



بخش هنر و معماری

گروه علمی مهندسی شهرسازی

مدل های کمی در شهرسازی

مهندس امیر اکبری مهام

مهندس خلیل گودرزی سروش





فهرست مطالب

عنوان	صفحه
مدل ها و کاربرد آنها در فرآیند برنامه ریزی	۱.....
انواع مدل ها	۲.....
مدل فیزیکی	۲.....
مدل انتزاعی	۲.....
مدل ریاضی	۳.....
مدل های توصیفی	۳.....
مدل های پیش بینی کننده	۴.....
مدل های برنامه ریزی	۴.....
کارکردهای مدل	۵.....
ارزیابی یک مدل	۵.....
ویژگی مدل های شهری	۵.....
تقسیم بندی مدل های شهری	۶.....
انواع مدل ها در برنامه ریزی شهری	۷.....
مدل های اقتصاد شهری	۸.....
مدل های برنامه ریزی حمل و نقل	۹.....
مدل های عمومی تولید و توزیع سفر	۱۰.....
مدل های غیر قطعی تولید و توزیع سفر	۱۰.....
مدل های انتخاب وسیله نقلیه	۱۰.....
مدل های انتخاب مسیر	۱۱.....
مدل های جغرافیایی	۱۲.....
مدل های مکان یابی	۱۲.....
مدل های پوششی	۱۳.....
مدل رشد خطی	۱۴.....
مدل رشد نمایی	۱۵.....
مدل رشد نمایی تعدیل شده	۱۶.....



.....۱۸	مدل نمایی مضاعف
.....۲۰	مدل منحنی لجستیک
.....۲۰	مدل مقایسه ای
.....۲۲	مدل رگرسیون خطی ساده
.....۲۴	مدل رگرسیون چندمتغیره
.....۲۶	بررسی و طبقه بندی مدل ها در برنامه ریزی مسکونی
.....۲۶	مدل های ساده
.....۲۷	مدل های پیچیده
.....۲۸	روش های برآورد نیاز به مسکن
.....۲۸	هدف از برآورد نیاز به مسکن
.....۲۸	تقاضای مسکن
.....۲۹	تقاضای بالقوه و تقاضای موثر
.....۲۹	تولید مسکن
.....۳۱	روش تعیین کمبود واحد مسکونی
.....۳۲	تشریح روش های برآورد نیاز به مسکن
.....۳۲	روش انبوهه
.....۳۲	روش شاخص
.....۳۳	روش نرخ سرپرستی
.....۳۳	روش کل
.....۳۳	روش خام
.....۳۴	مدل لجستیک
.....۳۵	برآورد نیاز به زمین برای تأمین مسکن
.....۳۵	روش استفاده از تراکم متوسط ساختمانی
.....۳۶	روش استفاده از گروه نما
.....۳۶	روش استفاده از سرانه مسکونی و تراکم خالص
.....۳۷	روش برآورد مساحت واحدهای مسکونی مورد نیاز خانوارها از دیدگاه فرهنگی
.....۴۰	مدل جاذبه ای تک قیدی
.....۴۴	مدل جاذبه دو قیدی برای پراکنش سفر



صفحه	عنوان
۴۵	تئوری نقطه جدایی
۴۶	مدل پیش بینی فضاهای گذران اوقات فراغت
۴۷	مدل های دسترسی
۴۹	مدل جاذبه ای/ پتانسیلی هنسن
۵۴	مدل آنتروپی شانون
۵۸	ضریب آنتروپی
۶۵	روش تعیین نخست شهر
۶۶	مرورمختصری بر روش های ارزیابی در برنامه ریزی شهری و منطقه ای
۶۷	چارچوب مفهومی فرایند تحلیل سلسله مراتبی
۷۰	مراحل فرایند تحلیل سلسله مراتبی در تجزیه و تحلیل داده ها
۷۰	ساختن سلسله مراتبی برای تجزیه و تحلیل داده ها
۷۱	تبیین ضریب اهمیت معیارها و زیر معیارها
۷۳	تعیین ضریب اهمیت گزینه ها
۷۸	تعیین امتیاز نهایی (اولویت) گزینه ها
۷۹	بررسی سازگاری در قضاوت ها
۸۱	ریشه و منشا و تفسیر های اولیه از توسعه
۸۲	توسعه چیست
۸۳	ضرورت استفاده از برنامه ریزی منطقه ای
۸۴	فرآیند سنجش سطوح توسعه مناطق
۸۴	الف - تعیین هدف مطالعه و تدوین چارچوب آن
۸۵	ب - تعیین سطوح مطالعه
۸۵	ج- شناخت نوع آمار قابل دسترسی
۸۵	د - انتخاب شاخص های توسعه
۸۶	تشریح روش تاکسونومی
۸۷	مراحل روش تاکسونومی
۸۹	تشریح روش تحلیل عاملی
۹۳	هدف از تحلیل عاملی



صفحه	عنوان
۹۴	پیشینه مدل تحلیل عاملی
۹۴	مراحل اجرای تحلیل عاملی
۹۵	تهیه ماتریس همبستگی
۹۵	استخراج عامل ها
۹۶	چرخش عامل ها
۹۷	تفسیر
۹۸	معرفی نرم افزار SPSS
۹۹	نوارهای SPSS
۱۰۰	پنجره ویرایش گر داده ها (DATA EDITOR)
۱۰۸	خروجی های SPSS (OUTPUTS)
۱۱۲	اجرای تحلیل عاملی در نرم افزار SPSS
۱۱۲	محاسبه آماره های توصیفی
۱۱۴	استخراج عوامل
۱۱۶	دوران عامل ها
۱۱۸	محاسبه نمره های عامل ها
۱۲۲	منابع و مأخذ:



فهرست جداول

عنوان

صفحه

جدول شماره (۱)، داده های جمعیتی شهر A طی سال های ۱۳۸۰-۱۳۴۰	۲۳
جدول شماره (۲): محاسبه جمعیت شهر A	۲۴
جدول شماره (۳): امتیاز و متغیرهای برآورد مساحت مورد نیاز خانوارها از دیدگاه فرهنگی	۳۸
جدول شماره (۴): آمار داده های زیربنا و هزینه در مثال فرضی	۴۲
جدول شماره (۵): محاسبه مرحله اول مدل جاذبه تک قیدی	۴۲
جدول شماره (۶): محاسبه مرحله دوم	۴۳
جدول شماره (۷): نمایش ارقام SJ	۴۴
جدول شماره (۸): داده های جمعیت و اشتغال برای منطقه فرضی	۵۲
جدول شماره (۹): ماتریس مساحت/ زمان سفر	۵۳
جدول شماره (۱۰): محاسبه شاخص دسترسی برای هر ناحیه	۵۳
جدول (۱۱): محاسبه پتانسیل توسعه	۵۳
جدول شماره (۱۲): محاسبه پتانسیل نسبی توسعه	۵۴
جدول شماره (۱۳): پیش بینی جمعیت و تخصیص آن در هر یک از نواحی	۵۴
جدول شماره (۱۴): محاسبه آنتروپی شانون برای سال ۱۳۷۵ در شهر فرضی	۵۵
جدول شماره (۱۵): محاسبه آنتروپی شانون برای سال ۱۳۸۵ در شهر فرضی	۵۶
جدول شماره (۱۶): محاسبات تغییرات ضریب آنتروپی در طبقات شهری استان آذربایجان غربی	۵۹
جدول شماره (۱۷): اندازه واقعی و تئوری مرتبه- اندازه شهرهای آذربایجان غربی ۱۳۷۵	۶۰
جدول شماره (۱۸): لگاریتم رتبه- اندازه شهرهای استان آذربایجان غربی در سال ۱۳۷۵	۶۳
جدول شماره (۱۹): درجه نخست شهری در نظام شهری بر پایه شاخص چهارشهر	۶۶
جدول شماره (۲۰): مقایسه زوجی یا دوجه دویی ال ساعتی	۶۹
جدول شماره (۲۱): ماتریس ارزیابی برای مکانیابی مورد نظر	۷۵
جدول شماره (۲۲): شاخص تصادفی بودن (R.I.)	۷۹
جدول شماره (۲۳): انواع تکنیک های استخراج عامل ها را بر حسب اکتشافی- تاییدی و توصیفی- استنباطی	۹۳



مدل ها و کاربرد آنها در فرآیند برنامه ریزی

به دلیل تنوع و پیچیدگی مسائل شهری، توصیف، تحلیل و پیش بینی رفتار آنها، امری بس دشوار می باشد به همین جهت در قرن اخیر خصوصاً در دهه های نیمه دوم قرن حاضر، برنامه ریزان، به منظور کاستن از پیچیدگی سیستم و قابل کنترل نمودن آن، به استفاده از مدل های کمی و ریاضی توسل جسته اند. اساساً مدل نمادی از واقعیت است. مدل مهمترین ویژگی های وضعیت دنیای واقعی را به صورتی ساده و کلی بیان می دارد. و برداشتی از واقعیت است که برای توضیح مفاهیم و تقلیل پیچیدگی جهان به نحوی که قابل درک بوده و ویژگی آن را براحتی مشخص شود به کار می رود.

مدل ها ابزارهایی عملی هستند که می توان به کمک آنها به درکی از واقعیت البته نه کل آن بلکه بخش مفید و قابل فهم آن دست یافت.

مدل ها از آن جهت که در شرایطی که امکان تجربه به دلایل تکنیکی، اقتصادی، سیاسی و اخلاقی وجود ندارد درک چگونگی رفتار یک سیستم را میسر سازند حائز اهمیت و ارزش هستند از طریق مدل می توان درباره دنیای واقعی به پیشگویی دست زد از این رو مدل ها ابزار آلات تفکر و تعمق نامیده می شوند مدل توانایی پیشنهاد فرضیه جدید و فرموله کردن فرضیه ها را دارد. یک مدل به طور ساده توصیف شماتیک، اما دقیق از سیستم است که ظاهراً با رفتار گذشته آن انطباق دارد و بنابراین این امید وجود دارد که بتوان از آن برای پیش بینی رفتار آینده نیز استفاده کرد. مدل ممکن است بسیار ساده باشد و همچنین از لحاظ محاسباتی بسیار پیچیده باشد. مدل ها در واقع پلی میان سطح مشاهده و سطح تئوریک به شمار می آید. آنها از نقشی منطقی برخوردارند که عبارتند از بیان چگونگی عملکرد یک پدیده که به طور واقعی خنثی هستند یعنی فی نفسه نه درست اند و نه نادرست.

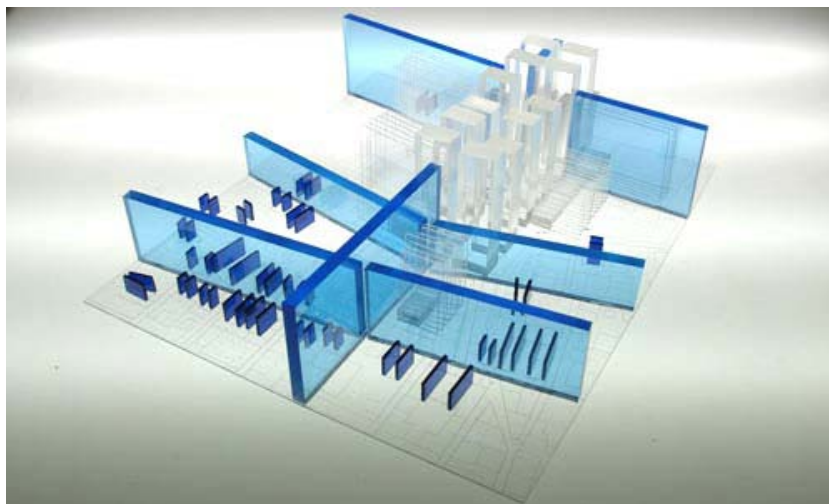
پیتر هاکت یکی از جغرافیدانان سرشناس، معتقد است دنیای واقعی خیلی پیچیده است، لذا جغرفیدانان برای سهولت در مطالعه به ساختن مدل ها پرداخته اند. در واقع مدل ها، نمایش های مطلوبی از واقعیت را بین بیان می دارند. تصاویر روشنی از اشکال ناهمواری های نواحی را بوجود می آورند. آنها بیان کننده پیچیده گی حقیقت هستند. اگر چه آنها نمی تواند تمام حقیقت را بین کنند اما ابزاری هستند که مسائل و قسمت های مفیدی از آن آشکار می سازند.

مدل تحلیلی در اصطلاح به یکی از مراحل انجام تحقیق تبیینی گفته می شود که در این مرحله رابطه میان متغیرهای وابسته و مستقل بررسی و ترسیم می شود. به منظور روشن تر شدن مفهوم مدل تحلیلی در ابتدا به بررسی انواع تحقیق و مراحل انجام آن که یکی از این مراحل، رسم مدل تحلیلی است، می پردازیم.

انواع مدل ها

مدل فیزیکی

ماکت کوچک شده از شیء مورد مطالعه می باشد. و به راحتی قالب درک بوده و اکثریت مردم با آن آشنا هستند. این نوع مدل ها ماکت کوچ شده شیئی مورد مطالعه می باشند به طوری که قادر به توصیف کامل رفتار سیستم مورد نظر برنامه ریزان نمی باشند.



مدل انتزاعی

مدلی است که به جای وسائل فیزیکی از نمادهای ریاضی برای نمایش موقعیت جهان واقعی و رفتار سیستم استفاده می نماید. زمانی که برنامه ریزان شهری توجه هاشان از جنبه های سه بعدی طراحی به نمایش روابط عملکردی و ، فرآیند بنیادی تحولات شهری معطوف گردد، مدل های انتزاعی از مدل های فیزیکی بسیار مفیدتر خواهند بود. مدل های فیزیکی قادر به توصیف کامل رفتار سیستم مورد نظر برنامه ریزان نمی باشند، در حالی که مدل های انتزاعی قادر به توصیف کامل رفتار این سیستم های می باشند. بنابراین مدل های انتزاعی با نمادی برای مقاصد برنامه ریزی از مهمترین نوع مدل های به شمار می روند.



مدل ریاضی

نوعی از مدل انتزاعی است که به صورت تجربه ایست که توسط یک عبارت جبری از متغیرهای ریاضی، پارامترها و ثابت ها، جایگزین اشیاء، نیروها، پدیده ها و غیره می شود. مدل های ریاضی یک جایگاه ویژه اصلی در تلاش های ما برای حل واقعیت های اجتماعی پیچیده دارند. مدل ریاضی در یک مسئله براساس تمام قوانین فیزیکی و هندسی حاکم بر مشاهدات و پارامترهای مسئله و آنچه تجربه نشان می دهد نوشته می شود (پرهیزکار، ۱۳۷۶: ۱۷).

مدل ها بر اساس اهداف و وظایف محوله آنها بایستی پاسخ گوی نیازهای مختلفی باشندلاری را بر آن داشت تا مدل های ریاضی را به سه دسته تقسیم کند. به گفته لاری مدل ها بر اساس خواسته های استفاده کنندگان و مدل ساز ممکن است به سه دسته تقسیم شوند. این دسته بندی مدل ها بترتیب پیچیدگی عبارتند از: مدل های توصیفی، مدل های پیش بینی کننده و مدل های برنامه ریزی.

مدل های توصیفی

مدل های توصیفی آنهایی هستند که بروی اوضاع موجود بیان یافته اند و تنها به بین وضع موجود پرداخته و با کم کردن تنوع روابط جهان واقعی و بیان آن به زبان ساده و به صورت روابط ریاضی، ابعاد ناشناخته پدیده ها را توصیف و نحوه تأثیر هر پدیده را بر دیگری بیان می کند.

