



Fuzzy Multiple Logistic Regression Under Fuzzy Random Errors

Maleki, M. , Nilisani, H. R.  Akbari, M. G. 

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

Corresponding author: H. R. Nili-Sani, hnilisani@birjand.ac.ir

Received: 27/1/2024 Revised: 26/9/2024 Accepted and Published Online: 27/9/2024.

Introduction

Data classification is one of the most critical subjects in machine learning and statistics. Several methods have been suggested for classification. One of the appropriate and efficient methods is logistic regression for crisp variables. But a big part of data in psychology, Sociology, clinical research, etc., is vague and random. In these cases, the classification is also vague, and the essential assumptions of logistic regression are invalid. In this paper, we generalized ordinary logistic regression to fuzzy random variables. We may consider response, error, or independent random variables or parameters non-precise. We tried to fit an appropriate logistic regression model for a case of fuzzy random variables and to determine the recurrent relations such that the possibility of updating the model was provided.

Material and Methods

It is a well-known fact that $\tilde{u}$ is a fuzzy number if and only if the 1-cut set isn't empty and for all $(0 \leq \alpha \leq 1)$, α -cut set is a closed and bounded interval (Goetschel and Voxman, 1986). An especial case of fuzzy numbers is L - R fuzzy numbers where L and R are non-increasing nonnegative real-valued function on $[0, \infty)$ (Taheri and Mashinchi, 2007). The set of fuzzy number is denoted by $F(\mathbb{R})$. The function $\tilde{X}^* : \Omega \rightarrow F(\mathbb{R})$ is called a fuzzy random variable, if is measurable (Joo and Kim, 2001). A triangular fuzzy random variable is denoted by $\tilde{X}^* = (X^L, X, X^U)$ if X^L, X and X^U are three random variables (Hesamian and Akbari, 2022).

Results and Discussion

In the remainder, Y is response random variable, $\mathbf{X}$ is a vector of independent

random variables, $\tilde{\epsilon}^*$ is fuzzy random variable and $\tilde{\mu} = E(Y|\mathbf{X}) = \beta' \mathbf{X} \oplus \tilde{\epsilon}$.

Lemma 1: Under the conditions of the above model if $\tilde{\epsilon}$ is a fuzzy number, $\tilde{\epsilon} = (m^*; a, b)_{LR}$, with the support of $[0, 1]$, then $\tilde{w} = \ln(\frac{\tilde{\mu}}{1-\tilde{\mu}}) = \beta' \mathbf{x} \oplus \tilde{\epsilon}$, is a fuzzy number with the following membership function

$$\tilde{w}(y) = \begin{cases} L(\frac{m - \frac{1}{1+e^{-y}}}{a}), & \frac{1}{1+e^{-y}} \leq m, \\ R(\frac{\frac{1}{1+e^{-y}} - m}{b}), & \frac{1}{1+e^{-y}} \geq m, \end{cases}$$

Theorem 1: Suppose $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$ are n independent observations of model $Y = \beta' \mathbf{X} \oplus \tilde{\epsilon}^*$, $\beta = (\beta_0, \dots, \beta_m)$, and $\tilde{\epsilon}_i^*$ are triangular fuzzy random variables, $\tilde{\epsilon}_i^* = (e_i; X_i^L, X_i^U)_T$, where e_i is the center of error(=0) and X_i^L and X_i^U are uniform random variables on $(0, 1)$. The least-square estimate of β of $\tilde{w} = \ln(\frac{\tilde{\mu}}{1-\tilde{\mu}}) = \beta' \mathbf{x} \oplus \tilde{\epsilon}$, and by Xu and Li's meter with $f(\alpha) = \alpha$ is given by $\beta = 2(\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i')^{-1} \sum_{i=1}^n z_i \mathbf{x}_i$, where

$$z_i = \int_0^1 \alpha (\ln \frac{h_{il}(\alpha - 1) + m_i}{1 - h_{il}(\alpha - 1) - m_i}) d\alpha + \int_0^1 \alpha (\ln \frac{h_{ir}(1 - \alpha) + m_i}{1 - h_{ir}(1 - \alpha) - m_i}) d\alpha.$$

Conclusion

In this study, we discussed logistic regression and assumed that the error and response are triangular fuzzy random variables, and the explanatory variables are crisp. Model parameters were estimated using the least squares approach, and a recurrence method for computing parameters was obtained. Two examples have been used to detect the applied aspect of our model.

Keywords: Multiple Logistic Regression, Fuzzy Random Variable.

Mathematics Subject Classification (2010): 62J12, 62H30.



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

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

رگرسیون لوژستیک چندگانه فازی تحت خطاهای تصادفی فازی

مریم مالکی، حمیدرضا نیلی‌ثانی، محمد قاسم اکبری

گروه آمار، دانشگاه بیرجند

نویسنده مسئول: حمیدرضا نیلی‌ثانی، hnilisani@birjand.ac.ir

تاریخ دریافت: ۱۴۰۲/۱۱/۷ تاریخ بازنگری: ۱۴۰۳/۷/۵ تاریخ پذیرش و انتشار: ۱۴۰۳/۷/۶

چکیده: در این مقاله، موضوع طبقه‌بندی داده‌ها مدنظر قرار داده می‌شود که در آن متغیر پاسخ به‌صورت دو یا چند ارزشی و متغیرهای پیشگو متغیرهای معمولی هستند اما، خطاها علاوه بر ماهیتی تصادفی، ماهیتی ابهامی نیز دارند. در این صورت متغیر پاسخ نیز متغیر تصادفی فازی است. بر این اساس مدلی بر پایه رگرسیون لوژستیک صورت‌بندی کرده و برآورد ضرایب با استفاده از روش کمترین توان‌های دوم بدست آورده می‌شود. با یک مثال نتایج حاصله برای حالت یک متغیر مستقل تشریح می‌گردند. در پایان روابط بازگشتی برای محاسبه برآورد پارامترها ارائه می‌شوند. این روابط بازگشتی می‌توانند در یادگیری ماشین و برای طبقه‌بندی داده‌های بزرگ مورد استفاده قرار گیرند.

واژه‌های کلیدی: رگرسیون لوژستیک چندگانه، متغیر تصادفی فازی.

کد موضوع‌بندی ریاضی (۲۰۱۰): 62J12 ، 62H30.

۱ مقدمه

طبقه‌بندی داده‌ها در آمار و یادگیری ماشین از اهمیت زیادی برخوردار است و روش‌های مختلفی برای آن پیشنهاد شده است. یکی از روش‌های مناسب و کارآمد طبقه‌بندی داده‌ها، رگرسیون لوژستیک است (آگریستی (۲۰۱۵)). از این روش برای رده‌بندی انواع متنوعی از داده‌ها استفاده می‌شود (مقیم‌بیگی (۱۴۰۱)). با گسترش فناوری اطلاعات امکان ذخیره‌سازی و پردازش مه داده‌ها، یعنی مجموعه داده‌هایی که اندازه آن فراتر از توانایی ابزارهای نرم افزارهای رایج برای پردازش و تحلیل داده است، از جمله داده‌های موجود در شبکه‌های اجتماعی، نجوم، ژنتیک و غیره، بیش



©نویسندگان). ناشر انجمن آمار ایران است.

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

از گذشته فراهم شده است. بخش قابل توجهی از این داده‌ها علاوه بر ماهیت تصادفی، ماهیت ابهامی یا نادقیق نیز دارند. در بسیاری از تحقیقات مرسوم و معمول، از جمله در تحقیقات بالینی، علوم اجتماعی و روانشناسی نیز چنین داده‌هایی مشاهده می‌شوند. اما، چنانچه داده‌ها نادقیق یا فازی باشند، مرزهای طبقه‌بندی نیز مبهم و نادقیق خواهند بود. به عنوان نمونه چنانچه هدف طبقه‌بندی داده‌ها به دو رده متمایز باشند، به دلیل نادقیق بودن مشاهدات، این طبقه‌بندی نیز نادقیق است و در چنین مواردی فرض‌های مورد نیاز برای رگرسیون برقرار نخواهد بود. مثلاً، در رگرسیون لوژیستیک معمولی متغیر پاسخ دو حالتی دارای توزیع برنولی است، اما اگر متغیر پاسخ نادقیق باشد، چنین فرضی در مورد متغیر پاسخ برقرار نیست. یک راهکار این است که مشاهدات نادقیق حذف شوند. این رویکرد غیر منطقی است، زیرا ممکن است همه مشاهدات نادقیق باشند. ناکارآمدی رگرسیون لوژیستیک معمولی برای مدل‌بندی مشاهدات رسته‌ای فازی و فراوانی اینگونه مشاهدات و ضرورت بررسی، مدل‌بندی و به خصوص طبقه‌بندی داده‌های نادقیق، از جنبه‌های مختلف و با روش‌های مختلف مورد توجه محققین قرار گرفته است.

