



Semnan University

Journal of Econometric Modelling

Journal homepage: <https://jem.semnan.ac.ir/?lang=en>



Research Article

Analysis of economic inequalities using bivariate total time on test transform and copula function

Mojtaba Esfahani (Corresponding Author)

Assistant Professor in statistics, Department of Statistics,
Velayat University of Iranshahr

m.esfahani@velayat.ac.ir

Gholam Reza Mohtashami Borzadaran

Professor in Statistics, Department of Statistics, Ferdowsi University of Mashhad

grmohtashami@um.ac.ir

Mohammad Amini

Professor in Statistics, Department of Statistics, Ferdowsi University of Mashhad

m-amini@um.ac.ir

PAPER INFO

Paper history:

Received: 25. 01. 2025

Revised: 07. 04. 2025

Accepted: 20. 03. 2025

JEL Classification:

D63, D31, D30, C16

Keywords:

Bivariate total time on test
transform,
Income inequality,
Copula function,
TTT Plot,
Lorenz curve

ABSTRACT

This study investigates the relationship between two key economic indicators: income and wealth. Despite apparent similarities in the behavior of these two variables, numerous instances highlight their differences and lack of complete correlation. The main objective of this paper is to analyze the dependency structure between individuals' income and wealth using advanced statistical tools. To model the dependency between income and wealth, the copula function—an innovative tool in probability theory—has been employed. In this regard, various copula functions are examined, and the Clayton copula family along with the Farlie-Gumbel-Morgenstern (FGM) copula family are identified as the most suitable choices for the studied data. Additionally, the study introduces a new bivariate index based on the Total Time on Test (TTT) transform in the bivariate case, utilizing copula functions. This index is applied to real-world data on the income and wealth of Iranian households, and the results demonstrate a significant relationship between the two variables at the societal level. Therefore, it can be concluded that using this index in economic data reveals that the proposed bivariate TTT index can effectively represent the dependency structure between income and wealth. Furthermore, the use of Clayton and FGM copulas also shows a good fit with the empirical data

© 2023 Published by Semnan University Press. All rights reserved.

تحلیل نابرابری‌های اقتصادی با استفاده از تبدیل زمان کل آزمون

دومتغیره و تابع مفصل

مجتبی اصفهانی (نویسنده مسئول)

استادیار گروه آمار، دانشکده علوم پایه، دانشگاه ولایت ایرانشهر، ایران

m.esfahani@velayat.ac.ir

غلامرضا محتشمی برزادران

استاد گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، ایران

grmohtashami@um.ac.ir

محمد امینی

استاد گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، ایران

m-amini@um.ac.ir

نوع مقاله: علمی - پژوهشی تاریخ دریافت: ۱۴۰۳/۱۱/۰۶ تاریخ پذیرش: ۱۴۰۴/۰۲/۳۰

چکیده

در این پژوهش، به بررسی ارتباط بین دو شاخص مهم اقتصادی یعنی درآمد و ثروت پرداخته شده است. با وجود شباهت‌های ظاهری در رفتار این دو متغیر، موارد متعددی وجود دارد که نشان‌دهنده تفاوت‌ها و عدم همبستگی کامل بین آنهاست. هدف اصلی مقاله، تحلیل ساختار وابستگی بین درآمد و ثروت افراد جامعه با بهره‌گیری از ابزارهای آماری پیشرفته است. برای مدل‌سازی وابستگی بین درآمد و ثروت، از تابع مفصل به عنوان یکی از ابزارهای نوین در نظریه احتمال استفاده شده است. در این راستا، ابتدا توابع مفصل مختلف مورد بررسی قرار گرفته و دو خانواده‌ی مفصل کلیتون و خانواده مفصل‌های فورلی-گامبل-مورگنستن (FGM) به عنوان مناسب‌ترین گزینه‌ها برای داده‌های مورد مطالعه شناسایی شده‌اند. همچنین، در ادامه پژوهش یک شاخص دو متغیره جدید بر پایه‌ی تبدیل زمان کل آزمون (TTT) در حالت دومتغیره و با استفاده از توابع مفصل معرفی گردیده است. این شاخص بر روی داده‌های واقعی مربوط به درآمد و ثروت خانوارهای ایرانی اعمال و نتایج حاصل، ارتباط معناداری را بین این دو متغیر در سطح جامعه نشان می‌دهد. بنابراین می‌توان نتیجه گرفت که استفاده از این شاخص در داده‌های اقتصادی نشان می‌دهد که شاخص پیشنهادی زمان کل آزمون دومتغیره می‌تواند ساختار وابستگی بین داده‌های درآمد و ثروت را به‌صورت مؤثری نمایش دهد. همچنین استفاده از توابع مفصل کلیتون و FGM نیز تطابق مناسب با داده‌های تجربی را نشان می‌دهد.

طبقه‌بندی JEL: D63، D31، D30، C16

کلید واژه‌ها: تبدیل زمان کل آزمون دومتغیره، نابرابری درآمد، تابع مفصل، نمودار TTT، منحنی لورنتس

۱. مقدمه و پیشینه تحقیق

نابرابری درآمد یکی از مسائل مهم اقتصادی و اجتماعی است که همواره توجه محققان و سیاست‌گذاران را به خود جلب کرده است. بررسی نحوه توزیع درآمد و سنجش نابرابری در یک جامعه، نقش مهمی در فهم بهتر مشکلات اقتصادی و ارائه راه‌حل‌های مناسب دارد. برای این منظور، تاکنون شاخص‌های زیادی معرفی شده و ویژگی‌های آنها مورد بررسی قرار گرفته‌است. از مهم‌ترین این شاخص‌ها می‌توان به منحنی لورنتس که توسط (لورنتس^۱، ۱۹۰۵) و ضریب جینی که توسط (جینی^۲، ۱۹۱۲) معرفی گردید، اشاره کرد. این ابزارها به تحلیل نابرابری درآمد کمک می‌کنند، اما دارای محدودیت‌هایی نیز هستند. با توجه به پیچیدگی‌های توزیع درآمد و وابستگی‌های موجود بین متغیرهای اقتصادی، نیاز به ابزارهای جدید و دقیق‌تر برای تحلیل نابرابری احساس می‌شود. علاوه بر منحنی لورنتس و ضریب جینی، شاخص‌های دیگری نیز برای سنجش نابرابری معرفی شده‌اند که هر کدام ویژگی‌ها و مزایای خاص خود را دارند. شاخص (بن‌فرونی^۳، ۱۹۳۰)، شاخص (زنگا^۴، ۱۹۸۶-۲۰۰۷)، شاخص لیمکوهلر^۵ و ضرایب (پترا^۶، ۱۹۱۵) و (تایل^۷، ۱۹۶۷) از جمله این ابزارها هستند. شاخص بن‌فرونی بر اساس نسبت درصد مشخصی از درآمد کل که در دست درصد مشخصی از جمعیت قرار دارد، محاسبه می‌شود و به تحلیل نابرابری در بخش‌های خاصی از جامعه می‌پردازد. شاخص زنگا نیز به عنوان اصلاحیه‌ای بر ضریب جینی، حساسیت بیشتری نسبت به تغییرات در بخش‌های مختلف توزیع درآمد دارد. شاخص لیمکوهلر نیز بر تحلیل تراکم تجمعی درآمد تمرکز دارد و نابرابری را از زوایای مختلفی بررسی می‌کند. ضرایب پترا و تایل نیز هر کدام با روش‌های خاص خود، نابرابری را اندازه‌گیری می‌کنند و ابزارهای مفیدی برای تحلیل اقتصادی به شمار می‌روند. با این حال، هر کدام از این شاخص‌ها نیز دارای محدودیت‌هایی هستند و نمی‌توانند به طور کامل پیچیدگی‌های نابرابری درآمد را

1. Lorenz

2. Gini

3. Bonferroni

4. Zenga

5. Leimkuhler

6. Pietra

7. Theil

نمایش دهند. برای تحلیل دقیق‌تر و کامل‌تر نابرابری، نیاز به ابزارهای جدیدی است که بتوانند وابستگی‌های بین متغیرهای اقتصادی را نیز در نظر بگیرند. یکی از این ابزارهای جدید، شاخص زمان کل آزمون^۱ (TTT) است که توسط (اصفهانی و همکاران، ۱۳۹۹) معرفی شده است. این شاخص بر اساس مفهوم زمان کل آزمون طراحی شده است. مزیت این شاخص نسبت به شاخص‌های سنتی، دقت بالای آن در نشان دادن نابرابری درآمد جوامع نزدیک به هم است. شاخص زمان کل آزمون قادر است با تحلیل وابستگی‌های پیچیده بین متغیرهای اقتصادی، تصویری دقیق‌تر از نابرابری ارائه دهد. برای مطالعه بیشتر درباره تبدیل TTT و ارتباط آن با مفاهیم اقتصادی می‌توان به (فام و تارکان^۲، ۱۹۹۴)، (پرز^۳ و همکاران، ۱۹۹۷)، (کیاج^۴، ۱۹۹۹) و (کاوزاک^۵ و همکاران، ۲۰۰۹) مراجعه کرد. از مطالعات دیگر در زمینه نابرابری‌های اقتصادی می‌توان به میرزایی و همکاران (۱۳۹۷)، (ابونوری و اسنانودی، ۱۳۸۴)، (ابونوری و همکاران، ۲۰۱۰) و (ابونوری و همکاران، ۲۰۱۳) اشاره کرد که با استفاده از ریزداده‌های ایران تحلیلی کامل برای برآورد و ارزیابی سازگاری شاخص‌های نابرابری اقتصادی ارائه دادند. همچنین مقایسه انواع روش‌های محاسبه ضریب جینی نیز توسط (ابونوری و مک کلوگان^۶، ۲۰۰۰) صورت گرفته است. (ابونوری، ۲۰۰۳) نیز یک روش جدید برای برآورد ضریب جینی ارائه کرده و مقایسه شاخص‌های نابرابری درآمد توسط (ابونوری و ذوقی، ۱۳۹۲) و ارتباط بین ساختار تولید و توزیع درآمد توسط (ابونوری و فراهتی، ۱۳۹۴) بیان شده است. تمام شاخص‌های معرفی شده برای سنجش نابرابری اقتصادی ابزارهای یک متغیره هستند یعنی متغیرها را به صورت جدا تحلیل می‌کنند و ساختار وابستگی بین آنها را نادیده می‌گیرند. برای این منظور تلاش‌هایی در جهت تعمیم منحنی‌های اقتصادی و شاخص‌های سنجش نابرابری اقتصادی به حالت چند متغیره انجام گرفت که می‌توان به (تاگوچی^۷، ۱۹۷۲ a,b)، (آرنولد^۸، ۲۰۱۲) و (آرنولد و سارابیا^۹، ۲۰۱۸) اشاره کرد.