هدف سازنده مدل توصیفی استفاده از کامپیوتر برای بازسازی ویژگیهای محیط شهری موجود و یا فرآیند تحولات شهری که قبلاً مشاهده گردیده می باشد. مدل های توصیفی مطلوب دارای ارزش علم می باشند زیرا این مدل ها با کم کردن تنوع روابط پیچیده جهان واقعی و بیان آن به زبان صریح و منسجم روابط ریاضی ابعاد ناشناخته ساختمان محیط شهری را توصیف می کند. این مدل ها نحوه تأثیر هر عنصر شهری بر دیگر عناصر شهری را بیان داشته و کمتر برنامه ریزی است که با نحوه کار چنین مدل هایی آشنا نباشد. این مدل ها همچنین می توانند ابزار کارآیی برای کار عملی باشند زیرا برای متغیرهایی که اندازه گیری آنها به سهولت میسر نیست می توان براساس اطلاعات موجود که مشتمل بر متغیرهای قابل اندازه گیری می باشند ارزش هایی معتبر استنتاج نمود.



مدل های پیش بینی کننده

مدل های پیش بینی کننده در عملکرد شبیه مدل های توصیفی است اما چون باید به به شبیه سازی آینده بپردازند به پیش بینی نیازهای دقیق تر و بیشتری نسبت به مدل های توصیفی دارند و در آنها به نقش صحیح پدیده علت و معلول توجه می شود. در مدل ها بایستی از امکان ارزیابی منطقی متغیرهای مدل تا زمان لازم در آینده اطمینان داشت.

در حالیکه مدل های شبیه ساز توصیفی قادر به استفاده از هر گونه رابطه ای و با هر روش برای توصیف یک وضعیت می باشند مدل های پیش بینی کننده تنها می توانند بر روابطی استوار باشند که در طی زمان بطور منطقی ثابت می ماند.

مدل های برنامه ریزی

مدل های برنامه ریزی یا مدل های نورمیتیو؛ این مدل ها بدنبال مدل های پیش بینی کننده بوجود آمده اند، اما تفاوت بسیار مهمی میان این مدل ها و مدل های پیش بینی کننده وجود دارد- این مدل ها به گونه ایی ساخته شده اند که فقط نتایج حاصله از فرضیاتی خاص را بیان نمی دارند، بلکه حیطه عملکرد ممکنه را در رابطه با اهداف مشخص تعیین می کنند. البته منظور این است که مدل برنامه ریزی باید در رابطه با هدفها و محدودیت های خاص ساخته شود. برای مثال می توان مدل را با اهدافی مانند ایجاد طرحی که کمترین هزینه را در بر داشته باشد، حداقل زمان سفر بکار را فراهم آورد و یا برای ایجاد حداکثر دسترسی به منطقه مرکزی شهر طراحی نمود. می توان محدودیت هایی مانند تراکم مجاز برای توسعه، مکان کاربری های خاص، میزان ترافیک ایجاد شده توسط یک منطقه یا حداکثر میزان سفر بکار را در ساختمان مدل های برنامه ریزی اعمال نمود. ستاده این مدل تا حد زیادی تجلی ایده آل های برنامه ریز است تا صرفاً پیش بینی ساده «نیروهای طبیعی» و به همین دلیل، این مدل نهایتاً با ارزش ترین نوع مدل محسوب می شود. هر چند تا به حال کاربرد مدل های نورمیتیو در برنامه ریزی محدودیت های تکنیکی شدیدی را دارا بوده است، ولی این محدودیتها کمتر به تکنیکهای صرفاً ریاضی راه حل مربوط بوده (هر چند تکنیکهای ریاضی پیچیده هستند، اما توسط کامپیوترهای کنونی قابل کنترل می باشند) و بیشتر در رابطه با تعیین ضوابط مطلوب برای استانداردها، هزینه ها و دامنه وسیعی از حالات و الگوهای رفتاری می باشند که باید در مدل گنجانده شوند.



کارکردهای مدل

برای مدل چهار کارکرد در نظر گرفته‌اند:

۱. سازماندهی؛ به معنای توانایی مدل در تنظیم و مرتبط ساختن داده‌ها و نشان دادن شباهت‌های میان داده‌ها است.
۲. اکتشافی؛ مدل‌ها ابزارهای اکتشافی می‌باشند که احتمالاً به کشف واقعیت‌ها و روش‌های ناشناخته منجر می‌شوند.
۳. پیش‌بینی کنندگی؛ یک مدل می‌تواند پیش‌بینی کند که چه پدیده‌های رخ خواهد داد. این پیش‌بینی‌ها در محدوده ساده بله یا خیر و تا پیش‌بینی‌های کمی مثل چه هنگام و چه اندازه هستند.
۴. سنجشی؛ داده‌هایی که توسط یک مدل به دست می‌آید می‌تواند مقیاسی برای سنجش یک پدیده باشد.

ارزیابی یک مدل

برای ارزیابی یک مدل باید به موارد زیر توجه کرد:

- الف) یک مدل تا چه اندازه کلی است؟
- ب) مدل تا چه اندازه سودمند است و چقدر در کشف روابط واقعیت‌ها یا روش‌های جدید مفید است؟
- ج) پیش‌بینی‌هایی که یک مدل به عمل می‌آورد چقدر دارای اهمیت است؟
- د) اندازه‌گیری‌هایی را که می‌توان با این مدل توسعه داد چقدر صحیح و درست است؟

ویژگی مدل‌های شهری

از نظر برنامه‌ریزان، مدلی مفید است که قادر به بازسازی پدیده یا مسایل موردنظر برنامه‌ریزی باشد چون رفتار سیستم رابطه نزدیکی با ساختار سیستم دارد لذا مدل‌های شهری موفق باید در خود ساختار سیستم شهری داشته باشند. ذیلاً ویژگی‌های اصلی مدل‌های شهری ذکر شده است.

الف- هر چند که در وهله اول مدلسازی ممکن است مدل دارای تعداد زیادی متغیر باشد، ولی با حذف تعداد بسیار زیادی از این متغیرها می‌توان پیچیدگی مسئله واقعی را از بین برد و تجزیه و تحلیل‌های ناشی از حل مدل را آسانتر انجام داد، ضمن آنکه کاهش متغیرها، باعث تسهیل جمع‌آوری اطلاعات می‌شود



که همواره در مدلسازی این مشکل وجود داشته است. در واقع کاهش تعداد متغیرها در مدل، به تحلیل گر سیستم این امکان را می دهد که بتواند بطور نسبتا ساده و دقیق تجزیه تحلیلهای مختلفی انجام دهد.

ب- بسیاری از متغیرها و روابطی که این متغیرها در مدل های برنامه ریزی شهری با یکدیگر دارند اساسا بصورت خطی هستند و یا برمبنای برآوردهای خطی بنا شده اند چرا که تجزیه و تحلیلهای ریاضی قادر به ارائه راه حل کلی برای سیستم های غیر خطی نیستند.

ج- مدل های شهری فقط بخشی از سیستم شهری را به نمایش می گذارند چرا که ساختن مدلی که بتواند زیر سیستم های آن را با یکدیگر پیوند دهد، بسیار مشکل است.

د- مدلهایی که مورد بحث قرار خواهند گرفت از نوع مدل های ایستا هستند. به عبارت دیگر چنین می توان گفت که عنصر زمان در این مدل ها نقشی ایفا نمی کنند.

تقسیم بندی مدل های شهری

در سیستم های شهری مدل های فراوانی وجود دارند که آنها را می توان اساسا به اقسام زیر تقسیم بندی نمود.

۱- مدل های اقتصاد شهری

۲- مدل های حمل و نقل

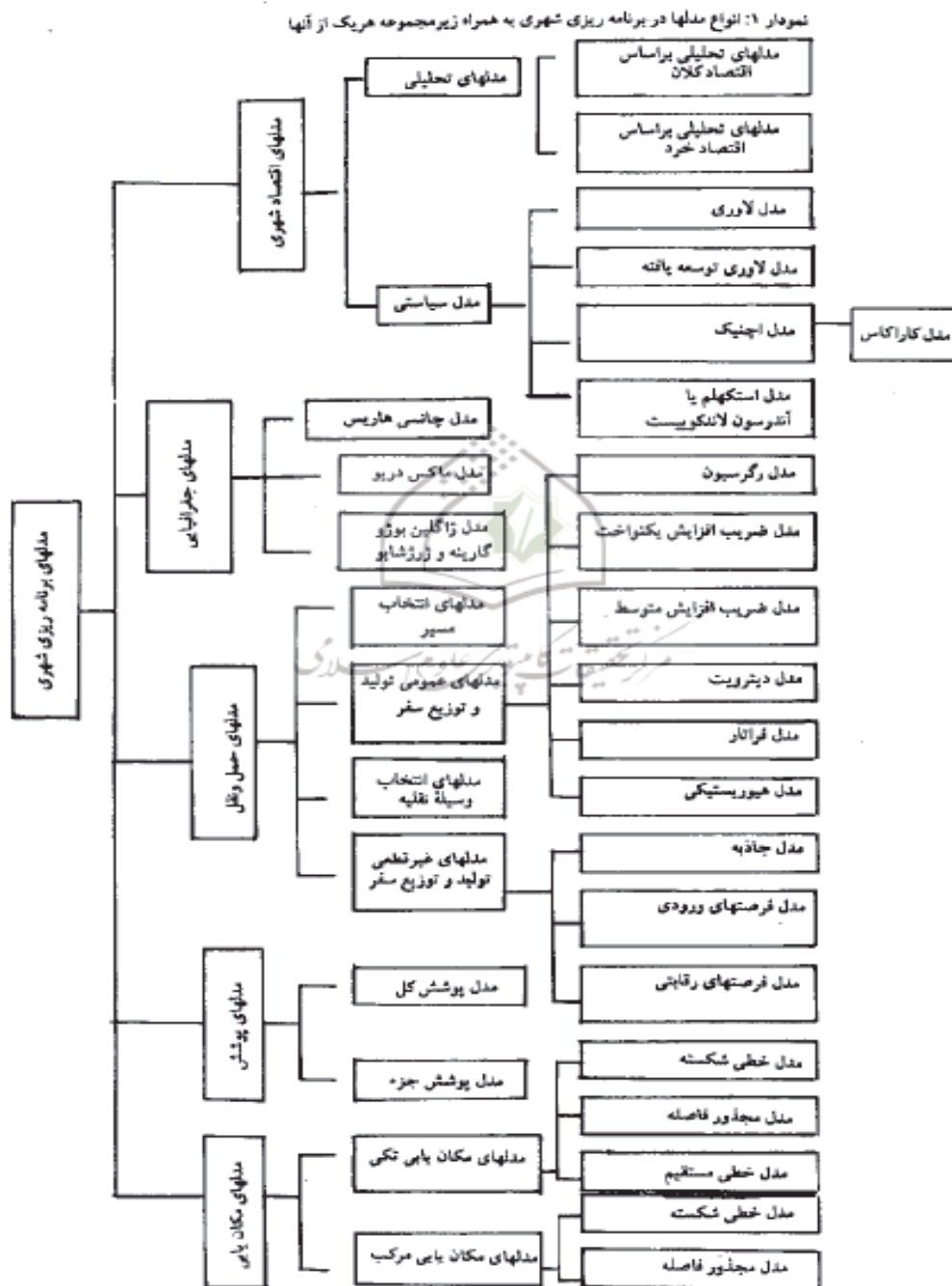
۳- مدل های جغرافیایی

۴- مدل های مکان یابی

۵- مدل های پوشش

علاوه بر مدل های فوق، مدل های دیگری در زمینه مسایل شهری وجود دارد.

در زمینه مدل های شهری، توجه به این نکته لازم است که اکثر این مدل ها در حالتی ایستا به شهر نگاه کرده و لذا این مدل ها، فاقد پویایی بین متغیرها هستند همچنین باید توجه داشت که در مدل های شهری، آمار و اطلاعات نقش عمده ای را ایفا می کنند و بنابراین مدلهایی قدرت کاربرد بیشتری دارند که بتوانند از آمار و اطلاعات موجود استفاده کرده و بهترین تجزیه و تحلیلها را انجام دهند.



انواع مدل ها در برنامه ریزی شهری

همانگونه که در قسمت ۲ عنوان شد مدل های برنامه ریزی شهری به ۵ دسته کلی تقسیم می شوند. نمودار ۱، ارتباط کلی این مدل ها را با زیرمجموعه های آنها نشان می دهد.



در این مقاله مدل های سطح یک نمودار شماره ۱ که شامل مدل های اقتصادی شهری، مدل های جغرافیایی، مدل های حمل و نقل، مدل های پوشش و مدل های مکان یابی می گردد به طور اختصار توضیح داده می شود و زیرمجموعه های آنها مورد بحث قرار نمی گیرد. علاقمندان می توانند برای جزئیات بیشتر به منابع و مآخذی که در انتهای این مقاله ارائه شده است مراجعه کنند.

مدل های اقتصاد شهری

به منظور توسعه و رشد شهرها، لازم است با توجه به تواناییهای بالقوه اقتصادی مطالعات اقتصادی همه جانبه ای انجام شود تا توسعه بهینه شهرها انجام پذیرد. بدیهی است که عوامل اقتصادی، زمینه اصلی ظهور و توسعه واحدهای جمعیتی را تشکیل می دهند که خود باعث توسعه شهرها می شوند. در تعیین چگونگی اثرگذاری تحولات اقتصادی به شکل گیری شهرها، برنامه ریزی شهری نه تنها باید اینگونه اثرات را پیش بینی کند بلکه باید میزان و اندازه اثرها را بشناسد.

برای شناسایی این اثرات، برنامه ریزی شهری به مدل های اقتصاد شهری نیاز دارد. در مورد مدل های اقتصاد شهری باید توجه داشت که روابط بین متغیرها نسبتاً پیچیده است. در واقع هدف مدل های اقتصادی، ساده کردن روابط، کمی کردن و قابل اندازه گیری نمودن آنها و بزبان ریاضی بیان کردن روابط اقتصادی در درون یک محدوده شهری می باشد. توجه به این نکته ضروری است که با رشد و توسعه اقتصادی شهرها پرجمعیت تر، و بزرگتر می شوند. اندازه شهر و رشد و توسعه اقتصادی آن عوامل موثری در تعیین مدل مناسب برای تحلیل ساخت اقتصادی شهر و تخمین دگرگونیهای آینده می باشد.

همانگونه که در نمودار شماره ۱ ملاحظه می شود، مدل های اقتصاد شهر به دو دسته مدل های تحلیلی و مدل های سیاسی تقسیم می شوند. در مدل های تحلیلی، چگونگی رشد و توسعه شهر در گذشته مورد تجزیه و تحلیل واقع می شود و با پیش بینی چگونگی رشد شهر در آینده، به کار خود خاتمه می دهد. مدل های سیاستی بیشتر در زمینه هایی نظیر کاربری زمین، خانوارهای ساکن و دسته بندی آنها برحسب گروههای اقتصادی - اجتماعی و تأثیر دگرگونیهای شبکه حمل و نقل بر سیستم شهری بکار برده می شود. مدل ها سیاستی بسیار پیچیده بوده و محصول این مدل ها، برنامه ریزی توزیع فعالیتها در شهر برحسب مناطق مختلف شهر است.



