ای دارد. در مطالعات بالینی که از نسبت شانس استفاده می شود، این شاخص شدت ارتباط بین هر یک از متغیرهای مستقل و وجود یا فقدان بیماری را فراهم می کند. در این مقاله هشت فاصله اطمینان نسبت شانس را که شامل فواصل اطمینان دقیق و مجانبی هستند، بر اساس طول فاصله اطمینان مورد بررسی قرار می دهیم.

پژوهش مذکور به بررسی انواع فاصله اطمینان نسبت شانس می پردازد و مقادیر آنها در مثالی تحت یک کار آزمایشی بالینی یک سوکور محاسبه می شوند. مشارکت کنندگان این کار آزمایشی بالینی ۸۰ نفر از بیماران مبتلا به سکته قلبی در بیمارستان های اراک در سال ۱۳۹۷ بودند. نمونه ها با استفاده از روش تخصیص تصادفی به دو گروه مداخله و کنترل تقسیم شدند. به منظور جمع آوری داده ها از چک لیست بررسی اختلال پس از سانحه (PTSD) استفاده شد. برای بیماران در گروه مداخله مراقبت معنوی به مدت (سه روز) انجام شد و یک ماه پس از ترخیص بیمار از وی خواسته شد تا چک لیست را تکمیل کند. تجزیه و تحلیل داده ها از طریق نرم افزار R نسخه ۳.۵.۲^۱ باشد. و به منظور تحلیل داده ها از آزمون های آماری کای اسکور، نسبت شانس، فاصله اطمینان ولف لوجیت، گارت، هموار مستقل، کارن فیلد دقیق، کارن فیلد مید-پی، اگرس-مین و باپتیستا-پیک دقیق و باپتیستا-پیک مید-پی استفاده شد.

در این مطالعه مقدار نسبت شانس ۰/۱۱ است. از سه فاصله اطمینان مجانبی بهترین فاصله اطمینان مجانبی گارت با طول فاصله ۲/۰۳۶ می باشد؛ و از بین پنج فاصله اطمینان دقیق بهترین فاصله اطمینان باپتیستا-پیک مید-پی با طول فاصله ۱/۹۷۳ است. در مطالعه حاضر انواع فاصله اطمینان برای نسبت شانس بررسی شدند. این فواصل شامل فواصل اطمینان دقیق و فواصل اطمینان مجانبی بودند. از سه فاصله اطمینان مجانبی ذکر شده بهترین فاصله اطمینان گارت است، که طول کمتر و دقت بیشتری دارد. اما در بین پنج فاصله اطمینان دقیق، فاصله اطمینان نسبت شانس باپتیستا-پیک مید-پی طول کمتر و دقت بیشتری دارد. لذا در این مطالعه بهترین فاصله اطمینان باپتیستا-پیک مید-پی است.

^۱مهشید عسکری: mahshidasgri70@gmail.com

واژه‌های کلیدی: نسبت شانس، فاصله اطمینان، سکت قلبی، PTSD، طول فاصله اطمینان

۱ مقدمه

نسبت شانس^۱ مقیاس مهمی برای دو نسبت مستقل دو جمله‌ای است. اغلب در مطالعات مورد شاهدهی و مقطعی و کوهورت مورد استفاده قرار می‌گیرد. همچنین در کار آزمایی بالینی یک معیار در متاآنالیزها محسوب می‌شود [۱]. نسبت شانس مقیاسی از ارتباط بین پیامد و مواجهه است و تقریب مناسبی برای بیماری‌های نادر می‌باشد. اگر $OR = 1$ مواجهه بر پیامد مؤثر نیست و اگر $OR > 1$ مواجهه بر پیامد تأثیر مثبت و افزایشی دارد و در صورتی $OR < 1$ مواجهه بر پیامد تأثیر منفی و کاهش‌دهنده دارد [۲]. آمار در پزشکی به‌کرات به مقایسه دو نسبت دو حالتی می‌پردازد. متغیر دو حالتی متغیری است که دارای دو نتیجه موفقیت و شکست است و موفقیت لزوماً نشان‌دهنده یک نتیجه مطلوب نیست. نتایج دو گروه مستقل را می‌توان در جدول 2×2 خلاصه کرد. جدول توافقی که اولین بار توسط لاندرسون^۲ ارائه شد، به صورت زیر است [۳].

جدول ۱: جدول توافقی 2×2			
	موفقیت	شکست	جمع
کنترل	n_{11}	n_{12}	n_{1+}
مداخله	n_{21}	n_{22}	n_{2+}
جمع	n_{+1}	n_{+2}	N

تعداد افراد را در گروه‌ها با (n_{1+}, n_{2+}) نشان می‌دهیم. در این جدول، حاشیه ثابت فرض شده است. جدول توافقی با حاشیه ثابت در مطالعات کار آزمایی بالینی و مورد شاهدهی و کوهورت بسیار کاربرد دارد. n_{11} موفقیت در گروه اول دارای توزیع دو جمله‌ای با پارامتر n_{1+} و p_1 است و n_{21} موفقیت در گروه دوم دارای توزیع دو جمله‌ای با پارامتر n_{2+} و p_2 است. لذا برآورد موفقیت‌ها را از نسبت درستمایی ماکزیمم به صورت مقابل به دست می‌آوریم.

$$p_1 = \frac{n_{12}}{n_{1+}}, \quad p_2 = \frac{n_{21}}{n_{2+}}$$

و نسبت شانس را به صورت مقابل برآورد می‌کنیم

$$\hat{\theta} = \frac{n_{11}/n_{12}}{n_{21}/n_{22}} = \frac{n_{11}n_{22}}{n_{21}n_{12}} \quad (1.1)$$

فاصله اطمینان برای نسبت شانس شامل فواصل اطمینان مجانبی و دقیق است. فاصله اطمینان مجانبی بر اساس تقریبی از توزیع نرمال ساخته می‌شود. فاصله اطمینان دقیق بر اساس توزیع دقیق مشاهدات ساخته می‌شود فاصله اطمینان دقیق به ویژه در حالت حجم نمونه کوچک و نسبت نزدیک به ۰ و ۱ کاربرد دارد.

۲ مواد و روشها

داده های این مطالعه از یک کار آزمایی بالینی یک سوکور بدست آمد. که بر بیماران سکت قلبی حاد بستری در بخش‌های CCU بیمارستان‌های امیرکبیر و امیرالمؤمنین اراک در دو گروه مداخله و کنترل اجرا شد. در گروه مداخله مراقبت معنوی در مدت بستری (سه

روز) انجام شد، که شامل امید افزایی معنوی یک مداخله محتوا محور در سه جلسه است و به ۴ محور اصلی می پردازد تا از این طریق امید بیماران را افزایش دهد. این چهار محور عبارت اند از: ارتباط با خدا، ارتباط با خود، ارتباط با دیگران و ارتباط با طبیعت است. یک ماه پس از ترخیص بیمار به منظور جمع آوری داده ها از چک لیست بررسی اختلال استرس پس از سانحه PCL^3 استفاده شد. در این مطالعه با متغیر دوحالتی روبرو هستیم که نتایج آن در جداول توافقی تنظیم می شود. یکی از مهم ترین شاخص هایی که در این نوع جداول برآورد می شود نسبت شانس است. برای یک جدول توافقی 2×2 با فراوانی مورد انتظار نسبت شانس به صورت مقابل بدست می آید.

$$\theta = \frac{p_1/(1-p_1)}{p_2/(1-p_2)} \quad (1.2)$$

و θ به صورت معادله (۱) برآورد می شود. برای برآورد دقیقتر نسبت شانس از فواصل اطمینان استفاده می کنیم. معیار بررسی عملکرد فواصل اطمینان، طول فاصله اطمینان است. هرچه طول کمتر باشد، فاصله اطمینان از دقت بالاتری برخوردار است. فواصل اطمینان OR در زیر اشاره شده است.

۱.۲ $logit Woolf$

فاصله اطمینان مجانبی ولف لوجیت ابتدا توسط ولف 4 [۴] ارائه شد. این فاصله اطمینان بر اساس تقریب توزیع نرمال از $\log(\hat{\theta})$ محاسبه شده و به صورت زیر تعریف می شود.

$$woolflogit = \log(\hat{\theta}) \pm z_{(\alpha/2)} \sqrt{\frac{1}{n_{11}} + \frac{1}{n_{12}} + \frac{1}{n_{21}} + \frac{1}{n_{22}}} \quad (2.2)$$

اگر یکی از سلول های جدول توافقی 2×2 صفر شود، انحراف معیار بی نهایت و فاصله اطمینان ناآگاهی بخش می شود و اطلاعات درستی نمی دهد.

۲.۲ $logitadjusted Gart$

گارت 5 [۵] با اضافه کردن 0.5 به هر سلول جدول توافقی 2×2 ، فاصله اطمینان زیر را ارائه کرد.

$$n_{\tilde{1}1} = n_{11} + 0.5, n_{\tilde{2}2} = n_{22} + 0.5, n_{\tilde{1}2} = n_{12} + 0.5, n_{\tilde{2}1} = n_{21} + 0.5$$

$$gartadjustedlogit = \log(\hat{\theta}) \pm z_{(\alpha/2)} \sqrt{\frac{1}{\tilde{n}_{11}} + \frac{1}{\tilde{n}_{12}} + \frac{1}{\tilde{n}_{21}} + \frac{1}{\tilde{n}_{22}}} \quad \tilde{\theta} = \frac{\tilde{n}_{11}\tilde{n}_{22}}{\tilde{n}_{12}\tilde{n}_{21}} \quad (3.2)$$

این فاصله اطمینان در زمانی که بیشتر سلول های جدول توافقی صفر هستند، بسیار مفید است. اگر $n_{11} = 0$ و $n_{21} \neq 0$ آنگاه $\hat{\theta} = 0$ و کران پایین فاصله اطمینان ($L > 0$) می شود و اگر $n_{11} = n_{1+}$ یا اگر $n_{21} = 0$ و $n_{11} \neq 0$ آنگاه $\hat{\theta} = \infty$ و کران بالای فاصله اطمینان تعریف نشده ($U = \infty$) است.

۳.۲ *logitIndependence – smooth*

اگرستی [۶] با اضافه کردن مقدار نامنفی $i, j = 1, 2$ $c_{ij} = \frac{2^{n_{i+j}}}{N^2}$ به هر سلول جدول توافقی 2×2 ، فاصله اطمینانی مشابه فاصله اطمینان گارت ارایه کرد. بنابراین در این فاصله اطمینان اگر $c = 0$ و $c = 0.5$ باشد به ترتیب فاصله اطمینان گارت و ولف به دست می آید.

۴.۲ *conditional exact Cornfield*

اگر تعداد موفقیت‌ها (n_{+1}) و تعداد شکست‌ها (n_{+2}) همچنین تمام مجموع توأم حاشیه‌ای در جدول (۱) ثابت باشد احتمال مشاهدات در جدول توافقی با x_{11} موفقیت دارای توزیع فوق هندسی غیر مرکزی است. که تابع احتمال آن به صورت زیر است:

$$f(x_{11}|\theta) = \frac{\binom{n_{1+}}{x_{11}} \binom{n_{2+}}{n_{+1}-x_{11}} \theta^{x_{11}}}{\sum_{i=n_{max}}^{n_{min}} \binom{n_{1+}}{i} \binom{n_{2+}}{n_{+1}-i} \theta^i} \quad (4.2)$$

و زمانی که $n_{min} = \min(n_{1+}, n_{+1})$ ، $n_{max} = \max(0, n_{1+} - n_{+1})$ کارن فیلد [۷] با معکوس کردن دو آزمون یک‌طرفه دقیق شرطی فاصله اطمینان شرطی دقیق (L, U) را معرفی کرد که (L, U) با حل روابط زیر می‌توان به دست می آیند.

$$\sum_{x_{11}=n_{11}}^{n_{min}} f(x_{11}|L) = \alpha/2 \quad \sum_{x_{11}=n_{max}}^{n_{11}} f(x_{11}|U) = \alpha/2 \quad (5.2)$$

فاصله اطمینان دقیق کارن فیلد همیشه با آزمون دقیق *Fisher – Irwin* مطابقت دارد [۶]. همچنین، کارن فیلد فواصل تقریبی مجانبی را به فواصل شرطی دقیق ترجیح می‌دهد.

۵.۲ *interval mid – pCornfield*

این فاصله اطمینان شبه دقیق ^۸ (براساس *pointmid* ساخته می‌شود) هستند [۸]؛ و با حل روابط زیر به دست می آید.

$$\sum_{x_{11}=n_{11}}^{n_{min}} f(x_{11}|L_{cm}) - \frac{1}{2} f(n_{11}|L_{cm}) = \alpha/2. \quad (6.2)$$

$$\sum_{x_{11}=n_{max}}^{n_{11}} f(x_{11}|U_{cm}) - \frac{1}{2} f(n_{11}|U_{cm}) = \alpha/2. \quad (7.2)$$

۶.۲ *conditional exact Batista – Pike*

باتیستا و پیک^۹ [۹] و [۱۰] بر اساس معکوس آزمون دوطرفه، فاصله اطمینان را پیشنهاد دادند که با حل روابط زیر به دست می آید:

$$\sum_{x_{11}=n_{max}}^{n_{min}} f(x_{11}|L).I(f(x_{11}|L) \leq f(n_{11}|L)) = \alpha \quad (۸.۲)$$

$$\sum_{x_{11}=n_{max}}^{n_{min}} f(x_{11}|U).I(f(x_{11}|L) \leq f(n_{11}|U)) = \alpha \quad (۹.۲)$$

۷.۲ *mid – p Batista – Pike*

این فاصله اطمینان از نوع فاصله اطمینان شبه دقیق است [۱۱]، به صورت زیر محاسبه می شود.

$$\sum_{x_{11}=n_{max}}^{n_{min}} f(x_{11}|L).I(f(x_{11}|L) \leq f(n_{11}|L)) - \frac{1}{2}f(n_{11}|L) = \alpha \quad (۱۰.۲)$$

$$\sum_{x_{11}=n_{max}}^{n_{min}} f(x_{11}|U).I(f(x_{11}|L) \leq f(n_{11}|U)) - \frac{1}{2}f(n_{11}|U) = \alpha \quad (۱۱.۲)$$

۸.۲ *unconditional exact Agresti – Min*

اگرستی و مین [۱۲] بر اساس معکوس یک آزمون *score* دوطرفه دقیق غیرشرطی فاصله اطمینان دقیق غیرشرطی را برای نسبت شانس پیشنهاد دادند:

$$T(n|\theta_*) = [n_{1+}(\hat{p}_1 - \tilde{p}_1)]^2 \left[\frac{1}{n_{1+}\tilde{p}_1(1 - \tilde{p}_1)} + \frac{1}{n_{2+}\tilde{p}_2(1 - \tilde{p}_2)} \right] \quad (۱۲.۲)$$

که در آن،

$$\theta_* = \frac{p_1/(1 - p_1)}{p_2/(1 - p_2)} \quad (۱۳.۲)$$

می باشد سپس روابط زیر را ارائه می شود :

$$R(T(n)|p_1, \theta_*) = \sum_{|T(x)| \geq |T(n)|} f(x|p_1, \theta_*) \quad (۱۴.۲)$$

$$R(T(n)|\theta_*) = \sup_{p_1} R(T(n)|p_1, \theta_*) \quad (۱۵.۲)$$

بنابراین با حل روابط زیر فاصله اطمینان اگرستی □ مین به دست می آید.

$$R(T(n)|L) = \alpha \quad R(T(n)|U) = \alpha \quad (۱۶.۲)$$

اگر $\theta_* < L$ در نتیجه $R(T(n)|\theta_*) < \alpha$ و در صورتی که $\theta_* > U$ می توان نتیجه گرفت $R(T(n)|\theta_*) > U$

۳ نتایج

در این پژوهش در مجموع ۸۰ نفر شرکت داشتند و افراد به طور تصادفی به دو گروه کنترل و مداخله تخصیص داده شدند. با توجه به جدول (۲) از ۴۰ نفر تحت گروه مداخله ۷ نفر دارای $PTSD > ۵۰$ و از ۴۰ نفر در گروه کنترل ۲۶ نفر دارای $PTSD > ۵۰$ بودند. در نتیجه در گروه مداخله ۶۵ درصد و در گروه کنترل ۱۷/۵ درصد افراد دارای اختلال استرس پس از سانحه بودند ($PTSD > ۵۰$).