تحلیل مدل‌های رگرسیونی لوژیستیک، وقتی متغیرهای پیشگو و وابسته به صورت فازی گزارش شوند، توسط **تاکمورا (۲۰۰۴)** مورد مطالعه قرار گرفته است. **تسینگ و لین (۲۰۰۵)** برای پیش‌بینی ورشکستگی در مؤسسات مالی انگلستان از رگرسیون لوژیستیک دو ارزشی بازه‌ای درجه دوم با ضرایب فازی استفاده کردند. آنها مسئله رگرسیونی را به عنوان یک مساله برنامه‌ریزی درجه دوم ریاضی مدل‌بندی کردند. در روش آنها، ابهام کلی بر مبنای پهنای باندها و مجموع توان دوم فاصله بین مراکز خروجی برآورد شده و مشاهدات ورودی نسبت به قیودی بر حسب الفا برش‌ها مینیمم می‌شوند. **ناگار و سریواستاوا (۲۰۰۸)** مدلی برای رگرسیون لوژیستیک فازی معرفی و مدل را برای پیش‌بینی یک نوع سرطان مورد بررسی قرار داده و نتایج را با شبکه‌های عصبی فازی مقایسه نمودند. **دام (۲۰۰۹)** به بررسی عوامل خطر مرتبط با سرطان دهان از طریق رگرسیون فازی پرداخت. بررسی رگرسیون لوژیستیک دو ارزشی بازه‌ای و رگرسیون لوژیستیک دو ارزشی امکانی توسط **طاهری و میرزایی (۲۰۰۹)** و **میرزایی و طاهری (۲۰۱۰)** انجام شده است. **پوراحمد و همکاران (b,a2011)** رگرسیون لوژیستیک دو ارزشی را برای حالتی که ورودی‌ها اعدادی دقیق و خروجی‌ها و ضرایب اعداد فازی باشند، ارائه داده‌اند و به روش کمترین توان دوم فازی پارامترها را برآورد کردند. **نامداری و همکاران (۲۰۱۳)** مدل‌های رگرسیونی لوژیستیک فازی را وقتی متغیرهای پیشگو به صورت معمولی و متغیرهای پاسخ به صورت متغیرهای زبانی بیان شوند، مورد بررسی و مطالعه قرار داده‌اند. **یو و تسینگ (۲۰۰۵)** مدل‌های رگرسیونی لوژیستیک فازی را برای پردازش نور در یک کارخانه تلویزیون به کار گرفتند. **نامداری و همکاران (۲۰۱۵)** مدل‌های رگرسیونی لوژیستیک فازی را بر اساس روش کمترین قدر مطلق انحرافات بررسی کرده‌اند. **حسامیان و اکبری (۲۰۱۷)** به بررسی مدل‌های رگرسیون لوژیستیک نیمه پارامتری با متغیر وابسته فازی شهودی و متغیرهای پیشگو دقیق پرداخته‌اند. **سلمانی و همکاران (۲۰۱۷)** مدل‌های رگرسیون لوژیستیک، وقتی متغیرهای مستقل و وابسته به صورت فازی گزارش شوند، مورد مطالعه قرار دادند.

رده بندی یک روش یادگیری با ناظر است. روش‌های رده‌بندی را می‌توان به دو دسته کلی تقسیم‌بندی کرد، روش‌های مبتنی بر یادگیری آماری، مانند روش‌های بیزی و روش‌های رگرسیون لوژیستیک و روش‌های مبتنی بر یادگیری ماشین، مانند روش‌های مبتنی بر درخت و روش‌های ماشین‌پرداز پشتیبان، برخی از روش‌ها نیز از تلفیق

روش‌های یادگیری آماری و یادگیری ماشین ایجاد می‌شوند. در روش‌های مبتنی بر یادگیری ماشین، که عموماً برای داده‌های بزرگ استفاده می‌شوند، هیچ فرض آماری پذیرفته نمی‌شود. در مقابل روش یادگیری آماری بر فرضیات آماری و ساختار احتمالی نمونه‌های تصادفی بنا شده و از این فرضیات و ساختارها در رده‌بندی، برآورد و آزمون فرض‌ها استفاده می‌شود. در تحقیقاتی که در زمینه رگرسیون لوژستیک فازی تاکنون انجام شده است (تا آنجایی که مؤلفین اطلاع دارند) هیچ ساختار احتمالی برای داده‌ها یا خطاهای تصادفی که ناشی از انتخاب تصادفی نمونه‌ها یا اندازه‌گیری‌ها است، در نظر گرفته نشده‌اند یا خطاهای موجود در مدل با فازی بودن ضرایب یا فازی بودن متغیرها تلفیق شده‌اند. در صورتی که این خطاها در بسیاری از موارد مشهود و با ماهیت فازی متفاوت است. به عنوان مثال هنگامی که وجود یا عدم وجود درد در بیماران خاص، مثلاً تحت یک عمل جراحی، مورد بررسی قرار می‌گیرد، میزان دردی که توسط فرد بیمار گزارش می‌شود از اهمیت خاصی برخوردار است. میزان درد هرچند که ماهیتی ابهامی دارد، آستانه و میزان درد از فردی به فرد دیگر می‌تواند متفاوت باشد، بنابراین ماهیتی تصادفی نیز دارد. در صورت نادیده گرفتن ماهیت تصادفی داده‌ها، همواره این امکان وجود دارد که مدل خوبی به داده‌ها برازش شود. پوراحمد و همکاران (۱۳۹۱) و نیکویی برازش با معیارهای مختلفی مانند قابلیت یا تعمیم مفاهیم مانند دقت اندازه‌گیری شود. اما، با توجه به اینکه در اکثر موارد اندازه نمونه کوچک است، در چنین شرایطی تضمینی وجود ندارد که مدل منتخب به خوبی به مجموعه دیگری از داده‌ها برازش گردد، چون مشخص نیست که ساختار احتمالی دو مجموعه داده‌ها یکسان باشند. از سوی دیگر، در صورت مواجهه با یک مجموعه از داده‌ها و مدل‌بندی آنها ضرورت دارد تا با تعیین روابط بازگشتی این امکان فراهم آید که با افزودن یک یا چند داده جدید، مدل به سهولت اصلاح گردد. در این مقاله تلاش می‌شود با ورود جمله خطا به مدل که ماهیتی تصادفی و فازی دارد، مدل مناسب به داده‌ها برازش و با تعیین روابط بازگشتی امکان به‌روزرسانی سریع مدل فراهم آید. در بخش ۲، مفاهیم مقدماتی که در بخش‌های بعدی مورد استفاده و استناد قرار خواهند گرفت، ارائه می‌شوند. بخش ۳ به ارائه نتایج اصلی اختصاص داده شده است. در بخش ۴ و به منظور تشریح نتایج اصلی یک مثال ارائه خواهد شد. جمع‌بندی و نتیجه‌گیری در بخش ۵ آورده شده است.

۲ مفاهیم مقدماتی

تعریف ۱. تابع $[0, 1] \rightarrow \mathbb{R}^n : \tilde{u}$ یک مجموعه فازی از مجموعه مرجع $\mathbb{R}^n$ نامیده می‌شود و α -برش مجموعه فازی $\tilde{u}$ به صورت

$$\tilde{u}_\alpha = \begin{cases} \{x \in \mathbb{R}^n : \tilde{u}(x) \geq \alpha\}, & 0 < \alpha \leq 1, \\ \text{supp } \tilde{u}, & \alpha = 0, \end{cases}$$

تعریف می‌شود، که در آن $\text{supp } \tilde{u} = \text{cl}\{x \in \mathbb{R}^n : \tilde{u}(x) > 0\}$. مجموعه فازی $\tilde{u}$ را نرمال نامند، اگر $\tilde{u}_1 \neq \emptyset$ و محدب نامند، اگر برای هر $0 \leq \alpha \leq 1$ ، $\tilde{u}_\alpha$ مجموعه محدب باشد. $\tilde{u}$ را تابع عضویت نیز می‌نامند.