مدل های برنامه ریزی حمل و نقل

جابجایی شدید و گروهی جمعیت شهری سبب تراکم و تاخیر در سفر و زمان سفر می شود. این جابجاییها به سبب جدایی جغرافیایی سکونتگاه و محل کار نیروی انسانی در شهرهای بزرگ می باشد. شکل کاملاً کلاسیک این جابجاییها و حرکات، مربوط به حرکت روزانه کارمندان ادارات و کارکنان شرکتها و واحدهای تجاری در مناطق مختلف شهر است.

جابجایی روزانه جمعیت شهری به ویژه طبقه فعال بزرگ یک حرکت موجی است که برای مقاصد مختلف زیر انجام می شود:

- جابجایی جمعیت فعال بین محل زندگی و محل اشتغال.

- جابجایی جمعیت برای خرید و انجام کارهای شخصی.

- جابجایی جمعیت برای گذراندن اوقات فراغت، گردشها، دیدوبازدیدها و غیره.

- جابجاییهایی که در محدوده و طبقه بندی فوق قرار نمی گیرند.

جابجایی روزانه کارمندان و کارکنان شرکتها، شکل ساده ای دارد. در هنگام صبح جهت حرکت رو به مرکز شهر و به هنگام عصر، جهت حرکت رو به خارج شهر است. این نوع جابجایی با ساعات کار ادارات، فروشگاهها و کوتاه و طولانی بودن روز با ساعات اشتغال هماهنگ است. در این حرکتها، عده ای بصورت انفرادی و با وسایل شخصی و گروهی به طور جمعی با استفاده از وسایل نقلیه عمومی جابجا می شوند. در این حرکات، علاوه بر طبقه فعال جامعه، دانش آموزان، دانشجویان دانشگاهها و مؤسسات آموزش عالی نیز شرکت دارند. حرکت روزانه کارگران از نظر جغرافیایی متفاوت است چرا که اصولاً سکونتگاههای کارگران در حول و حوش مناطق صنعتی قرار دارد و یا این سکونتگاهها، در فاصله کمی از واحد تولیدی استقرار یافته اند.

به هر شکل، تمام مطالب ذکر شده فوق، دلالت بر این مطلب دارد که وجود جابجاییها (بخصوص جابجایی با وسایل نقلیه) باعث اشغال حجم خیابنها و معابر می شود و با وجود تعداد معابر محدود، بدیهی است که ترافیک ایجاد می گردد. بدون تردید حمل و نقل شهری و ترافیک در زندگی مامی افراد ساکن در شهرها به نوعی تأثیر بر جای می گذارد.

اساساً مدل هایی که در برنامه ریزی حمل و نقل وجود دارند براساس موارد زیر تشکیل شده اند:



الف-براساس سه منظور:سفر به قصد کار از مبدأ خانه،سفر به قصد غیر کار از مبدأ خانه و سفر از مبدأ غیر خانه.

ب-براساس دو نوع سیستم حمل و نقل:سیستم حمل و نقل عمومی و سیستم حمل و نقل شخصی.

ج-براساس دو گونه زمان:ساعات تراکم و ساعات غیر تراکم.

د-براساس دو نوع رفت و آمد:رفت و آمد در حوزه مطالعاتی و رفت و آمد در خارج از حوزه مطالعاتی.همانگونه که از نمودار شماره ۱ مشاهده می شود،مدل های حمل و نقل به چهار دسته کلی تقسیم می شوند که ذیلا هریک از این مدل ها به اختصار توضیح داده می شود.

مدل های عمومی تولید و توزیع سفر

معمولترین روش مدلسازی در نظر گرفتن منظور از سفر است.الگوی سفرها با وسیله نقلیه که برای آنها مدلسازی انجام می شود برحسب منظور از سفر،به دو دسته تقسیم می شود.

الف-سفرهایی که مبدأ حرکت آنها خانه است.

ب-سفرهایی که مبدأ حرکت آنها،جایی بجز خانه است.

کلا مدل های عمومی تولید و توزیع سفر به ۵ مدل تقسیم می شوند که در نمودار ۱ آنها را ملاحظه می کنید.

مدل های غیر قطعی تولید و توزیع سفر

اساس مدل های غیر قطعی تولید و توزیع سفر براساس رابطه زیر قرار داد:

T-PB

در رابطه فوق، T یک ماتریس مربعی $n \times n$ با عناصر p, t_{ij} یک ماتریس قطری $n \times n$ است که عناصر آن مشخصات تولید سفر منطقه و B نیز یک ماتریس $n \times n$ با عناصر احتمالی b_{ij} است که این عناصر معرف احتمال سفر از منطقه i به منطقه j می باشد.از آنجا که تولید و جذب سفرها در مجموع مناطق باید باهم برابر باشند لذا تعدیل سازی اینگونه مدل ها ضروری است.

مدل های انتخاب وسیله نقلیه

در برنامه ریزی حمل و نقل یک شهر،بعد از تعیین میزان تولید و جذب سفر در هر یک از مناطق،باید به بررسی تعیین و انتخاب وسیله نقلیه پرداخت.در حقیقت مطالعات تفکیک وسایل سفر،سفرها را



برحسب نوع وسیله نقلیه مورد بررسی قرار می دهد. با استفاده از این مدل ها، درصد سفرهایی که توسط یکی از وسایل حمل و نقل صورت می پذیرد، تعیین می شود.

انتخاب بین انواع وسایل نقلیه برطبق این مدل ها، بستگی به چهار متغیر دارد. به عبارت دیگر، مسافران به منظور انتخاب نوع وسیله نقلیه، معمولاً چهار عامل زیر را در نظر می گیرند.

الف- مدت زمان نسبی سفر

ب- هزینه نسبی سفر

ج- موقعیت اقتصادی سفرکننده

د- میزان نسبی راحتی در سفر

علاوه بر عوامل فوق الذکر، عوامل دیگری نیز هستند که در تصمیم گیری سفرکننده جهت انتخاب نوع وسیله، موثر می باشند. از جمله این عوامل، طول مسیر سفر، تراکم جمعیت و تراکم استخدام را می توان نام برد.

با توجه به زمانهای مختلفی که سفرکننده در هر مقطع از سفرش صرف می کند و هزینه مربوطه بعلاوه موقعیت اقتصادی شخص و میزان راحتی سفر، وی می تواند با مدل هایی که وجود دارد وسیله نقلیه مورد نظر خود را انتخاب کند.

مدل های انتخاب مسیر

منظور از مدل های انتخاب مسیر، مدل هایی است که بتواند درصد سفرهایی را که از هر مسیر بین دو منطق وجود دارد تعیین کند. عمده ترین عواملی که در انتخاب هر مسیر توسط شخص سفرکننده در نظر گرفته می شود عبارتند از:

الف- مدت زمان سفر

ب- هزینه سفر

د- سطح سرویس (نسبت حجم ترافیک عبوری به ظرفیت مسیر)

از بین چهار عامل فوق، مدت زمان سفر نقش عمده ای را در مدل های انتخاب مسیر ایفا می کند. باید توجه داشت که سه عامل دیگر نیز عمدتاً بستگی به مدت زمان سفر دارند و از این رو، با در نظر گرفتن مدت زمان سفر، سایر عوامل نیز در نظر گرفته خواهند شد.



مدل های جغرافیایی

بررسیهای کلاسیک در قلمرو شناخت کار و نقش شهرها، عنوانی است که می تواند در این قسمت مورد ملاحظه قرار گیرد. روشها و مدل های پیشنهادی در این قسمت، مبتنی بر اصول ساده ای و در سطح بررسی رابطه ها، درصدها و منحنی های توزیع خلاصه می شود. در واقع داده ها باعث می شوند تا شناختی از الگوها و سیستم هایی که در یک شهر وجود دارند بدست آوریم و با توجه به آنها، مدل هایی را پیشنهاد کنیم. مدل های جغرافیایی اساسا مدل هایی می باشند که به تحلیل ساختار فعالیت های مختلف یک شهر با توجه به آمار و اطلاعاتی که وجود دارد، می پردازند.

مدل های مکان یابی

یکی از نخستین هدف های سیاست طرح ریزی سرزمین یک کشور در اکثر نقاط دنیا، جلوگیری از تمرکز مردم در یک شهر است، زیرا تمرکز اکثریت بیش از اندازه، باعث می شود تا توسعه شهرهای دیگر به مخاطره افتاده و یا اصلا انجام نشود. در بسیاری از ممالک، تضمین متعادل بین تمام شهرهای یک کشور، هدف اصلی به شمار می آید. در برخی از شهرهای بزرگ که به سبب صنعتی شدن با سرعت رشد می یابند، همکاری تنگاتنگی از لحاظ تبادل نیروی انسانی بین چند شهر ایجاد می شود و در نتیجه، جابجایی روزانه انسانها از مرز یک شهر فراتر می رود و حرکت روزانه جمعیت، شکل پیچیده تری به خود می گیرد. چنین مسایلی تنها بر دوش یک شهر معین سنگینی نمی کند بلکه پدیده ای است که در شهرهای مهاجرپذیر اتفاق می افتد. برای اینکه اینگونه حرکتها را محدود ساخت و در سطح معینی مهار کرد اقدام به ایجاد شهرک هایی می شود که محل کار و سکونت در فاصله دوری نسبت به یکدیگر قرار نداشته باشند. در این شهرکها، مکانهای سکونت تقریبا یکسان هستند و تنها شرایط مربوط به دسترسی آسان به وسایل حمل و نقل عمومی است که باعث می شود تا مکانی نسبت به مکان دیگر دارای ارجحیت باشد.

اولین مسئله ای که در ایجاد اینگونه شهرکها مد نظر قرار می گیرد انتخاب مکان آن می باشد. برای احداث شهرکها، عامل اساسی در انتخاب، مکانی است که خدمات را با کیفیتی درخور توجه ارائه دهد.



آنچه که در مدل های مکان یابی مورد بحث قرار می گیرد بیشتر فرم ریاضی مدل های مکان یابی است. اگر مبنا هزینه های رفت و آمد شهروندان به محل کار در نظر گرفته شود، می توان مکان بهینه شهرکها را تعیین کرد.

اینگونه مدل ها به دو دسته اساسی مدل های مکان یابی تکی (که استقرار یک شهرک را شامل می شود) و مدل های مکان یابی مرکب (که استقرار چند شهرک را شامل می شود) تقسیم می شوند و هرکدام از این مدل ها به نوبت خود به روشهای مختلف تقسیم می گردند. در این روشها مسافت بین مکان شهرک و محل اشتغال شاغلین، عامل اساسی است. این فواصل می توانند بصورت خطی شکسته (خطوط عمود بر هم مانند خیابانها و کوچه ها)، بصورت خط مستقیم (خطوطی که دو نقطه را بطور مستقیم به هم وصل می کنند) و بصورت مجذور فاصله در نظر گرفته شوند که این مورد آخری، در واقع روشی است که برای یافتن مکان بهینه، اگر مسافت بطور خط مستقیم باشد، بعنوان یک کاتالیزور بکار برده می شود.

در این مدل ها، فرض می شود که هزینه حمل و نقل بین شهرک و محل کار کارکنان، با بعد مسافت نسبت مستقیم دارد. در واقع، محل بهینه استقرار شهرک، بگونه ای تعیین می شود که کل مسافت پیموده شده بین مکان شهرک و محل کار شاغلین، حد اقل شود.

مدل های پوششی

از مدل های پوششی بیشتر به منظور انتخاب یک مکان از چند مکان موجود برای استقرار خدمات شهری استفاده می شود. برای مثال می توان از مکان استقرار مراکز تحویل کالا، شبکه های استقرار توزیع کالا، مکان یابی انبارها، مکان یابی ایستگاههای اورژانس و ایستگاههای آتش نشانی نام برد. یکی از مسایلی که می تواند در رابطه با مدل های پوششی مطرح شود مسئله مشخص کردن تعداد و مکانهای شهرکهای دانشجویی است که در آن دانشجویان بتوانند با استفاده از وسایل حمل و نقل عمومی، مثلاً خود را یک ساعته به محل دانشگاه برسانند و یا مسئله استقرار ایستگاههای آتش نشانی می باشد، به صورتی که مثلاً ۱۵ دقیقه بیشتر طول نکشد تا از کلیه نقاطی که ایستگاههای آتش نشانی قرار دارند به محل آتش سوزی رسید.

مدل های پوششی خود زیرمجموعه ای از مدل های برنامه ریزی صفر و یک می باشند. برای روشن ساختن این مدل ها، مدل برنامه ریزی صفر و یک زیر را در نظر بگیرید.



مقادیر aij در مدل فوق، ضرایب پوششی نام دارند و مقادیری که می‌توانند به خود اختصاص دهند، فقط صفر یا یک می‌باشند. اگر مشتری i توسط مکان j تحت پوشش قرار گیرد aij برابر یک و در غیر این صورت مقدار آن صفر خواهد بود. بعلاوه، اگر تسهیل به مکان j اختصاص یابد برابر عدد یک است و اگر غیر از این باشد مقدار آن صفر خواهد بود. هر یک از m مشتری، حد اقل توسط یکی از n تسهیلات که در محدودیتهای مدل فوق وجود دارد تحت پوشش قرار می‌گیرد. تابع هدف، تابعی است که مشتریها را تحت حد اقل هزینه پوشش می‌دهد و Cj تخصیص یک تسهیل به مکان j است.

برای روشنتر شدن معنی پوشش، تصور نمایید مشتریها باید توسط مکان j که محل استقرار ایستگاه آتش‌نشانی است و در فاصله ۵ دقیقه‌ای از مکان i قرار دارند تحت پوشش واقع شوند. مثال دیگر، می‌توان تسهیلات را بصورت کارخانه‌هایی فرض کرد که حد اقل یک کارخانه، باید مشتری i را مورد پوشش قرار دهد و یا به مشتری i خدمات ارائه کند اگر کارخانه در مکانهای ۱، ۲، یا ۳ قرار گیرد. بنابراین $ai1-ai2-ai3-1$ و مقادیر دیگر aik با شرط $k=1,2,3$ برابر صفر باشند.

همانگونه که از نمودار شماره ۱ مشاهده می‌شود، مدل‌های پوششی را می‌توان به مدل پوششی کل و مدل پوشش جزء تقسیم کرد. منظور از مدل پوشش کل، مشخص کردن تعداد حد اقل تسهیلات مورد نیاز است تا کل مشتریان پوشش داده شوند در حالیکه در مدل پوشش جزء، هدف مشخص کردن حداقل تعداد تسهیلات بگونه‌ای است که تعداد بیشتری از مشتریها تحت پوشش قرار گیرند و علت این امر، آن است که در این مدل، تعداد تسهیلات به اندازه‌ای نیست که بتواند به تمام مشتریان سرویس ارائه دهد.

مدل رشد خطی

این مدل، الگویی از رشد جمعیت را توصیه می‌کند که در آن میزان جمعیت همچنان با نرخ قبلی خود تغییر می‌کند. به همین ترتیب میزان تراکم جمعیت متناسب با زمان، افزایش یا کاهش خواهد داشت. اگر میزان افزایش یا کاهش جمعیت در طول برابر با مشخصه a باشد، متعاقب آن، سطوح جمعیتی $P_1, P_2, \dots, P_a$ در سال اول، دوم، ... و a ام برابر خواهد بود با:

$$P_1 = P_0 + a$$

$$P_2 = P_1 + a = (P_0 + a) + a = P_0 + 2a$$



$$P_3 = P_2 + a(P_2 + 2a) + a = P_2 + 3a$$

$$P_n = P_{n-1} + a = P_2 + (n-1)a + a = P_2 + na$$

بنابراین میزان جمعیت در زمان n برابر خواهد بود با:

$$P_n = P_2 + na \quad (۱)$$

P_n : میزان جمعیت در زمان n

P_2 : میزان جمعیت پایه،

a : میزان رشد در واحد زمان،

n : دوره زمانی برحسب (ماه، سال، نیمسال و ...)