1. Total Time on Test

2. Pham and Turkkan

3. Perez

4. Kiagh

5. Kawczak

6. McCloughan

7. Taguchi

8. Arnold

9. Sarabia

هدف اصلی این پژوهش معرفی یک تعمیم جدید از تبدیل زمان کل آزمون در حالت دومتغیره با استفاده از تابع مفصل است. این شاخص با استفاده از تابع مفصل^۱، می‌تواند به تحلیل وابستگی بین متغیرها و ارتباط بین آنها بپردازد. برای تجزیه و تحلیل این شاخص پیشنهادی، از ریز داده‌های درآمد و ثروت خانوارهای ایرانی در بازه زمانی سال‌های ۱۴۰۲-۱۴۰۰ با در نظر گرفتن چند تابع مفصل معروف استفاده شده است. در این پژوهش، ابتدا در بخش ۲ به معرفی شاخص‌های مختلف سنجش نابرابری درآمد و ویژگی‌های آنها و همچنین شاخص زمان کل آزمون به عنوان یک ابزار جدید و دقیق‌تر برای تحلیل نابرابری و ویژگی‌های آنها پرداخته شده است. سپس تعریف تابع مفصل به عنوان ابزاری برای مطالعه ارتباط بین متغیرهای مختلف معرفی و ویژگی‌های آن در ادامه این بخش مورد توجه قرار گرفته است. در ادامه در بخش ۳ با استفاده از مفهوم تابع مفصل و شاخص زمان کل آزمون یک شاخص جدید دومتغیره برای مطالعه و بررسی ارتباط بین دو متغیر درآمد و ثروت افراد معرفی و ویژگی‌های آن بیان می‌شود. بخش پایانی مقاله به مطالعه کاربردی و تجزیه و تحلیل داده‌های واقعی درآمد و ثروت خانواده‌های ایرانی پرداخته شده است.

۲. مبانی نظری تحقیق

در این بخش ابتدا شاخص‌های سنجش نابرابری درآمد را به طور مختصر مرور کرده و ارتباط آنها را با یکدیگر بیان می‌کنیم. سپس سایر مفاهیم و تعاریف مورد نیاز در مقاله را ارائه می‌دهیم.

۲-۱. شاخص‌های اقتصادی

یکی از شاخص‌های پرکاربرد در اقتصاد شاخص لورنتس است که توسط ماکس اوتو لورنتس در سال ۱۹۰۵ به صورت یک منحنی کمان شکل (محدب) که در زیر قطر مربع واحد قرار دارد، معرفی شد و در نابرابری‌های اقتصادی کاربرد فراوانی دارد به طوری که هر چه کمان بیشتر خم شود میزان نابرابری اقتصادی بیشتر است.

^۱. Copula Function

تعریف ۱-۲. فرض کنید درآمد افراد در جامعه یک متغیر تصادفی نامنفی با توزیع $F(x)$ و میانگین μ باشد. در این صورت

$$L(p) = \frac{1}{\mu} \int_0^p Q(u) du, \quad 0 \leq p \leq 1, \quad Q(u) = F^{-1}(u), \quad (1)$$

که در آن $L(p)$ را منحنی لورنتس توزیع $F(x)$ نامیده و $F^{-1}(u)$ نیز تابع معکوس تابع توزیع $F(x)$ است که تابع چندک نامیده می شود و μ نیز امید ریاضی متغیر تصادفی X است.

ضرایب جینی و پترا نیز که از شاخص های آماری برای سنجش میزان نابرابری توزیع درآمد هستند به ترتیب در (جینی، ۱۹۱۲) و (پیترا، ۱۹۱۵) معرفی و مورد بررسی مطالعه قرار گرفته اند. این دو ضریب که با استفاده از تبدیل لورنتس به دست می آیند به ترتیب عبارتند از:

$$G = \int_0^1 [u - L(u)] du = 1 - 2 \int_0^1 L(u) du, \quad (2)$$

$$P = \frac{1}{\mu} \int_0^{F(\mu)} (\mu - Q(p)) dp = F(\mu) - L(F(\mu)). \quad (3)$$

(بن فرونی، ۱۹۳۰) نیز برای بررسی نابرابری درآمد به کمک منحنی لورنتس معیاری را با عنوان منحنی بن فرونی معرفی نمود که با نماد $B(p)$ نشان داده می شود.

$$B(p) = \frac{1}{u\mu} \int_0^p F^{-1}(u) du = \frac{L(p)}{p}, \quad 0 \leq p \leq 1. \quad (4)$$

ضریب بن فرونی نیز به صورت $B = 1 - \int_0^1 B(p) dp$ بیان می شود و همواره $B \leq \frac{1+G}{2}$. منحنی بن فرونی معرفی شده در (۴) در مقایسه با منحنی لورنتس حساسیت بیشتری نسبت به درآمدهای پایین دارد و در بقیه ویژگی ها کاملاً مشابه منحنی لورنتس است. واضح است که همواره $L(p) \leq B(p)$. (ابونوری و مک کلوگان، ۲۰۰۰) و همچنین (ابونوری، ۲۰۰۳) نیز مطالعات جالبی در زمینه شاخص های نابرابری اقتصادی ارائه داده اند. در ادامه این بخش شاخص سنجش نابرابری اقتصادی زمان کل آزمون و ویژگی های آن بیان می شود.

۲-۲. شاخص نابرابری زمان کل آزمون

در این بخش ابتدا تبدیل TTT بر اساس (کوچار^۱ و همکاران، ۲۰۰۲) معرفی و سپس نشان داده می شود که می توان از نمودار این تبدیل برای بررسی نابرابری درآمد استفاده نمود. (اصفهانی و

¹. Kochar

همکاران، ۱۳۹۹) با بررسی ویژگی‌های این تبدیل نشان دادند که می‌توان از آن به عنوان یک شاخص سنجش نابرابری درآمد استفاده کرد. برای جزئیات بیشتر به (اصفهانی و همکاران، ۱۳۹۹) مراجعه شود.

تعریف ۲-۲. فرض کنید $F(x)$ یک توزیع درآمد با میانگین متناهی μ باشد. در این صورت تبدیل TTT مرتبط با توزیع $F(x)$ را با H_F^{-1} نشان داده و به صورت زیر بیان می‌شود:

$$H_F^{-1}(p) = \int_0^{F^{-1}(p)} \bar{F}(s) ds, \quad 0 \leq p \leq 1. \quad (5)$$

همچنین تبدیل TTT مقیاسی^۱ که در (بارلو و کامپو^۲، ۱۹۷۵) با نماد $\varphi(p)$ نشان داده شده است، به صورت زیر بیان می‌شود،

$$\varphi(p) = \frac{H_F^{-1}(p)}{\mu}, \quad 0 \leq p \leq 1. \quad (6)$$

که یک تابع صعودی و پیوسته در بازه $[0, 1]$ است و $\phi(0) = 0$ و $\phi(1) = 1$. برای درک بهتر مفهوم زمان کل آزمون و ارتباط آن با مفاهیم درآمد کافی است فرض کنید $x_1, x_2, \dots, x_n$ داده‌های درآمد حاصل از یک نمونه تصادفی با حجم n باشد. در این صورت $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$ داده‌های مرتب شده این نمونه می‌باشند. اگر هر کدام از $x_{(i)}$ را در نظر بگیریم بدیهی است که درآمد $n - i$ تا از اعضای نمونه از آن بیشتر و درآمد $i - 1$ تا از اعضای نمونه کمتر از آن خواهد بود. در این صورت آماره‌ی زمان کل آزمون در لحظه i ام یعنی $x_{(i)}$ نشان‌دهنده مجموع درآمد افراد تا رسیدن به مقدار $x_{(i)}$ است و تبدیل TTT در نقطه p میانگین درآمد افراد مورد بررسی، تا هنگامی که درآمد p درصد از کل افراد بررسی شده باشد، است. اگر $X_{(r)}$ آماره مرتب شده مرتبه r باشد آنگاه آماره زمان کل آزمون در بازه $(0, t)$ با نماد $T(t)$ نشان داده می‌شود که در آن $X_{(r)} < t \leq X_{(r+1)}$.

$$T(t) = nX_{(1)} + (n-1)(X_{(2)} - X_{(1)}) + \dots + (n-r+1)(X_{(r)} - X_{(r-1)}) + (n-r)(t - X_{(r)}), \quad (7)$$

^۱. Scale

^۲. Barlow and Campo

۲-۳. نمودار زمان کل آزمون به عنوان شاخصی برای سنجش نابرابری اقتصادی

نمودار زمان کل آزمون شامل یک منحنی صعودی است که از گوشه پایین سمت چپ شروع می‌شود و انتهای آن در گوشه بالا سمت راست مربع واحد قرار دارد. اگر تابع توزیع $F(x)$ به داده‌ها برازش داده شود آنگاه با افزایش حجم نمونه، نمودار TTT داده‌ها به نمودار تبدیل $\phi_X(p)$ مرتبط با تابع توزیع $F(x)$ همگرا می‌شود. یکی از ویژگی‌های مهم نمودار زمان کل آزمون تفسیر شهودی آن است و مانند منحنی لورنتس یک منحنی محدب است، با این تفاوت که این منحنی علاوه بر محدب بودن می‌تواند مقعر نیز باشد. نحوه رسم نمودار زمان کل آزمون در (برگمن و کلفسجو^۱، ۲۰۱۴) با جزئیات بیان شده است.

ابونوری و اسناوندی (۱۳۸۴) ویژگی‌های مورد نیاز یک تبدیل را برای اینکه به عنوان شاخصی برای سنجش نابرابری اقتصادی در نظر گرفته شود مورد مطالعه و ارائه قرار دادند. (اصفهانی و همکاران، ۱۳۹۹) با بررسی این ویژگی نتیجه گرفتند که نمودار زمان کل آزمون نیز ویژگی اصلی یک شاخص نابرابری یعنی محدب بودن، حساسیت نسبت به افزودن یک مقدار ثابت و عدم حساسیت نسبت به تغییر متناسب در درآمد افراد جامعه را دارد. بنابراین می‌توان از آن به عنوان معیاری برای سنجش نابرابری درآمد در جوامع مختلف استفاده کرد. همچنین با مقایسه این شاخص با سایر شاخص‌های شناخته شده در سنجش نابرابری درآمد ثابت کردند که شاخص زمان کل آزمون در مقایسه جوامع با نابرابری درآمد نزدیک به هم از دقت بیشتری نسبت به شاخص لورنتس برخوردار است.

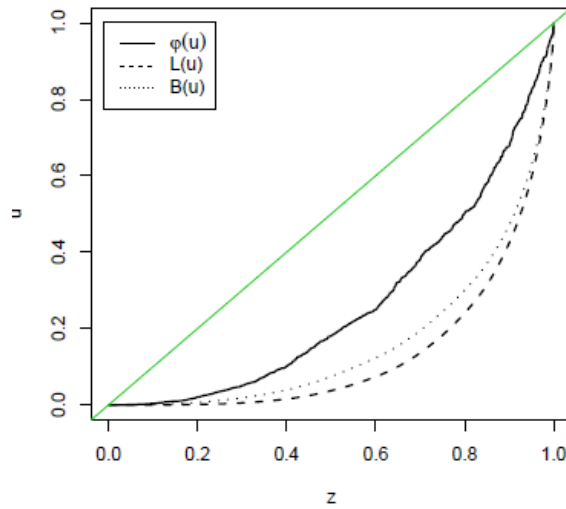
با توجه به تعریف تبدیل زمان کل آزمون در (۵) و همچنین روابط (۱) و (۴) به سادگی می‌توان نتیجه گرفت که همواره نمودار زمان کل آزمون از این دو نمودار بزرگتر است یعنی:

$$L(p) \leq B(p) \leq \phi(p). \quad (۸)$$

برقراری رابطه (۸) در شکل‌های ۱ و ۲ نشان داده شده است.