۴۰۰ رگرسیون لوژستیک چندگانه فازی

تعریف ۲. مجموعه فازی $\tilde{u}$ را بر $\mathbb{R}$ یک عدد فازی گویند، اگر $\tilde{u}$ مجموعه فازی محدب، نرمال، نیم‌پیوسته بالایی و $\text{supp } \tilde{u}$ فشرده باشد. خانواده همه اعداد فازی بر $\mathbb{R}$ با $F(\mathbb{R})$ نشان داده می‌شود.

بنابراین در حالتی که $\tilde{u}$ بر $\mathbb{R}$ تعریف شده باشد، $\tilde{u}$ یک عدد فازی است اگر و تنها اگر $\tilde{u}_\alpha \neq \emptyset$ و $\tilde{u}_\alpha$ برای هر $0 \leq \alpha \leq 1$ ، یک فاصله بسته و کراندار، $\tilde{u}_\alpha = [\tilde{u}_{\ell\alpha}, \tilde{u}_{r\alpha}]$ باشد. خواص بیشتری از اعداد فازی $\tilde{u} \in F(\mathbb{R})$ را می‌توانید در **گوتشل و واکسمن (۱۹۸۶)** بیابید. در ادامه فرض می‌شود $u_\ell(\alpha) = u_{\ell\alpha}$ و $u_r(\alpha) = u_{r\alpha}$ ، $0 \leq \alpha \leq 1$.

فرض کنید L و R توابع غیر صعودی از $\mathbb{R}^+$ به $[0, 1]$ باشند و $L(0) = R(0) = 1$. اگر ساختار تابع عضویت عدد فازی $\tilde{u}$ به صورت

$$\tilde{u}(x) = \begin{cases} L(\frac{m-x}{a}), & x \leq m, \\ R(\frac{x-m}{b}), & x \geq m, \end{cases}$$

باشد، $\tilde{u}$ را یک عدد فازی LR نامیده و با نماد $(m; a, b)_{LR}$ نشان داده می‌شود. عدد حقیقی m را مرکز و اعداد مثبت a و b به ترتیب پهنای چپ و راست می‌نامند. اگر $L = R$ و $a = b$ آنگاه عدد فازی را متقارن و اگر $L = R = \max\{0, 1 - |x|\}$ آن را عدد فازی مثلثی نامند و با نماد $(m; a, b)_T$ نمایش می‌دهند.

تعریف ۳. اصل گسترش، (طاهری و ماشین چی، ۱۳۸۷) فرض کنید n ، مجموعه مرجع، $X = X_1 \times \dots \times X_n$ حاصلضرب دکارتی آنها، $A_1, \dots, A_n$ ، مجموعه فازی به ترتیب از $X_1, \dots, X_n$ و $y = f(x_1, \dots, x_n)$ یک نگاشت از X به Y باشند. حاصل عمل f بر n مجموعه فازی $A_1, \dots, A_n$ به صورت مجموعه فازی B از Y با تابع عضویت

$$B(y) = f(A_1, \dots, A_n)(y) = \begin{cases} \sup_{X_1, \dots, X_n, y=f(X_1, \dots, X_n)} \min\{A_1(x_1), \dots, A_n(x_n)\}, & f^{-1}(y) \neq \emptyset, \\ 0, & f^{-1}(y) = \emptyset, \end{cases}$$

تعریف می‌شود، که در آن $f^{-1}(y)$ نگاشت معکوس y تحت f است.

فرض کنید $\tilde{\mu} = (m_1; a_1, b_1)_T$ و $\tilde{\nu} = (m_2; a_2, b_2)_T$ دو عدد فازی مثلثی باشند. تحت تعریف ۳ عملگر جمع، ضرب اسکالر و تفریق به شرح زیر تعریف می‌شوند (لی، ۲۰۰۵) :

الف: جمع $\tilde{\mu} \oplus \tilde{\nu} = (m_1 + m_2; a_1 + a_2, b_1 + b_2)_T$

ب: ضرب اسکالر

$$\lambda \otimes \tilde{\mu} = \begin{cases} (\lambda m_1; \lambda a_1, \lambda b_1)_T, & \lambda > 0, \\ (\lambda m_1; \lambda b_1, \lambda a_1)_T & \lambda < 0, \end{cases}$$

ج: تفریک $\tilde{\mu} \ominus \tilde{\nu} = \tilde{\mu} \oplus (-1 \otimes \tilde{\nu})$

تحت عملگرهای حسابی جمع و ضرب اسکالر اعداد فازی مثلثی یک فضای برداری خواهند بود.

تعریف ۴. (ژو و لی، ۲۰۰۱). فرض کنید $\tilde{u}$ و $\tilde{\nu}$ دو عدد فازی و $\tilde{\nu}_\alpha = [\nu_l(\alpha), \nu_r(\alpha)]$ و $\tilde{u}_\alpha = [u_l(\alpha), u_r(\alpha)]$ به ترتیب α -برش‌های $\tilde{u}$ و $\tilde{\nu}$ باشند. فاصله بین این دو عدد فازی بر پایه تابع وزنی $f(\alpha)$ به صورت

$$d(\tilde{u}, \tilde{\nu}) = \left[\int_0^1 f(\alpha) d^*(\tilde{u}_\alpha, \tilde{\nu}_\alpha) d\alpha \right]^{\frac{1}{\gamma}},$$

تعریف می‌شود، که در آن $d^*(\tilde{u}_\alpha, \tilde{\nu}_\alpha) = [u_l(\alpha) - \nu_l(\alpha)]^2 + [u_r(\alpha) - \nu_r(\alpha)]^2$ و $f(\alpha)$ تابعی صعودی روی بازه $[0, 1]$ است، به قسمی که $\int_0^1 f(\alpha) d\alpha = \frac{1}{\gamma}$ و $f(0) = 0$. در این حالت نرم عدد فازی $\tilde{u}$ به صورت $d(\tilde{u}, 0)$ تعریف می‌گردد.

فرض کنید (Ω, F, P) یک فضای احتمال باشد. پوری و والسکو (۱۹۸۶) تعریفی از متغیرهای تصادفی فازی را ارائه کردند که مبنای بیشتر مطالعات بعدی است.

تعریف ۵. تابع $\tilde{X}^* : \Omega \rightarrow F(\mathbb{R})$ را متغیر تصادفی فازی نامند، اگر اندازه‌پذیر (گات، ۱۹۸۶) باشد. متغیر تصادفی فازی $\tilde{X}^*$ را کراندار انتگرال‌پذیر گویند اگر $E\|\tilde{X}^*\| < \infty$.

لم ۱. (جو و کیم، ۲۰۰۱) فرض کنید $\tilde{X}^* : \Omega \rightarrow F(\mathbb{R})$ و $0 \leq \alpha \leq 1$ و $\tilde{X}^*(\omega) = \{(X_{\ell\alpha}(\omega), X_{r\alpha}(\omega))\}$ تحقق‌ی از $\tilde{X}^*$ باشد. $\tilde{X}^*$ یک متغیر تصادفی فازی است اگر و فقط اگر برای هر $0 \leq \alpha \leq 1$ کران $X_{\ell\alpha}$ پایین آلفا برش $\tilde{X}^*$ و $X_{r\alpha}$ کران بالای آلفا برش $\tilde{X}^*$ متغیرهای تصادفی معمولی باشند.

متغیرهای تصادفی فازی مثلثی را می‌توان با سه متغیر تصادفی و به شکل (X^L, X, X^U) تعریف کرد (حسامیان و اکبری، ۲۰۲۲).