مثال: اگر شهری در حال حاضر ۵۰۰۰ نفر جمعیت داشته باشد، و در عرض ۵ سال گذشته سالانه

۵۰۰۰ نفر افزایش یافته باشد، مطلوب است جمعیت این شهر را برای ۱۰ سال آینده محاسبه کنید؟

$$P_{10} = 5000 + 10 \times 5000 = 55000 \text{ نفر}$$

مدل رشد نمایی

در روش قبلی فرض بر این بود که تغییرات در سطوح جمعیتی متناسب با زمان سپری شده است. اما

در این روش مقدار افزایش جمعیت متناسب با میزان جمعیت موجود است. به طوری که نسبت بین

افزایش جمعیت و جمعیت کل، ثابت است ولی افزایش صعود می کند. به زبان ریاضی:

$$P_{t+n} = P(t)(1+r)^n$$

$P(t+n)$: جمعیت در سال $t+n$ (پایان دوره)،

$P(t)$: جمعیت در سال t (آغاز دوره)،

n : دوره زمانی (برحسب مال، سال، نیمسال و ...)،

r : نرخ رشد جمعیت سالانه

به طوری که r یعنی نرخ رشد سالانه جمعیت از طریق فرمول ذیل به دست می آید:

$$r = \sqrt[n]{\frac{P(t)}{P_0}} - 1 \times 100$$



مثال: اگر جمعیت شهر A در سال ۱۳۷۰ برابر ۲۰۰۰ نفر و در سال ۱۳۸۰ برابر ۳۲۰۰۰ نفر بوده است، مطلوب است جمعیت این شهر با استفاده از مدل رشد نمایی برای سال ۱۳۹۰ محاسبه نمایید.

$$r = \sqrt[10]{\frac{32000}{20000}} - 1 \times 100 = 4/8 \text{ سال}$$

$$P_{1390} = 32000 \left(1 + \frac{4/8}{100}\right)^{10} = 51140 \text{ نفر}$$

مدل رشد نمایی تعدیل شده

کاربرد مدل های رشد خطی و نمایی را می توان در موارد مختلفی تجسم نمود، لیکن از مشکلات بزرگ و محدودیت های اصلی این دو مدل آن است که در آن جمعیت قابلیت افزایش بی نهایت یا کاهش بی نهایت را دارد. فرض اصلی این مدل آن است که باقیمانده رشد جمعیت (تفاضل بین سطح جمعیت آخرین و سطح جمعیت فعلی) تابعی از نسبت ثابتی است از این تفاضل در یک دوره قبل. فرمول ریاضی آن عبارت است از:

$$P_n = P_\infty - V^n (P_\infty - P_0)$$

V: مقدار افزایش یا سرعت جمعیت،

P_∞ : تعداد جمعیت مطلوب،

P_0 : جمعیت حال حاضر،

P_n : جمعیت مورد انتظار،

N: تعداد سال برای جمعیت مورد انتظار

در این مدل پیش بینی جمعیت برای دوره n را از طرق تفاضل بین جمعیت نهایی یا (مطلوب) یا نسبت نهایی کاهش یا بنده ای از نرخ رشد کلی، به دست می آوریم. این نسبت قوه n عامل V می باشد. در نتیجه، هر قدر V بیشتر و بزرگ تر باشد، افزایش جمعیت کندتر می شود. نهایتاً درحاتی که $V=1$ باشد، سطح جمعیت ثابت مانده و هیچگونه افزایشی نخواهد داشت. برعکس اگر $V=0$ گردد سطح جمعیت فوراً و یک دفعه برابر با جمعیت نهایی یا مطلوب می شود. (در طی یک دوره زمانی).



مثال: اگر مقدار جمعیت مطلوب شهری ۲۰۰۰۰ نفر باشد و جمعیت کنونی آن ۱۰۰۰۰ نفر باشند، با اتخاذ سیاست منطقه ای، به طوری که این شهر طوری مهاجرپذیر باشد که هر ۱۰ سال فقط $\frac{1}{4}$ رشد سالانه قبلی به این شهر اضافه شود، در آن صورت چند سال طول می کشد جمعیت این شهر به ۱۵۰۰۰ نفر برسد.

$$P_{\infty}, 20000 :$$

$$P_n, 10000 :$$

$$10000 = 20000 - \frac{1}{4} (20000 - 10000)$$

$$P_{\infty} = 20000$$

$$\frac{1}{4} = V$$

$$15000 = 20000 - (0.25)^n (10000)$$

می توان معادله را به صورت ذیل نوشت:

$$0.25^n (10000) = 20000 - 15000$$

در آن صورت خواهیم داشت:

$$0.25^n (10000) = 5000$$

طرفین به ضریب مجهول تقسیم شده در آن صورت خواهیم داشت:

$$0.25^n = 0.5$$

دراین صورت می توانیم بنویسیم:

$$n = \frac{\log 0.5}{\log 0.25} \times 10$$

$$n = 5$$

$$P_1 = 20000 - (0.25)^1 (20000 - 10000) = 17500 \text{ نفر}$$

$$P_2 = 20000 - (0.25)^2 (20000 - 10000) = 19375 \text{ نفر}$$



$$P_7 = 20000 - (0/25)^7 (20000 - 10000) = 20156 \text{ نفر}$$

جمعیت در ۱۰ سال بعد

مدل نمایی مضاعف

مدل دیگری که بر مبنای یک سطح نهایی محدود جمعیتی استوار است، مدل نمایی مضاعف است:

$$P_t = P_\infty \cdot a^{b^t} \quad (1)$$

این مدل بر این فرض متکی است که نرخ رشد جمعیت متناسب با میزان جمعیت، همانند مدل نمایی فوق ولی با یک فاکتور نسبی که به جای ثابت بودن به صورت نمایی همراه با زمان، افزایش پیدا می کند. این امر، کندتر کردن نرخ رشد را در جهت محدودیت عنوان شده برای میزان جمعیت، فراهم می سازد. مزیت اصلی این مدل در این است که پس از تبدیل مناسب در واحدهای اندازه گیری، میزان جمعیت به صورت یک مدل خطی ساده از رشد، درمی آید. درواقع، اگر از طرفین معادله لگاریتم گرفته شود، خواهیم داشت:

$$\log P_t = \log P_\infty + b^t \log a$$

$$\log \left(\frac{P_\infty}{P_t} \right) = b^t \log a = b^t \log \frac{1}{a}$$

و اگر مجدداً لگاریتم گرفته شود:

$$\log \log \left(\frac{P_\infty}{P_t} \right) = \log \log \left(\frac{1}{a} \right) + t \log b \quad (2)$$

بدین ترتیب با استفاده از لگاریتم های مضاعف (P_∞ / P_t) و معکوس مشخصه a ، مدل به صورت یک مدل خطی مبدل می گردد. به هر حال، این جنبه از مدل نمایی مضاعف ما را قادر می سازد که معنی فیزیکی متغیرهای a ، b آن را مشخص کنیم. درواقع در زمان $t=0$ داریم:

$$\log \log \left(\frac{P_\infty}{P_0} \right) = \log \log \left(\frac{1}{a} \right) \quad \frac{P_\infty}{P_0} = \frac{1}{a}$$

یا سرانجام:

$$a = \frac{P_0}{P_\infty}$$

بنابراین، a نمایانگر نسبت میزان جمعیت پایه به میزان نهایی است.



همچنین $\log$ مشخصه b شیب خط مستقیم است، یعنی، نرخ تغییرات واحد جدید اندازه گیری میزان جمعیت: $\log \log \left(\frac{P_{\infty}}{P_t} \right)$ نسبت به زمان t ، لزوم یک سطح نهایی محدود برای این امر دلالت دارد که به سبب آن، $a = P_0 / P_{\infty}$ کوچک تر از یک است، b نیز باید کوچک تر از یک باشد. در واقع، در زمان $t = \infty$ ، P_t برابر با P_{∞} می گردد، یعنی $P_t = P_{\infty} = P_{\infty} a^{b^{\infty}}$ یا $1 = a^{b^{\infty}}$. بنابراین $b < 1$ است و در غیر این صورت b^{∞} بسیار بزرگ خواهد شد و P_{∞} نامحدود می گردد. اگر $b < 1$ باشد، b^{∞} مساوی صفر خواهد شد و $1 = P_{\infty} a^{b^{\infty}} = P_{\infty} a^0 = P_{\infty}$ و قید میزان محدود P_{∞} رعایت می گردد. از آنجا که b کوچک تر از یک است، لگاریتم آن منفی می باشد و شیب خط مستقیم به پایین خط خواهد بود.

به عنوان مثال: فرض می کنیم که می خواهیم یک مدل نمایی مضاعف را برای سیر تغییرات یک سطح جمعیتی که از مقدار پایه ۸۰۰۰۰ نفر شروع و سرانجام در مقدار ۲۰۰۰۰۰ نفر ثبت می شود، به کار ببندیم. این مدل به شکل $P_t = P_{\infty} = P_{\infty} a^{b^{\infty}}$ خواهد بود. با این حال باید آن را برای نیازهای خود سازگار کنیم، یعنی مشخصه های b, a مدل را تخمین زنیم. معادله (۳) نشان می دهد که a برابر است با P_{∞} / P_0 بنابراین:

$$a = \frac{80000}{200000} = 0.4$$

مقدار b از معادله (۲) بدست خواهد آمد. هرچند به سبب این که معادله وابسته به زمان است، نیاز به تعیین جمعیت در یک زمان معین داریم. بدین ترتیب، فرض کنید که انتظار داریم میزان جمعیت پس از ده سال به سطح ۱۵۰۰۰۰ نفر برسد.

$$\log \log \frac{200000}{150000} = \log \log \left(\frac{1}{0.4} \right) + \log \log b$$

$$-0.9 = -0.4 + \log \log b$$

$$\log b = \frac{-0.9 + 0.4}{10}$$

$$\log b = -0.5$$

$$b = 0.89$$

بنابراین، مدلی که منطبق بر شرایط موردنظر ماست، به صورت ذیل است:



$$P_t = (200000) \cdot e^{(0.08)t}$$

بنابراین برآورد جمعیت در سال موردنظر را می توان با جایگزین آن در رابطه فوق به دست آورد.

مدل منحنی لجستیک

در مدل های رشد نمایی فرض اصلی این بود که نرخ رشد جمعیت نسبت به سطح جمعیت ثابت است.

اگر این نرخ رشد را به جای ثابت بودن آن، تابع خطی کاهش یابنده ای از سطح جمعیت (برای مثال به

صورت $a-bP_t$) فرض نماییم نتیجه مدل لجستیک خواهد بود. یعنی:

$$P_t = \frac{1}{\left(\frac{1}{P_0} - \frac{b}{a}\right)e^{-at} + \frac{b}{a}}$$

P_t : برآورد جمعیت در سال t ،

P_0 : جمعیت سال پایه (اولیه)،

a, b : ضرایبی هستند که در فرمول رشد مضاعف طریقه محاسبه آن توضیح داده شده است،

C : عددی نپری برابر 2.718

در صورت نقطه گذاری اعداد در محور مختصات جمعیت و زمان، شکل این تابع به صورت "S" درمی

آید. در این مدل سطح نهایی جمعیت برابر با $\frac{a}{b}$ خواهد شد و درحالتی که $b=0$ باشد، همان مدل رشد

نمایی ساده به دست می آید.

مدل مقایسه ای

این مدل مبنای پیش بینی جمعیت را یک منطقه خاص قرار نمی دهد بلکه رشد جمعیت یک منطقه را

نسبت به مناطق دیگر درنظر می گیرد. برای مثال، ممکن است مشاهده کنیم که در یک شهر جمعیت و

رشد آن همواره رابطه با رشد استان مربوطه داشته باشد، در این حالت می توان از مدل مقایسه ای

استفاده نمود.

بنابراین اگر جمعیت در استان S در زمان t در دست باشد (P_t^s) ، جمعیت شهر مربوطه (P_t^c) از طریق

رابطه زیر بدست آورد:

$$P_t^c = K P_t^s \quad (1)$$



P_t^c : جمعیت شهر C در زمان t،

K: همان فاکتور نسبت،

P_t^s : جمعیت استان S در زمان t.

مثال: فرض کنیم جمعیت شهر C همواره ۱۰ درصد جمعیت استان مربوطه بوده است و جمعیت استان نیز براساس مدل خطی $P_t^s = 2 + 0/1t$ قابل پیش بینی است (P به میلیون نفر و t به سال) مطلوب است ابتدا جمعیت شهر C را در زمان صفر به دست آورده، سپس محاسبه کنید که جمعیت شهر C چند سال دیگر دوبرابر خواهد شد؟

$$P_0^C = 0/1(2 + 0/1(0)) = 0/2$$

جمعیت شهر C ۰/۲ میلیون نفر در زمان صفر است. حال اگر بخواهیم قسمت دوم مسئله را محاسبه کنیم عبارتند از:

$$P_0^C = 0/4 = 0/1(2 + 0/1t)$$

طرفین را به ضرب مجهول تقسیم می کنیم، آنگاه خواهیم داشت:

$$\frac{0/4}{0/1} = 2 + 0/1t$$

$$4 = 2 + 0/1t$$

آنگاه خواهیم داشت:

$$4 - 2 = 0/1t \Rightarrow 2 = 0/1t$$

$$t = \frac{2}{0/1} = 20 \text{ سال}$$

معادله فوق در صورتی بود که جمعیت رشد خطی داشته باشد. حال اگر جمعیت به صورت U رشد کند منتها با یک فاصله زمانی T، در این صورت خواهیم داشت:

$$P_t^u = P_t^s - T \quad (۱)$$



مثال: اگر سطح جمعیت در استان (P^s) تابعی است از سطح جمعیت در استان دیگر (PU) منتها با یک وقفه زمانی ۱۰ ساله و می دانیم که رشد جمعیت تابعی است نمایی که از میزان ۳ میلیون نفر شروع شده و رشدی معادل ۴ درصد را به صورت ذیل دارا است.

$$P_t^u = 3(1 + 0/04)^t$$

حال می خواهیم بدانیم در طی چندسال جمعیت استان S به ۵ میلیون نفر می رسد. می توان به صورت ذیل عمل کرد:

$$P_t^s = P_{t-10}^u 3(1 + 0/04)^{t-10}$$

$$5 = 3(1/04)^{t-10} \Rightarrow \frac{5}{3} = 1/04^{t-10}$$

$$1/67 = 1/04^{t-10}$$

معادله را می توان به صورت ذیل ساده نمود:

$$\log 1/67 = (t-10)\log 1/04$$

$$(t-10) = \frac{\log 1/67}{\log 1/04} = 13/12$$

$$(t-10) = 13/12$$

$$t = 13/12 + 10$$

$$t = 23/12$$

حدود ۲۳/۱۲ سال طول می کشد جمعیت استان S به سطح ۵ میلیون نفر برسد.