جدول ۲: نتایج اختلال استرس در بیماران سکته قلبی حاد

مراقبت معنوی	<i>PTSD</i>		جمع
	بیشتر از ۵۰ درصد	کمتر از ۵۰ درصد	
مداخله	۷ (۱۷/۵)	۳۳ (۸۲/۵)	۴۰ (۵۸/۸)
کنترل	۲۶ (۶۵)	۱۴ (۳۵)	۴۰ (۴۱/۳)
			۸۰

$$\chi^2 = ۱۸/۶۲$$

$$p - value = ۰/۰۰۰$$

با توجه به اینکه مقدار $OR = ۰/۱۱$ می باشد می توان نتیجه گرفت شانس استرس پس از سانحه در بیماران سکته قلبی حاد در گروه کنترل ۹/۱ برابر گروه مراقبت معنوی است ($p < ۰/۰۵$).

با توجه به نتایج جدول (۳) می توان نتیجه گرفت در بین سه فواصل اطمینان مجانبی، گارت با مقدار $۲/۰۳۶$ طول کمتر از ولف لوجیت و مستقل هموار است پس دقت بالاتری دارد. از میان پنج فاصله اطمینان دقیق طول فاصله اطمینان باپیتستا-پیک دقیق برابر با مقدار $۱/۹۷۳$ است چون طول کمتر، دارای دقت بیشتری می باشد. و بیشترین طول متعلق به فاصله اطمینان کارن فیلد دقیق برابر با مقدار $۲/۳۴۱$ است. بنابراین این فاصله اطمینان دقت کمتری نسبت به سایر فواصل اطمینان دارد.

جدول ۳: فواصل اطمینان برای OR در بیماران سکته قلبی حاد.

	فاصله اطمینان		
	کران پایین	کران بالا	طول
ولف لوجیت	۳/۰۹	۲۴/۸۴	۲/۰۸۴
گارت	۲/۹۵	۲۲/۶۰	۲/۰۳۶
مستقل هموار	۲/۹۶	۲۲/۷۵	۲/۰۳۹
کارن فیلد دقیق	۲/۷۹	۲۸/۹۸	۲/۳۴۱
کارن فیلد مید-پی	۳/۰۶	۲۵/۶۳	۲/۱۲۵
باپیتستا-پیک دقیق	۲/۸۵	۲۵/۳۹	۲/۱۸۷
باپیتستا-پیک مید-پی	۳/۱۷	۲۲/۸۰	۱/۹۷۳
اگرستی-مین	۰/۰۳	۰/۳۱	۲/۳۳۵

۴ بحث

در این پژوهش با توجه به اینکه حجم نمونه متوسط است. با مقایسه فواصل اطمینان مجانبی، طول فاصله اطمینان گارت برابر $۲/۰۳۶$ نسبت به فواصل اطمینان ولف لوجیت و هموار مستقل کمتر است. لذا در بین فواصل اطمینان مجانبی نسبت شانس، فاصله اطمینان گارت از دقت بالاتری برخوردار است. در بین فواصل اطمینان دقیق طول فاصله اطمینان بابتیستا-پیک-مید-پی برابر $۱/۹۷۳$ است. بنابراین دقت این فاصله اطمینان بیشتر از تمام فواصل اطمینان نسبت شانس است.

در مطالعاتی که اگرتی و مین در سال ۲۰۰۲ [۱۲] و فیگرلند^{۱۰} و همکاران در سال ۲۰۱۱ [۱۳] انجام دادند. فواصل اطمینان نسبت شانس از نظر پوشش احتمالی مورد بررسی قرار گرفتند. این مطالعات برای بررسی حجم نمونه‌های کوچک و بزرگ انجام شد. نتایج مطالعات نشان داد، که سه فاصله اطمینان مجانبی نسبت شانس عبارت‌اند از: ولف لوجیت، گارت و مستقل هموار عملکرد مشابهی در حجم نمونه کوچک و متوسط دارند. برای نمونه کوچک فاصله اطمینان گارت پوشش احتمالی محافظه‌کاری نسبت به فاصله اطمینان ولف لوجیت و هموار مستقل دارد و هر سه فاصله اطمینان برای نسبت نزدیک به صفر و یک محافظه‌کار عمل می‌کند. هنگامی که حجم نمونه کاهش می‌یابد هر سه فاصله اطمینان پوشش احتمالی مشابهی دارند. با افزایش نسبت شانس هر سه فاصله اطمینان واگرا می‌شوند و فاصله اطمینان گارت بین این سه فاصله اطمینان بهتر است. همچنین این فاصله اطمینان از دو فاصله اطمینان ولف لوجیت و هموار مستقل کوتاه‌تر است.

در این مطالعات بررسی پوشش احتمالی برای فواصل اطمینان دقیق، شامل کارن فیلد دقیق و کارن فیلد-مید-پی و بابتیستا-پیک دقیق و بابتیستا-پیک-مید-پی و اگرتی-مین نیز انجام شده، برای نمونه‌های کوچک هر دو فاصله اطمینان کارن فیلد بسیار محافظه‌کار عمل می‌کنند و با افزایش حجم نمونه تغییر نمی‌کنند. فاصله بابتیستا-پیک دقیق محافظه‌کاری ویژه برای حجم نمونه‌های کم در بین فواصل اطمینان دقیق دارد و فواصل اطمینان اگرتی-مین کمتر محافظه‌کار است. لذا بهترین فاصله اطمینان دقیق همان بابتیستا-پیک-مید-پی است.

در مطالعه‌ای دیگر که در سال ۲۰۱۴ [۱۴] تحت عنوان مقایسه فواصل اطمینان دو نسبت مستقل دوحالتی با استفاده از روش گرافیکی و میانگین متحرک توسط لود و دانه^{۱۱} انجام شد. از روش گرافیکی و حجم نمونه‌های کوچک و بزرگ استفاده شد. به‌منظور، سادگی محاسبات بدون نرم‌افزار، روش برون لی جفری^{۱۲} [۱۵] پیشنهاد شده است. در اکثر موارد پوشش احتمال یک‌طرفه پیشنهاد می‌شود. روش گارت-نم^{۱۳} [۱۶] تنها استثنا در روش تصحیح چولگی است. با توجه به تصحیح منظم که شامل یک خطای پیوستگی برای بهبود عملکرد نمونه‌های کوچک و انتخاب یک تصحیح پیوستگی برای بررسی پوشش احتمالی محافظه‌کار است. با استفاده از میانگین وزن‌دار شده یک روش ترکیبی جدید به دست می‌آید که تقریباً با گارت-نم عملکرد یکسانی دارد.

۵ نتیجه گیری

به‌طور کلی فاصله اطمینان بابتیستا-پیک-مید-پی برای نسبت شانس عملکرد بهتری نسبت به سایر فواصل اطمینان دقیق و مجانبی دارد و این نتیجه با یافته‌های فیگرلند و همکارانش مطابقت دارد.

^{۱۰}Fagerland

^{۱۱}DaneandLaud

^{۱۲}JeffreysBrown – Li

^{۱۳}Gart – Nam

مراجع

- [114] JJ. Deeks and DG. Altman, *Effect measures for metaanalysis of trials with binary outcomes.*, In: M. Egger , G. Davey Smith and DG. Altman (eds) *Systematic reviews in health care: meta-analysis in context* 2nd edn., BMJ Publishing Group, London, 2001, pp. 313–335.
- [115] M. Szumilas, Explaining odds ratios, *Journal of the Canadian academy of child and adolescent psychiatry.*, 19 (2010), 227.
- [116] S. Lydersen ,MW. Fagerland and P. Laake , Tutorial in biostatistics: recommended tests for association in 2×2 tables, *Stat Med.*, 28 (2009), 1159–1175.
- [117] B. Woolf , On estimating the relation between blood group and disease, *Ann Human Gene.*, 19 (1955), 251–253.
- [118] JJ. Gart , Alternative analyses of contingency tables, *J R Stat Soc B Stat Meth.*, 28 (1966), 164–179.
- [119] A. Agresti , On logit confidence intervals for the odds ratio with small samples, *Biometrics.*, 55 (1999), 597–602.
- [120] J. Cornfield , A statistical problem arising from retrospective studies, *Pro Third Berkeley Sym Math Stat Probab.*, 4 (1956), 135–148.
- [121] KF. Hirji , S-J. Tan and RM. Elashoff , A quasi-exact test for comparing two binomial proportions, *Stat Med.*, 10 (1991), 1137–1153.
- [122] J. Baptista and MC. Pike , Exact two-sided confidence limits for the odds ratio in a 2×2 table, *J R Stat Soc C Appl Stat.*, 26 (1977), 214–220.
- [123] TE. Sterne , Some remarks on confidence or fiducial limits, *Biometrika.*, 41 (1954), 275–278.
- [124] MW. Fagerland , Exacten and mid-p confidence intervals for the odds ratio, *Stata Journal.*, 12 (2012), 505.
- [125] A. Agresti and Y. Min , Unconditional small-sample confidence intervals for the odds ratio, *Biostatistics.*, 3 (2002), 379–386.
- [126] MW. Fagerland, S. Lydersen, and P. Laake, Recommended confidence intervals for two independent binomial proportions, *Statistical methods in medical research.*, 24 (2015), 224–254.
- [127] PJ. Laud, and A. Dane, Confidence intervals for the difference between independent binomial proportions: comparison using a graphical approach and moving averages, *Pharmaceutical statistics.*, 13 (2014), 294–308.

- [128] L. Brown, and Li. Xuefeng, Confidence intervals for two sample binomial distribution, *Journal of Statistical Planning and inference.*, 130 (2005), 359–375.
- [129] JJ. Gart, and JM. Nam, Approximate interval estimation of the ratio of binomial parameters: a review and corrections for skewness, *Biometrics.*, 44 (1988), 323–338.

برآورد طول دوره کمون در بیماران مبتلا به HIV در استان خراسان رضوی

نجمه نخعی راد^۱، گروه ریاضی و آمار، دانشگاه آزاد اسلامی واحد مشهد وحید فکور گروه آمار، دانشگاه فردوسی مشهد

چکیده: در این مقاله به برآورد طول دوره کمون در بیماران مبتلا به HIV در استان خراسان رضوی پرداخته شده است. دوره کمون یا دوره نهفتگی ویروس HIV به مدت زمان از ورود ویروس به بدن تا آشکار شدن ایدز اطلاق می‌شود. در ابتدا طول دوره کمون در حالت ناپارامتری با استفاده از جدول طول عمر و برآوردگر کاپلان-مه-یر برآورد شده است. سپس با استفاده از یک مدل پارامتری پیشنهادی برآورد پارامتری طول دوره کمون به دست آورده شده است. و در انتها عوامل مؤثر بر طول دوره کمون توسط روش رگرسیونی کاکس شناسایی شده اند.

واژه‌های کلیدی: سانسور راست، جدول طول عمر، برآوردگر کاپلان-مه-یر، رگرسیون کاکس.

۱ مقدمه

در سال ۱۹۸۱، در نیویورک، با مشاهده چندین بیماری وخیم شبیه سرطان و دیده شدن انواع متفاوت دیگری از عفونت‌های ناشناخته در ایالت متحده آمریکا، توجه دانشمندان به ویروس مشترکی جلب شد که با تغییر شکل مداوم، سیستم ایمنی بدن را از کار می‌اندازد و در نتیجه امکان ابتلا به انواع بیماری‌های خطرناک دیگر را فراهم می‌کند. اولین مواردی که مشاهده شد در بین معتادان تزریقی و همچنین همجنس‌گرایان مرد قرار داشت که به دلیل نامعلومی سیستم دفاعی بدن آن‌ها ضعیف شده بود و علائم بیماری التهاب ریه پنوموسیستیک کارینی (Pneumocystis Carinii) PC در آن‌ها مشاهده می‌شد. سپس نوعی سرطان پوست نادر با نام کاپوسی سارکوما (Kaposi's Sarcoma) KS در میان مردان همجنس‌گرا گزارش شد. گزارش موارد بیشتری از بیماری PC و KS زنگ خطری برای مرکز پیشگیری و کنترل بیماری بود. در آن اوایل مرکز کنترل بیماری هنوز نامی رسمی برای این بیماری انتخاب نکرده بود و معمولاً نام این بیماری را با بیماری که مریض با آن در ارتباط بود بیان می‌کردند. واژه ایدز در همایشی در ژوئیه ۱۹۸۲ معرفی شد و از سپتامبر ۱۹۸۲، مرکز کنترل بیماری از واژه ایدز برای نسبت دادن این بیماری استفاده کرد. همزمان در آفریقا نیز موارد متعددی از شیوع این علائم دیده شد که منجر به مرگ و میر فراوانی در این قاره گردید. دانشمندان بر این عقیده اند که شکل ابتدایی این ویروس در شامپانزه‌ها وجود داشته است که به دلیل تغذیه ساکنان گینه بیسائو و قصابی این حیوانات از طریق زخم‌های انسانی به انسان منتقل

^۱نجمه نخعی راد: najme.rad@gmail.com

شده و در طول یک قرن تکامل یافته و به دیگران انتقال یافته است. راه های عمده انتقال این ویروس عبارتند از: استفاده از سوزن و سرنگ مشترک (تزریق مواد مخدر با سرنگ مشترک، خالکوبی، وسایل آلوده کار به خصوص در بیمارستان ها و دندانپزشکی ها و ...)، رفتارهای جنسی پرخطر، خون و فرآورده های خونی و از مادر مبتلا به جنین. (بنت و همکاران، ۲۰۱۴)

HIV (Human Immune-deficiency Virus) یا ویروس نقص ایمنی انسانی، ویروسی است که سبب صدمه به دستگاه ایمنی بدن می شود و هنگامی که فرد مبتلا به HIV شده یا اصطلاحاً HIV مثبت می شود، ویروس به تدریج شروع به از بین بردن سیستم ایمنی بدن وی کرده به طوریکه بدن نمی تواند حتی با بیماری های جزئی مقابله کند و مستعد ابتلا به انواع بیماری ها می شود. ویروس HIV در دو نوع HIV^۱ و HIV^۲ است که HIV^۲ نادرتر و کم خطرتر می باشد. با پیشرفت ویروس و درگیری سلول های ایمنی موجود در خون (لنفوسیت ها) و همچنین سلول های ایمنی موجود در بافت ها مانند مغز استخوان، طحال، کبد و گره های لنفاوی این سلول ها دیگر توانائی تولید پادتن برای مقابله با بیماری ها را نداشته و فرد مبتلا به بیماری هایی نظیر انواع سرطان ها، توکسوپلاسمای مغزی و ... می گردد.

یک چنین وضعیتی ایدز (Acquired Immune Deficiency Syndrome) AIDS نامیده می شود که همراه با پیدایش علائم و نشانه های گوناگون در بدن می باشد. برخی از نشانه ها، همانطور که پیش تر گفته شد در سال ۱۹۸۱ در افراد بالغ جوانی دیده شد که دچار نقص ایمنی مادرزادی نبودند و این تعجب پزشکان را برانگیخت و به همین دلیل مجموعه علائم را نشانه های نقص ایمنی اکتسابی یا ایدز نام نهادند چون نمی توانستند آن را به بیماری مشخصی نسبت دهند و هنوز عامل آن ناشناخته بود. دوره کمون یا انکوباسیون در لغت به معنای تکوین بیماری عفونی از هنگام ورود پاتوژن تا پیدایش نشانه های بالینی می باشد. طبق همین تعریف، دوره کمون در بیماری ایدز از هنگام ورود ویروس HIV به بدن فرد تا پیدایش علائم ایدز ادامه دارد. این مدت زمان معمولاً سال ها به طول می انجامد و در طول این مدت فرد ناقل بیماری بوده بدون اینکه علائمی از بیماری را نشان دهد. شروع بیماری ایدز از دو طریق قابل تشخیص است:

۱. شمارش سلول های $Tcd4$ موجود در هر میلی متر خون: هر گاه تعداد این سلول ها به کمتر از ۲۰۰ عدد در هر میلی متر برسد اصطلاحاً بیمار وارد مرحله ایدز شده است.

۲. از طریق علائم و نشانگان: گاه ممکن است فرد تعداد سلول های $Tcd4$ بیش از ۲۰۰ داشته باشد ولی علائم و بیماری های از قبیل: قارچ دهانی، توکسوپلاسمای مغزی، سرطان پوست، PC، و یا پلاکت های مقاوم به درمان در وی مشاهده گردد، در این حالت نیز بیمار وارد مرحله ایدز شده است. (بنت و همکاران، ۲۰۱۴)

۲ اطلاعاتی درباره تحقیق:

اهداف تحقیق:

همانطور که پیش تر گفته شد فرد مبتلا به HIV در طول دوره کمون ناقل این ویروس بوده در صورتی که به هیچ طریقی به جز آزمایش خون قابل تشخیص نمی باشد. این افراد که معمولاً دارای رفتارهای پرخطر نیز می باشند با ظاهری سالم باعث انتقال ویروس به سایرین می گردند. لذا برآورد طول دوره کمون این بیماری دارای اهمیت می باشد.