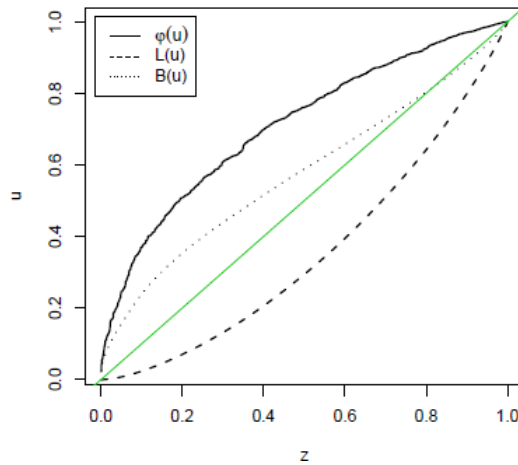
^۱. Bergman and Klefsjo

شکل (۱): نمودارهای زمان کل آزمون، لورنتس و بن فرونی برای داده‌های وایبول با پارامترهای ۰.۵ و ۵



منبع: محاسبات پژوهش با نرم افزار R

شکل (۲): نمودارهای زمان کل آزمون، لورنتس و بن فرونی برای داده‌های وایبول با پارامترهای ۲ و ۵



منبع: محاسبات پژوهش با نرم افزار R

۲-۴. تابع مفصل

مفهوم تابع مفصل برای اولین بار توسط (اسکلار^۱، ۱۹۵۹) معرفی شد و پس از آن محققان زیادی بر روی این مفهوم مطالعه و پژوهش انجام داده و نتایج مهمی را ارائه کرده‌اند. در این بخش ابتدا تعریف تابع توزیع توأم یک بردار از متغیرهای تصادفی و ویژگی‌های آن را بیان می‌شود. در ادامه تابع مفصل و ویژگی‌های آن را به اختصار مرور کرده و سپس ضابطه چند تابع مفصل معروف بیان می‌شود. همچنین قضیه اسکلار به عنوان یک ابزار مهم برای تولید تابع توزیع توأم با استفاده از تابع مفصل بیان می‌شود. برای جزئیات بیشتر به (موریللاس^۲، ۲۰۰۵) مراجعه شود.

تعریف ۲-۳. فرض کنید $X = (X_1, X_2)$ یک بردار تصادفی باشد. در این صورت تابع توزیع توأم این بردار تصادفی به صورت زیر بیان می‌شود.

$$F_{(X_1, X_2)}(x_1, x_2) = P(X_1 \leq x_1, X_2 \leq x_2)$$

برای همه x_1, x_2 که متعلق به اشتراک تکیه گاه متغیرهای تصادفی X_1 و X_2 هستند. یادآور می‌شود که دو متغیر تصادفی X_1 و X_2 از یکدیگر مستقل هستند اگر و تنها اگر

$$F_{(X_1, X_2)}(x_1, x_2) = F_{X_1}(x_1)F_{X_2}(x_2).$$

تعریف ۲-۴. تابع $C: [0,1]^2 \rightarrow [0,1]$ یک تابع مفصل است اگر برای هر

$u, v \in [0,1]$ در شرایط زیر صدق کند:

$$C(u, 0) = C(0, v) = 0 \quad (\text{الف})$$

$$C(u, 1) = u, \quad C(1, v) = v \quad (\text{ب})$$

ج) دارای خاصیت دوصعودی^۳ باشد یعنی برای $u_1 \leq u_2$ و $v_1 \leq v_2$ که $u_1, u_2, v_1, v_2 \in [0,1]$ داشته باشیم:

$$C(u_2, u_2) - C(u_1, v_2) - C(u_2, v_1) + C(u_1, v_1) \geq 0. \quad (۹)$$

اگر $C(u, v)$ یک مفصل پیوسته باشد، رابطه (۹) معادل با این است که $\frac{\partial^2 C(u, v)}{\partial u \partial v} \geq 0$.

1. Sklar

2. Morillas

3. Two Increasing

قضیه ۲-۵. (اسکلار-۱۹۵۹). فرض کنید X و Y دو متغیر تصادفی با تابع توزیع توام $H(x, y)$ و توزیع های حاشیه ای پیوسته $F(x)$ و $G(y)$ باشند. در این صورت تابع مفصل $C(u, v)$ به طور یکتا وجود دارد به طوریکه برای هر x و y ،

$$H(x, y) = C(F(x), G(y)).$$

برای همه $(u, v) \in [0, 1]^2$ هر تابع مفصل $C(u, v)$ از بالا و پایین به وسیله دو تابع مفصل کران دار است. این کران ها که به کرانهای فرشه-هافدینگ^۱ معروف هستند به ترتیب با نماد $M(u, v)$ و $W(u, v)$ نشان داده می شوند و به صورت $M(u, v) = \min\{u, v\}$ و $W(u, v) = \max\{u + v - 1, 0\}$ تعریف می شوند. برای هر تابع مفصل $C(u, v)$ می توان نوشت:

$$W(u, v) \leq C(u, v) \leq M(u, v).$$

در این مقاله از مفصل های کلیتون و گامبل^۲ استفاده خواهد شد که به ترتیب در روابط (۱۰) و (۱۱) نشان داده شده اند.

$$C_{\theta}(u, v) = [\max\{u^{-\theta} + v^{-\theta} - 1, 0\}]^{-\frac{1}{\theta}}, \quad \theta \neq 0, \quad \theta > -1, \quad (10)$$

$$C_{\theta}(u, v) = uv \exp\{-\theta \ln(u) \ln(v)\}, \quad \theta \geq 1. \quad (11)$$

یکی دیگر از مفصل های مهم خانواده مفصل های فارلی-گامبل-مورگنستن^۳ است که به اختصار با نماد FGM شناخته می شود. این خانواده به دلیل ویژگی های جبری ساده همواره مورد توجه محققان بوده است. خانواده توابع مفصل FGM به صورت زیر تعریف می شود:

$$C_{\theta}(u, v) = uv [1 + \theta(1 - u)(1 - v)], \quad \theta \in [-1, 1]. \quad (12)$$

اگر دو متغیر X_1 و X_2 از یکدیگر مستقل باشند برای نشان دادن ساختار وابستگی بین آن ها می توان از مفصل $\pi(u, v) = uv$ که به مفصل حاصلضرب معروف است، استفاده کرد.

¹. Frechet- Hoeffding bound

². Clayton & Gumbel

³. Farlie-Gumbel-Morgenstern

برای نشان دادن ارتباط خطی بین دو متغیر می‌توان از ضریب پیرسون^۱ استفاده کرد. هنگامی که ارتباط بین متغیرها غیرخطی باشد می‌توان از ابزارهای دیگری مانند ضریب اسپیرمن^۲ و کندال^۳ استفاده برد. با توجه به اینکه در این مقاله به مطالعه وابستگی بین متغیرها پرداخته می‌شود از ضریب وابستگی کندال استفاده می‌شود. ضریب وابستگی کندال که بر اساس تابع مفصل وابستگی بین دو متغیر X و Y محاسبه می‌شود، با نماد τ نشان داده می‌شود و به‌صورت زیر تعریف می‌شود:

$$\tau(C) = 4 \int_0^1 \int_0^1 C(x, y) dC(x, y) - 1$$

نکته ۲-۶. یادآور می‌شود که مفصل‌های FGM برای نشان دادن وابستگی ضعیف، مفصل کلیتون برای وابستگی قوی مثبت و همچنین مفصل گامبل برای وابستگی قوی منفی مناسب‌تر است.

۲-۵. توزیع گومپرتز^۴

متغیر تصادفی نامنفی X دارای توزیع گومپرتز با پارامترهای شکل α و مقیاس β است اگر تابع توزیع تجمعی و تابع چگالی آن به‌ترتیب به‌صورت زیر باشد،

$$G_X(x) = 1 - e^{-\alpha(e^{\beta x} - 1)}, \quad \alpha, \beta > 0, \quad x \geq 0,$$

$$g_X(x) = \alpha\beta e^{\alpha + \beta x - \alpha e^{\beta x}}, \quad x \geq 0,$$

و با نماد $X \sim G(\alpha, \beta)$ نمایش داده می‌شود.

از ویژگی‌های توزیع گومپرتز می‌توان به اریب بودن، دارای دم سمت راست بلند اشاره کرد. همچنین این توزیع یک توزیع نرخ خطر است، یعنی احتمال وقوع یک رویداد در هر زمان معین تنها به زمان سپری شده از آخرین رویداد بستگی دارد. این ویژگی توزیع گومپرتز، برای مدل‌سازی داده‌های درآمد مناسب است.

1. Pearson

2. Spearman

3. Kendall

4. Gompertz

۳. تبدیل زمان کل آزمون دومتغیره

همانطور که در بخش ۱ اشاره شد، منحنی لورنتس و ضریب جینی از مهمترین ابزارهای سنجش و تحلیل نابرابری درآمد و ثروت در یک جامعه هستند. با این حال، این ابزارها و سایر شاخص‌های معرفی شده تک متغیره هستند، یعنی متغیرها را به صورت جداگانه تحلیل می‌کنند و ساختار وابستگی آنها را نادیده می‌گیرند. یعنی در مطالعه همزمان توزیع متغیرهای درآمد و ثروت، تفاوت توزیع افراد ثروتمند و پردرآمد را نمی‌توان در نابرابری کلی مشاهده کرد. بنابراین لازم است تا ابزارهایی برای تحلیل توأم این دو متغیر معرفی گردد.

برای بررسی ارتباط بین متغیرهای درآمد و ثروت افراد در یک جامعه تاکنون شاخص‌های زیادی معرفی شده و ویژگی‌های آنها مورد بررسی قرار گرفته‌است. از مهم‌ترین این شاخص‌ها می‌توان شاخص‌های لورنتس چند متغیره را نام برد که تعمیم‌های زیادی از آن نیز معرفی شده است با توجه به مزیت شاخص سنجش نابرابری درآمد زمان کل آزمون نسبت به دیگر شاخص‌ها در نشان دادن نابرابری درآمد جوامع نزدیک به هم به نظر می‌رسد در حالت دومتغیره و مطالعه ارتباط بین متغیرهای درآمد و ثروت افراد نیز این شاخص نتایج مفیدی را نشان دهد.

(گروث^۱ و همکاران، ۲۰۲۲) توسیع‌هایی از شاخص لورنتس و ضریب جینی را برای مطالعه نابرابری متغیرهای X_1 و X_2 به‌طور همزمان ارائه کرده‌اند. در این توسیع‌ها ساختار وابستگی این متغیرها نیز به عنوان یک شاخص مهم مورد مطالعه قرار گرفته است.

در ادامه این بخش با توجه به روش ارائه شده در گروث و همکاران، (۲۰۲۲) و با در نظر گرفتن تبدیل زمان کل آزمون به عنوان معیار پایه برای سنجش نابرابری اقتصادی با استفاده از دو متغیر نامنفی درآمد و ثروت، یک توسیع جدید برای تبدیل زمان کل آزمون در حالت دو متغیره ارائه کرده و ویژگی‌های آن را مورد مطالعه و بررسی قرار می‌دهیم. برای این منظور فرض کنید $X = (X_1, X_2)$ نشان دهنده بردار متغیرهای تصادفی درآمد و ثروت است که هر کدام از X_i ها دارای تابع توزیع حاشیه ای $F_i(x_i)$ هستند.