مدل رگرسیون خطی ساده

در این مدل با استفاده از روابط میان دو متغیر می توان متغیری را از روی متغیر دیگری برآورد نمود. بدین صورت که متغیر Y در رابطه با زمان X تغییر می نماید. در این صورت می توان بین این دو گروه متغیر یک رابطه خطی به صورت ذیل نوشت:

$$y = ax + b \quad (۱)$$

مقادیر a, b ضرایبی هستند که از طریق فرمول های ذیل به دست می آیند:



$$a = \frac{\sum xy}{\sum x^2} \quad (2)$$

$$b = \frac{\sum y}{N} \quad (3)$$

مثال: جمعیت شهر A از سال ۱۳۴۰ تا ۱۳۸۰ به شرح جدول (۱-۲) می باشد خط رگرسیون جمعیت آن را در رابطه با زمان بدست آورده و برآورد نمایید که در سال ۱۴۰۰ جمعیت آن چقدر خواهد بود.

جدول شماره (1)، داده های جمعیتی شهر طي سال هاي 1340-1380

سال	جمعیت (y)
۱۳۴۰	۱۲۰۰۰
۱۳۵۰	۱۴۰۰۰
۱۳۶۰	۱۶۵۰۰
۱۳۷۰	۱۸۸۰۰
۱۳۸۰	۲۱۵۰۰
Σ	۸۲۸۰۰

حال می توان با استفاده از جدول ذیل مقادیر b را محاسبه می کنیم.

$$b = \frac{82500}{5} = 16560$$

N هم تعداد دوره ها است که هر ۱۰ سال به عنوان یک دوره در نظر گرفته شده است. برای بدست آوردن مقادیر a بایستی مقادیر X، X² را بدست آوریم. با در نظر گرفتن یکی از دوره ها به عنوان دوره مبدا می توان مقدار X را به دست آورد. در این صورت هر کدام از چند ده سال قبل از مبدا و بعد از آن مقادیر X خواهند بود، که در این مثال ۱۳۶۰ به عنوان سال مبدا که دوره های بعد و قبل نسبت به آن سنجیده می شود.



جدول شماره (2): محاسبه جمعیت شهر □

سال	y	x	x ²	xy
۱۳۴۰	۱۲۰۰۰	-۲	۴	-۲۴۰۰۰
۱۳۵۰	۱۴۰۰۰	-۱	۱	-۱۴۰۰۰
۱۳۶۰	۱۶۵۰۰	۰	۰	۰
۱۳۷۰	۱۸۸۰۰	۱	۱	۱۸۸۰۰
۱۳۸۰	۲۱۵۰۰	۲	۴	۴۳۰۰۰
Σ	۸۲۸۰۰	-	۱۰	۲۳۸۰۰

حال می توان مقدار a را بدست آورد:

$$a = \frac{23800}{10} = 2380$$

حال می توانیم با جایگزینی مقادیر a, b در معادله خطی، جمعیت را برای سال ۱۴۰۰ پیش بینی کنیم:

$$Y = 2380 \cdot X + 16560$$

باتوجه به این که در جدول شماره (۲-۲) هر ۱۰ سال به عنوان یک دوره نسبت سال مبداء در نظر گرفته شده است سال ۱۴۰۰ نسبت به سال ۱۳۶۰، ۴ دوره می شود که مقادیر X، ۴ خواهد بود. حال خواهیم داشت:

$$y = 2380 \times 4 + 16560$$

$$y = 26080 \text{ نفر}$$

جمعیت در سال ۱۴۰۰ به ۲۶۰۸۰ نفر خواهد رسید.

مدل رگرسیون چندمتغیره

تحلیل های رگرسیون چند متغیره براساس این فرضیه است که بین جمعیت و متغیرهای دیگر، رابطه ثابتی وجود دارد. این روش در ساده ترین شکل خود، تنها از یک متغیر مستقل به صورت زیر استفاده می کند:



$$P_{(t)} = a + bx_{(x)} \quad \text{رابطه (۱)}$$

در این فرمول، $x(t)$ متغیر مستقل است. گفتنی است که معادله های روش های خط روند و روش های نسبت، گونه ای دیگر از این فرمول کلی به شمار می آیند. با شناخت این معادله، می توان از آن برای پیش بینی استفاده کرد. معمولاً معادله یاد شده، به صورت زیر ارائه می شود:

$$\Delta P_t^{t+n} = a + b\Delta X_t^{t+n} \quad \text{رابطه (۲)}$$

که:

$$\Delta P_t^{t+n}: \text{تغییر در جمعیت بین زمان } t, t+n$$

$$\Delta X_t^{t+n}: \text{تغییر در متغیر مستقل بین زمان } t, t+n$$

گونه تغییر یافته تر، عامل فاصله زمانی را نیز در فرمول دخالت می دهد، یعنی می پذیرد که تغییر در متغیر مستقل برای موثر بودن، نیازمند زمان است، بنابراین:

$$\Delta P_t^{t+n} = a + b\Delta X_s^{s+m} \quad \text{رابطه (۳)}$$

در این فرمول، $s, s+m$ دو تاریخ هستند که برحسب ضرورت و تناسب گزینش شده اند؛ احتمال دارد که رابطه، غیرخطی باشد که در نتیجه پیش از محاسبه خط رگرسیون، باید انتقال لگاریتمی صورت گیرد:

$$\log P_{(t)} = a + bX_{(t)} \quad \text{رابطه (۴)}$$

$$P_{(t)} = a + b\log X_{(t)} \quad \text{رابطه (۵)}$$

معمولاً رگرسیون چندمتغیره مورد استفاده قرار می گیرد:

$$P_{(t)} = a + b_1x_1 + b_2x_2 + \dots + b_nx_n \quad \text{رابطه (۶)}$$

که به صورت فشرده می توان نوشت:

$$P_{(t)} = a + \sum_i b_i X_i \quad \text{رابطه (۷)}$$



متغیرهای به کاررفته ممکن است متغیرهای واقعی (X_i) تغییر در متغیرها (ΔX_i) یا بدون فاصله زمانی، یا انتقال لگاریتمی (لگاریتم X_i یا لگاریتم ΔX_i با تغییرات لگاریتم X_i) باشد. برای نمونه، می توان متغیرهای قابلیت دسترسی، هزینه مسکن، سطوح دستمزد، فرصت های خالی اشتغال، بیکاری و سرمایه گذاری صنعتی را نام برد. در برخی از مطالعات، به جای برازش خط رگرسیون براساس داده های قدیمی موجود در یک ناحیه، آن را بریک پایه داده های زمان حال یا داده های اخیر چند ناحیه برازش می دهند. سطح پیچیدگی این گونه الگوها سبب پنهان کردن نقص اساسی تئوریک می شود. ساختار علی اگو ضعیف است و ممکن است متغیرهای مستقل، تعداد زیادی از عوامل علی ویژه ای را پنهان کنند که به شیوه های گوناگون با متغیر مستقل مرتبط هستند. همچنین به نظر نمی رسد آن گونه که برای پیش بینی ضروری است، پارامترهای a , b_i ثابت باقی بمانند. مسائل دیگری وجود دارد که از نیازهای آمار کاربرد تکنیک سرچشمه می گیرند، متغیرهای مستقل باید دارای پراکنش نرمال بوده و مستقل باشند.

بررسی و طبقه بندی مدل ها در برنامه ریزی مسکونی

یکی از مهم ترین مراحل فرآیند برنامه ریزی مسکونی، آینده نگری از طریق شناسایی عناصر و روابط موجود در آن و فرا فکنی موقعیت آینده آن ها با استفاده از مدل ها و تکنیک های ریاضی و آماری است. با توجه به چنین اهمیتی است که به بررسی مدل ها در برنامه ریزی مسکونی اقدام می شود. اجزاء مدل ها شامل پیش بینی جمعیت، تعیین اندازه متوسط خانوار، تعیین تراکم در واحد مسکونی، پیش بینی تعداد واحد مسکونی، برآورد میزان پس انداز و نرخ بازدهی سرمایه گذاری می باشد. در یک تقسیم بندی کلی می توان مدل ها را به دو دسته اساسی به شرح زیر تقسیم نمود:

مدل ها ساده - مدل های پیچیده

مدل های ساده

در مدل های ساده برنامه ریزی مسکونی از طریق پیش بینی خطی پاره ای از متغیرها (عمدتاً جمعیت) و تدوین شاخص های استاندارد به عنوان هدف صورت می پذیرد. گام های متوالی برنامه ریزی مسکونی به روش ساده شامل پیش بینی یا تخمین اندازه متوسط خانوار، تعیین تراکم خانوار در واحد مسکونی و تعیین ویژگی های کالبدی فضای مسکونی کورد نیاز می باشد. ویژگی های این گام ها به شرح زیر است:

گام اول پیش بینی جمعیت



پیش‌بینی جمعیتی معمولاً بخش عمده مصالعات مدل‌های ساده را تشکیل می‌دهد. در این مدل‌ها، مطالعات جمعیتی به توصیف و تحلیل تعداد جمعیت طی دوره‌های گذشته بر حسب گروه‌های سنی، جنسی، سطح سواد، اشتغال، بی‌کاری صورت می‌گیرد و پیش‌بینی‌های جمعیتی از طریق تداوم نرخ رشد گذشته و یا نرخ بقا بر حسب گروه‌های سنی انجام می‌شود.

گام دوم - پیش‌بینی و یا تخمین اندازه متوسط خانوار

در مدل‌های ساده با فرض مطلوب بودن کمیت مشخص از بعد خانوار، تعداد خانوار براساس پیش‌بینی‌های جمعیتی برآورد می‌گردد.

گام سوم - تعیین تراکم خانوار در واحد مسکونی

در این گام نیز با فرض مطلوب بودن مشخصی از تراکم خانوار در واحد مسکونی، تعداد واحد مسکونی مورد نیاز در آینده تخمین زده می‌شود.

گام چهارم - تخمین نیاز به فضای مسکونی

در این گام ویژگی‌های کالبدی بر حسب سه متغیر تعیین می‌شوند که عبارتند از :

تعیین مساحت زیربنای هر واحد مسکونی، تعیین تراکم خالص مسکونی و تعیین مساحت فضای باز خصوصی (خانگی)

مدل های پیچیده

لازمه هر برنامه‌ریزی تخصصی مسکن، توجه به عوامل موثر در عرضه و تقاضای مسکن و تعیین میزان تاثیر آنها می‌باشد. در کلی‌ترین بیان مدل‌های پیچیده سعی در بیان این عوامل و تعیین میزان تاثیر آنها دارد تا بدین وسیله بتوان واقعیت‌های موجود را از طریق تغییر عوامل موثر و یا تعدیل تأثیرشان به اهداف مورد نظر هدایت کرد. مدل‌های پیچیده به دو دسته تقسیم می‌شوند:

الف) مدل‌های پیچیده با تمایل عمده نسبت به علم اقتصاد

ب) مدل‌های پیچیده با تمایل عمده نسبت به برنامه‌ریزی شهری و محلی

مدل‌های پیچیده ضمن استفاده از تمامی عناصر مدل‌های ساده، عناصر جدیدی را نیز وارد کرده‌اند. با این تفاوت که اکثر کمیت‌های مورد استفاده در مدل‌های پیچیده توسط تکنیک‌های لازم برای افق برنامه پیش‌بینی می‌گردد تا براین اساس واقعیت‌های محتمل طی سال‌های برنامه مشخص شود. در مدل‌های



پیچیده تقاضای موثر و یا امکانات عرضه مسکن به عنوان متغیر وابسته، تابعی از قدرت خرید و توان مالی خانوار، قیمت واحد مسکونی، تخمین تراکم های نفر و خانوار در واحد مسکونی و پتانسیل های توسعه فضایی در نظر گرفته می شود. قدرت خرید و توان مالی خانوار بخشی از درآمد قابل تصرف و همچنین میزان وام دریافتی و اندوخته پیشین هر خانوار می باشد. چنانچه قدرت مالی خانوار با قیمت واحد مسکونی منطبق باشد، در آن صورت متقاضی واحد مسکونی بر حسب توانایی خود صاحب مسکن شده و عرضه مسکن صورت می پذیرد. در صورتی که نرخ بازدهی سرمایه در بخش مسکن کمتر از بخش ها باشد وجود قابلیت توان خرید مسکن از سوی متقاضی، احتمال خارج شدن سرمایه از بخش مسکن به سوی بخش های اقتصادی وجود داشته و در نتیجه به دلیل کاهش منابع در بخش مسکن، مقدار عرضه مسکن کاهش می یابد (جوادی اشکلک، ۱۳۷۶: ۱۷).

روش های برآورد نیاز به مسکن

برای تنظیم و اجرای برنامه های مسکن، باید نیاز به مسکن را برحسب کمبود شرح داد. بدیهی است که برای تخمین نیاز به مسکن، باید عناصر اشلی این نیاز مورد بررسی قرار گیرد و یک روش کلی برای رسیدن به برآورد کلی نیز تنظیم شود.

هدف از برآورد نیاز به مسکن

برآورد مطمئن از نیاز به مسکن، عامل مهمی در تدوین سیاست و تنظیم و ارزیابی برنامه های مسکن است. برای تنظیم و اجرای برنامه ها، باید به مسکن در مناطق مختلف کشور مشخص شود و همچنین شناختی از خصوصیات اقتصادی و اجتماعی متقاضیان مسکن صورت گیرد که این باعث تعیین اولویت های تهیه مسکن خواهد شد. اولویت ها ممکن است بر مبنای کمبود مسکن، اهمیت مناطق معین، گروه های جمعیتی و هدف های برنامه شود (پورمحمدی، ۱۳۸۲: ۴۷).

تقاضای مسکن

تقاضای مسکن به عنوان تعداد مساکن قراردادی یا خانه های مناسب زندگی تعریف شده اند که در زمانی احتیاج به ساخته شدن یا ایجاد تغییراتی برای رسیدن به سطح استانداردهای قابل قبول ملی دارند. همچنین شامل تعداد مسکنی است که نیازمند به تعمیر و یا نگهداری هستند به طوری که در سطح استاندارد و در مدت زمان معین باقی بمانند. احتیاجات به طور کلی برای درجات مختلف و همه



موقعیت‌های مسکن برآورد شده است، یعنی احتیاج به تهیه خانه برای افراد بی‌خانمان، احتیاج به تهیه مسکن برای ساکنان زاغه‌ها و سایر مساکن موقتی، احتیاج به تهیه مسکن مجزا برای هر خانوار، احتیاج به نگهداری سطوح تراکم در یک سطح معین و برای آینده، احتیاج به تهیه مقرراتی برای افزایش پیش‌بینی شده تعداد خانوار و جایگزین کردن مساکن جدید با مساکن قدیم (زمانی که عمر آنها به پایان می‌رسد) همان: (۴۸).

تقاضای بالقوه و تقاضای موثر

منظور از تقاضای بالقوه مسکن، تقاضای کلیه کسانی است که تمایل به داشتن و تصاحب واحد مسکونی مناسب دارند که شامل کلیه افراد بی‌خانمان، ساکنان آلودگی‌ها، مستأجران و غیره است و منظور از تقاضای موثر تقاضای کلیه کسانی است که قادر و مایل به خریداری و تصاحب مسکن هستند، یعنی علاوه بر این که تمایل به خرید مسکن دارند، قدرت اقتصادی و یا اعتبار لازم خرید مسکن را نیز دارند. در امر برنامه‌ریزی باید نسبت به متقاضی آن شناخت کافی پیدا نمود که این افراد شامل چه کسانی هستند، تعداد آن‌ها چقدر است، چه فرهنگی دارند و شرایط مالی آن‌ها چگونه است (حمزه مصطفوی، ۱۳۷۴: ۳۳۹)

تولید مسکن

مبانی نظری و تئوریک تولید مسکن عمدتاً بر پایه اصل اقتصادی عرضه و تقاضا استوار گردیده است. علم اقتصاد همانند سایر علوم براساس تئوریه‌ها و قوانینی بنا شده است. یکی از مهم‌ترین و جامع‌ترین تئوریه‌ها و قوانینی بنا شده است. یکی از مهم‌ترین و جامع‌ترین تئوریه‌ها، تئوری رفتار بازار است. این تئوری بر بازار مسکن نیز احاطه دارد و چنین بیان می‌شود که:

تعیین قیمت در بازار به وسیله نیروهای عرضه و تقاضا انجام می‌گیرد. هر قدر قیمت بالاتر باشد، مقدار تقاضا کمتر است ولی مقدار بیشتری عرضه خواهد شد و هر قدر قیمت پایین‌تر باشد، مقدار تقاضا بیشتر می‌شود، اما مقدار عرضه کمتر خواهد بود. تعادل قیمت و مقدار در بازار زمانی برقرار می‌گردد که عرضه با مقدار تقاضا برابر شود.