هدف کلی:

برآورد طول دوره کمون در بیماران مبتلا به HIV در خراسان رضوی.

اهداف فرعی:

۱. برآورد طول دوره کمون به روش ناپارامتری (جدول طول عمر و برآوردگر کاپلان-مه-یر)

۲. برآورد طول دوره کمون براساس یک مدل پارامتری پیشنهادی

۳. یافتن عوامل مؤثر بر طول دوره کمون توسط روش رگرسیونی کاکس

نوع مطالعه:

از آنجایی که در حین اجرای تحقیق دخالتی در وضعیت موجود افراد صورت نپذیرفته، مطالعه به طور کلی یک مطالعه غیر تجربی یا مشاهده ای می باشد و طبق تعریف مطالعه مقطعی، مطالعه ای از این نوع نیز محسوب می گردد. در مطالعه مقطعی، تمام جامعه یا قسمتی از آن انتخاب شده و آزمایشگر داده های مورد نظرش را گردآوری می کند. چون اطلاعات جمع آوری شده، فقط آنچه را که در یک نقطه زمانی درباره ی پارامترهای جامعه وجود دارد نشان می دهد، به این نوع مطالعه مقطعی می گویند. همچنین این مطالعه با توجه به برخی سؤال ها که از سوابق افراد شده است، مشمول مطالعه گذشته نگر نیز می شود.

جامعه آماری:

جامعه آماری، کلیه بیماران مبتلا به HIV در استان خراسان رضوی را شامل می شود.

واحد آماری:

واحد آماری، فرد بیمار مبتلا به HIV می باشد.

روش نمونه گیری و حجم نمونه:

با توجه به مشخص نبودن چارچوب دقیق جامعه آماری از جمله حجم جامعه مورد بررسی یعنی تعداد افراد مبتلا به HIV در استان خراسان رضوی و محدودیت شناسایی و معرفی این افراد، نمونه مورد بررسی شامل کلیه افرادی است که به مرکز مشاوره بیماری های عفونی □ رفتاری وابسته به اداره بهداشت خراسان رضوی مراجعه کرده و یا توسط این مرکز و مراکزی از قبیل مراکز درمانی، انتقال خون، بهزیستی و زندان های استان شناسایی شده اند و در حال حاضر تحت پوشش این مرکز هستند.

روش و ابزار گردآوری اطلاعات :

برای جمع آوری اطلاعات از پرسش نامه استفاده شد که طبق اطلاعات موجود در پرونده های بیماران در مرکز بیماری های عفونی □ رفتاری تکمیل گردید و در مواردی که اطلاعات درج شده در پرونده ناقص بود از پزشک و افراد مشغول به کار در مرکز و همچنین بیماران نیز سؤالاتی پرسیده شد.

محدودیت های در تحقیق:

۱. حجم کم نمونه به دلیل شناسایی مشکل بیماران مبتلا به HIV و گاه در عین شناسایی عدم تمایل فرد به ادامه مراجعه به مرکز

۲. تعداد مشاهدات سانسور شده ی زیاد در مطالعه به دلیل نوبا بودن مرکز و فعالیت شناسایی بیماران مبتلا از طریق پایگاه های

دیده وری

۳. اطلاعات ثبت نشده در پرونده بیماران

پیش فرض ها در انجام تحقیق:

۱. از آنجائیکه زمان دقیق آلوده شدن فرد به ویروس HIV به طور دقیق قابل تشخیص نیست و ویروس HIV تنها از راه های گفته شده در قبل قابل انتقال است لذا از هر فرد با توجه به نوع رفتار پرخطر یا زمان اولین تزریق مشترک پرسیده شده و تاریخ آلوده شدن فرد به ویروس HIV نقطه وسط بازه ی زمان اولین رفتار پرخطر و زمان تشخیص ابتلا به HIV در نظر گرفته شده است.

۲. اطلاعات موجود در پرونده بیماران قابل اطمینان بوده و به طور دقیق ثبت شده اند.

۳. در طول دوره کمون هیچ دارویی که باعث تأثیر بر طول دوره گردد، مصرف نمی شود.

۴. بیمار آلوده به ویروس *HIV* حتماً روزی وارد مرحله ایدز می‌شود.

مدل تعریف شده برای داده ها:

مطالعه از سال ۱۳۷۵ آغاز شده و تا ۱۳۸۷ ادامه یافته است. در طول این مدت افراد در زمان های مختلف با مراجعه به مرکز و یا شناسایی، وارد مطالعه شده اما زمان پایان مطالعه برای تمام آن ها یکسان است. بیماران از زمان ورود به مطالعه تا شروع ایدز یا سانسور پیگیری شده اند.

متغیرهای تحقیق:

متغیر وابسته: طول دوره کمون به ماه متغیر وابسته است که متغیری کمی و مقیاسی است.

متغیرهای مستقل: متغیرهای مستقل بر اساس اطلاعات موجود در پرونده های بیماران در نظر گرفته شده اند چون به دلیل سختی دسترسی به بیماران (مثلاً به دلیل عدم آگاهی برخی خانواده ها امکان تماس تلفنی وجود نداشت، گذشته از اینکه اسم و آدرس این افراد محرمانه نیز می باشد و همچنین چون نوبت آزمایشات برای بیماران *HIV* مثبت، سالیانه و برای بیماران ایدزی هر سه ماه یکبار است ملاقات حضوری با تعداد زیادی از بیماران امکان پذیر نبود) امکان تعریف متغیرهای دیگر وجود نداشت. متغیرهای تعریف شده عبارتند از: جنسیت، سن، وضعیت تأهل، اعتیاد به مواد مخدر، اعتیاد به سیگار، میزان تحصیلات، شغل، درآمد، سابقه حضور در زندان، عامل ایجاد کننده بیماری، نوع ویروس، نحوه پیگیری، ابتلا به بیماری های هپاتیت *C(HCV)*، هپاتیت *B(HBS)*، سل *(TB)*، *VDRL* و ویروس های *CMV*، *HTLV1* و *HTLV2*.
روش های تجزیه و تحلیل:

۱. جدول طول عمر

۲. برآوردگر کاپلان-مه-یر

۳. رگرسیون کاکس

۴. استفاده از یک مدل پارامتری پیشنهادی برای برآورد زمان تا ایدز در داده های سانسور شده از راست تجزیه و تحلیل با نرم افزارهای *R* و *SPSS* انجام شده است.

۳ نتایج تحقیق:

۱.۳ اطلاعات توصیفی:

از مجموع ۶۴ فرد حاضر در مطالعه، ۵۳ نفر مرد و ۱۱ نفر زن هستند. میانگین سن این افراد ۳۶/۱۶ و انحراف معیار سن ۸/۴۷ سال می باشد. ۲۱ نفر متأهل و مابقی مجرد، مطلقه و بیوه هستند. ۸۷/۵ درصد افراد تحصیلات زیر دیپلم داشته و ۸۴/۴ درصد سابقه حضور در زندان را داشته اند. ۷۰ درصد افراد اعتیاد به سیگار و ۱۷/۲ درصد اعتیاد به مواد مخدر داشته و ۷۰/۴ درصد متادون مصرف کرده یا در ترک کامل هستند. ۸۰ درصد مبتلایان سابقه استفاده از سوزن و سرنگ مشترک، ۱۱ درصد رابطه جنسی پر خطر و ۹ درصد تجربه هر دو مورد را داشته اند. ۶۲ نفر از افراد آلوده به ویروس *HIV1* و ۲ نفر آلوده به ویروس *HIV2* هستند. تعداد مبتلایان به هر یک از بیماری های هپاتیت *C(HCV)*، هپاتیت *B(HBS)*، سل *(TB)*، *VDRL* و ویروس های *CMV*، *HTLV1* و *HTLV2* در جدول زیر آمده است:

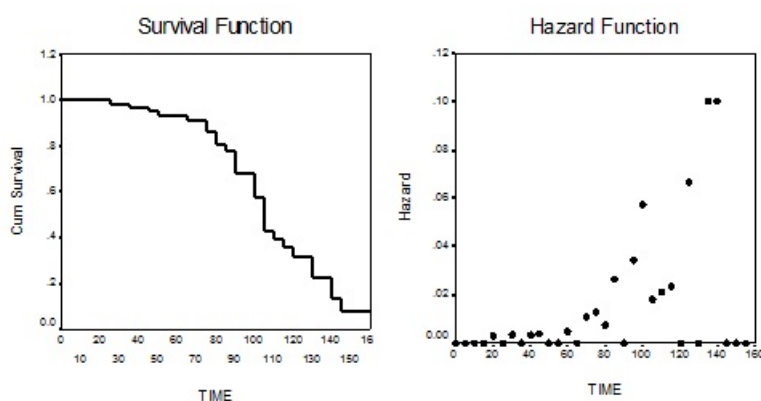
بیماری	HCV	HBS	TB	VDRL	CMV	HTLV ^۱	HTLV ^۲
دارد	۳۹	۹	۱۲	۰	۶	۵	۲
ندارد	۲۵	۵۵	۵۲	۶۴	۵۸	۵۹	۶۲

۲.۳ برآورد ناپارامتری طول دوره کمون:

همانطور که گفته شد مطالعه از سال ۱۳۷۵ آغاز شده و تا ۱۳۸۷ ادامه یافته است. در طول این مدت افراد در زمان های مختلف با مراجعه به مرکز و یا شناسایی وارد مطالعه شده اما زمان پایان مطالعه برای تمام آن ها یکسان است. بیماران از زمان ورود به مطالعه تا شروع ایدز یا سانسور پیگیری شده اند. به این ترتیب از مجموع ۶۴ فرد آلوده به ویروس HIV از آنجا که ۳۶ نفر تا پایان مطالعه پیشامد ایدز را تجربه نکرده اند سانسور می شوند. برای برآورد ناپارامتری طول دوره کمون از جدول طول عمر و برآوردگر کاپلان-مه-یر استفاده کرده ایم:

جدول طول عمر:

زمان های تشخیص ایدز را ۵ ماهه در نظر گرفته و احتمالات ۵ ماهه را محاسبه کرده ایم. نتایج نشان می دهد که میانه طول دوره کمون برای مشاهدات ۱۰۲/۵۸ ماه است. نمودار تابع بقا تجمعی و تابع مخاطره به دست آمده توسط این روش در شکل ۱ آمده است.



شکل ۱: نمودار تابع بقا و تابع مخاطره به دست آمده توسط جدول طول عمر

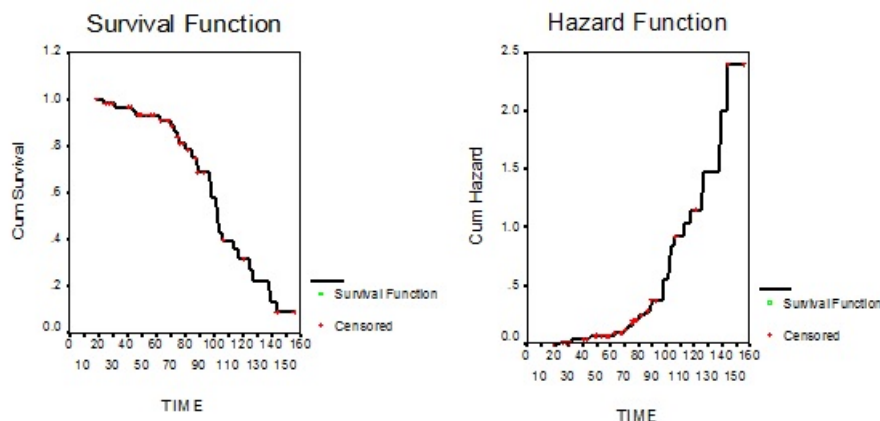
برآوردگر کاپلان-مه-یر: در این روش میانگین و میانه طول دوره کمون به ترتیب ۱۰۵ و ۱۰۲ ماه به دست آمده است (جدول ۱). نمودار تابع بقا تجمعی و تابع مخاطره به دست آمده توسط این روش در شکل ۲ آمده است.

جدول ۱: میانگین و میانه طول دوره کمون در روش کاپلان-مه-یر

	Survival Time	Confidence Interval ۹۵	Standard Error
Mean	۱۰۵	(۹۵, ۱۱۵)	۵
Median	۱۰۲	(۹۵, ۱۰۹)	۴

۳.۳ برآورد پارامتری طول دوره کمون با استفاده از یک مدل پیشنهادی:

y را یک متغیر برنولی تعریف می کنیم. به طوریکه $y = 1$ نشان دهنده وقوع ایدز و $y = 0$ بیانگر عدم رسیدن به مرحله ایدز باشد. فرض کنید T زمان تشخیص ایدز و X بردار متغیرهای توضیحی باشد که می تواند برای هر فرد تعریف شود. در این مدل احتمال وقوع



شکل ۲: نمودار تابع بقا و تابع مخاطره به دست آمده توسط روش کاپلان-میر

ایدز پس از آلوده شدن با ویروس *HIV* را برابر نسبت افراد آلوده در نمونه که به مرحله ایدز رسیده اند به کل افراد در نمونه در نظر گرفته و توزیع زمان تا پدیدار شدن ایدز را

$$f(t|y = 1, X) = \lambda \delta t^{\delta-1} \exp(-\lambda t^{\delta}),$$

فرض می‌کنیم که $\lambda, \delta > 0$.

اگر ایدز در زمان t_i برای نفر i ام مشاهده شود در این صورت سهم درستنمایی برای این شخص برابر است با

$$p(y_i = 1)f(t_i|y = 1, X)$$

و اگر نفر i ام تا زمان t_i بدون مشاهده ایدز پیگیری شود سهم درستنمایی وی

$$p(y_i = 0|X) + pr(y_i = 1|X) \int_{t_1}^{\infty} f(u|y = 1, X) du$$

می‌باشد. این مدل در سال ۱۹۸۸ توسط لویی و همکاران هنگامی که t معلوم یا سانسور شده باشد بحث شده است. برآورد درستنمایی ماکسیمم پارامترهای مدل در جدول ۲ آمده است ($-2 \log L = 359.833$). میانگین طول دوره کمون با استفاده از این مدل ۰۶/۱۰۰

جدول ۲: برآورد درستنمایی ماکسیمم پارامترهای مدل

Coefficients	Estimate	std. error
delta	۱/۲۵۹۷۴۲۸۴۶	۰/۰۷۵۲۸۰۹۰
lambda	۰/۰۰۲۷۵۶۳۳۹	۰/۰۰۰۸۰۶۵۳

ماه بدست آمده است که با میانگین بدست آمده در روش های ناپارامتری تفاوت چندانی ندارد.

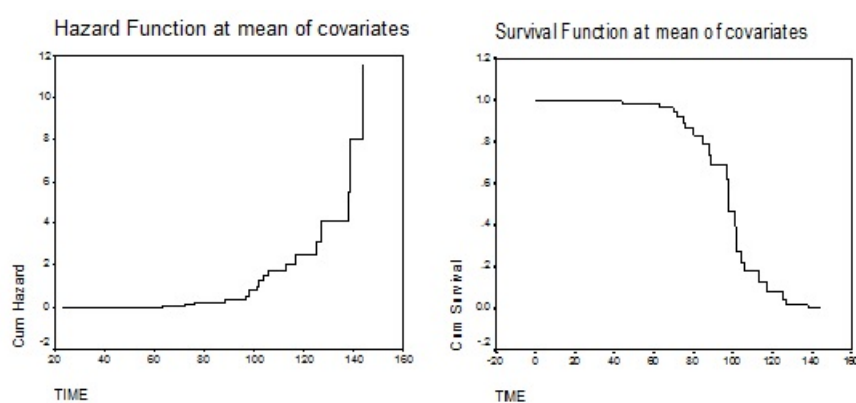
۴.۳ یافتن عوامل مؤثر بر طول دوره کمون توسط روش رگرسیونی کاکس:

رگرسیون کاکس در ۱۲ مرحله تکرار شده است به این ترتیب که ابتدا تمام متغیرهای توضیحی وارد مدل شده، سپس در طی ۱۲ مرحله از مدل خارج می‌شوند و در مرحله آخر با انتخاب خطای ۰/۰۵ متغیرهای جنسیت، اعتیاد به سیگار، اعتیاد به مواد مخدر (اعتیاد و مصرف متادون)، تحصیلات، راه انتقال (استفاده از سوزن و سرنگ مشترک)، ابتلا به هپاتیت C و ویروس HTLV۱ با ضرایب آمده در جدول زیر در مدل حضور دارند. نمودار تابع بقا تجمعی و تابع مخاطره به دست آمده توسط این روش در شکل ۳ آمده است.