1. Grothe

(لی و شیکد^۱، ۲۰۰۴) برای هر متغیر تصادفی نامنفی X ، مقدار تبدیل زمان کل آزمون در نقطه $F(X)$ را با نماد X_{ttt} نشان داده و آن را متغیر تصادفی مشاهده زمان کل آزمون^۲ نامیدند. با توجه به رابطه (۵) مقدار مشاهده شده X_{ttt} که با نماد x_{ttt} نشان می‌دهیم، به صورت زیر به دست می‌آید،

$$x_{ttt} = H_F^{-1}(F(x)) = \int_0^x \bar{F}(s) ds,$$

و در نتیجه متغیر تصادفی مشاهده زمان کل آزمون به صورت $X_{ttt} = \int_0^X \bar{F}(s) ds$ نشان داده می‌شود. در واقع هنگامی که متغیر تصادفی X مشاهده می‌شود متغیر X_{ttt} مشاهده زمان کل آزمون متناظر با X است. همانطور که مشاهده می‌شود مقدار مشاهده زمان کل آزمون در نقطه $F(x)$ برابر با مقدار تبدیل زمان کل آزمون در نقطه x است. در کاربرد اگر متغیر تصادفی X نشان دهنده درآمد یک فرد در جامعه باشد، آنگاه X_{ttt} نسبتی از کل درآمد آن بخشی از جامعه است که درآمد آنها کوچکتر یا برابر با x است و مقداری بین صفر و یک است. با توجه به اینکه در این پژوهش بیشتر از یک متغیر را مورد مطالعه قرار می‌دهیم برای سادگی در نوشتن از نماد X_i^{ttt} برای نشان دادن مشاهده زمان کل آزمون متغیر i -ام استفاده می‌کنیم که در آن $i = 1, 2$. برای متغیر تصادفی X_{ttt} ویژگی‌های جالبی را می‌توان به دست آورد که به طور خلاصه در ادامه به آن می‌پردازیم.

لی و شیکد (۲۰۰۴) نشان دادند که $E[X_i^{ttt}] = \int_0^\infty \bar{F}_i(x) dx$. همچنین با توجه به نتایج ارائه شده در لی و شیکد (۲۰۰۴) می‌توان نتیجه گرفت که تابع توزیع متغیر تصادفی X_i^{ttt} معکوس تبدیل زمان کل آزمون است. با توجه به اینکه در این پژوهش با برداری از دو متغیر تصادفی روبرو هستیم این نکته را در گزاره ۱.۳ به صورت کلی برای $i = 1, 2$ نشان داده‌ایم.

گزاره ۱.۳: فرض کنید $X = (X_1, X_2)$ یک بردار تصادفی با تابع توزیع توام $F_{(X_1, X_2)}(x_1, x_2)$ و توابع توزیع حاشیه‌ای $F_i(x_i)$ باشند و μ_i نشان دهنده امیدریاضی متغیر تصادفی X_i باشد، در این صورت تابع توزیع متغیر X_i^{ttt} معکوس تبدیل زمان کل آزمون یعنی $H_{F_i}(y_i)$ است.

اثبات: با استفاده از تعریف تابع توزیع برای متغیر تصادفی X_i^{ttt} و برای $p_1, p_2 \in [0, 1]$

^۱. Li and Shaked

^۲. Observed Total Time on Test

$$\begin{aligned}
 F_{X_i^{ttt}}(\mu_i p_i) &= P(X_i^{ttt} \leq \mu_i p_i) \\
 &= P(H_i^{-1}(F_i(X_i)) \leq \mu_i p_i) \\
 &= P(F_i(X_i) \leq H_i(\mu_i p_i)) \\
 &= P(U_i \leq H_i(\mu_i p_i)) \\
 &= H_i(\mu_i p_i).
 \end{aligned}
 \tag{۱۳}$$

قابل ذکر است که تساوی‌های رابطه (۱۳) به ترتیب با استفاده از تعریف تابع توزیع، جایگذاری مقدار مشاهده زمان کل آزمون با تبدیل زمان کل آزمون، معکوس‌گیری از طرفین نامساوی به وسیله تابع معکوس تبدیل زمان کل آزمون، هم توزیع بودن متغیر تصادفی یکنواخت استاندارد با تابع توزیع و خواص احتمالی متغیر تصادفی یکنواخت استاندارد حاصل شده است. رابطه (۱۳) نشان می‌دهد که تابع توزیع متغیر تصادفی X^{ttt} همان معکوس تبدیل زمان کل آزمون متغیر تصادفی X است. از رابطه (۱۳) همچنین می‌توان نتیجه گرفت که معکوس تبدیل زمان کل آزمون یک تابع توزیع است. این موضوع در کوچار و همکاران (۲۰۰۲) نیز بیان و ثابت شده است.

حال با در نظر گرفتن توزیع ارائه شده در (۱۳) و استفاده از تعریف تابع توزیع توام دومتغیره برای بردار تصادفی $\mathbf{X}^{ttt} = (X_1^{ttt}, X_2^{ttt})$ داریم:

$$F_{\mathbf{X}^{ttt}}(\mu_1 p_1, \mu_2 p_2) = P(X_1^{ttt} \leq \mu_1 p_1, X_2^{ttt} \leq \mu_2 p_2). \tag{۱۴}$$

با استفاده از قضیه اسکالر (قضیه ۴.۲) و تابع توزیع توام بردار $\mathbf{X}^{ttt} = (X_1^{ttt}, X_2^{ttt})$ می‌توان نوشت:

$$F_{\mathbf{X}^{ttt}}(\mu_1 p_1, \mu_2 p_2) = C(H_{F_1}(\mu_1 p_1), H_{F_2}(\mu_2 p_2)) \tag{۱۵}$$

که در آن $C(.,.)$ یک تابع مفصل است. بنابراین با توجه به اینکه مولفه‌های مفصل ارائه شده در رابطه (۱۵) هر کدام معکوس تبدیل زمان کل آزمون متناظر با متغیر مورد مطالعه می‌باشد، رابطه (۱۵) را می‌توان یک تعمیم جدید از معکوس تبدیل زمان کل آزمون در حالت دو متغیره در نظر گرفت که در تعریف ۲.۳ به بیان آن می‌پردازیم. برای جزئیات بیشتر به اصفهانی (۱۴۰۱) مراجعه شود.

تعریف ۳.۲. فرض کنید $X = (X_1, X_2)$ یک بردار تصادفی از متغیرهای تصادفی نامنفی با امید ریاضی $\mu = (\mu_1, \mu_2)$ باشد و X_i ها دارای تبدیل زمان کل آزمون $H_{F_i}^{-1}$ هستند. اگر $C(\cdot, \cdot)$ یک تابع مفصل و H_{F_i} معکوس تبدیل زمان کل آزمون باشد، آنگاه بایه کارگیری قضیه اسکالر برای توزیع های حاشیه ای H_{F_1} و H_{F_2} در نقاط $\mu_1 p_1$ و $\mu_2 p_2$ برای $(p_1, p_2) \in [0, 1]^2$ می توان نوشت:

$$T_{C, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) = C(H_{F_1}(\mu_1 p_1), H_{F_2}(\mu_2 p_2)), \quad (16)$$

که رابطه (۱۶) یک توسیع جدید از معکوس تبدیل زمان کل آزمون در حالت دو متغیره است که در آن $H_{F_i}(\mu_i p_i)$ تابع توزیع X_i^{ttt} است. برای استفاده از رابطه (۱۶) باید به این نکته توجه داشت که نابرابری اقتصادی برای هر کدام از متغیرهای تصادفی X_i توسط تبدیل $H_{F_i}^{-1}(p_i)$ اندازه گیری می شود که معکوس تبدیل $H_{F_i}(p_i)$ است، (اصفهانی و همکاران، ۱۳۹۹). همچنین نابرابری اقتصادی موجود در ساختار وابستگی $X = (X_1, X_2)$ به وسیله تابع مفصل $C = (u_1, u_2)$ نشان داده می شود.

با توجه به ویژگی های بیان شده در بخش ۲-۴ برای توابع مفصل، برخی از ویژگی های تبدیل زمان کل آزمون دومتغیره را می توان به صورت زیر بیان نمود.

با در نظر گرفتن توابع مفصل معرفی شده در بخش ۲-۴ به عنوان ساختار وابستگی^۱ بین متغیرها، تبدیل زمان کل آزمون دومتغیره بر اساس رابطه (۱۶) به ترتیب به صورت زیر بیان می شود:

$$T_{Gumbel, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) = H_{F_1}(\mu_1 p_1) H_{F_2}(\mu_2 p_2) \times \exp\{-\theta \ln(H_{F_1}(\mu_1 p_1)) \ln(H_{F_2}(\mu_2 p_2))\},$$

$$T_{FGM, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) = H_{F_1}(\mu_1 p_1) H_{F_2}(\mu_2 p_2) \times [1 + \theta (1 - H_{F_1}(\mu_1 p_1)) (1 - H_{F_2}(\mu_2 p_2))],$$

$$T_{\pi, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) = H_{F_1}(\mu_1 p_1) H_{F_2}(\mu_2 p_2).$$

^۱. Structure dependency

گزاره ۳.۳. فرض کنید $X = (X_1, X_2)$ یک بردار تصادفی از متغیرهای تصادفی نامنفی با امید ریاضی $\mu = (\mu_1, \mu_2)$ و تبدیل زمان کل آزمون دومتغیره $T_{C, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2)$ باشد. در اینصورت همواره

$$\max\{H_{F_1}(\mu_1 p_1) + H_{F_2}(\mu_2 p_2) - 1, 0\} \leq T_{C, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) \leq \min\{H_{F_1}(\mu_1 p_1), H_{F_2}(\mu_2 p_2)\}.$$

در ادامه این بخش با ارائه چند مثال رفتار تبدیل زمان کل آزمون دومتغیره را مورد بررسی و مطالعه قرار می‌دهیم.

مثال ۵.۳. فرض کنید $X = (X_1, X_2)$ یک بردار تصادفی از متغیرهای تصادفی وایبول^۱ با پارامترهای شکل و مقیاس به ترتیب α_i و β_i برای $i = 1, 2$ باشند که دارای تابع چگالی به فرم

$$f_{\alpha_i, \beta_i}(x_i) = \frac{\beta_i^{\alpha_i}}{\Gamma(\alpha_i)} x_i^{\alpha_i-1} \exp(-\beta_i x_i), \quad \alpha_i, \beta_i, x_i > 0$$

است و در آن $\Gamma(a) = \int_0^\infty t^{a-1} \exp(-t) dt$ تابع گاما است. در این صورت

$$H_{F_i}^{-1}(p_i) = \frac{\Gamma\left(\frac{1}{\alpha_i}\right)}{\alpha_i} P_{Gamma}(-\ln(1-p_i); \frac{1}{\alpha_i}, 1)$$

و در نتیجه

$$H_{F_i}(p_i) = 1 - \exp\left\{-F_{Gamma\left(\frac{1}{\alpha_i}, 1\right)}^{-1}\left(\frac{\alpha_i p_i}{\Gamma\left(\frac{1}{\alpha_i}\right) \beta_i}\right)\right\}$$

حال اگر ساختار وابستگی کلیتون را برای مولفه‌های بردار $X = (X_1, X_2)$ در نظر بگیریم، معکوس تبدیل زمان کل آزمون دومتغیره برای متغیرهای مورد نظر عبارت است:

¹ Weibull

$$T_{Clayton, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) \\ = [\max \left\{ (1 - \exp \left\{ -F_{Gamma(\frac{1}{\alpha_1}, 1)}^{-1} \left(\frac{\alpha_1 \mu_1 p_1}{\Gamma(\frac{1}{\alpha_1}) \beta_1} \right) \right\} \right\}^{-\theta} + \\ (1 - \exp \left\{ -F_{Gamma(\frac{1}{\alpha_2}, 1)}^{-1} \left(\frac{\alpha_2 \mu_2 p_2}{\Gamma(\frac{1}{\alpha_2}) \beta_2} \right) \right\}^{-\theta} - 1, 0 \right\}]^{-\frac{1}{\theta}}$$

با در نظر گرفتن مقادیر مختلف θ و حالت های مختلف وابستگی بین متغیرها می توان رفتار تبدیل جدید را مورد مطالعه قرار داد. در مثال های بعدی این موضوع را در قالب یک محاسبه عددی نیز مورد بررسی قرار می دهیم.

