

Bayesian Clustering of Spatially Varying Coefficients Zero-Inflated Survival Regression Models

S. Asadi, M. Mohammadzadeh*

Department of Statistics, Tarbiat Modares University, Tehran, Islamic Republic of Iran.

*Email: mohsen_m@modares.ac.ir

خوشه‌بندی بیزی مدل‌های رگرسیون بقای متورم صفر با ضرایب متغیر فضایی

سپیده اسعدی، محسن محمدزاده*

گروه آمار، دانشکده علوم ریاضی، دانشگاه تربیت مدرس، تهران، جمهوری اسلامی ایران

چکیده

این مطالعه به چالش‌های تحلیل داده‌های زمان تا وقوع رویداد می‌پردازد و به‌ویژه بر ماهیت گسسته مدت زمان تا وقوع پیشامد، مانند تعداد سال‌های تا طلاق، تأکید می‌کند. این وضعیت غالباً منجر به داده‌های بقا با توزیع صفر افزوده می‌شود، زیرا فراوانی مشاهدات صفر به طور قابل توجهی بالا است. برای مقابله با این مشکل، مطالعه از مدل رگرسیون وایبول گسسته صفر آماسیده استفاده می‌کند که به عنوان یک چارچوب مناسب برای ارزیابی تأثیر متغیرهای توضیحی در تحلیل بقا عمل می‌کند. با این حال، چالش‌هایی مانند عدم ایستایی در رابطه بین متغیرها و پاسخ‌ها و ناهمگونی فضایی در مناطق جغرافیایی می‌تواند منجر به مدلی با پارامترهای بسیار زیاد شود. برای کاهش این مشکل، ما یک روش خوشه‌بندی فضایی را برای خلاصه‌سازی فضای پارامترها پیشنهاد می‌کنیم. این مقاله از روش‌های بیز نا پارامتری برای بررسی ناهمگونی فضایی ضرایب رگرسیون بهره می‌برد و بر فرآیند رستوران چینی موزون جغرافیایی برای خوشه‌بندی پارامترهای مدل رگرسیون وایبول گسسته صفر آماسیده تمرکز می‌کند. از طریق مطالعات شبیه‌سازی نشان داده می‌شود که فرآیند رستوران چینی موزون جغرافیایی عملکرد بهتری نسبت به الگوریتم‌های خوشه‌بندی بدون نظارت مانند خوشه‌بندی K-میانگین و فرآیند رستوران چینی استاندارد دارد، به‌ویژه در سناریوهایی با اندازه‌های خوشه نامتعادل از دقت و کارایی محاسباتی بالاتری برخوردار است. شاخص‌های رند بالاتر و میانگین مربعات خطای پایین‌تر بر کارایی و دقت مدل پیشنهادی ما تأکید می‌کنند. اعمال این روش بر روی داده‌های زمان تا وقوع پیشامد طلاق، با انباشتگی در پنج سال اول زندگی زنانوشویی در ایران، منجر به یک خوشه‌بندی بهینه و برآوردهای کارآمد از ضرایب رگرسیونی متغیرهای تبیینی دخیل در این امر با تغییرپذیری فضایی شده است.

واژه‌های کلیدی: تحلیل بقا؛ ضریب متغیر فضایی؛ خوشه‌بندی فضایی

Modeling Some Repeated Randomized Responses

M. Tarhani, M. R. Zadkarami*, S. M. R. Alavi

Department of Statistics, Faculty of Mathematical Sciences and Computer, Shahid Chamran University of Ahvaz, Ahvaz,
Islamic Republic of Iran