Variables in the Equation^a

		B	SE	Wald	df	Sig.	Exp(B)
Step 12	SEX	3.587	1.411	6.466	1	.011	36.126
	AGE	-.075	.044	2.910	1	.088	.928
	SMOKING	-1.297	.507	6.556	1	.010	.273
	ADICATION			8.110	3	.044	
	ADICATION(1)	-1.323	1.281	1.067	1	.302	.266
	ADICATION(2)	-6.546	2.570	6.488	1	.011	.001
	ADICATION(3)	-2.137	1.181	3.274	1	.070	.118
	EDUCAT	.586	.267	4.815	1	.028	1.797
	FACTOR			14.634	2	.001	
	FACTOR(1)	1.649	2.049	.648	1	.421	5.204
	FACTOR(2)	5.431	1.641	10.958	1	.001	228.319
	HCV	-2.103	.774	7.381	1	.007	.122
	HTLV1	3.009	.772	15.179	1	.000	20.268

جدول ۳: متغیرهای موجود در مدل رگرسیون کاکس در مرحله آخر



شکل ۳: نمودار تابع بقا و تابع مخاطره به دست آمده توسط مدل رگرسیون کاکس

- [130] Bennett, J. E., Dolin, R., & Blaser, M. J. (2014). *Principles and practice of infectious diseases (Vol. 1)*. Elsevier Health Sciences.
- [131] Lui, K.- J., Darraw, W. W., Rutherford, G. W., (1988). *A model – based estimate of the mean incubation period for AIDS in homosexual men*. Science 240, 1333-1335

خواص مجانبی برآوردهای کاپلان-مه-یر و نلسون-آلن برای داده های سانسور شده ی وابسته نجمه نخعی راد^۱، گروه ریاضی و آمار، دانشگاه آزاد اسلامی واحد مشهد وحید فکور گروه آمار، دانشگاه فردوسی مشهد

چکیده: در مطالعات بقا، برآورد تابع بقا که در توصیف مدت زمان تا وقوع پیشامدی خاص به کار می رود و همچنین تابع تجمعی خطر از اهمیت بسیاری برخوردار می باشد. در حالتی که زمان های بقا مستقل از یکدیگرند، بیشتر برآوردهای این توابع معرفی و خواص مجانبی آن ها بررسی شده است. در این مقاله برقراری خواص مجانبی برآوردهای کاپلان-مه-یر و نلسون-آلن تابع تجمعی خطر برای همان برآوردهای پیشین، در صورت وابستگی زمان های بقا، در مدل سانسور راست نشان داده شده است که شامل سازگاری قوی و همگرایی ضعیف می باشد. پیش از بیان همگرایی ضعیف، ابتدا نشان می دهیم که تحت شرایط نظم می توان این برآوردها را به صورت میانگین متغیرهای تصادفی به اضافه یک جمله باقیمانده نمایش داد و در ادامه از آن در اثبات همگرایی ضعیف استفاده می کنیم.

واژه های کلیدی: α -آمیزنده، داده های سانسور شده ی وابسته، برآوردهای کاپلان-مه-یر، برآوردهای نلسون-آلن، سازگاری قوی، همگرایی ضعیف

۱ مقدمه

تحلیل بقا اصطلاحی است که برای بیان تحلیل داده های مربوط به زمان، از یک شروع خوش تعریف تا وقوع یک یا چند پیشامد مشخص یا نقطه پایان به کار می رود و در علوم کاربردی نظیر پزشکی، زیست شناسی، اپیدمی شناسی، مهندسی، اقتصاد و جمعیت شناسی مطرح است. در تحلیل داده های بقا به دو دلیل عمده نمی توان از روش های استاندارد تحلیل آماری استفاده نمود. دلیل اول اینکه داده های بقا به طور نرمال توزیع نشده اند و اکثر توزیع ها چوله به راست هستند. در نتیجه پذیرفتن فرض نرمال بودن این قبیل داده ها معقول نمی باشد که البته این مشکل را می توان با تبدیل داده ها برای رسیدن به توزیع متقارن تر برطرف نمود. هر چند که رویکرد بهتر، قبول یک مدل توزیعی جایگزین برای داده ها می باشد. دلیل دوم اینکه این داده ها به صورت هایی ظاهر می شوند که مشکلاتی را در تحلیل آن ها ایجاد می کند. از آن جمله وجود مشاهدات سانسور شده در بین آن ها است. مشاهدات سانسور شده زمانی رخ می دهند که وضعیت بقای واحدهای مورد مطالعه در رابطه با وقوع یا عدم وقوع پیشامدی خاص در دوره زمانی مشخصی، معلوم باشد. سانسور می توان به انواع سانسور راست، چپ و فاصله ای تقسیم بندی کرد. سانسور از راست که در این مقاله از آن صحبت شده است در سه

^۱نجمه نخعی راد: najme.rad@gmail.com

نوع کلی سانسور نوع I، سانسور نوع II و سانسور نوع III یا سانسور تصادفی می‌باشد:

سانسور نوع I: در سانسور نوع I پیشامد مطلوب تنها وقتی قابل مشاهده است که قبل از یک زمان از پیش تعیین شده رخ دهد. مطالعه ای را در نظر بگیرید که با تعداد ثابتی از اعضا (بیماران، حیوانات...) در زمانی مشخص آغاز می‌شود. به لحاظ صرفه جویی در زمان و هزینه مطالعه باید پیش از اینکه تمام اعضا به پیشامد مطلوب رسیده باشند خاتمه یابد و نتایج گزارش شوند. در این حالت اگر پیشامد مورد نظر برای عضوی رخ دهد، زمان تا پیشامد این عضو معلوم می‌باشد ولی اگر عضو تا پایان مطالعه بقا داشت و هیچ حادثه ای غیر از پیشامد مورد نظر باعث خارج شدن عضو از مطالعه نگردد در این صورت مشاهده، سانسور شده از راست با زمان بقای برابر با طول دوره مطالعه است.

حال مطالعه ای را در نظر بگیرید که در آن اعضا، زمان های سانسور مشخص ولی متفاوت از یکدیگر دارند. این شکل از سانسور نوع I سانسور فزاینده نوع I نامیده می‌شود.

نوع دیگری از سانسور که تعمیمی از سانسور نوع I می‌باشد به این صورت است که اعضا در زمان های متفاوتی وارد مطالعه می‌شوند اما زمان پایان مطالعه برای تمام اعضا یکسان است و از قبل توسط آزمایشگر تعیین می‌شود. زمان های سانسور اعضا هنگامی که وارد مطالعه می‌شوند مشخص است لذا در این نوع مطالعه هر عضو زمان سانسور ثابت و خاص خود را دارد. این شکل از سانسور، سانسور تعمیم یافته نوع I نامیده می‌شود.

سانسور نوع II: نوع دیگر سانسور راست، سانسور نوع II است که در آن اعضا در یک زمان مشابه وارد مطالعه می‌شوند و مطالعه تا زمانی که اولین r عضو، پیشامد مطلوب را تجربه کند ادامه می‌یابد. r عددی از پیش تعیین شده می‌باشد ($r \leq n$). یک چنین مطالعه ای صرفه جویی در زمان و هزینه را در پی دارد زیرا ممکن است زمان بسیار زیادی طول بکشد تا تمام اعضا پیشامد مطلوب را تجربه کنند. این نوع از سانسور معمولاً در تست عمر کالاها و تجهیزات استفاده می‌شود. نوعی سانسور که از سانسور نوع II نتیجه می‌شود سانسور فزاینده نوع II می‌باشد. این شکل از سانسور همانند سانسور فزاینده نوع I می‌باشد یعنی زمان در مطالعه ی (یا زمان سانسور) اعضا مشخص ولی متفاوت از یکدیگر است با این تفاوت که زمان های سانسور، در سانسور فزاینده نوع II برخلاف سانسور فزاینده نوع I تصادفی می‌باشد. در سانسور فزاینده نوع I زمان های سانسور از قبل توسط آزمایشگر مشخص می‌شوند و مقادیری ثابت می‌باشند. اما در سانسور فزاینده نوع II بدین گونه است که اعضای در مطالعه به $\square$ گروه به حجم های n_i تقسیم می‌شوند. خاتمه آزمایش در هر گروه وقتی است که تعداد r_i عضو آن گروه پیشامد مطلوب را تجربه کنند لذا اعضای باقیمانده که پیشامد مطلوب را تجربه نکرده اند سانسور می‌شوند. بنابراین زمان های سانسور تصادفی می‌باشند.

سانسور نوع III: گاهی اوقات ممکن است اعضای وارد شده به یک مطالعه پیشامدهایی غیر از پیشامد مطلوب در مطالعه را تجربه کنند و این باعث حذف آنها از ادامه مطالعه گردد. پیشامدهای نظیر مرگ های تصادفی یعنی مرگ در اثر هر عاملی غیر از پیشامد مطلوب در مطالعه، مهاجرت، عدم میل به شرکت در ادامه مطالعه و در چنین مواقعی پیشامد مورد نظر قابل مشاهده نیست. این نوع سانسور، سانسور تصادفی نامیده می‌شود.

در حالتی که مشاهدات سانسور شده ی تصادفی از راست و دو به دو مستقل از یکدیگرند، برآوردگر تابع بقا به برآوردگر کاپلان-مه-یر (۱۹۵۸) مشهور است و خواص مجانبی آن در طول سال های اخیر به طور وسیع توسط محققین بسیاری از جمله برسلو و کرولی (۱۹۷۴)، پترسون (۱۹۷۷)، گیل (۱۹۸۳، ۱۹۸۱، ۱۹۸۰)، لو و سینگ (۱۹۸۶)، ونگ (۱۹۸۷)، و اشتوت و ونگ (۱۹۹۳) مورد مطالعه قرار گرفته است. این در حالیکه نتایج اندکی برای مشاهدات بقا در حالتی که مستقل نیستند در دسترس است: ولکل و کرولی (۱۹۸۴) و یینگ و وی (۱۹۹۴) سازگاری و نرمال بودن مجانبی برآوردگر کاپلان-مه-یر را تحت وابستگی های نیمه مارکوفی و آمیختگی بیان کردند.

در ادامه ابتدا در بخش دوم متغیرها، تعاریف و پیش فرض های مورد نیاز در دنباله بحث را معرفی می‌کنیم. سپس در بخش سوم سازگاری قوی دو برآوردگر کاپلان-مه-یر تابع بقا و نلسون آلن تابع تجمعی خطر را نشان خواهیم داد. در بخش چهارم همگرایی ضعیف این دو برآوردگر را بررسی می‌کنیم.

۲ تعاریف، نمادها و پیش فرض ها

فرض کنید $X_1, \dots, X_n$ دنباله ای از زمان های بقا برای n فرد باشند. فرض کنید این دنباله دو به دو مستقل فرض نشده اند اما دارای یک تابع توزیع حاشیه ای پیوسته، مشترک و نامعلوم $F(t) = P(X_i \leq t)$ می باشند به طوریکه $F(\cdot) = 0$. و همچنین فرض کنید که متغیرهای تصادفی X_i سانسور شده از راست با متغیرهای تصادفی سانسور Y_i باشند در این صورت تنها (Z_i, δ_i) قابل مشاهده است که در آن

$$Z_i = X_i \wedge Y_i \quad \text{و} \quad \delta_i = I(X_i \leq Y_i) \quad i = 1, \dots, n$$

$I(A)$ متغیر تصادفی نشانگر پیشامد A و $\wedge$ نشانه مینیمم است.

در این مدل سانسور تصادفی، زمان های سانسور $Y_i, i = 1, \dots, n$ متغیرهای تصادفی مستقل و هم توزیع با تابع توزیع $G(y) = P(Y_i < y)$ فرض شده اند به طوریکه $G(\cdot) = 0$ ، که همچنین مستقل از متغیرهای تصادفی X_i می باشند.

هدف ارائه استنباط ناپارامتری درباره F بر پایه مشاهدات سانسور شده $(Z_i, \delta_i), i = 1, \dots, n$ می باشد برای این منظور دو فرآیند تصادفی $N_n(t)$ و $Y_n(t)$ را بر $[\cdot, \tau]$ به صورت زیر تعریف می کنیم:

$$N_n(t) = \sum_{i=1}^n I(Z_i \leq t, \delta_i = 1) = \sum_{i=1}^n I(X_i \leq t \wedge y_i)$$

$$Y_n(t) = \sum_{i=1}^n I(Z_i \geq t)$$

که $N_n(t)$ تعداد مشاهدات سانسور نشده کوچکتر یا مساوی t و $Y_n(t)$ تعداد مشاهدات سانسور شده و سانسور نشده ی بزرگتر یا مساوی t می باشد.

برآوردگر کاپلان مه -یر $(K - M)$ تابع بقا به صورت

$$1 - F_n(t) = \prod_{s \leq t} \left(1 - \frac{dN_n(s)}{Y_n(s)} \right)$$

است که در آن $dN_n(s) = N_n(s) - N_n(s^-)$. برای تابع توزیع F بر $[\cdot, \tau]$ ، تابع تجمعی خطر Δ درحالتی که F پیوسته است (گیل، ۱۹۸۰) به صورت:

$$\Delta(t) = \int_{\cdot}^t \frac{dF(s)}{1 - F(s^-)} = -\log(1 - F(t))$$

تعریف می شود که تابع تجربی آن، $\Delta_n(t)$ ، که به برآوردگر نلسون-آلن $\Delta(t)$ معروف است به صورت زیر می باشد.

$$\Delta_n(t) = \int_{\cdot}^t \frac{dN_n(s)}{Y_n(s)}$$

پیشتر فرض کردیم که زمان های بقا دو به دو مستقل نیستند و اکنون فرض می کنیم دارای وابستگی α -آمیزنده با تعریف زیر می باشند:

تعریف ۱.۲. دنباله $\{X_i, i = \cdot, \pm 1, \pm 2, \dots\}$ یک دنباله ی α -آمیزنده از متغیرهای تصادفی است هرگاه، برای دو مجموعه $A \in F_1^k(X), B \in F_{k+n}(X)$ و عدد صحیح مثبت n :

$$\alpha(n) = \sup_{A \in F_1^k(X), B \in F_{k+n}(X)} |P(A \cap B) - P(A)P(B)| \rightarrow 0; \quad n \rightarrow \infty$$

لم ۲.۲ (کای ۱۹۹۸). اگر $\{X_n; n \geq 1\}$ دنباله ای از متغیرهای تصادفی α -آمیزنده با ضریب آمیختگی $\alpha(n)$ و $\{Y_n; n \geq 1\}$ دنباله ای از متغیرهای تصادفی iid و مستقل از $\{X_n; n \geq 1\}$ باشد آن گاه $\{(X_n, Y_n); n \geq 1\}$ و همچنین $\{X_n \wedge Y_n; n \geq 1\}$ دنباله ای از متغیرهای تصادفی α -آمیزنده با ضریب آمیختگی $\alpha(n)$ می باشند.

فرض کنید H تابع توزیع Z_i ها به صورت زیر داده شده باشد:

$$\bar{H} = 1 - H = \bar{F}\bar{G} = (1 - F)(1 - G)$$

زمان های (احتمالاً نامتناهی) τ_H, τ_G, τ_F را به صورت

$$\tau_F = \inf\{y : F(y) = 1\}$$

تعریف می کنیم. به سادگی دیده می شود که $\tau_H = \tau_F \wedge \tau_G$ (ونگ و اشتوت ۱۹۹۳) با تعریف

$$F_*(t) = P(Z \leq t, \delta = 1)$$

داریم:

$$F_*(t) = \int_{\cdot}^{\infty} F(t \wedge z) dG(z) = \int_{\cdot}^t (1 - G(z)) dF(z)$$

قرار می دهیم:

$$\alpha = \alpha(F, G) = P(X \leq Y) = \int_{\cdot}^{\infty} F(z) dG(z) = \int_{\cdot}^{\infty} [1 - G(z)] dF(z)$$

فرض می کنیم $\alpha > 0$. واضح است که $F(t)/\alpha$ تابع توزیع شرطی Z به شرط $\delta = 1$ می باشد.