مهم‌ترین نیروهای فعال در بازار مسکن عبارتند از ۱- مصرف کنندگان یا متقاضیان مسکن ۲- عرضه کنندگان یا تولید کنندگان مسکن ۳- تأمین کنندگان منابع مالی ۴- دولت‌ها



هر یک از گروه های فوق در بازار مسکن منافی را دنبال می کنند و قانون حاکم بر رفتار آنها قانون "رفتار عقلایی" است. بر طبق این قانون مصرف کنندگان و تولید کنندگان دارای رفتار منطقی هستند. به عبارت دیگر هر دو گروه به دنبال کسب حداکثر ارضاء و مطلوبیت با حداقل هزینه می باشند. مصرف کنندگان سعی دارند بیشترین و بهترین کیفیت واحد مسکونی را با پایین ترین قیمت بدست آورند. تولید کنندگان نیز در اقتصاد بازار همواره تلاش می کنند با کمترین هزینه تولید، سود و منافع خود را حداکثر نمایند. در بازار مسکن عوامل دیگری مانند سرمایه داران و تأمین کنندگان منابع مالی نیز وجود دارند. آنها نیز سعی در افزایش سرمایه و به حداکثر رساندن سود خود دارند. دولت ها نیز در بازار مسکن اهدافی را دنبال می کنند. بخش مسکن به لحاظ تولید یکی از مهمترین نیازهای اساسی و ضروری خانواده ها و برطرف نمودن مشکل مسکن، از بخش های بسیار حساس اقتصاد جوامع مختلف بحساب می آید. دولت ها می توانند با برنامه ریزی دقیق و استفاده از ابزارهای سیاستگذاری در این بخش به اهداف خود نایل گردند.

عرضه مسکن تابع عواملی مانند موقعیت جغرافیایی، شرایط آب و هوایی، تکنولوژی ساخت، مهارت و ترکیب نیروی کار، ظرفیت و توان تولیدی جامعه، سرمایه گذاری در امور زیربنایی و غیره است. (شکرگزار، ۱۳۸۵ : ۴۰) عده ای از کارشناسان با استفاده از اصول اقتصادی عرضه و تقاضا به بررسی عوامل تولید، عرضه و تقاضای مسکن در جامعه پرداختند.

روش استفاده از مدل های اقتصاد سنجی برای پیش بینی عرضه و تقاضای مسکن: این مدل ها به طور کلی عبارت است از یک روابط کمی که بین متغیرهای اجتماعی و اقتصادی برقرار است. مدل پیش بینی عرضه مسکن به قرار زیر است.

$$H_s = (R, P_n, R_I, CRT, RGNP, GI, W, CONSM, UTIL, \dots V)$$

که در آن (F) ضریب ثابت، (H_s) عرضه مسکن (R) شاخص اجاره بها (PH) شاخص قیمت فروش، (R_I) نرخ کارمزد، (CRT) اعتبارات ساختمانی، (RGNP) نرخ رشد درآمد ملی، (GI) مخارج ساختمانی دولت، (W) شاخص دستمزد کارگر ساختمانی، (CONSM) شاخص بهای مصالح ساختمانی (VTIL) شاخص بهای آب و برق و غیره

مدل پیش بینی تقاضای مسکن:

$$H_d = F(POP, M, Y, NH, FA, \dots V)$$



که در آن (F) ضریب ثابت، (Hd) مسکن، (POP) جمعیت، (M) مهاجرت، (Y) درآمد، (NH) تعداد مساکن موجود، (FA) دسترسی به اعتبارات مالی و... است. معادلات عرضه و تقاضا به صورت معادلات مجرد یا به صورت معادلات متقارن با استفاده از روش های رگرسیونی مانند روش حداقل مربعات، گلرین ارات، 2SLS، 3SLS،... برآورد می گردد. (دلال پورمحمدی، ۳۷)

روش ها و مدل های مختلفی برای عرضه و تقاضای مسکن ارائه شده است. در مدل عرضه "لوکاس" تابع تولید، تابعی از سه عامل سرمایه، نیروی کار و سرمایه گذاری است. با تعمیم این نظریه برای تابع تولید مسکن خواهیم داشت.

$$O1 = F(Kt, Lt, It)$$

که در آن Kt سرمایه (زمین)، Lt نیروی انسانی شاغل در بخش مسکن و It سرمایه گذاری در بخش است. (خلیلی عراقی، ۱۳۷۹: ۱۷) مطالعات تابع تقاضای مسکن به وسیله افرادی چون دی جی اوت^۱، جی هاش، یو^۲ مطالعات کیرل^۳، مطالعات اولسن^۴ و میوت^۵ انجام شده است. اساس تئوریک تمامی این مطالعات بر تئوری رفتار مصرف کننده است. (شکرگزار، ۱۳۸۵: ۴۱).

روش تعیین کمبود واحد مسکونی

برای تعیین کمبود واحد مسکونی در برنامه ریزی مسکن شهری از روش زیر استفاده شده و ابتدا کمبود واحد مسکونی تعریف می شود. (سازمان مدیریت و برنامه ریزی استان فارس، ۱۳۷۰: ۲۱). **حاصل تفاضل واحدهای مسکونی معمولی موجود از تعداد خانوارهای معمولی را اصطلاحاً کمبود واحد مسکونی می نامند.** به عبارت دیگر با مقایسه تعداد خانوارها با تعداد مسکن موجود و براساس یک استاندارد معین (غالباً یک مسکن برای هر خانوار) می توان کمبود مسکن را برآورد نمود. برای مقایسه بهتر این شاخص می توان از درصد کمبود واحدهای مسکونی استفاده نمود که از طریق رابطه زیر به دست می آید (زیاری و دهقان، ۱۳۸۲: ۶۷):

^۱ Ott.D.J
^۲ YOO.H.J
^۳ Kearn
^۴ Olsen
^۵ Mute



$$\text{درصد کمبود واحد مسکونی} = \frac{\text{واحد مسکونی} - \text{تعداد خانوار}}{\text{تعداد خانوار}} \times 100$$

تشریح روش های برآورد نیاز به مسکن

برای برآورد نیاز به مسکن از روش های مختلفی استفاده می شود مهمترین روش های برآورد نیاز به

مسکن به شرح زیر هستند

✓ روش انبوهه،

✓ روش شاخص،

✓ روش های نرخ سرپرستی،

✓ روش کلی،

✓ روش خام،

✓ روش لجستیک.

روش انبوهه

در این روش، جمعیت پیش بینی شده (P) به میانگین بعد خانوار (S) تقسیم می گردد تا تعداد خانوار

پیش بینی شده (H) به دست آید، بنابراین خواهیم داشت $H = \frac{P}{S}$ و اگر K ضریب خانوار را در واحد

مسکونی دخالت دهیم. تعداد مساکن مورد نیاز (E) بدست می آید. $E = \frac{P}{K(S)}$. (حکمت نیا و موسوی،

۱۳۸۵: ۱۱۸).

روش شاخص

این روش یکی دیگر از روش هایی است که برای برآورد تعداد واحدهای مسکونی مورد نیاز می توان

استفاده کرد و در این روش ابتدا تعداد خانوار پیش بینی می گردد، سپس با دخالت دادن شاخص خانوار

به مسکن می توان تعداد واحدهای مسکونی مورد نیاز را برآورد نمود. این روش از رابطه زیر به دست

می آید: که در آن (P) مقدار پیش بینی جمعیت، (H) تعداد خانوار، (K) شاخص خانوار به مسکن و

(E) تعداد واحد مسکونی مورد نیاز (همان: ۱۱۹).

رابطه (۱) $H = PR$



$$E_{(t)} = \frac{H}{K} \quad \text{رابطه (۲)}$$

روش نرخ سرپرستی

در این روش، جمعیت پیش‌بینی شده معمولاً بر حسب سن، جنس و وضع ازدواج به گروه‌هایی تقسیم می‌شوند، که در هر گروه جمعیتی، سهم معینی از افراد رئیس خانوار خواهند بود و از طریق تعداد افراد رئیس خانوار و جمعیت سال پایه می‌توان نرخ های سرپرستی در هر گروه را محاسبه کرد. سپس با جمع کردن تعداد افراد هر گروه سنی به نرخ های سرپرستی، تعداد خانوارها پیش‌بینی می‌شود. از روی تعداد خانوارهای پیش‌بینی شده تعداد خانوارهای بالقوه محاسبه شده و واحدهای مسکونی مورد نیاز برآورد می‌گردد (پور محمدی، ۱۳۸۲: ۵۶).

روش کل

در این روش از رابطه زیر استفاده می‌کنند:

$$E_{(t)} = K(E_1 + E_2 + E_3 + E_4 + E_{6(t)}) + E_5 + E_{7(t)}$$

$E_{(t)}$: واحد مسکونی مورد نیاز تا زمان □

□ : ضریب خانه های خالی

E_1 : کمبود فعلی مسکن

E_2 : واحدهای مسکونی زیر استاندارد

E_3 : نیاز ناشی از واحد مسکونی به ازای هر خانوار

E_4 : نیاز ناشی از حذف یا کاهش تراکم

$E_{6(t)}$: نیاز ناشی از افزایش جمعیت در زمان □

E_5 : واحدهای مسکونی که نیاز به تخریب و تجدید بنا دارند

$E_{7(t)}$: واحدهای مسکونی که در طول مدت برنامه‌ریزی نیاز به تخریب و تجدید بنا دارند

روش خام

ساختار کلی این روش به شرح زیر است :



$$E(t) = H - U + H_{(t)} + rU_{(t)}$$

در این رابطه :

$E_{(t)}$: واحد مسکونی مورد نیاز تا زمان □

□ : تعداد خانوارها در شروع دوره برآورد

□ : تعداد واحدهای مسکونی قابل قبول در شروع دوره برآورد.

$H_{(t)}$: تعداد خانوارها در پایان دوره برآورد

$rU(t)$: درصد واحدهای مسکونی که تا زمان □ نیاز به تخریب و تجدید بنا دارند.

برای بدست آوردن میزان تخریب واحدهای مسکونی می توان از رابطه زیر استفاده نمود (توفیق، ۱۳۶۹:

۴۹).

$$\text{درصد تخریب در یک دوره} = \frac{\text{مسکن موجود ۱۳۸۵} + \text{مسکن موجود ۱۳۷۵} + \text{واحدهای ساخته شده طی (۸۵-۱۳۷۵)}}{\text{مسکن موجود ۱۳۷۵}} \times ۱۰۰$$

مدل لجستیک

مدل لجستیک روش مناسبی برای برآورد تعداد واحد مسکونی منتسب در یک شهر است این روش

براساس روند گذشته به پیش بینی آینده مسکن مبادرت می کند، که بر فرض های زیر استوار است:

- میانگین وسعت خانوار تا افق برنامه ریزی تغییر نمی کند و ثابت باقی می ماند.
- حد نهایی تراکم مطلوب برابر یک واحد مسکونی به ازای یک خانوار است. بدین ترتیب کرانه زیرین شاخص مسکن به ازای نفر، معادل معکوس خانوار به دست می آید.

ساختار کلی مدل به شرح زیر است:

$$HP_{(t)} = \frac{K}{1 + e^{(a+abt)}}$$

$HP(t)$: شاخص واحد مسکونی به ازاء نفر در سال □،

□: عدد ثابت اولر،

□، □: پارامترهای ثابت،

□: رقم نهایی واحد مسکونی / نفر یا معادل معکوس وسعت خانوار که از رابطه زیر به دست می آید:



$$K = \frac{1}{\text{رأین اصح عو}}$$

با در دست داشتن مقدار $\square$ و اطلاعات مربوط به شاخص واحد مسکونی به ازای نفر برای سال های گذشته، می توان مقدار $\square$ و $\square$ را از رابطه (۳) و (۴) به دست آورد:

$$a = \text{Ln} \frac{K - HP_0}{HP_0} \quad \text{رابطه (۳)}$$

$$a = \text{Ln} \frac{HP_0(K - HP_1)}{HP_1(K - HP_0)} \quad \text{رابطه (۴)}$$

در رابطه های بالا :

$\square$: زمان (سال)،

HP : شاخص مسکن در سال مبدا،

HP1 : شاخص مسکن در سال مقصد (حکمت نیا و موسوی، ۱۳۸۵: ۱۲۲).

برآورد نیاز به زمین برای تأمین مسکن

با استفاده از روش های زیر می توان زمین مورد نیاز برای فعالیتهای خانه سازی را پیش بینی کرد:

- روش استفاده از تراکم متوسط ساختمانی،
- روش استفاده از گروه نما
- روش استفاده از سرانه مسکونی و تراکم خالص،
- روش برآورد مساحت واحدهای مسکونی از دیدگاه فرهنگی (شاخص فرهنگی مسکن)

روش استفاده از تراکم متوسط ساختمانی

در این روش، برای محاسبه میزان زمین مورد نیاز از رابطه بین قیمت و تراکم استفاده می شود. ابتدا گرایش متوسط به تراکم از طریق رابطه زیر به دست می آید، سپس با ضرب آن در بعد خانوار، مساحت زمین مورد نیاز برای شهر به دست می آید (یزدانی و الیاسی، ۱۳۸۱: ۱۵۹).

$$۱۰۰ \times \text{قیمت زیر بنا} / \text{قیمت زمین} \times \frac{3}{2} = \text{تراکم متوسط ساختمانی}$$

نتیجه به دست آمده را ضرب به در بعد خانوار می کنیم تا میزان زمین مورد نیاز در ده سال آینده به

دست آید.



روش استفاده از گروه نما

یکی از روش های برآورد نیاز به زمین، جهت تأمین مسکن مورد نیاز شهروندان، استفاده از روش گروه نما است. در این روش، ابتدا واحدهای مسکونی شهر از نظر مساحت زمین به طبقاتی تقسیم می شوند. مثلاً بناهای با مساحت کمتر از ۵۰ متر مربع، ۵۰ تا ۹۹ متر مربع، ۱۰۰ تا ۱۴۹ متر مربع، ۱۵۰ تا ۱۹۹ متر مربع و سپس درصد واحدهای مسکونی (فراوانی) هر طبقه محاسبه می شود. هر طبقه از واحدهای مسکونی که بالاترین درصد کمی از مساکن شهر (فراوانی) را به خود اختصاص دهد، طبقه نما اطلاق می شود. بعد از مشخص شدن طبقه نما، نیاز به زمین مورد نیاز از طریق رابطه زیر پیش بینی می شود:

$$M_0 = L_0 \frac{(F^0 + F_{-1})}{(F_0 - F_{-1}) + (F_0 - F_{+1})} \times i$$

در رابطه بالا:

□: نمای مساحت واحدهای مسکونی،

□: کرانه پایین طبقه نما،

□: فراوانی در طبقه نما،

(۱- □) فراوانی طبقه پیش از نما،

(□ + ۱) فراوانی طبقه بعد از نما،

□ : فاصله طبقات.