* Email: zadkarami@yahoo.co.uk

مدل سازی چند روش پاسخ تصادفیده مکرر

مهتاب طرهانی، محمدرضا زادکرمی*، سید محمدرضا علوی

گروه آمار دانشگاه شهید چمران اهواز، اهواز، ایران

چکیده

در برخی تحقیقات اجتماعی محقق با موضوعاتی مواجه است که به دلیل حساسیت بالا، پاسخهای قابل اعتمادی به دست نمی‌آید. روشهای پاسخ تصادفیده برای افزایش سطح محرمانگی و ارائه پاسخ صادقانه به کار می‌روند. با این وجود برآوردهای به دست آمده از این روش واریانس بزرگتری دارند. تکرار پاسخ تصادفیده به ازای هر عضو نمونه باعث افزایش تعداد مشاهدات شده و استفاده از میانگین مشاهدات هر شخص برای برآوردیابی واریانس برآورد را کاهش می‌دهد و مقادیر برآورد را به واقعیت نزدیکتر می‌نماید. در این مقاله با استفاده از پاسخهای تصادفیده‌ی جمعی پیوسته مکرر، میانگین پاسخهای هر شخص در نظر گرفته شده است و در قالب یک مدل خطی به برآورد میانگین متغیر حساس و پارامترهای رگرسیون پیرداخته شده است. برای این منظور داده‌های درآمد خانوار جمع آوری شده از دانشجویان به روش تصادفیده‌ی جمعی مکرر بررسی شده است که از هر شخص تقاضا شده پاسخ خود را پنج بار تصادفیده کند. نتایج به دست آمده با دو روش، یک بار برای مشاهده اول و بار دیگر با در نظر گرفتن میانگین مشاهدات هر شخص در قالب مدل رگرسیون به دست آمده است. نتایج نشان می‌دهد برآوردهای به دست آمده از روش دوم به واقعیت نزدیکتر بوده و واریانس کمتری دارند.

واژه های کلیدی: پاسخ تصادفیده؛ پاسخ تصادفیده مکرر؛ مدل رگرسیون خطی؛ متغیر پیوسته حساس؛ تکرار مشاهدات

Bounds for the Varentropy of Basic Discrete Distributions and Characterization of Some Discrete Distributions

F. Goodarzi*

Department of Statistics, Faculty of Mathematical Sciences, University of Kashan, Kashan, Islamic Republic of Iran.

* Email: f-goodarzi@kashanu.ac.ir

کران‌هایی برای ورن‌آنترپی توزیع‌های گسسته پایه و مشخص‌سازی برخی توزیع‌های گسسته

فرانک گودرزی*

گروه آمار، دانشکده علوم ریاضی، دانشگاه کاشان، کاشان، جمهوری اسلامی ایران

چکیده

با توجه به اهمیت ورن‌آنترپی در نظریه اطلاع و از آنجا که شکل بسته‌ای برای ورن‌آنترپی برخی توزیع‌های گسسته نمی‌توان یافت، هدف ما تعیین کران‌هایی برای ورن‌آنترپی این توزیع‌ها و معرفی ورن‌آنترپی گذشته برای متغیرهای تصادفی گسسته است. در این مقاله، ابتدا کران‌های بالا و پایینی را برای ورن‌آنترپی توزیع‌های پواسون، دوجمله‌ای، دوجمله‌ای منفی و فوق هندسی به دست آورده‌ایم. با توجه به آن‌که کران‌های بالای حاصل به صورت امید ریاضی توانهای دوم لگاریتمی بیان می‌شوند، یک عبارت معادل با استفاده از ضرایب تفاضل لگاریتمی ارائه می‌کنیم. به همین ترتیب، کران‌های پایینی را نیز بر حسب ضرایب تفاضل لگاریتمی ارائه می‌دهیم. علاوه بر این، کران بالایی برای واریانس تابعی از تابع باقیمانده عمر معکوس گسسته به دست می‌آید. همچنین، نامساوی‌هایی شامل گشتاورهای توابعی منتخب از طریق نرخ شکست معکوس را بررسی کرده و برخی توزیع‌های گسسته را با استفاده از نامساوی کُشی شوارتز مشخص‌سازی می‌کنیم.