اگر $\tau_{F_*} = \inf\{t : F(t) = \alpha\}$ به آسانی مشاهده می شود که $\tau_{F_*} = \tau_F \wedge \tau_G$ پس $\tau_H = \tau_{F_*}$

به این ترتیب برای تابع پیوسته F ، تابع تجمعی خطر را می توان به صورت

$$\Delta(t) = \int_{\cdot}^t \frac{F_*(s)}{\bar{H}(s)}$$

نوشت. قرار می دهیم:

$$\bar{N}_n(t) = N_n(t)/n, \quad \bar{Y}_n(t) = Y_n(t)/n$$

پیش از اثبات خواص مجانبی برآوردگرهای کاپلان مه-یر و نلسون-آلن فرض های زیر را می پذیریم.

K_1 : $\{X_j; j \geq 1\}$ دنباله ای از متغیرهای تصادفی α -آمیزنده با تابع توزیع پیوسته است.

K_2 : متغیرهای زمان سانسور $\{Y_j; j \geq 1\}$ با تابع توزیع پیوسته G و مستقل از $\{X_j; j \geq 1\}$ می باشند.

K_3 : $\alpha(n) = O(n^{-v})$ برای یک $v > 3$

۳ بررسی سازگاری قوی برآوردگرهای کاپلان مه-یر و نلسون-آلن

در این بخش به بررسی سازگاری قوی دو برآوردگر کاپلان مه-یر و نلسون-آلن می پردازیم. قضیه زیر سازگاری قوی را برای برآوردگر نلسون آلن نشان می دهد:

قضیه ۱.۳. تحت فرض های $(k_1) - (k_3)$

$$\sup_{r \geq \cdot} |\bar{Y}_n(t) - \bar{H}_n(t)| = O(a_n) \quad a.s. \quad (1.3)$$

$$\sup_{r \geq \cdot} |\bar{N}_n(t) - F_*(t)| = O(a_n) \quad a.s. \quad (2.3)$$

که $a_n = \left(\frac{\log \log n}{n}\right)^{\frac{1}{\tau}}$ و در نتیجه برای هر $0 < \tau < \tau_H$

$$\sup_{\cdot \leq r \leq \tau} |\Delta_n(t) - \Delta(t)| = O(a_n) \quad a.s. \quad (3.3)$$

پیش از اثبات قضیه ۱.۳ لم زیر را بیان می کنیم:

لم ۲.۳ (کای و روساس، ۱۹۹۲). اگر دنباله $\{X_n, n \geq 1\}$ دنباله ای مانا و α -آمیزنده از متغیرهای تصادفی با تابع توزیع F و ضریب آمیختگی $\alpha(n) = O(n^{-v})$, $v > 3$ و F_n تابع توزیع تجربی بر پایه $X_1, \dots, X_n$ باشد. آنگاه

$$\sup_{x \in R} |F_n(x) - F(x)| = O(a_n) \quad a.s.$$

که در آن $a_n = \left(\frac{\log \log n}{n}\right)^{\frac{1}{\tau}}$

اثبات. قضیه ۱.۳

از لم ۲.۲ به آسانی مشاهده می شود که $\{(Z_i, \delta_i); i \geq 1\}, \{Z_i; i \geq 1\}$ دو دنباله از متغیرهای تصادفی α -آمیزنده مانا می باشند. پس قسمت اول و دوم قضیه با استفاده از لم ۲.۳ و این حقیقت که $\bar{Y}_n$ و $\bar{N}_n$ تابع های تجربی هستند، نتیجه می شود. برای اثبات قسمت سوم، از لم ۲ در گیل (۱۹۸۱) و دو قسمت اول قضیه، نتیجه می شود که:

$$\sup_{\cdot \leq r \leq \tau} |\Delta_n(t) - \Delta(t)| \leq \frac{\rho(\bar{N}_n, F_*)}{\bar{Y}_n(\tau)} + \frac{\rho(\bar{Y}_n, \bar{H})[\bar{N}_n(\tau) + \rho(\bar{N}_n, F_*)]}{\bar{Y}_n(\tau)[\bar{Y}_n(\tau) - \rho(\bar{Y}_n(\tau), \bar{H})]} = O(a_n) \quad a.s.$$

□

که ρ سوپریمم بر $[\cdot, \tau]$ است و به این ترتیب حکم ثابت می شود.

پس از اثبات سازگاری قوی برآوردگر نلسون-آلن حال، در قضیه بعد سازگاری قوی را برای برآوردگر کاپلان مه-یر نشان می دهیم.

قضیه ۳.۳. تحت فرض های $(k_1) - (k_3)$ برای هر $0 < t < \tau_H$

$$\sup_{\cdot \leq r \leq \tau} |F_n(t) - F(t)| = O(a_n) \quad a.s.$$

اثبات. با نوشتن بسط تیلور برای تابع $e^{-\Delta(t)}$ در نقطه $\Delta_n(t)$ و تابع $e^{\log(1-F_n(t))}$ در نقطه $\Delta_n(t)$ داریم:

$$\begin{aligned} F_n(t) - F(t) &= 1 - F(t) - (1 - F_n(t)) \\ &= e^{-\Delta(t)} - e^{\log(1-F_n(t))} \\ &= e^{-\Delta_n^*(t)}[\Delta_n(t) - \Delta(t)] + e^{-\Delta_n^{**}(t)}[-\log(1 - F_n(t)) - \Delta_n(t)] \\ &= e^{-\Delta_n^*(t) + \Delta(t)} e^{-\Delta(t)} [\Delta_n(t) - \Delta(t)] \\ &\quad + e^{-\Delta_n^{**}(t) + \Delta_n(t)} e^{-\Delta_n(t) + \Delta(t)} e^{-\Delta(t)} [-\log(1 - F_n(t)) - \Delta_n(t)] \\ &= I + II \end{aligned} \quad (4.3)$$

که Δ^* و Δ^{**} در روابط زیر صدق می‌کنند:

$$\rho(\Delta_n^*, \Delta) \leq \rho(\Delta_n, \Delta) \quad (5.3)$$

$$\rho(\Delta_n^{**}, \Delta_n) \leq \rho(-\log(1 - F_n), \Delta_n) \leq C(\tau)/n \quad (6.3)$$

از طرفی می‌دانیم

$$|e^{-x} - e^{-y}| \leq |x - y|; \quad x, y \geq 0$$

پس:

$$\begin{aligned} I &= (e^{-\Delta_n^*(t) + \Delta(t)} - 1)e^{-\Delta(t)}[\Delta_n(t) - \Delta(t)] + e^{-\Delta(t)}[\Delta_n(t) - \Delta(t)] \\ &\leq |e^{-\Delta_n^*(t)} - e^{-\Delta(t)}| |\Delta_n(t) - \Delta(t)| + e^{-\Delta(t)} |\Delta_n(t) - \Delta(t)| \\ &= O(\rho^\gamma(\Delta_n, \Delta)) + e^{-\Delta(t)} |\Delta_n(t) - \Delta(t)| \end{aligned} \quad (7.3)$$

و

$$\begin{aligned} II &= [1 + (-\Delta_n^{**}(t) + \Delta_n(t)) + (-\Delta_n^{**}(t) + \Delta_n(t))^\gamma / \gamma + \dots] \\ &\quad [1 + (-\Delta_n(t) + \Delta_n(t)) + (-\Delta_n(t) + \Delta_n(t))^\gamma / \gamma + \dots] e^{-\Delta(t)} (-\log(1 - F_n(t) - \Delta_n(t))) \\ &= O(n^{-1}) \quad a.s \end{aligned} \quad (8.3)$$

و به این ترتیب از (5.3)–(8.3) نتیجه می‌شود:

$$F_n(t) - F(t) = \bar{F}(t)[\Delta_n(t) - \Delta(t)] + O(a_n^\gamma) \quad (9.3)$$

در نتیجه با استفاده از قضیه قبل حکم ثابت می‌شود.

□

۴ بررسی همگرایی در توزیع برآوردگرهای کاپلان مه-یر و نلسون-آلن

پیش از بیان قضایای مربوط به همگرایی ضعیف ابتدا نشان می‌دهیم که برآوردگر کاپلان مه-یر و برآوردگر نلسون-آلن را می‌توان به صورت میانگین متغیرهای تصادفی به اضافه باقیمانده ای از مرتبه $O(n^{-\frac{1}{\gamma}}(\log n)^{-\lambda})$ ، $\lambda > 0$ نوشت و سپس از آن در اثبات همگرایی ضعیف این دو برآوردگر استفاده می‌کنیم. این مطلب در قضیه ۱.۴ بیان شده است. برای این منظور قرار می‌دهیم:

$$g(x) = \int_{\cdot}^x (\bar{H}(s))^{-\gamma} dF_*(s)$$

و برای x, z حقیقی مثبت و $\delta = 0$ یا ۱

$$\zeta(z, \delta, x) = g(z \wedge x) - I(z \leq x, \delta = 1) / \bar{H}(z)$$

مشاهده می‌شود که

$$E(\zeta(Z_i, \delta_i, x)) = 0$$

$$Cov(\zeta(Z_i, \delta_i, x), \zeta(Z_i, \delta_i, t)) = g(s \wedge t)$$

قضیه ۱.۴ (کای ۱۹۹۸). تحت فرض های $(k_1) - (k_3)$

$$i) \Delta_n(t) - \Delta(t) = -\frac{1}{n} \sum_{i=1}^n \zeta(Z_i, \delta_i, x) + r_n(t)$$

$$ii) F_n(t) - F(t) = -\frac{\bar{F}(t)}{n} \sum_{i=1}^n \zeta(Z_i, \delta_i, x) + r_n(t)$$

که برای هر $\tau < \tau_H$

$$\sup_{0 \leq t \leq \tau} |r_n(t)| = O(b_n) \quad a.s$$

و

$$b_n = n^{-\frac{1}{\gamma}} (\log n)^{-\lambda}$$

حال با استفاده از نمایش برآوردگرهای کپلان مه-یر و نلسون-آلن به صورت میانگین متغیرهای تصادفی، همگرایی ضعیف آن ها را در قضیه های زیر اثبات می کنیم.

قضیه ۲.۴. تحت فرض های $(k_1) - (k_3)$ برای $t \in [0, \tau]$ و هر $\tau < \tau_H$

$$\sqrt{n}[\Delta_n(t) - \Delta(t)] \xrightarrow{L} N(0, \sigma^2(t))$$

که

$$\sigma^2(t) = \text{Var}(\zeta(Z_i, \delta_i, t)) + 2 \sum_{j=2}^{\infty} \text{Cov}(\zeta(Z_1, \delta_1, t), \zeta(Z_j, \delta_j, t))$$

اثبات. طبق لم ۲.۲ $\{\zeta(Z_i, \delta_i, t), i = 1, 2, \dots\}$ دنباله ای پایا، α -آمیزنده و کراندار از متغیرهای تصادفی است. به این ترتیب با استفاده از قضیه ۱.۴ و قضیه ۴۰۵۰۱۸ در ابرایگیمف و لی نیک (۱۹۷۱) حکم به سادگی ثابت می شود. $\square$

قضیه ۳.۴. تحت فرض های $(k_1) - (k_2)$ برای $t \in [0, \tau]$ و هر $\tau < \tau_H$

$$\sqrt{n}[F_n(t) - F(t)] \xrightarrow{L} N(0, \sigma^2(t) \bar{F}^2(t))$$

اثبات. با استفاده از قضیه ۲.۴ و رابطه (۹.۳) حکم به سادگی ثابت می شود. $\square$

مراجع

- [132] Breslow, N., Crowley, J. (1974). *A large sample study of the life table and product limit estimator under random censorship*. Ann. Statist. 2, 437-453.
- [133] Cai, Z. (1998). *Asymptotic properties of Kaplan- Meier estimator for censored dependent data*. Statistics & Probability Lett. 37, 381-389.

- [134] Cai, Z., Roussas, G. G. (1992). *Uniform strong estimation under α -mixing, with rate*. Statist. probab. Lett. 15, 47-55.
- [135] Csorgo, M., Revesz, P. (1981). *Strong Approximation in probability and Statistics*. Academic Press, New York,
- [136] Gill, R. D. (1980). *Censoring and Stochastic Integrals*. Mathematical Center Tracts No. 124, Mathematisch Centrum, Amsterdam.
- [137] Gill, R. D. (1981). *Testing with replacement and the product limit estimator*. Ann. Statist. 9, 853-860.
- [138] Gill, R.D. (1983). *Large sample behavior of the product limit estimator on the whole line*. Ann. Statist. 11, 49-56.
- [139] Ibragimov, I. A., Linnik Yu, V. (1971). *Independent and Stationary Sequences of Random Variables*. Walters Noordhoff, Groningen, the Netherlands.
- [140] Kaplan, E. L., Meier, P. (1958). *Non parametric estimation from incomplete observation*. J. Amer. Statist. Assoc. 53, 457-481
- [141] Lo, S. H., Singh, K. (1986). *The product limit estimator and the bootstrap: some asymptotic representations*. Probab. Theory Rel. Fields 71, 455-465.
- [142] Peterson, A. V. (1977). *Expressing the Kaplan-Meier estimator as a function of empirical sub-survival functions*. J. Amer. Statist. Assoc. 72, 854-858.
- [143] Stute, W., Wang, J.L. (1993). *A strong law under random censorship*. Ann. Statist. 21, 1591-1607.
- [144] Voelkel, J., Crowley, J. (1984). *Nonparametric inference for a class of semi-Markov process with censored observation*. Ann. Statist. 12, 142-160.
- [145] Wang, J. G. (1987). *A note on the uniform consistency of Kaplan – Meier estimator*. Ann. Statist. 15, 1313-1316.
- [146] Ying, Z., Wei, L. J. (1994). *The Kaplan – Meier estimate for dependent failure time observations*. J. Multivariate Anal. 50, 17- 29.

کاربرد تکنیک مدل‌سازی معادلات ساختاری در شناسایی عوامل خطر بیماری فشار خون

راضیه یوسفی^۱، کارشناسی ارشد آمارزیستی، گروه اپیدمیولوژی و آمارزیستی، دانشگاه علوم پزشکی مشهد، آزاده ساکی، استادیار آمارزیستی، گروه اپیدمیولوژی و آمارزیستی، دانشگاه علوم پزشکی مشهد

چکیده: در ساختاربندی مدل‌های چندسطحی تک معادله‌ای، معمولاً یک متغیر پاسخ بر اساس یک یا چندین متغیر تبیینی مدل‌بندی می‌شود. با این حال، مثال‌های بیشمار مرتبط با برخی از پدیده‌های اجتماعی و اقتصادی وجود دارد که مدل‌های تک معادله‌ای به تنهایی نمی‌توانند تشریح کاملی از ارتباط بین متغیرهای پاسخ و تبیینی را فراهم کنند. مدل‌های چندسطحی به ظاهر نامربوط می‌توانند چنین کمبودهایی را برطرف کنند. همانند سایر مدل‌های خطی، مدل‌های چندسطحی به ظاهر نامربوط نیز دربرگیرنده پیش فرض‌هایی مانند نرمال بودن توزیع خطاها است که ممکن است فرض مذکور در مسائل واقعی برقرار نباشند. در این صورت، استفاده از توزیع‌های نامتقارن مانند توزیع چوله-نرمال می‌تواند مثمرتر باشد. در مقاله حاضر، با در نظر گرفتن فرض چوله-نرمال برای خطای سطح دوم مدل‌های چندسطحی به ظاهر نامربوط، نحوه بکارگیری این دست مدل‌ها مورد مطالعه قرار می‌گیرد. سپس، با استفاده از رویکرد بیزی برآوردهای مدل و استنباط‌های آماری مرتبط انجام می‌شود. علاوه بر این، با برازش مدل‌های پیشنهادی در مورد داده‌های هزینه-درآمد سال ۱۳۹۰ خانوارهای روستایی ایرانی، نحوه عملکرد آن‌ها مورد بررسی قرار می‌گیرد.

واژه‌های کلیدی: مدل‌های چندسطحی به ظاهر نامربوط، توزیع چوله-نرمال، استنباط بیزی، هزینه و درآمد.

مقدمه

با توجه به رشد و گسترش جمعیت و به طبع آن افزایش شیوع بیماری‌های مختلف، امروزه مطالعات زیادی در خصوص بیماری‌ها و شناسایی ریسک فاکتورهای موثر بر آن‌ها انجام می‌شود. در سال‌های اخیر، مزایای مدل‌های معادلات ساختاری به کمک تکنیک‌های معمول آماری مورد ملاحظه قرار گرفته و موجب افزایش بکارگیری تکنیک مدل‌سازی معادلات ساختاری در تحلیل داده‌ها و آزمون اینگونه مدل‌ها گردیده است. در این مطالعه سعی شد در مدل‌سازی فشارخون و ریسک فاکتورهای مرتبط با آن، برآورد پارامترها براساس روش درست‌نمایی ماکزیمم محاسبه گردد.