پس از به دست آوردن طبقه نما مساحت آن را به تعداد واحدهای مسکونی مورد نیاز ضرب می کنیم تا مقدار زمین مورد نیاز برای تأمین واحدهای مسکونی به دست آید (توفیق، ۱۳۶۹: ۵۳). به عبارت دیگر:

تعداد واحد مسکونی مورد نیاز * مساحت نما = زمین مورد نیاز

روش استفاده از سرانه مسکونی و تراکم خالص

شاخص های سرانه مسکونی که از تقسیم کردن کل مساحت زمین های مسکونی شهر (به متر مربع) به جمعیت آن، و تراکم خالص که از تقسیم کردن جمعیت شهر به مساحت زمین های مسکونی (به هکتار) به دست می آید مثل دو روی یک سکه اند. از نظر تشخیص وضع نسبی شهرها و نیز آینده نگری شاید این دو مهمترین شاخص باشند.



در این روش، سرانه مسکونی به جمعیت پیش‌بینی شده ضرب می‌شود تا مقدار زمین مورد نیاز به دست آید. روش تراکم خالص براساس نفر یا خانوار در هکتار محاسبه می‌شود و از رابطه زیر به دست می‌آید (حکمت نیا و موسوی، ۱۳۸۵: ۱۲۶):

$$A = \frac{H}{D}$$

در این رابطه:

A: مقدار زمین مورد نیاز،

H: تعداد مساکن مورد نیاز،

D: تراکم مسکونی.

روش برآورد مساحت واحدهای مسکونی مورد نیاز خانوارها از دیدگاه فرهنگی

این روش برای برآورد مساحت واحدهای مسکونی مورد نیاز خانوارها با توجه به ویژگی‌های فرهنگی پیشنهاد شده است. در این روش متغیرهای زیادی از جمله بعد خانوار، شغل، تحصیلات، نوع معیشت، سابقه شهرنشینی، میزان درآمد و ویژگی‌های اجتماعی مورد نظر است و به هر یک امتیاز داده می‌شود. بنابراین موقیت اجتماعی - اقتصادی سرپرست خانوار با توجه به جدول زیر امتیازاتی را کسب خواهند نمود که در مجموع امتیازات، میزان مساحت مسکن خانوار مشخص خواهد کرد (همان: ۱۲۸).



جدول شماره (3): امتیاز و متغیرهای برآورد مساحت مورد نیاز خانوارها از دیدگاه فرهنگی

مساحت زمین به متر مربع									
۴۰ -۴۵	۶۰	۸۰	۱۰۰	۱۲۰	۱۴۰	۱۶۰	۱۸۰	۲۰۰	۴۰۰
۴ امتیاز	کارمندان (ساده و فنی)		۲ - شغل (امتیاز ۴)		۱ امتیاز		۱ نفر	۱ - بعد خانوار (امتیاز ۵)	
۶ امتیاز	پزشکان، استادان، مهندسان				۲ امتیاز		۲ نفر		
۲ امتیاز	کشاورزان، باغداران، دامداران				۴ امتیاز		۳-۴ نفر		
۳ امتیاز	تولید کنندگان، عمده فروشان				۴ امتیاز		۵-۶ نفر		
۴ امتیاز	فرهنگیان				۴ امتیاز		۷ نفر به بالا		
۳ امتیاز	مشاغل خدماتی								
۵ امتیاز	شهری		۴ - نوع زندگی (امتیاز ۳)		۱ امتیاز	کمتر از سیکل		۳ - تحصیلات (۶ امتیاز)	
					۳ امتیاز	دیپلم و فوق دیپلم			
۲ امتیاز	روستایی				۴ امتیاز	لیسانس			
					۵ امتیاز	فوق لیسانس			
۱ امتیاز	ایلی				۶ امتیاز	دکتری و متخصص			
۱ امتیاز	فاقد شغل		۶ - ویژگی‌های اجتماعی (امتیاز ۲)		۱ امتیاز	کمتر از ۵ سال		۵ - سابقه شهرنشینی (امتیاز ۳)	
۴ امتیاز	کار در محل				۲ امتیاز	۵-۸ سال			
۳ امتیاز	کار در خارج از شهر				۳ امتیاز	۸ - ۱۰ سال			
۲ امتیاز	اشتغال متحرک				۴ امتیاز	۱۰ - ۱۵ سال			
					۵ امتیاز	۱۵ - ۲۰ سال			
					۶ امتیاز	۲۰ سال به بالا			
					۱ امتیاز	کمتر از ۲۰ هزار تومان		۷ - میزان درآمد (امتیاز ۴)	
					۲ امتیاز	۲۰ - ۴۰ هزار تومان			
					۳ امتیاز	۴۰ - ۶۰ هزار تومان			
					۴ امتیاز	۶۰ - ۸۰ هزار تومان			
					۵ امتیاز	۸۰ - ۱۰۰ هزار تومان			
					۶ امتیاز	۱۰۰ هزار تومان به بالا			

منبع : پور محمدی ۱۳۸۲: ۱۱۱

تئوری مدل جاذبه^۱

اولین نظریه که واکنش متقابل تعدادی از فعالیتهای انسانی را در سازمان فضایی سرزمین مورد بررسی قرار می‌دهد تئوری مدل جاذبه است. نام مدل جاذبه از قانون جاذبه نیوتون در فیزیک گرفته شده است. جغرافیدانان یکی از مهمترین امانت هایی که از علوم فیزیک گرفته اند همین مدل جاذبه است (*Haggett, 1968, P.35*). بر اساس این قانون، واکنش متقابل دو جسم به جرم های M_i و M_j که با



فاصله d از هم قرار دارند با حاصل ضرب جرم آنها نسبت مستقیم و با مجذور فاصله آنها نسبت معکوس دارد. این رابطه به زبان ریاضی به صورت زیر بیان می گردد: (Chisholm, 1979, PP.150-151)

$$F_{ij} = a \frac{M_i \cdot M_j}{d_{ij}^2}$$

F_{ij} : نیروی کششی یا جاذبه بین جسم i و j ،

a : ضریب ثابت،

M_i : جرم جسم i ،

M_j : جرم جسم j ،

d_{ij} : فاصله بین جسم i و جسم j .

در مطالعات جغرافیایی این مدل توسط جغرافیدانان بسیاری مورد تعدیل قرار گرفت. به طوری که جهت استفاده از تأثیرات متقابل فضایی مراکز خرید و روابط متقابل اقتصادی بین شهر و سکونتگاه ها مورد استفاده قرار گرفت.

مدل تحلیل جاذبه (گرانشی)

این مدل بر پایه تئوری جاذبه نیوتن استوار است. در این روش تأکید بر جریان های بالقوه بین مراکز است. در این مدل فرض می شود که تأثیرات متقابل دو مرکز جمعیتی دارای نسبتی مستقیم با توده این مراکز و نسبتی معکوس با فاصله بین آنها است. متغیرهایی که به عنوان شاخص توده و فاصله به کار می روند بستگی به ماهیت مسئله مورد مطالعه و آمار و ارقام در دسترس دارند. برای نشان دادن مقدار توده هر مرکز می توان از ارقام مربوط به حجم جمعیت، اشتغال، درآمد و مصرف یاری گرفت و برای نشان دادن فاصله بین مراکز می توان از فاصله به کیلومتر، زمان سفر یا هزینه حمل و نقل و هزینه فرصت های از دست رفته، استفاده کرد. ساختار کلی مدل به شرح زیر است: (رفیعی، ۱۳۷۱: ۱۶)

(۱)

$$\tau_{ij} = K \left[\frac{P_i P_j}{d_{ij}^2} \right]$$

که در آن:

τ_{ij} : قدرت جاذبه بین شهر،



P_i و P_j : توده دو مرکز i و j ،

d_{ij}^2 : فاصله بین i و j ،

K : عدد ثابت.

مفهوم توان های جمعیتی یا گرانشی مفهومی است که در اثر کاربرد مدل های گرانشی در تعیین محدوده مناطق به دست آمده است. توان جمعیتی P_i ، که در اثر توده مرکز i ایجاد می شود، بدین گونه نشان داده می شود:

(۲)

$$P_i P_j = P \left[\frac{P_j}{d_{ij}} \right]$$

که در آن:

P : توان جاذبه مرکز i از مرکز j ،

حال اگر مرکز i توسط تعدادی واحدهای جمعیتی پیرامونی (P_i) احاطه شده باشد، در این صورت کل توان جذب مرکز i عبارتست از:

(۳)

$$V_j = K \sum_{i=1}^n \left[\frac{P_i}{d_{ij}} \right]$$

بدین ترتیب، با محاسبه توان جذب جمعیتی در منطقه مورد بررسی، می توان محدوده حوزه نفوذ را مشخص نمود.

مدل جاذبه ای تک قیدی

یکی دیگر از مدل هایی که بر پایه قانون جاذبه نیوتن استوار است مدل جاذبه ای تک قیدی است که کاربرد اصلی آن مکانیابی خرده فروشی در مناطق شهری است. به این علت تک قیدی خوانده می شود که شرط یا قیدی را مراعات می کند، ساختار کلی مدل به شرح زیر است:

$$\sum T_{ij} = O_i$$

اگر سفر از هر منطقه مبدا i به تمام مناطق j با یکدیگر جمع گردند، مجموع آنها برابر تعداد مبدأهای هر منطقه که به عنوان داده به مدل داده شده است، O_i می باشد از طرف دیگر اگر مدل در وضعیتی به



کار رود که هم مبدأ ها و هم مقصد ها (D_j, O_i) معلوم باشند این مدل قادر به بازسازی مقدار معلوم D_j نخواهد بود ، یعنی

$$\sum T_{ij} \neq O_i$$

یا اگر سفرها از تمام مناطق مبدأ به هر منطقه مقصد با یکدیگر جمع شوند، تعداد کل برابر مقصدهای معلوم D_i نخواهد بود با وجود این ، مدل در مواقعی مفید خواهد بود که ارزش D_i (مقصد های پيس بينی شده) معلوم نباشد. بنابراین از این مدل می توان برای پيس بينی یا برآورد مقصد استفاده نمود. بنابراین از این مدل می توان برای پيس بينی یا برآورد مقصد استفاده نمود. در مثال زیر، مدل توصیف جریان پرداخت ها مورد استفاده قرار گرفته است تا جریان افراد، اما شکل مدل به همان صورت است. مدلی که مورد استفاده قرار می گیرد شبیه مدلی است که لکشمین ابداع نموده است و بطور وسیعی در مطالعات برنامه ریزی مورد استفاده قرار می گیرد. مدل، جریان پرداخت ها میان نواحی مسکونی و مراکز خرید را تعیین می کند و مقدار فروش هر مرکز را با جمع کردن جریان پرداخت ها از تمام نواحی به مراکز خاص را به نوبت برآورد می کند. مدل این نکته را بیان می سازد که فروش یا بازده خرده فروشی یا یک مرکز خرید نسبت مستقیمی با اندازه یا جاذبه مرکز و نسبت معکوس یا مسافت از نواحی مسکونی و رقابت با دیگر مناطق دارد.

(۳)

$$S_{ij} = C_i A_j F_j^a d_{ij}^b$$

در رابطه (۳):

S_{ij} : پرداخت از ناحیه مسکونی i به مرکز خرید j ،

C_i : کل پرداختها در ناحیه مسکونی i ،

F_j : اندازه یا جاذبه مکز خرید j ،

(۴)

$$\left(\sum F_j d_{ij}^{-b} \right)^{-1} = A_i$$

فاصله از ناحیه مسکونی i و b توان هستند.

 d_{ij}



بنابراین کل فروش در مرکز خرید از حاصل جمع هزینه کردن تمام نواحی مسکونی در آن مرکز خرید به دست می آید :

(۵)

$$S_j = \sum_j C_i A_i F^a_j d_{ij}^{-b} = \sum_i S_{ij}$$

توان a در ارتباط با جاذبه مراکز خرید است که در اکثر مطالعات ضروری می باشد. زیرا که مراکز خرید بزرگ داد و ستد بیشتری نسبت به اندازه شان جذب می کنند. در این مدل علاوه بر ماتری زمان / مسافت به آمار پرداخت ها و سطح زیر بنا احتیاج است.

مثال : چهار منطقه شهری داریم که به دنبال آنیم که کل خرید های مناطق مسکونی را از مراکز خود با استفاده از مدل تک قیدی جاذبه محاسبه کنیم . به صورتی که زمان / مسافت از مثال مدل هسنن استفاده شده است. مقدار $a=1$ و $b=2$ در نظر گرفته شده است. و سطح زیر بنا و هزینه در ناحیه i در جدول زیر نشان داده شده است. مطلوب است خرید مناطق مسکونی از مراکز خرید را محاسبه کنید.

جدول شماره (4): آمار داده های زیربنا و هزینه در مثال فرضی

داده ها نواحی	جاذبه (زیربنا) فوت مربع	هزینه در ناحیه j (پوند)
۱	۴۰۰۰	۱۴۰۰۰۰
۲	۵۰۰۰	۷۰۰۰۰
۳	۱۰۰۰۰	۱۰۰۰۰
۴	۸۰۰۰	۶۰۰۰

مرحله اول : محاسبه $F^1_j d^2_{ij}$

جدول شماره (5): محاسبه مرحله اول مدل جاذبه تک قیدی

ناحیه	$D=1$	$D=2$	$D=3$	$D=4$	
$D=1$					۲۳۷۴۹/۹۹
$D=2$					۲۲۷۷۷/۷۷
$D=3$					۱۸۴۷۲/۲۱
$D=4$					۱۶۵۲۷/۷۷



$$\text{مرحله دوم: محاسبه } A_i F_j^1 d_{ij}^{-2} \text{ که آن برابر است با } \Pr_{ij} = \frac{F_i^1 d_{ij}^{-2}}{\sum_j F_j^1 d_{ij}^{-2}}$$

جدول شماره (6): محاسبه مرحله دوم

ناحیه	D-۱	D-۲	D-۳	D-۴
D-۱				
D-۲				
D-۳				
D-۴				

مرحله سوم: محاسبه $S_{ij} = C_i A_i F_j^1 d_{ij}^{-2}$ در مرحله اول و دوم محاسبه شده حال در

مرحله سوم رابطه (۵) مدل مورد محاسبه قرار می گیرد.

$$S_{11} = C_1 \Pr_{11} = 1400000 \times / 421 = 5894400$$

$$S_{12} = C_1 \Pr_{12} = 1400000 \times / 058 = 81200$$

$$S_{13} = C_1 \Pr_{13} = 1400000 \times / 468 = 655200$$

$$S_{14} = C_1 \Pr_{14} = 1400000 \times / 052 = 72800$$

$$S_{21} = C_1 \Pr_{21} = 700000 \times / 049 = 34300$$

$$S_{22} = C_1 \Pr_{22} = 700000 \times / 244 = 170800$$

$$S_{23} = C_1 \Pr_{23} = 700000 \times / 488 = 341600$$

$$S_{24} = C_1 \Pr_{24} = 700000 \times / 22 = 1540000$$

$$S_{31} = C_1 \Pr_{31} = 100000 \times / 24 = 24000$$

$$S_{32} = C_1 \Pr_{32} = 100000 \times / 3 = 30000$$

$$S_{33} = C_1 \Pr_{33} = 100000 \times / 338 = 33800$$

$$S_{34} = C_1 \Pr_{34} = 100000 \times / 12 = 12000$$

$$S_{41} = C_1 \Pr_{41} = 600000 \times / 037 = 22200$$

$$S_{42} = C_1 \Pr_{42} = 600000 \times / 189 = 112400$$



$$S_{43} = C_1 \Pr_{43} = 600000 \times / 168 = 100800$$

$$S_{44} = C_1 \Pr_{44} = 600000 \times / 655 = 362000$$

اگر این نتایج را در جدولی نمایش دهیم، به راحتی میتوان مجموع جریان هزینه از هر ناحیه به مرکز را به دست آورد.