مثال ۳.۶. فرض کنید $\mathbf{X} = (X_1, X_2)$ یک بردار تصادفی تابع بقای دو متغیره گامبل^۱ به فرم $\bar{F}(x_1, x_2) = \exp\{-\lambda_1 x_1 - \lambda_2 x_2 - \delta x_1 x_2\}$ ، $x_1, x_2, \lambda_1, \lambda_2, \delta > 0$ باشد. با انجام محاسباتی ساده تبدیل زمان کل آزمون برای متغیرهای X_1 و X_2 به ترتیب $H_{X_1}^{-1}(p_1) = \frac{p_1}{\lambda_1}$ و $H_{X_2}^{-1}(p_2) = \frac{p_2}{\lambda_2}$ و در نتیجه معکوس این تبدیل ها نیز به ترتیب $H_{X_1}(t) = \lambda_1 t$ و $H_{X_2}(t) = \lambda_2 t$ به دست می آید. همچنین $\mu_i = \frac{1}{\lambda_i}$ برای $i = 1, 2$ است. حال با در نظر گرفتن ساختار وابستگی کلیتون برای بردار تصادفی $\mathbf{X}$ ، معکوس تبدیل زمان کل آزمون دو متغیره به صورت زیر به دست خواهد آمد.

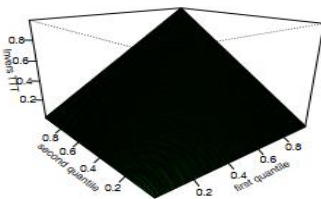
$$T_{C, H_{F_1}, H_{F_2}}^{-1}(\mu_1 p_1, \mu_2 p_2) = C(H_{F_1}(\mu_1 p_1), H_{F_2}(\mu_2 p_2)) \\ = [\max\{p_1^{-\theta} + p_2^{-\theta} - 1, 0\}]^{-\frac{1}{\theta}} \\ = (p_1^{-\theta} + p_2^{-\theta} - 1)^{-\frac{1}{\theta}} \quad (17)$$

نمودار تبدیل به دست آمده در رابطه (۱۷) برای مقادیر مختلف پارامتر θ در شکل ۳ نشان داده شده است. همانطور که مشاهده می شود برای مقدار وابستگی منفی زیاد، شکل (آ) نمودار به صورت مخروط ناقص است و در شکل (ب) که نمودار برای مقدار وابستگی صفر یعنی حاشیه ای های مستقل را نشان می دهد شکل نمودار یک مخروط کامل با توزیع متناسب داده ها

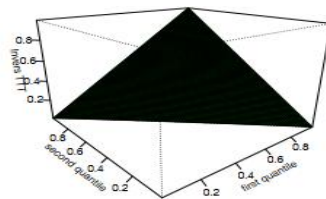
^۱. Gumbel

است. در شکل های (ج) و (د) هم که به ترتیب نمودار را برای وابستگی مثبت کم (ج) و وابستگی مثبت قوی (د) نشان می دهد چگالی نمودار به سمت نقاط بالاتر افزایش می یابد. بنابراین می توان نتیجه گرفت که، با تغییر وابستگی از منفی به مثبت، نمودار به صورت مخروطی شکل می گیرد و با افزایش وابستگی مثبت بین متغیرها، چگالی نمودار و در نتیجه تراکم آن در بالاترین نقاط آن افزایش می یابد.

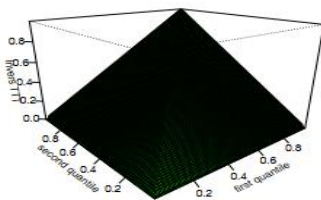
شکل (۳): نمودار معکوس تبدیل زمان کل آزمون برای توزیع دو متغیره گامبل



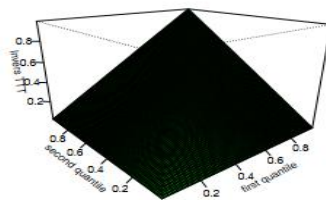
(ب) $\tau = 0$



(آ) $\tau = -0.99$



(د) $\tau = 0.99$



(ج) $\tau = 0.5$

منبع: محاسبات پژوهش با نرم افزار R

مثال ۷.۳. فرض کنید $X = (X_1, X_2)$ یک بردار تصادفی از متغیرهای تصادفی هم توزیع با تابع بقای دو متغیره پارتو^۱ به فرم

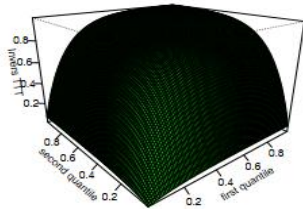
$$\bar{F}(x_1, x_2) = (1 + \alpha_1 x_1 + \alpha_2 x_2)^{-\lambda}, x_1, x_2, \alpha_1, \alpha_2, \lambda > 0,$$

باشد. حال با در نظر گرفتن ساختار وابستگی کلیتون برای متغیرها با انجام محاسباتی شبیه مثال قبل می توان تبدیل های زمان کل آزمون یک متغیره و همچنین معکوس آنها را برای متغیرهای پارتو به دست آورد. سپس با استفاده از رابطه (۱۶) تبدیل معکوس دومتغیره زمان کل آزمون را با در نظر گرفتن ساختار همبستگی کلیتون بین متغیرهای بردار $X = (X_1, X_2)$ بدست آورد. شکل ۳ نمودار معکوس تبدیل زمان کل آزمون دومتغیره را برای مشابه مثال قبل رفتار تبدیل معکوس زمان کل آزمون دومتغیره برای بردار $X = (X_1, X_2)$ و مقادیر مختلف پارامتر θ مفصل کلیتون نشان می دهد.

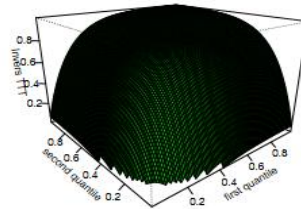
شکل (۴) رفتار تبدیل را نسبت به تغییر وابستگی بین متغیرها از منفی به مثبت نشان می دهد، یعنی حالت وابستگی منفی زیاد (آ)، حاشیه های مستقل (ب)، وابستگی مثبت کم (ج) و وابستگی مثبت قوی (د). همانطور که مشاهده می شود، با تغییر وابستگی از منفی به مثبت، به ترتیب شاهد یک گنبد ناقص، یک گنبد متوازن و برای حالت وابستگی منفی و عدم وابستگی هستیم. همچنین برای وابستگی های مثبت هم، با افزایش وابستگی تراکم داده ها در بالا و اطراف سقف گنبد بیشتر می شود و شکل گنبندی نمودار کامل می شود.

^۱. Pareto

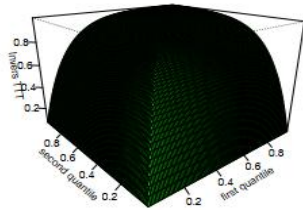
شکل (۴): نمودار معکوس تبدیل زمان کل آزمون برای توزیع دو متغیره پارتو



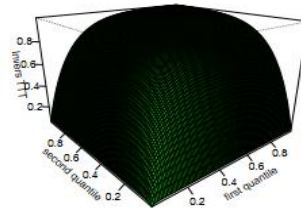
(ب) $\tau = 0$



(ا) $\tau = -0.99$



(د) $\tau = 0.99$



(ج) $\tau = 0.5$

منبع: محاسبات پژوهش با نرم افزار R

۴. تحلیل داده‌های ثروت و درآمد

در این بخش ابتدا به مروری کوتاه بر مفاهیم درآمد و ثروت پرداخته و سپس با استفاده از نتایج ارائه شده در بخش ۴ به تحلیل داده‌های درآمد و ثروت خانواده‌های ایرانی می‌پردازیم. میزان درآمدی که در یک دوره زمانی، مانند یک ماه یا سال، از یک شغل، سرمایه‌گذاری یا منابع دیگر کسب می‌شود، درآمدی موقتی است و می‌تواند بر اساس وضعیت شغلی، شرایط بازار یا تغییر

در بازده سرمایه‌گذاری در نوسان باشد. درآمد به‌طور کلی به پولی اشاره دارد که فرد یا خانوار در یک دوره زمانی مشخص به‌دست می‌آورد.

ثروت نیز ارزش کل دارایی‌های یک فرد منهای بدهی‌های آنها است که به‌عنوان ارزش خالص نیز شناخته می‌شود. ثروت در طول زمان انباشته می‌شود و می‌تواند امنیت مالی بلندمدتی ایجاد کند. بنابراین ثروت به مجموع دارایی‌ها و منابع مالی یک فرد اشاره دارد.

درآمد و ثروت هر دو شاخص مهم امنیت مالی هستند، اما ویژگی‌های متفاوتی دارند و همیشه یکسان نیستند. درآمد را می‌توان برای افزایش ثروت از طریق پس‌انداز و سرمایه‌گذاری استفاده کرد و ثروت می‌تواند در مواقع بی‌ثباتی درآمد به‌عنوان یک تکیه‌گاه در نظر گرفته شود. علیرغم اینکه به نظر می‌رسد ارتباط مستقیمی بین این دو متغیر وجود داشته باشد، اما رابطه آماری بین درآمد و ثروت به میزان قابل توجهی پایین است. این مفهوم در علم اقتصاد با استفاده از مباحثی مانند اقتصاد رفتاری، سواد مالی، و غیره توضیح داده می‌شود. در اینجا با مرور مختصر این عوامل به کنکاش بیش‌تری در داده‌های در دسترس خواهیم پرداخت و مفاهیم کمی بیش‌تری را توضیح خواهیم داد.

یکی از کلیدهای اصلی برای افزایش ثروت، پس‌انداز منظم از درآمد است. افرادی که بخشی از درآمد خود را به‌طور منظم پس‌انداز می‌کنند، می‌توانند به مرور زمان ثروت خود را افزایش دهند. پس‌انداز تنها بخشی از معادله است. سرمایه‌گذاری هوشمندانه می‌تواند به رشد سریع‌تر ثروت کمک کند. همچنین ایجاد یک بودجه دقیق و پایبندی به آن می‌تولند به افراد کمک کند تا از درآمد خود به بهترین نحو استفاده کنند. با بودجه‌بندی مناسب، افراد می‌توانند هزینه‌های غیرضروری را کاهش دهند و بیش‌تر پس‌انداز کنند. کاهش بدهی‌ها نیز می‌تواند تأثیر مثبتی بر روی ثروت داشته باشد. بنابراین پرداخت بدهی‌ها به معنای افزایش دارایی خالص است.

آموزش مالی یکی از عوامل کلیدی در مدیریت بهتر درآمد و ثروت است. افرادی که با اصول مالی آشنا هستند، معمولاً تصمیمات بهتری در زمینه سرمایه‌گذاری و مدیریت پول اتخاذ می‌کنند. عوامل اقتصادی مانند تورم، نرخ بهره، و شرایط بازار نیز می‌توانند بر درآمد و ثروت تأثیر بگذارند. برای مثال اگر درآمد فرد نتولند با نرخ تورم همگام شود، قدرت خرید او کاهش خواهد یافت. تغییرات در نرخ بهره نیز از عوامل تاثیرگذار بر هزینه‌های وام و همچنین بازده سرمایه‌گذاری‌ها است. نسبت درآمد به ثروت می‌تواند نشان‌دهنده وضعیت مالی یک فرد باشد.