مواد و روش‌ها

در این مطالعه مقطعی مبتنی بر جمعیت، مدل‌سازی معادلات ساختاری به منظور مدل‌سازی ارتباطات بین عوامل مستعدکننده بیماری فشارخون بکارگرفته شد. اطلاعات مربوط به نمونه تصادفی شامل ۹۷۶۵ نفر از مطالعه مشهد مورد استفاده قرار گرفت. داده‌ها به کمک نرم افزار Amos22 مورد تحلیل قرار گرفتند.

نتایج

^۱راضیه یوسفی : o.akhgari@gmail.com

$$(CFI = ۰/۹۳۹, TLI = ۰/۹۰۸, NFI = ۰/۹۳۵, RMSEA = ۰/۰۴, SRMR = ۰/۰۳۷)$$

مورد تایید قرار گرفت. متغیر سن و پس از آن چاقی و کم تحرکی بیشترین تاثیر را روی فشار خون نشان دادند.

نتیجه گیری

یافته ها شواهدی را برای تایید مدل مفهومی در نظر گرفته شده در مورد عوامل مستعدکننده بیماری فشارخون نشان دادند که می توانند در بکارگیری سیاست های پیشگیرانه در جلوگیری از ابتلا به این بیماری مفید باشند.

۱ مقدمه

با توجه به رشد و گسترش جمعیت و به طبع آن افزایش شیوع بیماری های مختلف، امروزه مطالعات زیادی در خصوص بیماری ها و شناسایی ریسک فاکتورهای موثر بر آن ها انجام می شود. از جمله روش هایی که به منظور انجام اینگونه تحلیل ها کاربرد دارد مدل های رگرسیونی هستند. اما از آنجاییکه در اغلب این بیماری ها بیشتر عوامل مرتبط همبستگی دارند بکارگیری مدل های رگرسیون معمول که با فرض عدم هم خطی و عدم همبستگی بین پیشگوها برازش داده می شوند مناسب نیست. بنابراین نیاز به مدل های کامل تری وجود دارد که بتواند ساختارهای همبستگی را تشخیص دهد و عوامل اصلی و مستقل را در غالب متغیرهای پنهان به همراه سایر متغیرهای آشکار مورد بررسی قرار دهد. بدین منظور مدلسازی معادلات ساختاری می تواند مورد توجه قرار گیرد. مدل یابی معادله ساختاری یک تکنیک تحلیل چند متغیری کلی و بسیار نیرومند از خانواده رگرسیون چند متغیری و به بیان دقیق تر بسط مدل خطی کلی است که به پژوهشگر امکان می دهد مجموعه ای از معادلات رگرسیون را بصورت همزمان مورد آزمون قرار دهد. مدلسازی معادلات ساختاری می تواند به راحتی برای تحلیل های پیچیده که شامل روابط مفروض بین یک یا چند متغیر پنهان مستقل و یک یا چند متغیر پنهان وابسته باشد، بکار گرفته شود ((۱۴۷)–(۱۴۹)). از مزایای اصلی مدل های معادلات ساختاری میتوان به انعطاف پذیری، مدلسازی متغیرهای پنهان، در نظر گرفتن خطا و آزمون مدل های مفروض اشاره نمود (۱۴۸). مباحث عملی مملو از متغیرهای پنهانی است که نمی توانند برای هر فرد اندازه گیری شوند. در سال های اخیر مدل های معادلات ساختاری به طور فزاینده ای در مسایل مهمی همچون علوم رفتاری، پزشکی و علوم اجتماعی بکار گرفته شدند و بنابر این امروزه تحلیل مدل های معادلات ساختاری بخش مهمی از تحقیقات آماری را به خود اختصاص می دهد. نرم افزارهای گوناگون تحلیل این معادلات را بر مبنای مدل های خطی که در آنها رابطه بین متغیرهای پنهان خطی است انجام می دهند (۱۵۰).

با توجه به گسترش کاربرد مدل های معادلات ساختاری، در این مقاله سعی شده است به معرفی این روش و بیان کاربرد آن در شناسایی عوامل خطر بیماری ها پرداخته شود. بنابراین پس از شرح این تکنیک، شناسایی عوامل خطر بیماری فشار خون در یک نمونه واقعی با استفاده از مدلسازی معادلات ساختاری انجام گرفت. امید است با تعیین مدل مناسب گامی در جهت پیشگیری اولیه از بروز این بیماری و گسترش آن در جامعه برداشته شود.

۲ روش کار

در این مطالعه به بررسی عوامل خطر بیماری فشار خون بر روی داده های مطالعه مشهد پرداخته شده است. مطالعه مشهد که از سال ۱۳۸۷ در شهر مشهد با هدف تعیین شیوع عوامل خطر ساز بیماری های غیر واگیردار و بهبود شیوه ی زندگی برای کاهش این عوامل خطر ساز آغاز شد، شامل ۹۷۶۵ فرد بین سن ۶۵–۳۵ سال می باشد. روش نمونه گیری طبقه ای-خوشه ای بود. به طوری که کل شهرستان مشهد به سه مرکز بهداشت ۳،۲،۱ تقسیم بندی شده بود و خوشه ها مناطق تحت پوشش این مرکز بودند. توضیحات بیشتر در مورد نحوه

استخراج نمونه و فرایند جمع آوری اطلاعات در مطالعه غیور مبرهن و همکاران آمده است. (۱۵۱) متغیرهای مورد استفاده در این مطالعه عبارت بودند از: متغیر پنهان موقعیت اقتصادی و اجتماعی با نشانگرهای میزان درآمد خانوار در ماه، میزان تحصیلات سرپرست خانوار (برحسب تعداد سال‌های تحصیل) و تعداد اعضای خانوار، متغیر پنهان فشارخون براساس دو متغیر فشارخون سیستولیک و دیاستولیک افراد، متغیر پنهان تغذیه ناسالم شامل میزان اسید چرب اشباع، کلسترول، ساکاروز و سدیم مصرفی افراد در رژیم روزانه آنها، متغیر پنهان چاقی براساس دو شاخص مربوط به چاقی BMI و WHR، متغیرهای مشاهده شده دیگری که در این تحلیل وارد شدند عبارت بودند از سن، میزان فعالیت فیزیکی افراد، میزان قند خون و تری گلیسیرید خون.

یک مدل معادلات ساختاری برای برآورد اثرات فاکتورهای خطر موثر بر بالا رفتن فشارخون تعیین گردید. پس از تعریف مدل مفهومی شناسایی پذیر بودن و مناسب بودن مدل مورد توجه قرار گرفت. همانطور که اشاره شد هدف مدل‌های معادلات ساختاری آزمون این فرضیه است که آیا ماتریس وارینانس مشاهده شده برای یک مجموعه از متغیرهای اندازه‌گیری شده، برابر با ماتریس کوواریانس باز تولید شده از طریق مدل مفروض است یا خیر و این فرضیه به صورت زیر بیان میشود:

$$\Sigma = \Sigma(\theta)$$

که در آن Σ : ماتریس کوواریانس تولید شده از جامعه توسط متغیرهای مشاهده شده و $\Sigma(\theta)$: ماتریس کوواریانس باز تولید شده از مدل مفروض براساس θ است. مدل‌های معادلات ساختاری پایه، شامل مدل اندازه‌گیری و مدل ساختاری هستند. مدل اندازه‌گیری جزئی از مدل ساختاری است که نحوه سنجش یک متغیر پنهان را با استفاده از دو یا تعداد بیشتری متغیر مشاهده شده تعریف می‌کند. مدل اندازه‌گیری عبارت است از:

$$x_i = \Lambda\omega_i + \varepsilon_i \quad i = 1, 2, \dots, n$$

که در آن x_i یک بردار $1 \times p$ از متغیرهای نشانگر توضیحی برای $1 \times q$ بردار تصادفی متغیرهای پنهان ω_i و Λ یک ماتریس $p \times q$ از بارهای ضرایب بدست آمده از رگرسیون x_i بر ω_i است و ε_i یک بردار تصادفی $1 \times p$ از خطاهای اندازه‌گیری است که دارای توزیع $N(0, \psi_\varepsilon)$ هستند. فرض می‌شود که برای $i = 1, 2, \dots, n$ ω_i ها مستقلند و از توزیع $N(0, \phi)$ پیروی می‌کنند، و با بردار تصادفی ε_i ناهمبسته هستند. مدل ساختاری جزئی از مدل‌های معادلات ساختاری است که نشان می‌دهد متغیرهای پنهان (و گاهی اوقات آشکار) چگونه بر یکدیگر اثر می‌گذارند. این مدل به ترتیب زیر است:

$$\eta_i = \beta\eta_i + \Gamma\xi_i + \delta_i \quad i = 1, \dots, n$$

وقتی β یک ماتریس $m \times m$ از پارامترهای ساختاری است که بیان‌کننده رابطه بین متغیرهای پنهان درون‌زا است و قطر آن صفر است. Γ ماتریس پارامتر رگرسیون $m \times n$ برای ارتباط دادن متغیرهای درون‌زا و برون‌زا است. δ_i بردار $1 \times m$ که فرض می‌شود از توزیع $N(0, \psi_\delta)$ پیروی می‌کند و ψ_δ یک ماتریس کوواریانس قطری است. همچنین فرض می‌شود که δ_i با ξ_i ناهمبسته اند ((۱۴۷)، (۱۵۲)). هر مدل اندازه‌گیری متشکل از سه نوع متغیر است که شامل متغیر پنهان، متغیر مشاهده شده و متغیر خطا می‌شوند. محقق همواره به معیارهایی نیاز دارد تا با کمک گرفتن از آنها و تفسیرشان، به ارزیابی مدل تدوین شده بر مبنای چارچوب نظری و پیشینه تجربی دست زند (۱۵۳). برخی از این شاخص‌ها عبارتند از:

★ شاخص کای اسکوئر

شاخص کای اسکوئر را می‌توان به عنوان عمومی‌ترین و پرکاربردترین شاخص برازش در مدل‌سازی معادله ساختاری تلقی کرد. اولین نکته‌ای که درباره تفسیر مقدار این شاخص می‌توان اظهار کرد این است که هرچه مقدار آن کوچکتر باشد برازش داده‌ها به مدل بهتر است. فرمول محاسبه این شاخص عبارت است از:

$$\chi^2 = (n - 1) \times F_{ML}$$

که در آن:

$$F_{ML} = \ln |\Sigma| - \ln |S| + tr [(S) (\Sigma^{-1})] - K$$

$\ln |S|$ لگاریتم طبیعی دترمینال ماتریس کوواریانس مشاهده شده و $\ln |\Sigma|$ لگاریتم طبیعی دترمینال ماتریس کوواریانس باز تولید شده است.

به دلیل حساس بودن این شاخص به حجم نمونه از شاخص اصلاح شده آن نسبت به درجه آزادی استفاده می‌شود. هرچند در صورتی که حجم نمونه بسیار بالا باشد شاخص کای اسکوتر اصلاح شده نیز مقادیر بزرگی خواهد داشت و قابل اتکا نیست (۱۵۴).

★ شاخص NFI

این شاخص به افزودن پارامتر به مدل حساس است به نحوی که هرچه پارامتر به مدل افزوده شود مقدار این شاخص نیز افزایش می‌یابد.

$$NFI = \frac{\chi_n^2 - \chi_m^2}{\chi_n^2}$$

مقدار قابل قبول برای این شاخص حداقل ۰/۹۰ و مقداری که نشان دهنده ی یک برازش خوب است حداقل ۰/۹۵ می باشد.

★ شاخص TLI

این شاخص تلاش می‌کند تا نقطه ضعف شاخص NFI را در به حساب نیاوردن جریمه شاخص برای افزودن پارامتر مرتفع کند. این شاخص بر مبنای متوسط ضرایب همبستگی بین متغیرها در مدل قرار دارد. هرچه این ضرایب کوچکتر باشند شاخص توکر لويس نیز مقدار کوچکتری را نشان خواهد داد. این شاخص بین صفر تا یک تغییر می‌کند و مقدار ۰/۹۵ یا بیشتر منعکس کننده ی یک مدل خوب است.

$$TLI = \frac{(\chi_n^2/df_n) - (\chi_m^2/df_m)}{(\chi_n^2/df_n) - 1}$$

★ شاخص CFI

این شاخص نیز بر مبنای همبستگی بین متغیرهای حاضر در مدل قرار دارد به نحوی که ضرایب بالای همبستگی بین آن ها به مقادیر بالای این شاخص می‌انجامد. مقدار آن از فرمول بدست می‌آید. مقدار بیشتر از ۰/۹۰ برای این شاخص بیانگر برازش مناسب مدل است.

$$CFI = \frac{(\chi_n^2 - df_n) - (\chi_m^2 - df_m)}{(\chi_n^2 - df_n)}$$

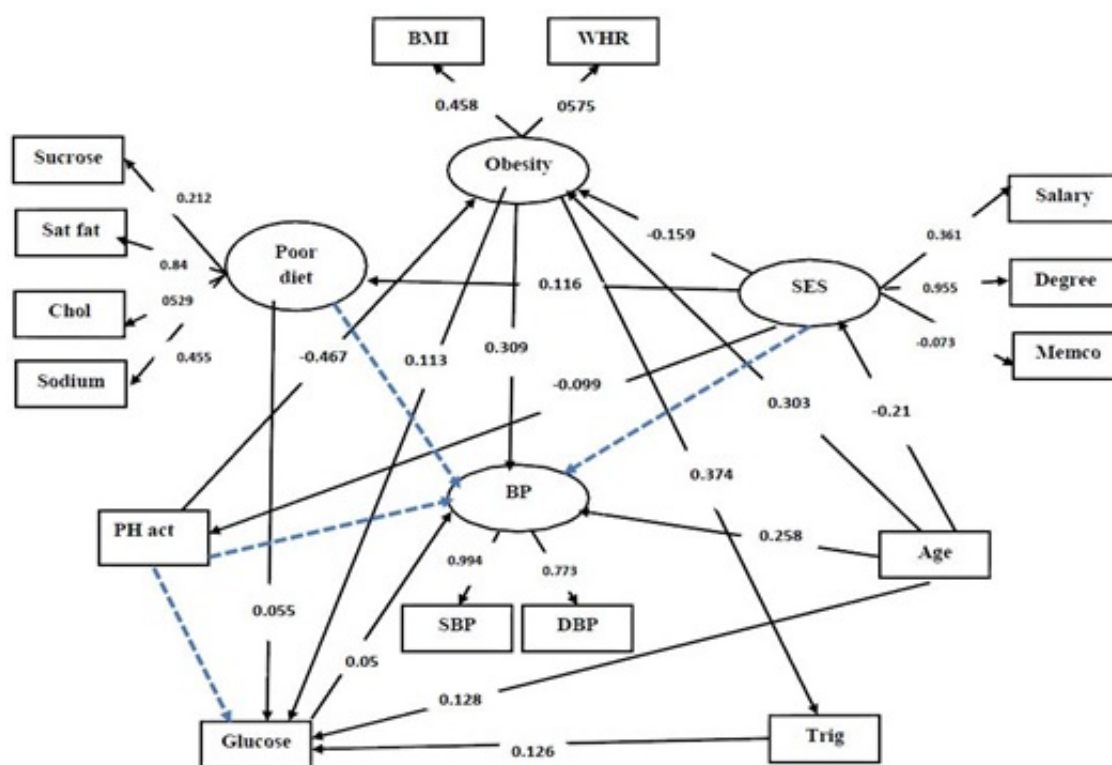
★ شاخص RMSEA

این شاخص بر مبنای تحلیل ماتریس باقیمانده قرار دارد. مدل‌های قابل قبول دارای مقدار $0/50$ یا کوچکتر برای این شاخص هستند. برازش مدل‌هایی که دارای مقادیر بالاتر از $0/1$ هستند ضعیف برآورد می‌شود.

$$RMSEA = \sqrt{\frac{\chi_m^2 - df_m}{(n - 1) \times df_m}}$$

۳ نتایج

از ۹۷۶۵ فرد شرکت داده شده در تحلیل، حدود ۴۰ درصد مرد و ۶۰ درصد زن بودند. میانگین سنی افراد $1/48$ بود. ۲۳ درصد از کل جمعیت به بیماری پر فشاری خون مبتلا بودند. (فشارخون سیستولیک بالای ۱۴۰ و فشارخون دیاستولیک بالای ۹۰). شیوع این بیماری در بین مردان بیشتر از زنان بود. (۲۳/۳ درصد در مقابل ۲۴/۴ درصد). برای مشخصات توصیفی نمونه مورد مطالعه به مطالعه غیور و همکاران مراجعه شود (۱۵۱)



شکل ۱: ضرایب برآورد شده از برازش مدل مفهومی

مدل مفهومی براساس شاخص های نیکویی برازش، برازش مناسبی به داده‌ها نشان داد

$$(CFI = 0/939, TLI = 0/908, NFI = 0/935, SRMR = 0/037, RMSEA = 0/04)$$

. نمودار ۱ ضرایب مسیر و بارهای عاملی بدست آمده را نشان می‌دهد. اثرات مستقیم سن، چاقی و قندخون بر فشارخون معنی دار بودند. همانطور که ملاحظه می‌شود، متغیر موقعیت اقتصادی و اجتماعی با اینکه اثر مستقیم بر فشارخون ندارد اما اثر غیر مستقیم بر آن دارد. بررسی اثرات کل متغیرها نشان می‌دهد متغیر سن بیشترین تاثیر بر روی فشارخون را دارد. متغیر فعالیت بدنی اثر مستقیم بر روی

فشارخون ندارد اما اثر کل نسبتاً بالایی دارد. منفی بودن این مقدار نشان می‌دهد که با افزایش فعالیت بدنی، فشارخون کاهش می‌یابد. اثرات کل متغیرها در جدول ۱ آمده اند. فعالیت بدنی به شدت موجب کاهش احتمال بروز چاقی می‌گردد و همچنین تری گلیسرید خون را نیز کاهش می‌دهد. چاقی بر فشارخون سیستمیک نسبت به فشارخون دیاستولیک تأثیر گذارتر است. در مجموع کلیه متغیرها تأثیر بیشتری روی فشارخون سیستمیک نشان داده اند.