تعریف ۶. امید ریاضی متغیر تصادفی فازی کراندار انتگرال‌پذیر $\tilde{X}^*$ یک عدد فازی است که با

$$E(\tilde{X}^*) = \left\{ \left[\int X_{\ell\alpha} dp, \int X_{r\alpha} dp \right] \mid 0 \leq \alpha \leq 1 \right\},$$

تعریف می‌شود. در این صورت $\left\{ \alpha \in [\circ, \mathbf{1}] : x \in E(\tilde{X}_\alpha^*) \right\}$. $E(\tilde{X}^*)(x) = \sup$
فرض کنید $\mathbf{X}$ بردار پیشگو و Y متغیر پاسخ نامنفی دو یا چند حالتی باشند. در بسیاری موارد ارتباط بین متغیر
پاسخ با بردار پیشگو، $\mathbf{X}$ ، مدل‌بندی می‌شود. برای اینکه

$$Y = \beta' \mathbf{X} \oplus \tilde{\epsilon}^*, \quad \beta = (\beta_\circ, \dots, \beta_m), \quad \mathbf{X} = (\mathbf{1}, X_1, \dots, X_m), \quad (1)$$

که در آن β بردار پارامترهای مجهول و $\tilde{\epsilon}^*$ متغیر تصادفی فازی (مثلی) است، یک الگوی مناسب برای مدل‌بندی
ارتباط بین متغیر پاسخ و بردار پیشگو باشد، لازم است پارامترهای مدل برآورد شوند.

لم ۲. تحت شرایط (۱) اگر $\tilde{\epsilon}^*$ متغیر تصادفی فازی مثلی با مرکز باشد، Y به شرط $\mathbf{X}$ نیز متغیر تصادفی فازی
مثلی است.

برهان: به ازای هر $\omega \in \Omega$ و بنا به اصل گسترش $Y|\mathbf{X}$ یک عدد فازی مثلی است. با توجه به خاصیت خطی
بودن توابع اندازه‌پذیر، $Y|\mathbf{X}$ اندازه‌پذیر و لذا یک متغیر تصادفی فازی مثلی است.

در مدل رگرسیون لوژیستیک مقادیر متغیرهای پیشگو می‌توانند هر مقدار حقیقی باشند اما مقادیر متغیر پاسخ
برچسب‌های رده‌بندی نامنفی (متغیرهای تصادفی فازی نامنفی) هستند، لذا به جای پیشگویی متغیر تصادفی پاسخ،
که ماهیتی ابهامی نیز دارد، امید ریاضی آن، یعنی

$$\tilde{\mu} = E(Y|\mathbf{X}) = \beta' \mathbf{X} \oplus \tilde{\epsilon}, \quad (2)$$

مدل‌بندی می‌شود، که در آن $\tilde{\epsilon} = E(\tilde{\epsilon}^*)$ یک عدد فازی است. در این صورت $\tilde{\mu}$ نیز یک عدد فازی است.

۳ نتایج اصلی

تابع خطی سمت راست (۱) یک تابع بیکران است. متغیرهای پیشگو می‌توانند هر مقدار حقیقی باشند، لذا برای
به الگو درآوردن مقدار متناهی (سمت چپ) مناسب نخواهد بود. در ادامه از تعمیم نظریه رگرسیون لوژیستیک برای
مدل‌بندی ارتباط بین $\tilde{\mu}$ و $\mathbf{X}$ استفاده خواهد شد. فرض می‌شود که $\tilde{\epsilon}$ یک عدد فازی LR ، $\tilde{\epsilon} = (m^*; a, b)_{LR}$ ،
، $a, b > \circ$ باشد.

گزاره ۱. اگر در مدل (۱)، $a, b > \circ$ باشد، آنگاه تحت شرط $\mathbf{X} = \mathbf{x}$ ، $\tilde{\mu} = (m = m^* + \beta' \mathbf{x}; a, b)_{LR}$ ،
، $a, b > \circ$ است.

برهان: بنا به اصل گسترش، $\tilde{\mu}$ یک مجموعه فازی است. چون $\tilde{\mu}$ تابع یک به یک است،

$$\tilde{\mu}(y) = \sup_{e, y = \beta'x + e} \tilde{\epsilon}(y - \beta'x) = \tilde{\epsilon}(e). \quad (۳)$$

و لذا از ویژگی جمع اعداد مثلثی، نتیجه مورد نظر حاصل می شود. بدیهی است که در معادله (۲) تکیه گاه طرف چپ بازه صفر و یک، یا در حالت کلی یک بازه مثبت و متناهی، و طرف راست $\mathbb{R}$ است. لذا مدل بندی برای تبدیل (۲) بر نمونه تصادفی به حجم n اعمال می شود:

$$\tilde{w}_i = \ln\left(\frac{\tilde{\mu}_i}{1 - \tilde{\mu}_i}\right) = \beta'x_i \oplus \tilde{\epsilon}_i. \quad i = 1, \dots, n. \quad (۴)$$

لم ۳. تحت شرایط مدل (۲) اگر $\tilde{\epsilon}$ یک عدد فازی LR ، $\tilde{\epsilon} = (m^*; a, b)_{LR}$ با تکیه گاه $[0, 1]$ باشد، آنگاه $\tilde{w} = \ln\left(\frac{\tilde{\mu}}{1 - \tilde{\mu}}\right) = \beta'x \oplus \tilde{\epsilon}$ ، یک عدد فازی با تابع عضویت زیر خواهد بود.

$$\tilde{w}(y) = \begin{cases} L\left(\frac{m - \frac{1}{1+e^{-y}}}{a}\right), & \frac{1}{1+e^{-y}} \leq m, \\ R\left(\frac{\frac{1}{1+e^{-y}} - m}{b}\right), & \frac{1}{1+e^{-y}} \geq m, \end{cases}$$

برهان: بنا به اصل گسترش

$$\tilde{w}(y) = \sup_{x \in [0, 1], y = \ln \frac{x}{1-x}} \tilde{\mu}(x) = \tilde{\mu}\left(\frac{1}{1+e^{-y}}\right),$$

لذا بنا به گزاره ۱ حکم ثابت است.

فرع ۱. اگر $\tilde{\epsilon} = E(\tilde{\epsilon}^*) = (0; \frac{1}{\alpha})_T$ ، آنگاه $\tilde{\mu} = (m = \beta'x; \frac{1}{\alpha})_T$

$$\tilde{\mu}(x) = \begin{cases} 1 - \alpha|m - x|, & m - \alpha \leq x \leq m, \\ 1 - \alpha|x - m|, & m < x < m + \alpha, \end{cases}$$

و $\tilde{\mu}_\alpha = [m + \alpha(\alpha - 1), m + \alpha(1 - \alpha)]$ اگر $\tilde{\epsilon} = (m^*; a, b)_T$ و $a, b > 0$ آنگاه

$$\tilde{\mu}(x) = \begin{cases} 1 - \frac{|m-x|}{a}, & m - a \leq x \leq m, \\ 1 - \frac{|x-m|}{b}, & m < x \leq m + b. \end{cases}$$

فرع ۲. اگر $\tilde{\epsilon} = E(\tilde{\epsilon}^*) = (\circ; \frac{1}{\tau})_T$ ، آنگاه آلفا برش‌ها $\tilde{w}$ به صورت

$$\tilde{w}_\alpha = \{y : \tilde{w}(y) \geq \alpha\} = \{y : \tilde{\mu}(\frac{1}{1+e^{-y}}) \geq \alpha\}.$$

نتیجه می‌شود. با فرض آنکه $\tilde{\mu}$ عددی فازی با تکیه‌گاه مقادیر نامنفی است، $m + \circ \cdot \delta(\alpha - 1) \geq \circ$ و

$$\begin{aligned} \tilde{w}_\alpha &= \{y : m + \circ \cdot \delta(\alpha - 1) \leq \frac{1}{1+e^{-y}} \leq m + \circ \cdot \delta(1 - \alpha)\} \\ &= \{y : \ln \frac{m + \circ \cdot \delta(\alpha - 1)}{1 - m - \circ \cdot \delta(\alpha - 1)} \leq y \leq \ln \frac{m + \circ \cdot \delta(1 - \alpha)}{1 - m - \circ \cdot \delta(1 - \alpha)}\}, \end{aligned}$$

یعنی

$$\tilde{w}_\alpha = [\ln \frac{(\alpha - 1) \circ \cdot \delta + m}{1 - ((\alpha - 1) \circ \cdot \delta + m)}, \ln \frac{(1 - \alpha) \circ \cdot \delta + m}{1 - ((1 - \alpha) \circ \cdot \delta + m)}] = [A_\ell, A_r]. \quad (5)$$