واژه‌های کلیدی: ورن‌آنترپی؛ نرخ شکست معکوس؛ تبدیل دو جمله‌ای؛ نامساوی کُشی شوارتز

Evaluating Feature Selection Methods for Macro-Economic Forecasting, applied for Iran's macro-economic variables

M. Goldani*

Department of Political Science and Economics, Faculty of Literature and Humanities, Hakim Sabzevari University, Sabzevar, Islamic Republic of Iran

* Email: m.goldani@hsu.ac.ir

ارزیابی روش‌های انتخاب ویژگی برای پیش‌بینی کلان‌اقتصادی: کاربرد برای متغیرهای کلان‌اقتصادی ایران

مهدی گلدانی*

گروه اقتصاد و علوم سیاسی، دانشکده علوم انسانی، دانشگاه حکیم سبزه‌واری، سبزوار، جمهوری اسلامی ایران

چکیده

این مطالعه به ارزیابی روش‌های مختلف انتخاب ویژگی برای بهبود دقت پیش‌بینی مدل‌های کلان‌اقتصادی می‌پردازد، با تمرکز بر شاخص‌های کلان‌اقتصادی ایران که از داده‌های بانک جهانی استخراج شده‌اند. یک مقایسه جامع با استفاده از ۱۴ تکنیک انتخاب ویژگی انجام شد که در دسته‌های روش‌های فیلتری، روش‌های مبتنی بر پوسته، روش‌های تعبیه‌شده، و روش‌های مبتنی بر شباهت طبقه‌بندی شده‌اند. چارچوب ارزیابی شامل معیارهای میانگین مربعات خطا و میانگین مطلق خطا در یک فرآیند اعتبارسنجی متقاطع ۱۰-بخشی بود. یافته‌های کلیدی نشان داد که روش انتخاب گام‌به‌گام، روش‌های مبتنی بر درخت و تکنیک‌های مبتنی بر شباهت، به‌ویژه فاصله‌های هاسدورف و اقلیدسی، به‌طور مداوم دقت پیش‌بینی بالاتری از خود نشان دادند. میانگین مقادیر میانگین مطلق خطا برای روش انتخاب گام‌به‌گام ۳۲۰۳ و برای فاصله هاسدورف ۶۲۰۶۹ بود. در مقابل، روش‌های حذف بازگشتی ویژگی و آستانه‌گذاری واریانس عملکرد نسبتاً ضعیفی با مقادیر میانگین مطلق خطا به مراتب بالاتر داشتند. روش‌های مبتنی بر شباهت، میانگین رتبه ۹۰۱۲۵ را در بین مجموعه داده‌ها کسب کردند که نشان‌دهنده استحکام آن‌ها در مواجهه با داده‌های کلان‌اقتصادی با ابعاد بالا است. این نتایج بر پتانسیل ادغام معیارهای شباهت با روش‌های سنتی انتخاب ویژگی برای بهبود دقت و کارایی مدل‌های پیش‌بینی تأکید دارد. این مطالعه بینش‌های ارزشمندی برای پژوهشگران و سیاست‌گذاران ارائه می‌دهد که به دنبال توسعه ابزارهای پیش‌بینی اقتصادی قابل اعتماد هستند.

واژه‌های کلیدی: انتخاب ویژگی؛ دقت پیش‌بینی؛ شاخص‌های بانک جهانی؛ تحلیل کلان‌اقتصادی؛ روش‌های شباهت

Comparison of Adaptive Neural-Based Fuzzy Inference System and Support Vector Machine Methods for the Jakarta Composite Index Forecasting

A. Mutmainnah*, S. A. Thamrin, G. M. Tinungki

Department of Statistics, Faculty of Mathematics and Natural Sciences, Hasanuddin University, Makassar, 90245 Indonesia
Makassar, Indonesia

* Email: amutmainnah8@gmail.com

مقایسه سیستم استنتاج فازی مبتنی بر عصبی تطبیقی و روش‌های ماشین بردار پشتیبان برای پیش‌بینی شاخص ترکیبی جاکارتا

آیو موتمیننه*، سری آستوتی تامرین، جورجینا ماریا تینونگی

گروه آمار، دانشکده ریاضیات و علوم طبیعی، دانشگاه حسن الدین، ماکاسار، ۹۰۲۴۵ اندونزی ماکاسار، اندونزی