افرادی که درآمد بالایی دارند اما ثروت کمی جمع‌آوری کرده‌اند، ممکن است نیاز به بررسی عمیق‌تری از عادات مالی خود داشته باشند. در نهایت، رابطه بین درآمد و ثروت پیچیده است و نیازمند مدیریت دقیق و برنامه‌ریزی مالی هوشمندانه است. صحیح نیست که گفته شود با درآمد بالاتر می‌توان به ثروت بیش‌تری رسید و یا بالعکس. افرادی که تولنایی مدیریت درآمد خود را دارند و از آن به خوبی استفاده می‌کنند، معمولاً در جمع‌آوری ثروت موفق‌تر خواهند بود. این موضوع نه تنها در وضعیت مالی افراد بلکه در رفاه کلی جامعه نیز تأثیرگذار است. با توجه به اینکه داده‌های نقدی مربوط به ثروت خانوارها همواره مجهول است و نمی‌توان به طور دقیق بیان نمود باید از یک شاخص جایگزین که با استفاده از دارایی‌های غیرنقدی خانوارها به دست آمده است، برای متغیر ثروت استفاده نمود. در ادامه روش ساخت این شاخص و ارزیابی اعتبار آن بیان می‌شود.

۴-۱. ساخت شاخص ترکیبی ثروت خانوار

در بسیاری از کشورها، از جمله ایران، داده‌های مستقیم و رسمی مربوط به ثروت خانوارها، شامل مجموع دارایی‌های نقدی و غیرنقدی، به صورت عمومی و قابل دسترسی منتشر نمی‌شوند. با این حال، تحلیل نابرابری اقتصادی بدون در نظر گرفتن عامل «ثروت» در کنار «درآمد»، تصویری ناقص از ساختار رفاه در جامعه ارائه می‌دهد. از این رو، در پژوهش حاضر، با هدف بررسی وابستگی میان دو مؤلفه کلیدی رفاه خانوار (درآمد و ثروت) و ارزیابی عملکرد شاخص نوین TTT در شرایط داده‌ای واقعی، ساخت یک شاخص ترکیبی ثروت خانوار بر پایه اطلاعات جانبی صورت گرفته است. این شاخص با بهره‌گیری از متغیرهای جانبی موجود در پرسشنامه هزینه و درآمد خانوار و مبتنی بر روش‌های رایج در پژوهش‌های اقتصادسنجی و علوم اجتماعی طراحی شده است.

در ادبیات توسعه و سنجش رفاه، استفاده از متغیرهای جانشین^۱ برای برآورد سطح ثروت خانوار، به‌ویژه در کشورهایی با داده‌های ناقص، امری متداول است. مطالعاتی نظیر گروث و همکاران (۲۰۲۲)، فیلمر و پریچت^۲ (۲۰۰۱) و سان و استیفل^۳ (۲۰۰۳) در زمینه تحلیل فقر چند

۱. Proxy Variables

۲. Filmer and Pritchett

۳. Sahn and Stifel

بعدی، بر ترکیب اطلاعاتی چون مالکیت وسایل با دوام، وضعیت مسکن و دسترسی به خدمات زیربنایی برای ساخت شاخص ثروت تأکید کرده‌اند.

در همین راستا، در این پژوهش، شاخص ترکیبی ثروت خانوار با استفاده از داده‌های موجود در بخش‌های دوم و سوم پرسش‌نامه هزینه و درآمد خانوار مرکز آمار ایران (سال‌های ۱۴۰۰ تا ۱۴۰۲) استخراج شده است. مهم‌ترین متغیرهای مورد استفاده در ساخت این شاخص عبارتند از:

- نوع مالکیت مسکن (ملکی، اجاره‌ای، سایر)
- ارزش تقریبی ملک محل سکونت
- مالکیت وسایل نقلیه شخصی و تعداد آن
- برخورداری از وسایل بادوام مانند یخچال، ماشین لباس‌شویی، تلویزیون و رایانه
- نوع مصالح به کار رفته در ساختمان به عنوان شاخص کیفیت سکونت
- سطح تحصیلات سرپرست خانواده

با استفاده از یک رویکرد امتیازدهی ساده، برای هر خانوار یک نمره نسبی ثروت محاسبه شده است. به عنوان مثال، خانوارهایی که در خانه‌ی ملکی با ارزش بالا زندگی می‌کنند، دارای خودرو و چند وسیله بادوام هستند، امتیاز بالاتری نسبت به خانوارهای فاقد این امکانات دریافت کرده‌اند. این امتیاز به عنوان شاخص ترکیبی ثروت وارد مدل گردیده است. این روش توسط گروث و همکاران (۲۰۲۲) نیز مورد استفاده قرار گرفته است.

شایان ذکر است که این شاخص، ثروت را به صورت نسبی و غیرنقدی می‌سنجد و از دقت سنج‌های مالی دقیق (مانند تراز دارایی‌ها و بدهی‌ها) برخوردار نیست. اما به عنوان یک تقریب ساختاریافته، می‌تواند برای تحلیل وابستگی میان ثروت و درآمد با استفاده از روش‌های تابع مفصل، مبنایی قابل قبول فراهم آورد. همچنین، اعتبار شاخص در این پژوهش از طریق تحلیل رفتار آماری آن در دهک‌های مختلف درآمدی، و همبستگی آن با درآمد خانوارها، به صورت ضمنی ارزیابی شده است.

در مطالعات آتی، استفاده از تکنیک‌های تحلیل عاملی یا تحلیل مؤلفه‌های اصلی^۱ (PCA) برای ساخت نسخه‌های پیچیده‌تر و دقیق‌تر از شاخص ثروت خانوار پیشنهاد می‌شود.

^۱. Principal Component Analysis

۴-۲. ارزیابی روایی و پایایی شاخص ثروت

روایی و پایایی شاخص‌های ترکیبی در پژوهش‌های علوم اجتماعی به ویژه زمانی که متغیر مورد نظر به صورت مستقیم در دسترس نیست از اهمیت بالایی برخوردار است. در این پژوهش برای ارزیابی روایی و پایایی شاخص ترکیبی ثروت خانوار از چند روش مکمل استفاده شده است.

۴-۲-۱. روایی همگرا^۱

به منظور ارزیابی روایی همگرا، همبستگی بین شاخص ثروت با چند متغیر نظری مرتبط بررسی شد. متغیر درآمد خانوار یکی از این متغیرها بود که ضریب همبستگی بین این دو متغیر برابر با ۰/۴۱ به دست آمد که نشان دهنده رابطه مثبت و معنادار بین این دو شاخص (درآمد و ثروت) است. میانگین شاخص ثروت در خانوارهای شهری به طور معناداری بالاتر از خانوارهای روستایی بود که مطابق با انتظار نظری درباره تمرکز بیشتر دارایی‌ها در نواحی شهری است. شاخص ثروت در میان خانوارهایی که سرپرست دارای تحصیلات دانشگاهی هستند به طور معناداری بالاتر از سایر سطوح تحصیلی بوده است که تائیدی بر همگرایی نظری شاخص با ظرفیت انباشت انسانی و اقتصادی می باشد.

۴-۲-۲. پایایی^۲

به منظور ارزیابی پایایی درونی شاخص از ضریب آلفای کرونباخ^۳ برای متغیرهای تشکیل دهنده استفاده شده است. مقدار این ضریب برابر با ۰/۷۲ به دست آمده که نشان دهنده سطح پایایی قابل قبول برای این نوع شاخص ترکیبی در مطالعه مورد نظر است. در مجموع نتایج حاصل از آزمون‌های روایی و پایایی حاکی از آن است که شاخص پیشنهادی توانایی قابل قبولی در بازنمایی مفهوم ثروت خانوار دارد. هر چند با در نظر گرفتن محدودیت‌های ذاتی ناشی از نبود داده‌های نقدی باید در تحلیل نتایج با دقت بیشتر تفسیر شود. جدول‌های ۱ و ۲ نتایج مربوط به آزمون‌های پایایی و روایی را نشان می دهد. تحلیل‌های آماری مربوط به مطالعه پایایی و روایی با استفاده از نسخه ۲۲ نرم افزار SPSS انجام شده است.

1. Convergent Validity

2. Reliability

3. Cronbach's alpha

جدول (۱): نتایج تحلیل عاملی شاخص ترکیبی ثروت خانوار

نوع دارایی	اشتراک پذیری (Communality)	بار عاملی (Factor Loading)	متغیر جانشین ثروت	ردیف
غیرنقدی	۰/۷۱	۰/۸۴	مالکیت مسکن	۱
غیرنقدی	۰/۶۷	۰/۸۲	ارزش تقریبی مسکن و نوع مصالح	۲
غیرنقدی	۰/۵۸	۰/۷۶	مالکیت خودرو	۳
غیرنقدی	۰/۴۸	۰/۶۹	وسایل با دوام	۴
شبه نقدی	۰/۴۱	۰/۶۴	مالکیت سهام یا اوراق	۵
غیرنقدی	۰/۳۵	۰/۵۹	سطح تحصیلات	۶

منبع: یافته‌های پژوهش

جدول (۲): نتایج تحلیل عاملی شاخص ترکیبی ثروت خانوار

تفسیر	مقدار	شاخص
پایایی قابل قبول برای شاخص ترکیبی ثروت	۰/۷۲	آلفای کرونباخ
نشان دهنده کفایت نمونه‌گیری مطلوب	۰/۷۸	شاخص KMO ^۱

منبع: یافته‌های پژوهش

۳-۴. تحلیل داده واقعی

در این بخش بدنبال تحلیل رابطه‌ی بین متغیر درآمد و شاخص ثروت خانوارهای ایران در بازه زمانی سال‌های ۱۴۰۰ تا ۱۴۰۲ هستیم. برای داشتن تحلیلی جامع و با توجه به تورمی که در ایران رایج است، بهتر است که این اعداد براساس تورم به اعداد قابل مقایسه‌ای تبدیل شوند. بر این اساس، باید تاریخ پایه‌ای در نظر گرفت که اعداد بر اساس آن تاریخ بازتنظیم شوند. با توجه به اینکه تاریخ نگارش این پژوهش، آبان ماه ۱۴۰۳ است، تمامی اعداد براساس نرخ تورم اعلامی مرکز آمار ایران به لینک (<https://amar.org.ir/prices>) بازتنظیم شده است و برای درک بهتر و راحتی کار با داده‌ها از مقیاس تقسیم بر ده میلیون تومان استفاده شده است.

جدول ۳ یک خلاصه مقدماتی از داده‌ها بر اساس دامنه تغییرات و چارک‌های اصلی نشان می‌دهد. با توجه به داده‌های جدول ۳ و فاصله بین کمترین و بیشترین مقدار درآمد و همچنین مقادیر چارک‌های داده‌ها می‌توان نتیجه گرفت که پراکندگی داده‌ها زیاد است و همچنین تمرکز بر داده‌های قبل میانه است، یعنی توزیع داده‌ها چوله به راست به نظر می‌رسد. برای به

^۱. Kaiser-Meyer-Olkin Measure

دست آوردن یک مدل آماری مناسب برای این داده‌ها مدل معروف وایبل به آن‌ها برازش داده شد که البته نتیجه خوبی نداشت.

جدول (۳): خلاصه داده‌های درآمد ماهانه تقسیم بر ده میلیون تومان

بیشترین	چارک سوم	میانه	میانگین	چارک اول	کمترین
۱۶/۹۷۳	۵/۴۶۲	۳/۱۰۷	۳/۶۹۸	۲/۱۵۶	۰/۷۳۸

منبع: یافته‌های پژوهش

در ادامه جستجو برای یافتن مدل مناسب برای داده‌های جدول ۳، خروجی نتایج منجر به انتخاب یکی از مدل‌های معروف آماری یعنی توزیع گومپرتز شد. جدول ۴ نتایج مربوط به این مورد را ارائه می‌دهد.