در حالت کلی‌تر تحت شرایط لم ۳، $\tilde{\epsilon} = E(\tilde{\epsilon}^*) = (m^*; a, b)_T$ ، آنگاه آلفا برش‌های $\tilde{w}$ نتیجه می‌شود

$$\begin{aligned} \tilde{w}_\alpha &= \{y : \tilde{w}(y) \geq \alpha\} = \{y : \tilde{\mu}(\frac{1}{1+e^{-y}}) \geq \alpha\} \\ &= \{y : \ln \frac{m + b(\alpha - 1)}{1 - m - b(\alpha - 1)} \leq y \leq \ln \frac{m + a(1 - \alpha)}{1 - m - a(1 - \alpha)}\}, \end{aligned}$$

فرع ۳. تحت شرایط مدل (۱) و با این فرض که $\tilde{\epsilon}^*$ متغیر تصادفی فازی مثلثی با مرکز $\circ$ باشد، مقدار مشاهده‌شده y به شرط x مقدار مشاهده‌شده از یک متغیر تصادفی فازی مثلثی به مرکز $\beta'x$ و پهنای متناظر با مقدار مشاهده‌شده از متغیرهای تصادفی X^U و X^L ، مثلاً h_r و h_ℓ ، است. یعنی در این حالت

$$\tilde{w}_\alpha = [\ln \frac{(\alpha - 1)h_\ell + m}{1 - ((\alpha - 1)h_\ell + m)}, \ln \frac{(1 - \alpha)h_r + m}{1 - ((1 - \alpha)h_r + m)}] = [A_\ell, A_r]. \quad (6)$$

تحت شرایط مدل (۱) و با این فرض که $\tilde{\epsilon}^*$ متغیر تصادفی فازی مثلثی با مرکز e باشد، مقدار مشاهده‌شده y به شرط x مقدار مشاهده‌شده از یک متغیر تصادفی فازی مثلثی به مرکز $\beta'x + e$ است.

قضیه ۱. فرض کنید $(x_1, y_1), \dots, (x_n, y_n)$ ، n مشاهده مستقل از مدل (۱) و $\tilde{\epsilon}_i^*$ ها متغیرهای تصادفی فازی مثلثی $\tilde{\epsilon}_i^* = (e_i; X_i^L, X_i^U)_T$ باشند، که در آن مرکز خطا e_i (و برابر صفر) و پهنای X_i^L و X_i^U متغیرهای

مالکی، م. و همکاران ۴۰۵

تصادفی دارای توزیع یکنواخت در بازه $(۰, ۱)$ است. در این صورت

$$\beta = \left(\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i' \right)^{-1} \sum_{i=1}^n z_i \mathbf{x}_i, \quad (۷)$$

برآورد کمترین توانهای دوم پارامتر β در مدل (۴) و نسبت به متر (۱) به ازای $f(\alpha) = \alpha$ است، که در آن m_i مرکز $\tilde{\mu}_i$ است و

$$\begin{aligned} z_i &= \int_0^1 \alpha \left(\ln \frac{h_{i\ell}(\alpha - 1) + m_i}{1 - h_{i\ell}(\alpha - 1) - m_i} \right) d\alpha + \int_0^1 \alpha \left(\ln \frac{h_{ir}(1 - \alpha) + m_i}{1 - h_{ir}(1 - \alpha) - m_i} \right) d\alpha \\ &= z_{i1} + z_{i2}. \end{aligned} \quad (۸)$$

برهان: تحت شرایط قضیه عدد فازی $E(\tilde{\epsilon}^*)$ به صورت $E(\tilde{\epsilon}^*) = (0; \frac{1}{\gamma}, \frac{1}{\gamma})_T$ با آلفا برش‌های $(E(\tilde{\epsilon}^*))_\alpha = [e_i; h_{i\ell}, h_{ir}]_T$ خواهد بود. کران‌های آلفا برش مقادیر مشاهده‌شده متغیر تصادفی فازی $\tilde{\epsilon}_i^* = (e_i; h_{i\ell}, h_{ir})_T$ به ترتیب $\tilde{\mu}_i = (m =$ در این صورت $\tilde{\epsilon}_{r\alpha} = (1 - \alpha)(1 - h_{ir}) + e_i$ ، و $\tilde{\epsilon}_{\ell\alpha} = (-1 + \alpha)h_{i\ell} + e_i$ است. اگر $\mathbf{X}_i = \mathbf{x}_i$ آنگاه تحت شرایط مدل (۴)، $\beta' \mathbf{x}_i + e_i; \frac{1}{\gamma}, \frac{1}{\gamma})_T$

$$\begin{aligned} \tilde{w}_i &= \beta' \mathbf{x}_i + E(\tilde{\epsilon}^*) \\ \tilde{w}_{i\alpha} &= [\beta' \mathbf{x}_i + e_i + \frac{(\alpha - 1)}{\gamma}, \beta' \mathbf{x}_i + e_i + \frac{(1 - \alpha)}{\gamma}] = [B_{i\ell\alpha}, B_{ir\alpha}]. \end{aligned} \quad (۹)$$

پارامترهای مدل به روش کمترین مجموع توان دوم خطاها، به گونه‌ای برآورد می‌شود که فاصله بین بازه مشاهده‌شده $\tilde{w}_{i\alpha} = [B_{i\ell\alpha}, B_{ir\alpha}]$ و بازه برآورد شده $\tilde{w}_{i\alpha} = [A_{i\ell\alpha}, A_{ir\alpha}]$ (که به اختصار از ذکر اندیس α خودداری

می‌گردد) مینیمم شود. در این مرحله از فاصله (۱) استفاده می‌شود. در این صورت

$$\begin{aligned} SSE &= \sum_{i=1}^n \int_0^1 \alpha d^2(\tilde{w}_{i\alpha}, \tilde{\hat{w}}_{i\alpha}) d\alpha \\ &= \sum_{i=1}^n \int_0^1 \alpha \left([A_{i\ell} - B_{i\ell}]^2 + [A_{ir} - B_{ir}]^2 \right) d\alpha \\ &= \sum_{i=1}^n \left[\int_0^1 \alpha [A_{i\ell}^2 + A_{ir}^2] d\alpha + \int_0^1 \alpha (2(\beta' \mathbf{x}_i)^2 + \frac{\alpha^2 - 2\alpha + 1}{2}) d\alpha \right. \\ &\quad \left. - 2 \left[\overbrace{k_{i1} + k_{i2}}^{k_i} + \overbrace{(z_{i1} + z_{i2})}^{z_i} (\beta' \mathbf{x}_i) \right] \right] \\ &= \sum_{i=1}^n \left[\int_0^1 \alpha [A_{i\ell}^2 + A_{ir}^2] d\alpha + ((\beta' \mathbf{x}_i)^2 + \frac{1}{2}) - 2[k_i + z_i(\beta' \mathbf{x}_i)] \right]. \end{aligned}$$

اکنون با مشتق‌گیری از SSE نسبت به β و برابر صفر قرار دادن آن رابطه $\beta = 2(\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i')^{-1} \sum_{i=1}^n z_i \mathbf{x}_i$ را به دست می‌آوریم. در حالت یک متغیره برآورد ضرایب به صورت زیر خواهند بود.

$$\hat{\beta}_0 = \bar{z} - \hat{\beta}_1 \bar{x}, \quad \hat{\beta}_1 = \frac{\sum_{i=1}^n z_i x_i - n \bar{x} \bar{z}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2}. \quad (10)$$

تذکر ۱. قضیه ۱ را می‌توان به هر تابع چندجمله‌ای f و در حالت کلی‌تر هر تابع f که برای آن محاسبه z_i ها امکان‌پذیر باشد تعمیم داد.

فرع ۴. اگر در مفروضات قضیه ۱ فرض پهنای H_i متغیرهای تصادفی دارای توزیع یکنواخت در بازه $(0, 1)$ به فرض $H_i \sim U(\theta_1, \theta_2)$ تغییر یابد، با تکرار اثبات می‌توان نشان داد نتایج قضیه همچنان معتبر است. کافی است کران‌های انتگرال‌ها از $(0, 1)$ به (θ_1, θ_2) تغییر یابند.

۴ مثال‌های کاربردی

در این بخش داده‌ها در مورد بیماران قلبی است و متغیر زبانی به صورت دوحالتی (خیلی کم و خیلی زیاد)، متوسط و بیشتر از آن خیلی زیاد و کمتر از متوسط، خیلی کم در نظر گرفته می‌شود.