چکیده

شاخص ترکیبی جاکارتا (JCI) یک معیار اساسی برای ارزیابی عملکرد تمام سهام‌های فهرست شده در بورس اوراق بهادار اندونزی (IDX) است. با توجه به پیچیدگی ذاتی، غیرخطی و غیرثابت بودن داده‌های بازار سهام، انتخاب روش‌های پیش‌بینی قوی ضروری است. این مطالعه عملکرد سیستم استنتاج عصبی فازی تطبیقی (ANFIS) و ماشین بردار پشتیبان (SVM) را در پیش‌بینی حرکات JCI مقایسه می‌کند. محقق دقت پیش‌بینی را با استفاده از ریشه میانگین مربعات خطا (RMSE) و میانگین درصد مطلق خطا (MAPE) ارزیابی کرد. مرحله آموزش نشان داد که مدل ANFIS بهینه از تابع عضویت زنگ تعمیم‌یافته استفاده می‌کند و از جایگزین‌های دوزنقه‌ای و گاوسی بهتر عمل می‌کند. به طور همزمان، بهترین پیکربندی SVM از یک هسته خطی (هزینه = ۱۰) استفاده می‌کند که عملکرد برتر را در مقایسه با تابع پایه شعاعی (RBF) و هسته‌های سیگموئید نشان می‌دهد. در مرحله آزمایش، ANFIS به $RMSE=39.894$ و $MAPE=0.4647$ دست یافت، در حالی که SVM $RMSE=38.728$ و $MAPE=0.4516$ را ثبت کرد. این نتایج بر قابلیت‌های پیش‌بینی برتر SVM تاکید می‌کند و آن را به عنوان ابزاری قابل اعتماد برای پیش‌بینی بازار سهام قرار می‌دهد. یافته‌های این مطالعه بینش‌های ارزشمندی را برای سرمایه‌گذاران و سیاست‌گذاران در جهت‌یابی عدم قطعیت‌های بازار و بهینه‌سازی استراتژی‌های سرمایه‌گذاری فراهم می‌کند.

واژه‌های کلیدی: پیش‌بینی؛ ماشین بردار پشتیبان؛ شاخص ترکیبی جاکارتا؛ سیستم استنتاج فازی مبتنی بر عصبی تطبیقی

Second-ordered Characterization of Generalized Convex Functions and Their Applications in Optimization Problems

M. T. Nadi^{1*}, J. Zafarani²

¹ School of Mathematics, Institute for Research in Fundamental Sciences (IPM), P. O. Box 19395-5746, Tehran, Islamic Republic of Iran

² Department of Mathematics, Sheikhabaee University and University of Isfahan, Isfahan, Islamic Republic of Iran

* Email: mt_nadi@yahoo.com

مشخص سازی مرتبه دوم توابع محدب تعمیم یافته و کاربردهای آنها در مسائل بهینه سازی

محمد تقی نادی^{۱*}، جعفر زعفرانی^۲

^۱ پژوهشکده ریاضیات، پژوهشگاه دانشهای بنیادی (IPM)، تهران، جمهوری اسلامی ایران
^۲ گروه ریاضی، دانشگاه شیخ بهایی و دانشگاه اصفهان، اصفهان، جمهوری اسلامی ایران

چکیده

این مطالعه، برخی از توسعه های مشخص سازی مرتبه دوم توابع محدب تعمیم یافته را با استفاده از هم مشتق نگاشت زیر مشتق بررسی می کند. بطور دقیقتر، مشخص سازی مرتبه دوم توابع شبه محدب، محدب نما و محدب پایا ارائه می شوند. علاوه بر این بعضی از کاربردهای زیر مشتق های مرتبه دوم در مسائل بهینه سازی، مانند برنامه ریزی غیر خطی مقید و نامقید را ارائه می دهد.

واژه های کلیدی: زیر مشتق مرتبه دوم؛ ویژگی نیم معین مثبت؛ زیر مشتق مرتبه دوم منظم؛ شرایط بهینه مرتبه دوم