۴ بحث

در چند دهه گذشته محققان زیادی بر روی عوامل موثر بر بیماری فشارخون مطالعاتی را انجام داده اند. در این مطالعات مدل‌های گوناگونی نیز برای شناسایی این عوامل بکارگرفته شد. از جمله زنگ و همکاران در مطالعه خود عوامل خطر مربوط به فشارخون را به کمک یک مدل رگرسیون لجستیک بررسی کردند (۱۵۵). در مطالعه ون نیز این مدل برای شناسایی ارتباط بین فشارخون و عوامل خطر آن بکارگرفته شد (۱۵۵)، (۱۵۶)). اما در کلیه این مدل‌ها تنها اثرات عوامل در مقایسه بین افراد مبتلا به این بیماری و افراد سالم مورد توجه قرار گرفته است، در صورتیکه تأثیر این عوامل در بین افراد در معرض خطر و حتی در میان سالمین به منظور پیشگیری از ابتلا نیز اهمیت ویژه‌ای دارد. بنابراین در این مطالعه سعی شده است با در نظر گرفتن یک مدل آماری به بررسی عواملی بپردازیم که در همه‌ی افراد جامعه از سالم گرفته تا افراد مبتلا به پرفشاری خون در مرحله‌ی حاد، موجب افزایش خطر بیماری می‌گردد. در این شرایط نیاز به مدلی بود که بتواند عوامل مختلف را در یک تحلیل چند متغیره وارد کند و از طرفی همبستگی این عوامل را در نظر بگیرد. بنابراین این مطالعه براساس یک مدل پیشبینی با استفاده از یک مطالعه‌ی مقطعی برای شناسایی عوامل مستعدکننده‌ی بیماری فشارخون در افراد بالای ۳۵ سال صورت گرفت. متغیر فشارخون به عنوان متغیر پاسخ بصورت پنهان و متغیرهای چاقی، موقعیت اقتصادی و اجتماعی و تغذیه‌ی ناسالم به شکل متغیرهای پنهان تعریف شدند. بعلاوه متغیر سن متغیر برونزای این مدل بود. متغیرهای میانجی قندخون و تری‌گلیسرید نیز به مدل اضافه شدند. انتخاب این متغیرها با مطالعات زیادی هماهنگ بود (۱۵۷). مطالعات زیادی از مدلسازی معادلات ساختاری جهت شناسایی عوامل موثر بر بیماری‌ها و از جمله فشار خون بهره گرفته اند. ((۱۶۰)–(۱۵۸)).

پس از تعیین متغیرها و مدل مفهومی، مدل کلاسیک معادلات ساختاری با استفاده از روش برآورد MLR برازش یافت و سپس شاخص‌های نیکویی برازش بررسی گردید. مقدار شاخص CFI و TLI و NFI بیشتر از ۰.۹۰، شاخص RMSEA و SRMR کمتر از ۰.۰۵۰ همگی گواهی بر مناسب بودن مدل کلاسیک در این مطالعه بودند. مدل مفهومی با یافته‌های پژوهش برازش مناسبی داشت. سن، چاقی و کمبود فعالیت بدنی عواملی هستند که موجب افزایش ریسک ابتلا به بیماری پرفشاری خون می‌گردند. این یافته‌ها نیز با پژوهش‌های دیگر همسو است. در مطالعه برملاگ و همکاران رابطه قوی میان افزایش فشارخون سیستمیک و دیاستولیک در هر دو گروه دارای فشارخون بالا و تحت درمان و گروه افراد با سطح پایین فشارخون مشاهده شد. (۱۶۱). ککینوس و همکاران در مطالعه‌ی اذعان داشتند که می‌توان پذیرفت که افزایش فعالیت بدنی با شدت و دوره مناسب با عث کاهش فشارخون می‌گردد. بنابراین برای پیشگیری اولیه در یک برنامه سلامت ملی باید برنامه‌های ورزشی مورد توجه ویژه قرار گیرد (۱۶۲). موسویان و همکاران تأثیرات فعالیت بدنی مشخصی را بر فشارخون زنان یائسه دارای اضافه وزن بررسی کرده اند. در این مطالعه این فعالیت موجب کاهش میانگین فشارخون سیستمیک و دیاستولی شده است (۱۶۳).

۵ نتیجه گیری

تکنیک مدلسازی معادلات ساختاری این امکان را فراهم نمود تا به طور همزمان به ارزیابی کیفیت سنجش متغیرها و مقبولیت اثرات مستقیم و غیرمستقیم و همچنین تعامل‌های تعریف شده میان متغیرها بپردازیم. براساس نتایج بدست آمده به نظر می‌رسد انتخاب شیوه‌ی زندگی سالم شامل تحرک بدنی کافی و در نتیجه جلوگیری از بروز چاقی برای پیشگیری و مدیریت افزایش فشارخون امری ضروری به نظر می‌رسد.

جدول ۱: ضرایب بدست آمده از مدل معادلات ساختاری

فشارخون	قندخون	تغذیه ناسالم	تری گلیسرید	چاقی	فعالیت بدنی	موقعیت اقتصادی اجتماعی	سن	
-	-	-	-	-	-	-	-۰/۲۱	موقعیت اقتصادی اجتماعی
-	-	-	-	-	-	-۰/۰۹۹	۰/۰۰۲۱	فعالیت بدنی
-	-	-	-	-	-۰/۴۶۷	-۰/۱۱۳	۰/۳۲۷	چاقی
-	-	-	-	۰/۳۷۴	-۰/۱۷۵	-۰/۰۴۲	۰/۱۲۲	تری گلیسرید
-	-	-	-	-	-	۰/۱۱۶	-۰/۰۲۴	تغذیه ناسالم
-	-	۰/۰۵۵	۰/۱۲۶	۰/۱۶۱	-۰/۰۷۵	-۰/۰۱۲	۰/۱۷۹	قندخون
-	۰/۰۵	۰/۰۰۳	۰/۰۰۶	۰/۳۱۷	-۰/۱۴۸	-۰/۰۳۵	۰/۳۶۸	فشارخون
-	-	۰/۴۵۵	-	-	-	۰/۰۵۳	-۰/۰۱۱	سدیم
-	-	۰/۲۱۲	-	-	-	۰/۰۲۵	-۰/۰۰۵	ساکاروز
-	-	-	-	-	-	-۰/۰۷۳	۰/۰۱۵	جمعیت خانوار
-	-	۰/۵۲۹	-	-	-	۰/۰۶۱	-۰/۰۱۳	کلسترول
-	-	۰/۸۴۱	-	-	-	۰/۰۹۷	-۰/۰۰۲	اسیدچرب اشباع
-	-	-	-	-	-	۰/۹۵۵	-۰/۲۰۱	تحصیلات
-	-	-	-	-	-	۰/۳۶۱	-۰/۰۷۶	درآمد
-	-	-	-	۰/۵۷۵	-۰/۲۶۹	-۰/۰۶۵	۰/۱۸۸	شاخص دورکمره دور هیپ
-	-	-	-	۰/۴۵۸	-۰/۲۱۴	-۰/۰۵۲	۰/۱۵	شاخص توده بدنی
۰/۹۹۴	۰/۰۴۹	۰/۰۰۳	۰/۰۰۶	۰/۳۱۵	-۰/۱۴۷	-۰/۰۳۵	۰/۳۶۶	فشارخون سیستول
۰/۷۷۳	۰/۰۳۸	۰/۰۰۲	۰/۰۰۵	۰/۲۴۵	۰/۱۱۴	۰/۰۲۷	۰/۲۸۵	فشارخون دیاستول

- [147] Ullman JB. *Structural equation modeling: Reviewing the basics and moving forward*. Journal of personality assessment. 2006;87 (1):35-50.
- [148] Altindag I, Genc A. *Bayesian Nonlinear Structural Equation Modeling*. Journal of Selcuk University Natural and Applied Science. 2015;4(1):174-92.
- [149] Song XY, Xia YM, Lee SY. *Bayesian semiparametric analysis of structural equation models with mixed continuous and unordered categorical variables*. Statistics in medicine. 2009;28(17):2253-76.
- [150] Lee S-Y, Song X-Y. *Model comparison of nonlinear structural equation models with fixed covariates*. Psychometrika. 2003;68(1):27-47.
- [151] Ghayour-Mobarhan M, Moohebbati M, Esmaily H, Ebrahimi M, Parizadeh SMR, Heidari-Bakavoli AR, et al. *Mashhad stroke and heart atherosclerotic disorder (MASHAD) study: design, baseline characteristics and 10-year cardiovascular risk estimation*. International journal of public health. 2015;60(5):561-72.
- [152] Kaplan D. *Structural equation modeling: Foundations and extensions*: Sage Publications; 2008.
- [153] Sivo SA, Fan X, Witta EL, Willse JT. *The search for "optimal" cutoff properties: Fit index criteria in structural equation modeling*. The Journal of Experimental Education. 2006;74(3):267-88.
- [154] Kline RB. *Principles and practice of structural equation modeling*: Guilford publications; 2015.
- [155] Ahmad WMAW, Nawli MABA, Aleng NA, Halim NA, Mamat M, Hamzah MP, et al. *Association of Hypertension with Risk Factors Using Logistic Regression*. Applied Mathematical Sciences. 2014;8(52):2563-72.
- [156] Van Onna M. *Bayesian estimation and model selection in ordered latent class models for polytomous items*. Psychometrika. 2002;67(4):519-38.
- [157] Bardenheier BH, Bullard KM, Caspersen CJ, Cheng YJ, Gregg EW, Geiss LS. *A Novel Use of Structural Equation Models to Examine Factors Associated With Prediabetes Among Adults Aged 50 Years and Older*. Diabetes Care. 2013;36(9):2655-62.
- [158] Cois A, Ehrlich R. *Analysing the socioeconomic determinants of hypertension in South Africa: a structural equation modelling approach*. BMC public health. 2014;14(1):1.

- [159] Shen B-J, Todaro JF, Niaura R, McCaffery JM, Zhang J, Spiro III A, et al. *Are metabolic risk factors one unified syndrome? Modeling the structure of the metabolic syndrome X*. American Journal of Epidemiology. 2003;157(8):701-11.
- [160] Nock NL, Wang X, Thompson CL, Song Y, Baechle D, Raska P, et al., editors. *Defining genetic determinants of the Metabolic Syndrome in the Framingham Heart Study using association and structural equation modeling methods*. BMC proceedings; 2009: BioMed Central.
- [161] Bramlage P, Pittrow D, Wittchen H-U, Kirch W, Boehler S, Lehnert H, et al. *Hypertension in overweight and obese primary care patients is highly prevalent and poorly controlled*. American journal of hypertension. 2004;17(10):904-10.
- [162] Kokkinos PF, Giannelou A, Manolis A, Pittaras A. *Physical activity in the prevention and management of high blood pressure*. Hellenic J Cardiol. 2009;50(1):52-9.
- [163] Mousavian As, SHaherian S, Namvar f, GHanbarzadeh m. *Comparing effect of walking and selected aerobic exercise on blood pressure in overweight postmenopausal women*. Journal of Nursing and Midwifery Urmia University of Medical Sciences. 2014;12(16):427-35. eng @ 2008-6326, 2014.

مدلسازی معادلات ساختاری بیزی و کاربرد آن در مدلسازی عوامل خطر بیماری ها

راضیه یوسفی^۱، کارشناسی ارشد آمارزیستی، گروه اپیدمیولوژی و آمارزیستی، دانشگاه علوم پزشکی مشهد آزاده ساکی استادیار آمارزیستی، گروه اپیدمیولوژی و آمارزیستی، دانشگاه علوم پزشکی مشهد

چکیده: مدلسازی معادله ساختاری، ابزاری قدرتمند در دست پژوهشگر است که وی را در چگونگی تدوین مبانی و چارچوب نظری پژوهش در قالب مدل های اندازه گیری و ساختاری یاری می رساند. اما این تکنیک نیاز به برقراری پیش فرض هایی دارد که مشکلاتی ایجاد می کند. به منظور حل این مشکلات رویکرد بیزی درمورد تحلیل بسیاری از مدل های معادله ی ساختاری پیشرفته پیشنهاد شد. در این مطالعه سعی شده است ضمن معرفی رویکرد بیزی مدلسازی معادلات ساختاری به بیان نتایج این روش در داده های واقعی نیز پرداخته شود.

مواد و روش ها

در این مطالعه مقطعی مبتنی بر جمعیت، مدلسازی معادلات ساختاری بیزی با استفاده از روش نمونه گیری گیس معرفی و در داده های واقعی مطالعه مشهد بکار گرفته شد. مدل های بیزی با پیشین های فاقد اطلاع نرمال، یکنواخت و پیشین حاوی اطلاع به داده ها برازش یافتند.

نتایج

در این مطالعه با توجه به نزدیک بودن برآوردهای دو مدل فاقد اطلاع این نتیجه حاصل شد که انتخاب توزیع فاقد اطلاع تاثیر چندانی در برآورد نداشت و در اینحالت برآوردها براساس اطلاعات موجود در دادهها استخراج شدند. مدل بیزی حاوی اطلاع زمان بیشتری برای رسیدن به همگرایی نیاز داشت. اما بررسی آماره DIC مدل ها نشان داد که مدل های بیزی فاقد اطلاع با پیشینهای یکنواخت و نرمال تعریف شده یکسان بودند اما با بکارگیری اطلاعات پیشین شاخص DIC به اندازهی ۱۳۹ واحد کاهش یافت.

نتیجه گیری

نتایج این مطالعه نشان داد که بکارگیری رویکرد بیزی معادلات ساختاری نتایجی حداقل در سطح رویکرد کلاسیک ارائه می کند و در صورتی که اطلاعات پیشین مناسبی وجود داشته باشد منجر به نتایج مناسب تر و دستیابی به برآوردها در شرایط دشوار می گردد.

واژه های کلیدی: مدلسازی معادلات ساختاری، تکنیک بیزی، نمونه گیر گیس

^۱راضیه یوسفی:

۱ مقدمه

مدلسازی معادله ساختاری، ابزاری قدرتمند در تدوین مبانی و چارچوب نظری پژوهش در قالب مدل‌های اندازه‌گیری و ساختاری با فرضیات معین است. ((۱۶۷)-(۱۶۴)). اما اغلب ملاحظه می‌شود که در داده‌های جمع‌آوری شده در یک مطالعه این فرضیات برقرار نیستند و بنابراین در این شرایط نمیتوان روش درست‌نمایی ماکزیمم را بکار برد زیرا شاخصهای برازش مدل، برآورد پارامترها و خطاهای استاندارد با افزایش عدم نرمالیتی، هرچه بیشتر اریب خواهند شد ((۱۶۸)). اثرات تخلف از فرضیات نرمالیتی باعث بوجود آمدن مقادیر بزرگ کیدو و در نتیجه رد بسیاری از مدلها خواهد شد. ((۱۶۹)).