مثال ۱. بیماری قلبی: داده‌های بیماری قلبی، حاوی ۷۶ متغیر و ۳۰۳ مشاهده است، اما به دلایل مختلف تنها ۱۴ متغیر و در پایگاه داده‌ای کلیوند منتشر شده‌اند. این مجموعه داده از پایگاه داده‌ای UCI دریافت شده است. نام و

شماره بیمه بیماران از مجموعه داده‌ها حذف و مقادیر ساختگی جایگزین آن شده است. از مجموعه متغیرهای سن، جنسیت، نوع درد، فشارخون در حال استراحت، و وجود بیماری قلبی متغیر هدف یا پاسخ پنج متغیر پیشگو سن، فشارخون در حال استراحت، کلسترول سرم، ماکسیمم ضربان قلب، افسردگی و وجود بیماری قلبی مورد بررسی قرار می‌گیرند. اطلاعات مربوط به ۵ مشاهده نمونه در جدول ۱ ارائه شده است. متغیر پاسخ از ۰ (نداشتن بیماری قلبی) تا ۴ (داشتن بیماری قلبی شدید) ارزش‌گذاری شده است. برای سادگی محاسبات متغیر هدف دوتایی در نظر گرفته می‌شود. در این متغیر ۱ نشان‌دهنده (بیماری قلبی شدید) و ۰ (نداشتن بیماری قلبی شدید) است. در ۱۳

جدول ۱. نمایش ۵ مشاهده اول

Heart	Maximum	Cholesterol	Pressure	Blood_A	Type	Pain	Chest	Sex	Age
۰	۶	۱۵۰	۲	۱	۲۳۳	۱۴۵	۱	۱	۶۳
۲	۳	۱۰۸	۲	۰	۲۸۶	۱۶۰	۴	۱	۶۷
۱	۷	۱۲۹	۲	۰	۲۲۹	۱۲۰	۴	۱	۶۷
۰	۳	۱۸۷	۰	۰	۲۵۰	۱۳۰	۳	۱	۳۷

تا از مشاهدات مقدار متغیر هدف ۱ (گروه ۱) و در مابقی مشاهدات (۲۹۰ مشاهده) صفر (گروه ۲) است. یعنی داده‌ها خیلی نامتعادل هستند. برای ایجاد تعادل در مجموعه داده‌های آموزشی از روش نمونه‌گیری متناسب استفاده شده است. قاعده مشخصی برای نسبت نمونه‌های منتخب از رده‌های مختلف وجود ندارد اما معمولاً تلاش می‌شود که نسبت نمونه از هر رده (در حالت دو رده‌ای) تقریباً و حداقل ۲۰ درصد باشد. بنابراین از گروه اول ۱۲ نمونه و از گروه دوم ۵۸ نمونه انتخاب شده است. در این صورت نسبت نمونه در گروه اول برابر $P_1 = 0.8$ و در گروه دوم برابر $P_2 = 0.2$ است. در مجموعه داده‌های نمونه منتخب یک داده گم‌شده (مقدار متغیر الکتروکاردیوگرافی متغیر ۵۶) وجود دارد که با توجه به فراوانی مقادیر متغیر مقدار ۰ را به عنوان برآورد مقدار گم‌شده در نظر گرفته شده است. بر اساس اطلاعات هر یک از ۷۰ نمونه، از پزشک خواسته شده است تا امکان بیمار بودن آن‌ها را با توجه به علائم و وضعیت عمومی آن‌ها به صورت یکی از واژه‌های زبانی "بیمار بودن (۱) و بیمار نبودن (۰) بیان کند. پذیرفته می‌شود که نظر پزشک ماهیت تصادفی نیز دارد به گونه‌ای که در درجه عضویت واژه‌های زبانی تأثیرگذار است. معادله (۴) را می‌توان به صورت

$$\tilde{\mu}_i = \frac{1}{1 + \exp(-(\beta_0 + \beta_1 X_i \oplus \tilde{\epsilon}_i))} = \frac{\exp(\beta_0 + \beta_1 X_i \oplus \tilde{\epsilon}_i)}{1 + \exp(\beta_0 + \beta_1 X_i \oplus \tilde{\epsilon}_i)}, \quad i = 1, \dots, n.$$

نیز نوشت. شکل تعمیم‌یافته مانسکی و مک فادن (۱۹۸۱) مشخص‌سازی معادله (۴) را می‌توان به صورت

$$\tilde{\mu}_i = \frac{P_1}{P_1 + P_2 \exp(-(\beta_0 + \beta_1 X_i \oplus \tilde{\epsilon}_i))}, \quad i = 1, \dots, n.$$

بیان کرد. اگر $\gamma = \frac{P}{P_1}$ و $\bar{\gamma}_1 = \ln(\frac{P}{P_1})$ آنگاه

$$\bar{\mu}_i = \frac{\exp(\gamma_1 + \beta_0 + \beta_1 X_i \oplus \bar{\epsilon}_i)}{1 + \exp(\gamma_1 + \beta_0 + \beta_1 X_i \oplus \bar{\epsilon}_i)}, \quad i = 1, \dots, n.$$

اگر $(\bar{\epsilon}; \frac{1}{\gamma}, \frac{1}{\gamma}) = E(\bar{\epsilon}^*)$ ، آنگاه از تعمیم نتیجه (۲) حاصل می‌شود:

$$\begin{aligned} \bar{\mu}_\alpha &= \{y : \bar{\mu}(y) \geq \alpha\} = \{y : \bar{\epsilon}_1(\ln \frac{y}{1-y}) \geq \alpha\} \\ &= \{y : \frac{1}{\exp(-(\frac{1}{\gamma}(\alpha - 1) + \gamma_1 + \beta'_1 \mathbf{x})) + 1} \leq y \leq \frac{1}{\exp(-(\frac{1}{\gamma}(1 - \alpha) + \gamma_1 + \beta'_1 \mathbf{x})) + 1}\}. \end{aligned}$$

فرض می‌شود که میزان بیماری قلبی اعداد فازی مثلثی به صورت

نداشتن بیماری قلبی شدید:

$$(\circ 3; u(\circ, \circ 3); u(\circ 3, \circ 6))_T \quad \text{متغیر تصادفی فازی حدوداً } \circ 3 \text{ تابع عضویت:}$$

بیماری قلبی شدید: متغیر تصادفی فازی حدوداً $\circ 8$ با تابع عضویت $(\circ 8; u(\circ 6, \circ 8); u(\circ 8, 1))_T$ است. در این صورت بردارهای m ، x_ℓ و x_u از خواص اعداد فازی مثلثی به دست می‌آیند که به ترتیب عبارت‌اند از

$$\begin{aligned} m &= \{\circ 3, \circ 3, \circ 3, \dots, \circ 8, \circ 8\}, \\ x_\ell &= \{\circ 4, \circ 23, \circ 5, \dots, \circ 71, \circ 67\}, \\ x_r &= \{\circ 60, \circ 48, \circ 40, \dots, \circ 93, \circ 86\}, \end{aligned}$$

از تفریق بردارهای فوق، پهنای چپ و راست حاصل می‌شوند.

$$\begin{aligned} h_\ell &= m - x_\ell = \circ 26, \circ 57, \circ 25, \dots, \circ 59, \circ 13 \\ h_r &= x_r - m = \circ 30, \circ 18, \circ 10, \dots, \circ 13, \circ 56. \end{aligned}$$

در پرونده داده‌ها نخست نمونه‌های فاقد بیماری قلبی شدید و سپس داده‌های دارای بیماری قلبی شدید آورده شده است. پس از برآورد پارامترها، مدل به صورت

$$\tilde{w}_i = -\circ 495 + \circ 23X_1 + \circ 09X_2 + \circ 01X_3 - \circ 24X_4 + \circ 235X_5 \oplus \bar{\epsilon}.$$

جدول ۲. جدول مربوط به پهنایها

h_r	h_ℓ	x_r	x_ℓ	m
۰/۳۰	۰/۲۶	۰/۵۶	۰/۵۴	۰/۳۰
۰/۱۸	۰/۵۷	۰/۴۸	۰/۲۳	۰/۳۰
۰/۱۰	۰/۲۵	۰/۴۰	۰/۵۵	۰/۳۰
۰/۱۷	۰/۵۴	۰/۴۷	۰/۲۶	۰/۳۰
۰/۵۱	۰/۱۷	۰/۳۱	۰/۱۳	۰/۳۰