جدول (۴): خلاصه داده‌های تولید شده از توزیع گومپرتز با پارامترهای مختلف

بیشترین	چارک سوم	میانه	میانگین	چارک اول	کمترین	توزیع گومپرتز
۱۹/۷۸۴	۸/۶۳۴	۶/۱۱۵	۵/۹۸۸	۳/۳۶۲	۰/۰۰۰	$G(۰/۰۶ و ۰/۲)$
۸/۲۸۱	۴/۴۹۵	۳/۳۵۶	۳/۴۴۸	۲/۲۵۵	۰/۰۰۰	$G(۰/۰۶ و ۰/۶)$
۱۴/۸۴۰	۵/۴۶۴	۳/۷۶۵	۳/۴۴۰	۱/۷۲۱	۰/۰۰۰	$G(۰/۱۴ و ۰/۲)$
۶/۸۳۳	۳/۲۲۶	۲/۳۱۱	۲/۲۹۸	۱/۳۳۹	۰/۰۰۰	$G(۰/۱۴ و ۰/۰۶)$

منبع: یافته‌های پژوهش

جدول ۴ خلاصه داده‌ها را برای چهار نوع توزیع گومپرتز نشان می‌دهد. داده‌های این جدول پس از انتخاب پارامترهای مناسب در مدل گومپرتز پیشنهادی و برازش آن به داده‌ها حاصل شده است. با توجه به داده‌های این جدول می‌توان نتیجه گرفت که داده‌های درآمد از توزیع گومپرتز با پارامترهایی در حدود پارامترهای داده شده در جدول ۴ پیروی خواهند کرد.

برای پیدا کردن یک مدل مناسب برای داده‌های درآمد، با استفاده از برآوردگر ماکزیمم درست‌نمایی^۱ به توزیع $G(۰/۱۰ و ۰/۳۹۱)$ می‌رسیم. برای نشان دادن اینکه این مدل به خوبی عمل می‌کند، می‌توان از نمودارهای کانتور^۲ و چندک-چندک^۳ استفاده کرد. شکل‌های ۵ و ۶ به ترتیب دو نمودار کانتور و چندک-چندک مدل برازش شده را نشان می‌دهد. زیرا نمودار

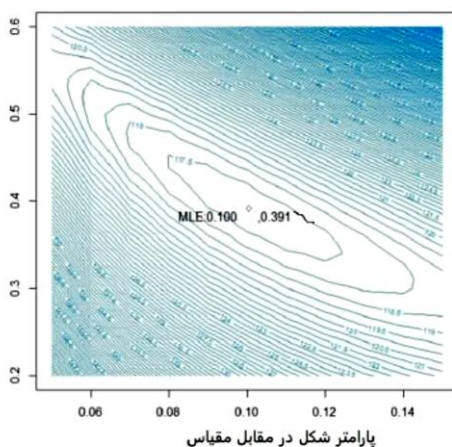
۱. Maximum Likelihood Estimator

۲. Contour

۳. Q-Q plot

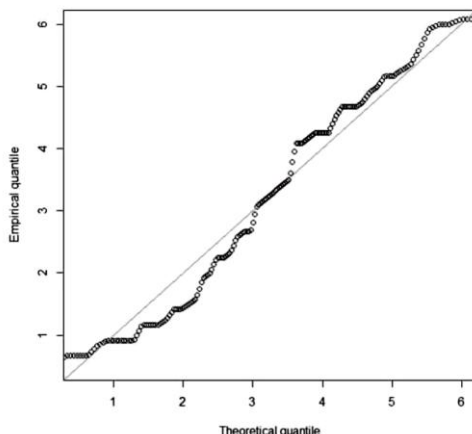
کانتور داده‌ها به‌طور متناسب حول مقدار برآورد درست‌نمایی ماکزیمم داده‌ها پراکنده شده است و نشان دهنده پراکندگی متناسب حول این نقطه است.

شکل (۵): نمودار کانتور داده‌های درآمد



منبع تحقیق: یافته‌های پژوهش

شکل (۶): نمودار چندک-چندک داده‌های درآمد



منبع تحقیق: یافته‌های پژوهش

همچنین نمودار چندک-چندک نیز که نحوه پراکندگی داده‌ها حول نیمساز مربع فرضی را نشان می‌دهد بیانگر مناسب بودن توزیع پیشنهادی است. زیرا داده‌ها حول نیمساز پراکنده

شده اند و فاصله معناداری از آن وجود ندارد. با توجه به این شکل می‌توان استنباط کرد که مدل بدست آمده برای این داده‌ها مناسب خواهد بود.

براساس این مدل بدست آمده می‌توان شبیه‌سازی‌های بیش‌تری از داده‌های درآمد را انجام داده و برپایه‌ی آن‌ها درک مناسب‌تری از رابطه‌ی بین داده‌های درآمد و ثروت به دست آوریم.

در ادامه به دنبال یافتن مدل مناسب برای داده‌های ثروت هستیم که البته نیازی به تکرار این فرآیند نیست و فقط مدل برازش‌شده را بیان خواهیم کرد. انجام فرآیند قبل برای داده‌های ثروت به مدل $G(0/01 و 0/01)$ برای این داده‌ها نتیجه می‌شود. برازش این داده‌ها بدلیل پراکندگی و چوله بودن زیاد آن‌ها سخت‌تر و از دقت کم‌تری برخوردار است، ولی با این وجود می‌توان استنباط روشنی از این نوع داده‌ها نیز داشت. خلاصه این داده‌ها در جدول ۵ نشان داده شده‌اند. داده‌های این جدول نیز میزان اختلاف زیاد بین کمترین و بیشترین مقدار داده ثروت را نشان می‌دهد. همچنین نشان دهنده اختلاف زیاد چارک‌های داده‌ها با بیشترین مقدار داده‌ها است که این موضوع پراکندگی خیلی زیاد داده‌ها و در نتیجه چوله بودن آنها را نتیجه می‌دهد.

جدول (۵): خلاصه داده‌های ثروت تقسیم بر ده میلیون تومان

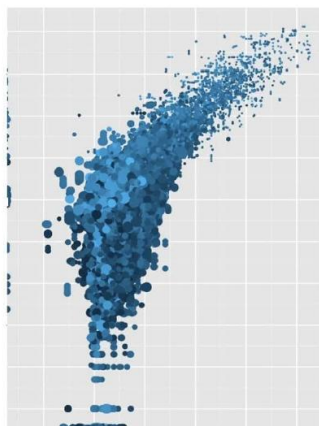
بیشترین	چارک سوم	میانه	میانگین	چارک اول	کمترین
۱۹۵۰/۶۵۴	۸۶/۴۳۸	۵۹/۰۴۳	۵۲/۰۰۱	۲۵/۳۲۱	۱۲/۵۶۴

منبع تحقیق: یافته‌های پژوهش

شکل ۷ پراکندگی داده‌های ثروت را نشان می‌دهد.

براساس این شکل، تمرکز ثروت بین ۱/۵ تا ۵ میلیارد می‌باشد (هر سلول نمودار نشان دهنده ۱ میلیارد است). همچنین از شکل ۷ می‌توان نتیجه گرفت که بدون در نظر گرفتن متغیر درآمد، حدود ۷۰ درصد ثروت بین ۲ تا ۵ میلیارد قرار دارند. بنابراین دوباره می‌توان گفت که در کل داده‌ها ارتباط چندانی بین داده‌های درآمد و ثروت وجود ندارد. اما در مقادیر کمتر می‌توان ارتباط مثبت را بین دو متغیر مشاهده نمود و هر چه به مقدار این دو متغیر افزوده می‌شود از میزان ارتباط بین آنها نیز کاسته می‌شود.

شکل (۷): نمودار پراکندگی داده‌های درآمد در مقابل داده‌های ثروت



منبع تحقیق: یافته‌های پژوهش

با توجه به این موضوع و برای درک بهتر این مورد، هر دو داده‌ی درآمد و ثروت را به ده قسمت تقسیم کرده که اصطلاحاً به آن دهک می‌گوییم. جدول ۶ مقدار ضریب وابستگی τ و τ^1 کندال بین داده‌های دهک‌ها و همچنین تابع مفصل مناسب برای توصیف ساختار وابستگی بین دو متغیر درآمد و ثروت در هر دهک ارائه گردیده است. از مقادیر وابستگی به دست آمده برای دهک‌های مختلف جامعه می‌توان نتیجه گرفت که میزان وابستگی بین دو متغیر ثروت و درآمد به تفکیک دهک‌های مختلف شرایط متفاوتی نسبت به کلیت جامعه دارد و مقدار آن در دهک‌های پایین بیشتر است. تابع مفصل‌های کلیتون و FGM پیشنهاد شده برای توصیف ساختار وابستگی بین متغیرها در دهک‌های مختلف با استفاده از روابط بین مقدار ضریب وابستگی τ و مقدار پارامترهای این مفصل‌ها که به ترتیب برای این دو مفصل کلیتون و FGM روابط $\theta_C = \frac{\tau\tau^1}{1-\tau}$ و $\theta_{FGM} = \frac{4}{5}\tau$ است، به دست آمده است. در مقدار وابستگی کلی داده‌ها حدود ۰/۲۵۶ محاسبه شده است که مقداری بسیار پایین‌تر از حد انتظار است، ولی بررسی جزئی‌تر میزان وابستگی‌ها در هر دهک می‌تواند ذهنیت مناسب‌تری ارائه دهد. با توجه به ضریب وابستگی τ به دست آمده برای هر دهک می‌توان نتیجه گرفت که برای دهک‌های پایین جامعه که مقدار τ زیاد است مفصل کلیتون را می‌توان به عنوان ساختار وابستگی مناسب در نظر

^۱. Tau

گرفت. همچنین برای دهک‌های بالا چون مقدار τ کم است و نشان‌دهنده وابستگی ضعیف بین دو متغیر ثروت و درآمد در این دهک‌ها است، مفصل FGM را می‌توان به عنوان یک ساختار وابستگی مناسب برای متغیرهای تحت مطالعه در نظر گرفت. قابل ذکر است هنگامی که وابستگی بین متغیرها زیاد باشد مفصل کلیتون مناسب‌ترین تابع مفصل برای توصیف ساختار وابستگی بین متغیرها می‌باشد. همچنین برای حالتی که مقدار وابستگی کم و ضعیف باشد خانواده مفصل‌های FGM بهترین توصیف را از ساختار وابستگی بین متغیرها ارائه می‌دهند.

جدول (۶): مفصل مناسب برای ساختار همبستگی داده‌های درآمد

و ثروت در دهک‌های مختلف جامعه

دهک	مقدار ضریب τ	مفصل مناسب
یک	۰/۴۱۱	Clayton(۱/۳۹۵۵۸۵۷)
دو	۰/۳۶۸	Clayton(۱/۱۶۴۵۵۷۰)
سه	۰/۲۷۴	Clayton(۰/۷۵۴۸۲۰۹)
چهار	۰/۲۲۷	Clayton(۰/۵۸۷۳۲۲۱)
پنج	۰/۲۰۹	Clayton(۰/۵۲۸۴۴۵۰)
شش	۰/۱۸۲	Clayton(۰/۴۴۴۹۸۷۸)
هفت	۰/۱۶۳	Clayton(۰/۳۸۹۴۸۶۳)
هشت	۰/۱۵۷	FGM(۰/۷۰۶۵)
نه	۰/۱۴۷	FGM(۰/۶۶۱۵)
ده	۰/۱۳۶	FGM(۰/۶۱۲۰)

منبع تحقیق: محاسبات پژوهش با نرم افزار R

اگر دهک‌های یک، دو و سه را افراد با درآمد پایین، دهک‌های چهار تا هفت را افراد با درآمد متوسط و بقیه را افراد ثروتمند بدانیم، مجدداً می‌توان به این نتیجه جالب رسید برای افراد با درآمد پایین رابطه‌ی بین داده‌های درآمد و شاخص ثروت مثبت و جالب توجه است. همچنین این رابطه برای افراد با درآمد متوسط کم اما معنادار است. و بالاخره در بین افراد با درآمد بالا رابطه بین داده‌های درآمد و ثروت علی‌رغم انتظار معنادار نیست و به سمت صفر میل می‌کند. جدول ۶ ارتباط بین دهک‌ها را نیز به صورت شفاف‌تری نشان می‌دهد. با توجه به جدول ۶ مقدار ضریب وابستگی بین متغیر درآمد و شاخص ثروت در افراد فقیر جامعه بالاست و هر چه

سطح رفاه عمومی افزایش می‌یابد، به سرعت از مقدار ضریب وابستگی کاسته می‌شود. ضریب وابستگی کل داده‌ها نیز عددی بین مقدار ضریب وابستگی دهک‌های پنجم و ششم است که البته با توجه به اینکه میانه داده‌های درآمد نیز در این بازه است، منطقی به نظر می‌رسد.