به منظور حل این مشکلات رویکرد بیزی با چند مزیت در مورد تحلیل بسیاری از مدل‌های معادله‌ی ساختاری پیشرفته پیشنهاد شد. اول اینکه روش‌های بیزی دامنه فرضیات مورد آزمون را گسترش می‌دهد و نتایج می‌توانند به طرق شهودی و بدون اتکا به آزمون‌های معناداری فرض صفر تفسیر شوند. دوم، برآورد بیزی میتواند یافته‌های پیشین موجود در تحقیقات مشابه قبلی یا بدست آمده از متآنالیز را با داده‌های جدید ترکیب کند. سوم، چون برآورد بیزی میتواند یافته‌های قبلی را بکار گیرد، داده‌های بدست آمده از مطالعات با حجم نمونه‌ی کم کمتر دچار مشکل خواهند شد. چهارم، روشهای بیزی امکان برآورد مدلها را وقتی روشهای سنتی برآورد به علت پیچیدگی مدل ناکارا هستند، فراهم می‌کند ((۱۶۶)، (۱۷۰)، (۱۷۱)). توسعه و کاربرد رویکرد بیزی مدلسازی معادلات ساختاری، نسبتاً کند بوده است اما با تکنولوژی جدید و روش‌هایی چون نمونه گیر گیس، که امروزه به کمک نرم افزارها به آسانی قابل دستیابی است، امکان پذیر گشته است. ((۱۷۶)-(۱۷۲)). در این مطالعه سعی شد پس از شرح مختصری از مفاهیم مربوط به رویکرد بیزی معادلات ساختاری به بیان نتایج بدست آمده از این روش در یک داده واقعی پرداخته شود.

۲ تعیین مدل معادلات ساختاری بیزی

فرض کنید مدل اندازه گیری به صورت زیر باشد:

$$y = \alpha + \Lambda\eta + Kx + \varepsilon$$

وقتی y بردازی از متغیرهای آشکار باشد. α برداری از عرض از مبدا مدل اندازه گیری، Λ ماتریس بارهای عاملی، η برداری از متغیرهای پنهان، ماتریسی از ضرایب رگرسیون مربوط به متغیرهای آشکار Λ به مشاهده شده x ، ε بردار خطا و Σ ماتریس کوواریانس قطری می‌باشد. مدل ساختاری عامل‌های مشترک را به یکدیگر و به متغیرهای آشکار x مربوط می‌کند و بصورت زیر بیان می‌شود:

$$\eta = \nu + B\eta + \Gamma x + \zeta$$

وقتی ν برداری از عرض از مبدا مدلساختاری، B و Γ ماتریسهای ضرایب ساختاری، و ζ برداری از خطاهای ساختاری با ماتریس کوواریانس Ψ می‌باشد.

۳ نمونه گیری زنجیره مارکوف مونت کارلو برای مدل‌های معادلات ساختاری بیزی

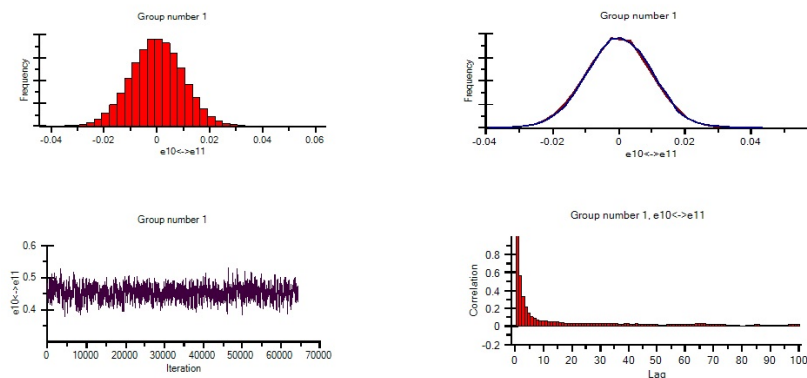
روش‌های زنجیره مارکوف مونت کارلو عمومی ترین روش برای دستیابی به برآورد بیزی در مسائل است ((۱۷۷)). نمونه گیر گیس یکی از این روشهاست که ایده‌ی اولیه آن توسط متروپلیس و نلر در سال ۱۹۵۳ ارائه و سپس با مقاله ای که گیمن و گیمن در مورد مدل‌های پردازش تصویر ارائه دادند وارد مرحله جدیدی شد ((۱۷۸)). رویکرد بیزی با در نظر گرفتن η بعنوان داده گمشده شروع می‌شود. فرض کنیم $\eta = (\eta_1, \eta_2, \dots, \eta_n)$ سپس، داده های مشاهده شده $Y = (y_1, y_2, \dots, y_n)$ با η در تحلیل پسین ترکیب می‌شوند.

مدل بیزی	تعداد نمونه تولید شده	تعداد زنجیره	PSR	$PPP - value$	DIC	سرعت همگرایی
پیشین فاقد اطلاع یکنواخت	۵۵۱۹۲	۲	۱/۰۰۱۸	۰/۵	۶۵۶۲۵/۳۶	۱
پیشین فاقد اطلاع نرمال	۵۶۵۰۱	۲	۱/۰۰۱۷	۰/۵۲	۶۵۶۲۶/۱۷	۲
پیشین حاوی اطلاع	۵۶۷۳۱	۶۴	۱/۰۰۱۸	۰/۵۵	۶۵۴۸۶/۴۲	۳

جدول ۲: میانگین‌های پسین با پیشین‌های متفاوت

مسیر	پیشین ۱	میانگین پسین	پیشین ۲	میانگین پسین
فشارخون	←	فشارخون دیاستول	$N(0, 10^{-10})$	۰/۴۷۰
چاقی	←	شاخص دورکمر به دور هیپ	$N(0, 10^{-10})$	۰/۰۳۸
موقعیت اقتصادی اجتماعی	←	درآمد	$N(0, 10^{-10})$	۰/۵۴۱
سن	←	چاقی	$N(0, 10^{-10})$	۰/۴۱۰
سن	←	فشارخون	$N(0, 10^{-10})$	۰/۶۰۷
چاقی	←	فشارخون	$N(0, 10^{-10})$	۰/۴۳۱
فعالیت بدنی	←	چاقی	$N(0, 10^{-10})$	-۱/۹۴۳
قندخون	←	فشارخون	$N(0, 10^{-10})$	۰/۲۳۹
موقعیت اقتصادی اجتماعی	←	فعالیت بدنی	$N(0, 10^{-10})$	-۰/۰۳۵
چاقی	←	قندخون	$N(0, 10^{-10})$	۰/۰۳۵
تغذیه ناسالم	←	کلسترول	$N(0, 10^{-10})$	۱/۰۹۹
تغذیه ناسالم	←	قندخون	$N(0, 10^{-10})$	۰/۱۲۰
موقعیت اقتصادی اجتماعی	←	تغذیه ناسالم	$N(0, 10^{-10})$	-۰/۰۴۰
موقعیت اقتصادی اجتماعی	←	جمعیت خانوار	$N(0, 10^{-10})$	-۰/۲۳۳
موقعیت اقتصادی اجتماعی	←	چاقی	$N(0, 10^{-10})$	۰/۰۶۴
تری گلیسرید	←	قندخون	$N(0, 10^{-10})$	-۰/۲۲۸
سن	←	موقعیت اقتصادی اجتماعی	$N(0, 10^{-10})$	۰/۰۶۳
سن	←	قندخون	$N(0, 10^{-10})$	۰/۲۲۸
چاقی	←	تری گلیسرید	$N(0, 10^{-10})$	۰/۳۵۱
تغذیه ناسالم	←	ساکاروز	$N(0, 10^{-10})$	۰/۹۲۵
تغذیه ناسالم	←	سدیم	$N(0, 10^{-10})$	۰/۰۲۷
کوواریانس خطاها				
فشارخون سیستمیک	و	دیاستولیک	-	۰/۰۰۱
درآمد	و	تحصیلات	-	۰/۰۱۶
درآمد	و	جمعیت خانوار	-	۰/۰۰۲
تحصیلات	و	جمعیت خانوار	-	۰/۰۰۱
اسیدچرب اشباع	و	کلسترول	-	-
اسیدچرب اشباع	و	سدیم	-	۰/۰۰۴
اسیدچرب اشباع	و	ساکاروز	-	۰/۰۰۱
کلسترول	و	سدیم	-	۰/۰۰۵
کلسترول	و	ساکاروز	-	۰/۰۰۱
سدیم	و	ساکاروز	-	۰/۰۰۵

نمودار ۲ نمودارهای همگرایی مربوط به یکی از مسیرهای موجود را نشان می‌دهد. در این نمودارها همگرایی به وضوح دیده می‌شود. از آنجاییکه مقادیر نمونه‌ها هم مثبت و هم منفی هستند و در اطراف صفر نیز بصورت متقارن قرار گرفته‌اند بنابراین می‌توان نتیجه گرفت به لحاظ آماری فرضیه برابری این پارامترها با صفر می‌تواند پذیرفته شود. بنابراین پیش فرض مدلسازی معادلات ساختاری مبنی بر مستقل بودن خطاهای متغیرهای نشانگر به این ترتیب تأمین شده است.



شکل ۲: نمودار کوواریانس خطای بین متغیرهای فشارخون سیستولیک و دیاستولیک

۵ بحث

در این مطالعه با توجه به اینکه توزیع متغیرها غیرنرمال شناخته شد و قبل از بکارگیری برآوردهای درستنمایی ماکزیمم نیاز به برقراری پیش فرض‌های زیادی وجود داشت تا برآوردهای ناریب بدست آیند، بنابراین رویکرد بیزی معادلات ساختاری مورد توجه قرار گرفت. مودن (۱۸۲) و مودن و آسپارو (۱۶۵) در مطالعات خود بیان داشتند که پارامترهایی چون کوواریانس خطاها که در رویکرد کلاسیک شناسایی‌پذیر نیستند و امکان برآورد آنها وجود ندارد به کمک تکنیک بیزی قابل برآورد هستند، به کمک تکنیک بیزی امکان آزمون مستقیم احتمال اینکه یک پارامتر در یک محدوده قرار گیرد، فراهم می‌گردد و از طرفی بکارگیری توزیع‌های پیشین به دستیابی به برآوردهای دقیق‌تری منجر می‌شود. در این مطالعه با توجه به نزدیک بودن برآوردهای دو مدل فاقد اطلاع این نتیجه حاصل شد که انتخاب توزیع فاقد اطلاع تأثیر چندانی در برآورد نداشت و در این حالت برآوردها براساس اطلاعات موجود در داده‌ها استخراج شدند. مطالعات مودن (۱۸۲)، جلمن (۱۸۳) و اشبی (۱۸۴) نیز بر این نکته اذعان داشته‌اند که بکارگیری توزیع پیشین فاقد اطلاع موجب می‌شود برآوردهای پسین براساس تابع درستنمایی استخراج شود.

اما در مرحله بعد برازش مدل بیزی حاوی اطلاع به این ترتیب صورت گرفت که پیش فرض تکنیک معادلات ساختاری مبنی بر صفر بودن کوواریانس خطاهای مربوط به نشانگرهای یک عامل به عنوان پیشین بکار رفت. در رویکرد کلاسیک مدلسازی معادلات ساختاری امکان در نظر گرفتن فرض مستقل بودن خطای نشانگرها به علت شناسایی ناپذیر شدن مدل وجود ندارد، در حالیکه در مسائل واقعی ممکن است بین نشانگرهای یک عامل با توجه به ارتباطشان همبستگی وجود داشته باشد. به این منظور برای کوواریانس بین نشانگرهای یک عامل توزیع نرمال با میانگین صفر و واریانس ۰/۰۱ در نظر گرفته شدند. زیفورو و اسوالد نیز در مطالعه خود که در زمینه علوم رفتاری انجام شده است این پیشین را بکار گرفته و به مدل مناسبتری دست یافته‌اند (۱۷۱). مدل بیزی حاوی اطلاع زمان بیشتری برای رسیدن به همگرایی نیاز داشت. اما بررسی آماره DIC مدل‌ها نشان داد که مدل‌های بیزی فاقد اطلاع با پیشینه‌های یکنواخت و نرمال تعریف شده یکسان بودند اما با بکارگیری اطلاعات پیشین شاخص DIC به اندازه‌ی ۱۳۹ واحد کاهش یافت. این کاهش نشان داد که وجود این همبستگی‌ها باعث مقدار هرچند جزئی کمبود برازش در مدل بوده است.

۶ نتیجه گیری

نتایج این مطالعه نشان داد که بکارگیری رویکرد بیزی معادلات ساختاری نتایجی حداقل در سطح رویکرد کلاسیک ارائه می‌کند و در صورتی که اطلاعات پیشین مناسبی وجود داشته باشد منجر به نتایج مناسبتر و دستیابی به برآوردها در شرایط دشوار می‌گردد. بنابراین با توجه به ماهیت رخدادهای طبیعی که منجر به بروز شرایط دشوار در تحلیل می‌گردد به نظر می‌رسد رویکرد بیزی میتواند به عنوان راه حلی در این شرایط مورد توجه قرار گیرد.

مراجع

- [164] Hoyle RH. *Handbook of structural equation modeling*. Guilford Press; 2012.
- [165] Muthen LK, Muthén BO. 2010 *Mplus user's guide*. Muthén and Muthén, Los Angeles. 1998.
- [166] Altindag I, Genc A. *Bayesian Nonlinear Structural Equation Modeling*. Journal of Selcuk University Natural and Applied Science. 2015;4(1):174-92.
- [167] Browne MW. *Generalized least squares estimators in the analysis of covariance structures*. ETS Research Bulletin Series. 1973;1973(1):i-36.
- [168] Yanuar F, editor *The Estimation Process in Bayesian Structural Equation Modeling Approach*. Journal of Physics: Conference Series; 2014: IOP Publishing.
- [169] Hancock G, Mueller R. *Structural Equation Modeling: A Second Course (Quantitative Methods in Education and the Behavioral Sciences)*. IAP-Information Age Publishing Inc, Charlotte. 2006.
- [170] Dunson DB. *Bayesian latent variable models for clustered mixed outcomes*. Journal of the Royal Statistical Society: Series B (Statistical Methodology). 2000;62(2):355-66.
- [171] Zyphur MJ, Oswald FL. *Bayesian Estimation and Inference A User's Guide*. Journal of Management. 2013;0149206313501200. English.
- [172] Arminger G, Muthén BO. *A Bayesian approach to nonlinear latent variable models using the Gibbs sampler and the Metropolis-Hastings algorithm*. Psychometrika. 1998;63(3):271-300.
- [173] Lee S-Y, Song X-Y. *Model comparison of nonlinear structural equation models with fixed covariates*. Psychometrika. 2003;68(1):27-47.
- [174] Ansari A, Jedidi K. *Bayesian factor analysis for multilevel binary observations*. Psychometrika. 2000;65(4):475-96.
- [175] Van Onna M. *Bayesian estimation and model selection in ordered latent class models for polytomous items*. Psychometrika. 2002;67(4):519-38.

- [176] Song X-Y, Lee S-Y. *Analysis of structural equation model with ignorable missing continuous and polytomous data*. Psychometrika. 2002;67(2):261-88.
- [177] Efron B. *The bootstrap and markov-chain monte carlo*. Journal of biopharmaceutical statistics. 2011;21(6):1052-62.
- [178] Geman S, Geman D. *Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images*. IEEE Transactions on pattern analysis and machine intelligence. 1984 (6):721-41.
- [179] Lee S-Y. *Structural equation modeling: A Bayesian approach*. John Wiley & Sons; 2007.
- [180] Lee S-Y, Song X-Y. *Basic and advanced Bayesian structural equation modeling: With applications in the medical and behavioral sciences*. John Wiley & Sons; 2012.
- [181] Ghayour-Mobarhan M, Moohebbati M, Esmaily H, Ebrahimi M, Parizadeh SMR, Heidari-Bakavoli AR, et al. *Mashhad stroke and heart atherosclerotic disorder (MASHAD) study: design, baseline characteristics and 10-year cardiovascular risk estimation*. International journal of public health. 2015;60(5):561-72.
- [182] Muthén B. *Bayesian analysis in Mplus: A brief introduction*. Unpublished manuscript www.statmodel.com/download/IntroBayesVersion. 2010;203.
- [183] Gelman A, Carlin JB, Stern HS, Rubin DB. *Bayesian data analysis*. Chapman & Hall/CRC Boca Raton, FL, USA; 2014.
- [184] Ashby D. *Bayesian statistics in medicine: a 25 year review*. Statistics in medicine. 2006;25(21):3589-631.