حاصل شده است. برای اولین داده آزمونی، مقدار پیش‌بینی به صورت

$$\tilde{w}_i = -0.4949 + 0.0235 \times 69 + 0.0093 \times 145 + 0.0010 \times 233 \\ - 0.0236 \times 150 + 0.2350 \times 2/3 \oplus \tilde{\epsilon} = 0.682 \oplus \tilde{\epsilon},$$

و آلفا برش‌ها به صورت

$$\tilde{\mu}_\alpha = \left[\frac{1}{1 + \exp(-((\alpha - 1) \cdot 0.5 - 0.682))}, \frac{1}{1 + \exp(-((1 - \alpha) \cdot 0.5 - 0.682))} \right] [\tilde{A}_L^\alpha, \tilde{A}_U^\alpha],$$

خواهند بود. با استفاده از معیار غیر فازی‌ساز $E_{\tilde{A}} = \int_0^1 (1 - 2^{-\alpha})(\tilde{A}_L^\alpha + \tilde{A}_U^\alpha) d\alpha$ ، خروجی به عدد حقیقی تبدیل می‌شود. در حالت فوق مقدار خروجی ۰/۴۵۵ خواهد بود. اگر هنگامی که مقدار خروجی کمتر از ۰/۵ باشد، رده صفر در نظر گرفته شود، الگوریتم فوق به درستی رده اولین مشاهده آزمونی را تعیین می‌کند. با استفاده از این الگوریتم رده هر یک از مشاهدات آزمونی، شامل ۲۳۳ مشاهده، تعیین شدند. نرخ خطای کل یا به طور ساده‌تر نرخ خطا که برابر است با مجموع منفی‌ها و مثبت‌های کاذب تقسیم بر تعداد کل مشاهدات برابر است با $\frac{9+1}{233}$. به همین صورت نرخ منفی کاذب، تعداد منفی‌های کاذب تقسیم بر تعداد کل طبقه‌بندی‌های منفی، طبقه‌بندی صفر، برابر $\frac{9}{233}$ است. جدول ۳ نرخ خطای کل، نرخ منفی‌های کاذب و نرخ مثبت‌های کاذب را نشان می‌دهد. کد برنامه در <https://github.com/maryammaleki68/Fuzzy-logit/edit/main/README.md> آدرس قابل دسترس است.

جدول ۳. خلاصه نتایج برای داده آزمونی

تشخیص مدل			
واقعی	عدم ابتلا	ابتلا	جمع
عدم ابتلا	۲۲۳	۹	۲۳۲
ابتلا	۱	۰	۱

مثال ۲. شبیه سازی: در این مثال متغیرهای سن، فشارخون در حال استراحت، کلسترل سرم، ماکسیمم ضربان قلب و افسردگی، هر یک با استفاده از توزیع یکنواخت در بازه مقادیر متغیرهای پیشگو که در زیر آمده است، تولید شده و برآورد پارامترها محاسبه شده است. همچنین مقدار کمترین مربعات خطا را برای مقادیر β برآورد شده از داده‌های اصلی و داده‌های شبیه سازی شده محاسبه و برابر ۰/۰۸۵۲۷ است. کدهای مربوط به برنامه شبیه سازی در بخش پیوست آمده است. متغیر سن دارای توزیع $U(۳۳, ۸۰)$ ، متغیر فشارخون در حال استراحت دارای توزیع $U(۱۰۰, ۲۰۰)$ ، متغیر کلسترل سرم دارای توزیع $U(۱۴۰, ۵۷۰)$ ، متغیر ماکسیمم ضربان قلب دارای توزیع $U(۱۱۰, ۲۰۰)$ و متغیر افسردگی دارای توزیع $U(۰, ۶/۵)$ است.

در ادامه روابطی بازگشتی برای برآورد پارامترها به دست آورده شده است. این روابط بازگشتی می‌توانند ابزار مهمی در یادگیری ماشین و برای داده‌های بزرگ باشند. در ادامه از نماد $\bullet$ به جای $n + ۱$ استفاده می‌شود. مثلاً $\bar{z}^\bullet$ نشان‌دهنده میانگین z_i به ازای حجم نمونه برابر $n + ۱$ است. در این صورت در رابطه $\hat{\beta}_\bullet^\bullet = \bar{z}^\bullet - \hat{\beta}_\bullet^\bullet \bar{x}^\bullet$ خواهیم داشت:

$$\bar{z}^\bullet = \frac{n}{n+1} \bar{z} + \frac{z^\bullet}{n+1}, \quad \bar{x}^\bullet = \frac{n}{n+1} \bar{x} + \frac{x^\bullet}{n+1}.$$

در نتیجه

$$\hat{\beta}_\bullet^\bullet = \frac{n}{n+1} (\bar{z} - \hat{\beta}_\bullet^\bullet \bar{x}) + \frac{z^\bullet}{(n+1)} - \hat{\beta}_\bullet^\bullet \frac{x^\bullet}{(n+1)}.$$

که در آن $z^\bullet$ ، $x^\bullet$ مربوط به داده جدید است. همچنین برآورد $\hat{\beta}_\bullet^\bullet$ نیز به صورت

$$\begin{aligned} \hat{\beta}_\bullet^\bullet &= \frac{\sum_{i=1}^{n+1} z_i x_i - (n+1) \bar{x}^\bullet \bar{z}^\bullet}{\sum_{i=1}^{n+1} x_i^2 - (n+1) \bar{x}^{\bullet 2}} \\ &= \frac{\sum_{i=1}^n z_i x_i - \left(\frac{n}{(n+1)} n \bar{x} \bar{z} + \frac{n}{(n+1)} \bar{x} z^\bullet + \frac{n}{n+1} x^\bullet \bar{z} + \frac{z^\bullet x^\bullet}{n+1} \right) + x^\bullet z^\bullet}{\sum_{i=1}^n x_i^2 - (n+1) \left(\frac{n}{(n+1)^2} n \bar{x}^2 + 2 \frac{n}{(n+1)^2} \bar{x} x^\bullet + \frac{x^{\bullet 2}}{(n+1)^2} \right) + x^{\bullet 2}}, \end{aligned}$$

است. برای مقادیر بزرگ n ، $\frac{n}{(n+1)} \cong 1$ و $\frac{1}{(n+1)} \cong 0$ ، لذا

$$\hat{\beta}_1^\bullet \cong \frac{\overbrace{\left(\sum_{i=1}^n z_i x_i - n\bar{x}\bar{z}\right)}^k - \overbrace{\left(\bar{x}z^\bullet + x^\bullet\bar{z} - x^\bullet z^\bullet\right)}^\theta}{\underbrace{\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger}_m - \underbrace{\left(\bar{x}x^\bullet + x^\bullet\bar{x}\right)}_\Omega} \cong \hat{\beta}_1 - \frac{k\Omega + m\theta}{m(m+\Omega)}.$$

بنابراین

$$\hat{\beta}_1^\bullet \cong \hat{\beta}_1 - \frac{(\sum_{i=1}^n z_i x_i - n\bar{x}\bar{z})(-\bar{x}x^\bullet + x^\bullet\bar{x}) + (\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger)(\bar{x}z^\bullet + x^\bullet\bar{z} - x^\bullet z^\bullet)}{(\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger)(\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger - \bar{x}x^\bullet + x^\bullet\bar{x})}.$$

به‌طور مشابه و با جایگذاری تقریب $\hat{\beta}_1^\bullet$ در (۱۱) برای پارامتر $\hat{\beta}_2^\bullet$ خواهیم داشت:

$$\begin{aligned} \hat{\beta}_2^\bullet &= \frac{n}{(n+1)}(\bar{z} - \hat{\beta}_1^\bullet \bar{x}) + \frac{z^\bullet}{(n+1)} - \hat{\beta}_1^\bullet \frac{x^\bullet}{(n+1)} \\ &\cong \hat{\beta}_2 \\ &- \frac{(\sum_{i=1}^n z_i x_i - n\bar{x}\bar{z})(-\bar{x}x^\bullet + x^\bullet\bar{x}) + (\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger)(\bar{x}z^\bullet + x^\bullet\bar{z} - x^\bullet z^\bullet)}{(\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger)(\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger - \bar{x}x^\bullet + x^\bullet\bar{x})} \bar{x} \\ &+ \frac{z^\bullet}{(n+1)} - (\hat{\beta}_1 \\ &- \frac{(\sum_{i=1}^n z_i x_i - n\bar{x}\bar{z})(-\bar{x}x^\bullet + x^\bullet\bar{x}) + (\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger)(\bar{x}z^\bullet + x^\bullet\bar{z} - x^\bullet z^\bullet)}{(\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger)(\sum_{i=1}^n x_i^\dagger - n\bar{x}^\dagger - \bar{x}x^\bullet + x^\bullet\bar{x})}) \\ &\frac{x^\bullet}{(n+1)}. \end{aligned}$$