شکل (۸): نمودار داده‌های درآمد خانوارهای ایرانی در سال ۱۴۰۲



منبع تحقیق: یافته‌های پژوهش

اگر ضریب وابستگی را برآورد خوبی برای مقدار ضریب تاو کندال (τ) بدانیم، با توجه به روابط $\theta = \frac{\tau}{1-\tau}$ و $\theta = \frac{4}{5}\tau$ که به ترتیب بین پارامترهای مفصل کلایتون و FGM و مقدار ضریب تاو کندال آن‌ها برقرار است، می‌توان مقدار مناسب برای پارامترهای این مفصل‌ها را به دست آورد، که البته این مهم در جدول ۶ و در ستون مفصل مناسب نشان داده شده است. به عنوان نکته پایانی در این بخش می‌توان تاو کندال را از روابط $\tau = \frac{\theta}{\theta+2}$ و $\tau = \frac{\theta}{9}$ نیز به دست آورد که در واقع فرم تغییر یافته فرمول‌های θ است.

۵. بحث و نتیجه‌گیری

در این پژوهش به بررسی و تحلیل سطح رفاه افراد جامعه با تکیه بر داده‌های درآمد و شاخص ثروت آن‌ها پرداخته شده است. هرچند در نگاه اولیه به نظر می‌رسد که بین این دو متغیر رابطه‌ای

مستقیم و پایدار برقرار باشد، اما در مواردی می‌توان شواهدی یافت که این ارتباط را رد می‌کند. به عنوان مثال، ممکن است دو فرد با درآمدهای بسیار متفاوت، سطح ثروت نسبتاً مشابهی داشته باشند. چنین مواردی نشان می‌دهد که تحلیل‌های سطحی نمی‌توانند به‌درستی ساختار وابستگی میان این دو شاخص اقتصادی را تبیین نمایند و لازم است با بهره‌گیری از ابزارهای آماری پیشرفته، به تحلیل دقیق‌تری پرداخت. در همین راستا، ساختار وابستگی میان متغیر درآمد و شاخص ثروت با استفاده از توابع مفصل مورد مطالعه قرار گرفت؛ روشی که از انعطاف‌پذیری بالایی در مدل‌سازی وابستگی برخوردار بوده و حالت استقلال کامل را نیز به عنوان حالتی خاص شامل می‌شود. در گام نخست، با تعریف شاخص زمان کل آزمون در حالت دومتغیره و بررسی سایر شاخص‌های نظری مرتبط، چارچوب مفهومی پژوهش تبیین شد. در ادامه، داده‌های مربوط به درآمد خانوارها که توسط مرکز آمار ایران منتشر شده‌اند، همراه با داده‌های مربوط به شاخص ثروت در بازه زمانی سال‌های ۱۴۰۰ تا ۱۴۰۲ مورد استفاده قرار گرفت تا شاخص دومتغیره معرفی‌شده در مقاله به‌صورت تجربی بررسی شود. برای این منظور از متغیر X_1 برای نشان دادن داده‌های شاخص ثروت افراد و از متغیر X_2 برای نشان دادن درآمد افراد جامعه استفاده شده است. در پایان باید به این نکته اشاره کرد که استفاده از داده‌های دوره ۳ ساله می‌تواند به عنوان یک محدودیت در تحلیل داده‌های کاربردی و استفاده از شاخص معرفی‌شده تلقی شود. به همین دلیل استفاده از داده‌های طولی در بازه‌های زمانی بلندمدت و در شرایط اقتصادی متنوع‌تری برای بررسی پایداری و تغییرپذیری ساختار وابستگی میان درآمد و ثروت می‌تواند در تحقیقات آینده مورد توجه قرار گیرد. همچنین یادآور می‌شود که داده‌های مربوط به شاخص ثروت از یک شاخص ترکیبی ارائه شده برای شاخص ثروت به دست آمده‌اند که محدود به داده‌های غیر نقدی خانوارها بوده است و دارایی‌های نقدی مانند سپرده‌های بانکی، اوراق بهادار و سرمایه‌گذاری‌های مالی را شامل نمی‌شود. لذا نتایج به دست آمده باید با در نظر گرفتن این محدودیت با دقت بیشتری تفسیر و بیان شود.

References

Abounoori, E. and McCloughan, P. (2000). Measuring the Gini coefficient: An empirical assessment of non-parametric and parametric methods (No. 2000_06).

- Abounoori, E. (2003). Modeling the income distribution and Gini coefficient using the Log-Logistic distribution. *Journal of Social Sciences and Humanities of Shiraz university*.19(2), 13-24.
- Abounoori, E. and Asnavandi, E. (2004), Estimation and assessment of the consistency of economic inequality indicators using microdata in Iran. *Economic Reaserch*, 71, 171-210. (In Persian)
- Abounoori, E., Khoshkar, A., & Davoudi, P. (2010). An Analysis of Theil Inequality Index in Terms of Different Provinces in Iran. *Economics Research*, 10(36), 201-222. (In Persian)
- Abounoori, E., Khoshkar, A., & Davoudi, P. (2013). Gini Coefficient Decomposition in Iran in Terms of Urban and Rural Areas. *Journal of Economic Research (Tahghighat-E-Eghtesadi)*, 48(3), 1-12. (In Persian)
- Abounoori, E and Zoghi, E. (2013). Estimating and comparing income distribution inequality using parametric and nonparametric methods. *Macroeconomics Research Letter(MRL)*, 8(16), 13-30. (In Persian)
- Abounoori, E. and Farahati, M. (2016). The Structure of production and income distribution in Iran. *Journal of Economic Modelling* 9(432), 1-23. (In Persian)
- Arnold, B. C. and Sarabia, J. M. (2018). Analytic expressions for multivariate Lorenz surfaces. *Sankhya A*, 80, 84-111.
- Arnold, B. C. (2012). Majorization and the Lorenz order: A brief introduction (Vol. 43). Springer Science and Business Media.
- Barlow, R. E., Bartholomew, D. J., Bremner, J. M. & Brunk, H. D. (1972), *Statistical Inference under Order Restrictions*. John Wiley & Sons, New York.
- Barlow, R. E. and Campo, R. (1975), Total time on test processes and applications to failure data analysis. In: *Reliability and Fault Tree Analysis*. SIAM, Philadelphia, PA.
- Bergman, B. and Klefsjö, B. (2014), Total time on test plots. *Wiley Stats Ref: Statistics Reference Online*.
- Bonferroni, C. (1930), *Elementi di Statistica Generale*. Seeber - Firenze.
- Esfahani, M., Mohtashami Borzadaran, G. R., & Amini, M. (2020). The Total time on test transform in measuring income inequality. *Journal of Econometric Modelling*, 5(4), 89-120. (In Persian)

- Esfahani, M. (2022). Extensions of Total Time on Test Transform. PHD Thesis, Faculty of Mathematical Sciences, Ferdowsi University of Mashhad, Iran. (In Persian)
- Filmer, D. and Pritchett, L. H. (2001). Estimating wealth effects without expenditure data—or tears: an application to educational enrollments in states of India. *Demography*, 38, 115-132.
- Gini, C. (1912), Variabilitia e mutabilitia. Reprinted in *Memorie di metodologica statistica* (Ed. Pizetti E, Salvemini, T). Rome: Libreria Eredi Virgilio Veschi.
- Grothe, O., Kächele, F. and Schmid, F. (2022). A multivariate extension of the Lorenz curve based on copulas and a related multivariate Gini coefficient. *The Journal of Economic Inequality*, 20(3), 727-748.
- Kawczak, J., Kulperger, R. and Yu, H. (2009), Equivalent processes of total time on test, Lorenz and inverse Lorenz processes. *Statistics and Probability Letters*, 79(1), 125-130.
- Kaigh, W. D. (1999), Total time on test function principal components. *Statistics and Probability Letters*, 44(4), 337-341.
- Kochar, S. C., Li, X. and Shaked, M. (2002), The total time on test transform and the excess wealth stochastic orders of distributions. *Advances in Applied Probability*, 34(4), 826-845.
- Li, X. and Shaked, M. (2004), The observed total time on test and the observed excess wealth. *Statistics and Probability Letters*, 68(3), 247-258.
- Lorenz, M. O. (1905), Methods of measuring the concentration of wealth. *Publications of the American Statistical Association*, 9(70), 209-219.
- Mirzaei, S., Mohtashami Borzadaran, G. R and Amini, M. (2018), Zenga Index in Measuring Income Inequality. *Journal of Econometric Modelling*, 3(4), 113-133. (In Persian)
- Morillas, P. M. (2005), A method to obtain new copulas from a given one. *Metrika* 61(2), 169–184.
- Pietra, G. (1915), Delle relazioni tra gli indici di variabilitā. C. Ferrari.
- Perez-Ocon, R., Gámiz-Pérez, M. L. and Ruíz-Castro, J. E. (1997), A study of different ageing classes via total time on test transform and Lorenz curves. *Applied Stochastic Models in Business and Industry*, 13(3-4), 241-248.
- Pham, T. G. and Turkkan, N. (1994), The Lorenz and the scaled total-time-on-test

transform curves: a unified approach. IEEE Transactions on Reliability, 43(1), 76-84.

Sahn, D. E. and Stifel, D. (2003). Exploring alternative measures of welfare in the absence of expenditure data. Review of income and wealth, 49(4), 463-489.

Sklar, M. (1959). Fonctions de répartition à n dimensions et leurs marges. In Annales de l'ISUP (Vol. 8, No. 3, pp. 229-231).

Taguchi, T. (1972a). On the two-dimensional concentration surface and extensions of concentration coefficient and pareto distribution to the two dimensional case—I: On an application of differential geometric methods to statistical analysis. Annals of the Institute of Statistical Mathematics, 24(1), 355-381.

Taguchi, T. (1972). On the two-dimensional concentration surface and extentions of concentration coefficient and pareto distribution to the two dimensional case—II: On an application of differential geometric methods to statistical analysis. Annals of the Institute of Statistical Mathematics, 24(1), 599-619.

Theil, H. (1967), Economics and Information Theory, North Holland Publishing Company, Amsterdam.

Zenga, M. (1984), Proposta per un indice di concentrazione basato sui rapporti fra quantili di popolazione e quantili di reddito. Giornale Degli Economisti e Annali Di Economia, 301-326.

Zenga, M. (2007), Inequality curve and inequality index based on the ratios between lower and upper arithmetic means. Statistica and Applicazioni, 5, 3-28.