بحث و نتیجه‌گیری

رگرسیون لوژستیک یکی از روش‌های یادگیری بر پایه آمار و احتمال و ابزار مناسبی برای رده بندی مشاهدات است. در این روش و بر مبنای متغیرهای پیشگو مقدار متغیر وابسته که متغیری رسته‌ای یا اسمی است، برآورد می‌شود. تعیین متغیر وابسته معادل تعیین رده‌ای است که مشاهده به آن رده تعلق دارد. اما در مطالعات اجتماعی، روانشناسی، بالینی،... این امکان وجود دارد که متغیرهای پیشگو یا وابسته یا پارامترهای مدل ماهیتی ابهامی داشته باشند. حالتی که متغیرهای وابسته، مستقل یا پارامترهای مدل فازی باشند، توسط مؤلفین مختلف و تحت شرایط

متنوعی مورد بررسی قرار گرفته‌اند. در این مقاله حالتی بررسی شده است که متغیر پاسخ علاوه بر ماهیت فازی بودن ماهیت تصادفی نیز داشته باشد. در این مقاله فرض شده که متغیر پاسخ به دلیل وجود خطای تصادفی فازی، متغیر تصادفی فازی مثلثی باشد و پارامترهای مدل رگرسیون لوژیستیک برآورد شدند. انتظار می‌رود که این مقاله چشم انداز گسترده‌ای را برای مطالعه و کاربرد رگرسیون لوژیستیک هنگامی که برخی یا همه متغیرها تصادفی فازی باشند، ایجاد کند. اگر حجم نمونه بزرگ باشد، محاسبه برآوردها مشکل خواهد بود. در انتها روشی بازگشتی برای محاسبه پارامترها به دست آمده است. این روابط، کاربرد روش پیشنهادی برای داده‌های بزرگ را امکان‌پذیر می‌سازند. با یک مثال به تشریح نتایج به دست آمده پرداخته شده است. در این مثال داده‌ها با استفاده از نمونه‌گیری متناسب به دو دسته، داده‌های آزمایشی و داده‌های آزمونی تقسیم شده‌اند. بر مبنای داده‌های آزمایشی پارامترها برآورد شدند. با استفاده از پارامترهای برآورده شده رده اولین مشاهده از داده‌های آزمونی (و به درستی) تعیین گردید.

تقدیر و تشکر

مؤلفین از داوران، سردبیر و ویراستار محترم که پیشنهادات ارزشمند ایشان موجب بهبود مقاله گردید کمال تشکر و سپاسگزاری را دارند.

مراجع

- طاهری، س. م. و ماشین چی، م. (۱۳۸۷)، مقدمه‌ای بر احتمال و آمار فازی، دانشگاه شهید باهنر کرمان.
- مقیم بیگی، م. (۱۴۰۱)، رگرسیون لوژیستیک چند جمله‌ای نیمه پارامتری برای رده‌بندی داده‌های شکل، مجله علوم آماری، ۱۶(۲)، ۴۴۹-۴۶۸.
- Agresti, A. (2015), *Foundations of Linear and Generalized Linear Models*, Hoboken: Wiley.
- Fauci, A., Braunwald, S. E., Kasper, D. L., Hauser, S. L., Longo, D. L., Jameson, J. L., and Loscalzo, J. (2008), *Harrison's Principals of Internal Medicine*, New York: Wiley. 2275- 2279.
- Goetschel, R., and Voxman, W. (1986), Elementary Fuzzy Calculus, *Fuzzy Sets and Systems*, 18, 31-43.
- Gut, A. (2005), *Probability, A Graduated Course*, Springer Science+Business Media, Inc.

- Hesamian, G. and Akbari, M. G. (2017), Semi-Parametric Partially Logistic Regression Model with Exact Inputs and Intuitionistic Fuzzy Output Applied, *Soft Computing*, **58**, 517-526.
- Hesamian, G. and Akbari, M. G. (2022), A Fuzzy Empirical Quantile-Based Regression Model Based on Triangular Fuzzy Numbers, *Computational and Applied Mathematics*, **41**(6), 1-26
- Joo, S. Y., and Kim, Y. K. (2001), Kolmogorov's Strong Law of Large Numbers for Fuzzy Random Variables, *Fuzzy Sets and Systems*, **120**(3), 499-503.
- Lee, K. H. (2005), *First Course on Fuzzy Theory and Applications*, Berlin: Springer.
- Mirzaei Yeganeh, S., and Taheri, S.M. (2010), Possibilistic Logistic Regression by Linear Programming Approach, *Proc of 7th Probability and Random Process*, Isfahan University of Technology, Isfahan, Iran, 139-143.
- Mohd Dom, R. (2009), A Fuzzy Regression Model for The Prediction of Oral Cancer Susceptibility. Ph.D. Thesis, *Computer Science and Information Technology*, University of Malaya, Kuala Lumpur.
- Manski, C., and McFadden, D. (1981), Alternative Estimators and Sample Designs for Discrete Choice Analysis. In: Manski, McFadden, Editors, *Structural Analysis of Discrete Data with Econometric Applications*, London: MIT Press.
- Nagar, P., and Srivastava, S. (2008), Adaptive Fuzzy Regression Model for The Prediction of Dichotomous Response Variables Using Cancer Data: A Case Study, *Journal of Applied Mathematics Statistics and Informatics*, **4**, 183-191.
- Namdari, M., Taheri, S.M., Abadi, A., Rezaei, M., and Kalantari, N. (2013), Possibilistic Logistic Regression for Fuzzy Categorical Response Data, *International Conference on Fuzzy Systems*, **8**, 1-6.
- Namdari, M., Yoon, J. H., Abadi, A., Taheri, S. M., and Choi, S. H. (2015), Fuzzy Logistic Regression with Least Absolute Deviations Estimators, *Soft Computing*, **19**, 909-917.

- Pourahmad, S., Ayatollahi, S. M. T., Taheri, S. M., and Agahi, Z. H. (2011.a), Fuzzy Logistic Regression Based on The Least Squares Approach with Application in Clinical Studies, *Computer Mathematics Applications*, **62**, 3353-3365.
- Pourahmad, S., Ayatollahi, S. M., and Taheri, S. M. (2011.b), Fuzzy Logistic Regression, A New Possibilistic Regression and Its Application in Clinical Vague Status, *Iranian Journal of Fuzzy Systems*, **8**, 1-17.
- Puri, M. L., and Ralescu, D. A. (1986), Fuzzy Random Variables, *Journal of Mathematical Analysis and Applications*, **114**, 409-422.
- Salmani, F., Taheri, S. M., Yoon, J. H., Abadi, A., Alavi Majd, H., and Abbaszadeh, A. (2017), Logistic Regression for Fuzzy Covariates: Modeling, Inference and Applications, *International Journal of Fuzzy Systems*, **19**, 1635-1644.
- Taheri, S. M., and Mirzaei Yeganeh, S. (2009), Logistic Regression with Non-Precise Response, *Proceedings of The 57th ISI Conference, Durban, South Africa*, 98-101.
- Takemura, K. (2004), Fuzzy Logistic Regression Analysis for Fuzzy Input -Output Data, *International Conference on Soft Computing and Intelligent Systems and the 5th International Symposium on Advanced Intelligent Systems, Japan*, 1-6.
- Tseng, F., and Lin, L. (2005), A Quadratic Interval Logit Model for Forecasting Bankruptcy, *Omega*, **33**, 85-91.
- Yu, J. R., and Tseng, F. M. (2014), Fuzzy Piecewise Logistic Growth Model for Innovation Diffusion: A Case Study of The TV Industry, *International Journal of Fuzzy Systems*, **18**, 1-12.
- Xu, R., and Li, C. (2001), Multidimensional Least-Squares Fitting with A Fuzzy Model, *Fuzzy Sets and Systems*, **119**(2), 215-223.