جدول شماره (7): نمایش ارقام Sj

	۵۸۹۴۰۰	۸۱۲۰۰	۶۵۵۲۰۰	۷۲۸۰۰	۱۳۹۸۶۰۰
	۳۴۳۰۰	۱۷۰۸۰۰	۳۴۱۶۰۰	۱۵۴۰۰۰	۷۰۰۷۰۰
	۲۴۰۰۰	۳۰۰۰۰	۳۳۸۰۰	۱۲۰۰۰	۹۹۸۰۰۰
	۲۲۲۰۰	۱۱۳۴۰۰	۱۰۰۸۰۰	۳۶۳۰۰۰	۵۹۹۴۰۰
	۸۸۵۹۰۰	۶۶۵۴۰۰	۱۴۳۵۶۰۰	۷۰۹۸۰۰	۳۶۹۶۷۰۰

مدل جاذبه دو قیدی برای پراکنش سفر

مدل جاذبه تک قیدی، به دنبال مکان یابی مراکز خرده فروشی بود. اما در مدل جاذبه دو قیدی در حالی که مکان سکونت کارگران و مشاغل آنها معلوم است برای تعیین تعداد سفرهای موجود میان دو ناحیه مورد استفاده قرار می گیرد. ساختار کلی مدل به شرح زیر است:

$$T_{ij} = A_i B_j O_i D_j d_{ij}^{-b}$$

در رابطه (۱) :

T_{ij} : سفرهای میان مناطق i و j

O_i : کل سفرهایی که از ناحیه i آغاز گردیده،

D_j : کل سفرهایی که به ناحیه j ختم می گردد.



$$A_i = (\sum_j B_j D_j d_{ij}^{-b})^{-1} \quad (2)$$

$$B_i = (\sum_j A_j O_j d_{ij}^{-b})^{-1} \quad (3)$$

افزودن جمله B_j به مدل رعایت دو قید را تضمین می کند.

$$\sum_j T_{ij} = O_i \quad (4)$$

$$\sum_j T_{ij} = D_i \quad (5)$$

مبداء ها (O_i) و مقصدهای (D_j) معلوم را می توان به طور دقیق از اطلاعات رابطه متقابل به دست آورد. این نکته حائز اهمیت است که عبارت A_i شامل جمله B_j می باشد، در حالی که B_j نیز شامل A_i می باشد. این مفهوم است که A_i و B_j باید به عملی مکرر محاسبه شوند. یعنی با استفاده از ارزش A_i در یک مرحله ارزش B_j بدست می آید. سپس ارزش محاسبه شده B_j برای محاسبه ارزش جدید A_i به کار می رود، این ارزش جدید با ارزش اولیه A_i متفاوت است. ارزش جدید A_i را در معادله برای محاسبه B_j قرار داده، ارزش جدیدی برای B_j بدست می آید. سپس این ارزش جدید B_j در معادله ای که برای محاسبه A_i به کار می رود قرار داده و ارزش جدیدی برای A_i محاسبه می گردد. این فرآیند تا آنجا که تفاوتی میان ارزش های قبلی A_i و B_j ارزش جدیدی که محاسبه گردیده وجود نداشته باشد ادامه می یابد. البته بایست توجه نمود که به دست آوردن ارزش A_i و B_j به صورتی دستی خسته کننده و دوشوار است که بهتر است از برنامه رایانه ای استفاده گردد.

تئوری نقطه جدایی^{۱۱}

این تئوری اولین تغییر و تعدیل در تئوری تأثیر متقابل است. این مدل سعی می کند خط مرز منطقه تجاری بین دو شهر را از هم مشخص و جدا کند. ساختار کلی مدل به شرح زیر است: (Ibid,

(PP.449 -450)

$$B.P.D = \frac{d}{1 + \sqrt{\frac{P_L}{P_S}}}$$



: فاصله نقطه جدایی بین دو شهر،

: فاصله بین دو شهر،

$$B.P.D = d_{p_L p_s} - B.P.D$$

: جمعیت شهر بزرگتر،

: جمعیت شهر کوچکتر.

مثال: جمعیت شهر □ برابر ۱۰۰۰۰ نفر جمعیت و شهر □ برابر ۳۰۰۰۰ نفر جمعیت است. فاصله بین این

دو شهر ۱۰۰ مایل است فاصله نقطه جدایی بین دو شهر را مشخص نمایید؟

$$B.P.D_y = \frac{1001}{\sqrt{3000010000} + \sqrt{3}} = \frac{1001}{\sqrt{3000010000} + \sqrt{3}}$$

$$= \frac{1001}{+1/73} = \frac{1002}{/73} = 6/36 \text{ مایل}$$

$$B.P.D_2 = 100 - 36/6 = 63/3 \text{ مایل}$$

بنابراین در بین دو شهر □ و □ تا ۳۶/۶ مایل منطقه تجاری شهر □ بوده و ۶۳/۴ مایل در منطقه تجاری

شهر □ می باشد.

مدل پیش بینی فضاهای گذران اوقات فراغت

از کاربردهای دیگر مدل جاذبه، پیش بینی مراکز تفریح است. ابتداء با استفاده از نرخ های مشارکت

می توان تعداد شرکت کنندگان را به دست آورد سپس تسهلات لازم برای مراکز تفریحی فراهم آورد.

ساختار کلی مدل به شرح زیر است: (زیاری، ۱۳۸۱، صص ۱۵۶-۱۵۵)

(۱)

$$\square = \frac{N}{P}$$

در رابطه (۱):

□: نرخ مشارکت،



□: جمعیت،

□: تعداد مشارکت کنندگان.

از رابطه (۱)، رابطه (۲) را با طرفین وسطین کردن می توان به دست آورد:

(۲)

$$N = rP$$

آنگاه تعداد مشارکت کنندگان، به نیاز به تسهیلات تبدیل می شود.

کاربرد مدل جاذبه در پیش بینی مراکز تفریحی به شرح زیر است: (همان منبع)

(۳)

$$V_{ij} = \frac{GP_i A_i}{D_{ij}^b}$$

در رابطه (۳):

V_{ij} : تعداد بازدیدکنندگان از ناحیه به تسهیلات □،

P_i : جمعیت ناحیه □،

A_i : جذب تسهیلات □،

D_{ij}^b : فاصله بین □ و □،

□ و □: ضرایب ثابت.

مدل های دسترسی

مدل های دسترسی تقریباً شکل تکامل یافته مدل جاذبه اند که بوسیله اندیشمندان مختلف اصلاح، توسعه

و تکامل یافته اند. (پورمحمدی، ۱۳۸۲: ۶۷). جغرافیدانان و برنامه ریزان از دسترسی به عنوان یک

شاخص نام می برند. اسمیت[□] ۱۹۷۷ از دسترسی به عنوان مبنا در برنامه ریزی مکان یاد می کند، ضمن

اینکه پرد[□] ۱۹۷۷ و ناکس[□] هر دو معتقدند که اهمیت اندازه گیری «کیفیت زندگی»[□] در دسترسی به



خدمات و سرویس ها، عامل کلیدی است. (پرهیزگار، ۱۳۷۷، ص ۱۲۱). اشنایدر و سیمونز[□] (۱۹۷۱) شاخص فرصت دسترسی را به شرح زیر مطرح می کنند:

$$O_i = \sum_j \frac{S_j}{b_{ij}^t} \quad (۱)$$

که در این فرمول O_i شاخص فرصت، S اندازه ای از خدمات منطقه j و i و t زمان سفر و b شاخص کاهش فاصله است. (همان منبع، ص ۱۲۱).

ناکس با وارد کردن پارامترهایی چون سرعت مسافرت بوسیله سواری یا حمل و نقل عمومی O_i را به صورت زیر در می آورد: (همان منبع)

$$(۲)$$

$$T_{Aj} = C_i \left(\frac{A_i}{S_a} \right) + (100 - C_i) \left(\frac{A_i}{S_t} \right)$$

که در آن TA شاخص جدید دسترسی برای منطقه i ، C_i درصد اتومبیل شخصی در منطقه i و S_a و S_t به ترتیب زمان سفر در یک فاصله معین توسط سواری و وسیله نقلیه عمومی است. اینگرام[□] (۱۹۷۱) آن را با توجه به سنجه گوسی[□] به صورت زیر بیان می کند:

(۳)

$$A_i = \sum_j S_j^{exp} \left(\frac{-d_{ij}^\gamma}{p_\gamma} \right)$$

متغیرها در فرمول، همانند مدل های فوق الذکرند و γ مقداری ثابت است. ابرگ (Oberge) سنجه فرصت های تجمعی را به شرح زیر توسعه داده است:

(۴)

$$A_i = O_i(D) \left(D - \sum \right)$$

[□] Schneider and Symons

[□] Ingram

[□] Gaussian Measure



که در آن $O_i(D)$ مجموع فرصتهای قابل دسترسی به خانوار Δ در فاصله Δ مقیاسی از بیشینه (Max) فاصله پیاده روی محسوب می گردد. گای (Guy) مثالی از خرده فروشی را به صورت زیر بیان کرده است:

(۵)

$$A_i = \frac{\sum d^{min}_{ij} (K)}{\sum KEK}$$

که در آن d^{min}_{ij} حداقل فاصله مستقیم تا فروشگاه در j است که در آن کالای K وجود دارد و EK هزینه متوسط هر خانوار برای کالای K است.

اسمیت با استفاده از کار اشنایدر و سیمونز از طریق تقسیم مجموع حاصلضربهای AO مربوط به هر محل در جمعیت آنها N بر جمعیت کل به یک شاخص کارایی به شرح زیر دست یافت.

(6)

$$\square = \frac{\sum (AO_i \times N_i)}{\sum N_i}$$

کونینگ (Koenig) روش رفتاری سنجی های دسترسی را که بر مفهوم سودمندی استوار است، شرح می دهد. وی رابطه مناسب برای دسترسی را بر مبنای رفتارگرایی برای انتخاب مقصد بکار برده است:

(۷)

$$U^{t}_{ij} = \gamma^{t}_{ij} - C^{t}_{ij} + \varepsilon^t$$

که در آن U^{t}_{ij} سودمندی فرد $\square$ است که در $\square$ زندگی می کند و مقصدش در $\square$ قرار دارد، γ^{t}_{ij} سود ناخالص رسیدن فرد $\square$ به مقصد $\square$ است، C^{t}_{ij} یعنی هزینه سفر تولید شده یا وقت صرف شده فرد $\square$ برای سفر از $\square$ به $\square$ جمله تصادفی است. آنگاه مسئله، تشخیص یک تابع احتمالی مناسب برای متغیر تصادفی ε^t است. (همان منبع، صص ۱۲۴-۱۲۶).

مدل جاذبه ای / پتانسیلی هنسن

یکی از نخستین نمونه های کاربرد مدل جاذبه ای برای برنامه ریزی مدلی است که هنسن ابداع نمود. مدل هنسن مدلی مکانی است برای پیش بینی مکان جمعیت طراحی شده است. این مدل بر این فرض



استوار است که دسترسی به اشتغال عامل مهمی در تعیین مکان جمعیت است. در واقع این مدل یک مدل صرفاً جاذبه ای نیست زیرا بر اساس روابط متقابل میان مناطق ساخته نشده است. صحیح تر آن است که این مدل را به عنوان مدل پتانسیلی توصیف نمود، زیرا مدل هسن به روابط متقابل پتانسیلی یا دسترسی نسبی مناطق توجه دارد (لی، ۱۳۶۶: ۹۵).

هسن^۱ مطرح نمود که برای بیان رابطه میان مکان جمعیت و اشتغال می توان از شاخص دسترسی استفاده نمود. این شاخص دسترسی به اشتغال را برای هر منطقه تعیین می کند و به صورت زیر محاسبه می گردد.

(۱)

$$A_{ij} = \frac{E_j}{d_{ij}^b}$$

A_{ij} : شاخص دسترسی منطقه i در رابطه با منطقه j .

E_j : کل اشتغال در j .

d_{ij} : مسافت میان i و j .

b : ضریب یا توان d_{ij} می باشد.

این عبارت نماینگر دسترسی منطقه i در رابطه با یک منطقه j می باشد. شاخص کلی برای منطقه i مجموع تمام تک شاخص هاست. بنابراین:

(۲)

$$A_i = \sum_j \frac{E_j}{d_{ij}^b}$$

همچنین هسن متوجه شد که علاوه بر دسترسی، یکی از عوامل مهمی که جمعیت را به یک منطقه خاص جذب می کند، مقدار زمین موجود برای کاربرد مسکونی می باشد. او این زمین های خالی را امکان رشد^۲ یک منطقه نامید. بنابراین پتانسیل توسعه را به این صورت می توان نوشت:



(۳)

$$D_i = A_i H_i$$

H_i : امکان رشد منطقه i می باشد.

پتانسیل توسعه را می توان به عنوان میزان جاذبه هر منطقه بر اساس اشتغال و مقدار زمین مناسب برای توسعه مناطق مسکونی تصور نمود. جمعیت بر اساس پتانسیل نسبی توسعه هر منطقه به مناطق تخصیص داده می شود یعنی، پتانسیل توسعه هر ناحیه تقسیم بر پتانسیل توسعه تمام مناطق:

(۴)

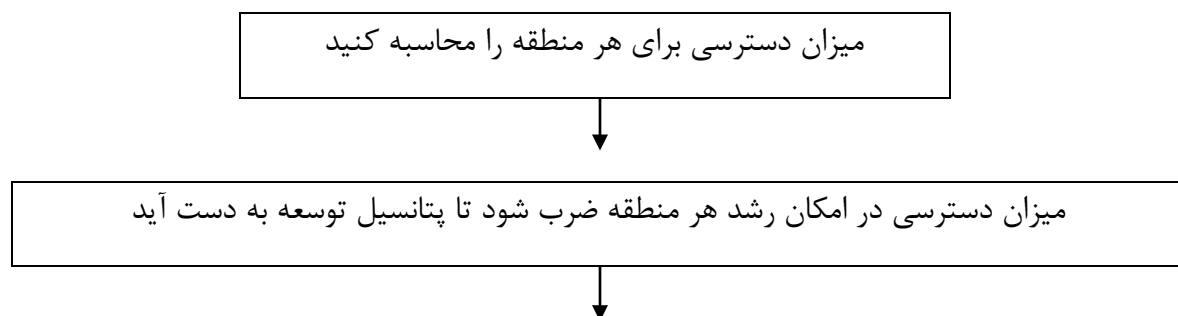
$$\frac{A_i H_i}{\sum_i A_i H_i}$$

به عبارتی، هسن این مسئله را مطرح نمود که سهم رشد کل جمعیت هر منطقه به میزان جاذبه آن منطقه در رابطه با تمام مناطق رقیب می باشد. اگر رشد کل جمعیت C_t باشد، میزان سهم رشد در هر منطقه i برابر است با:

(۵)

$$G_i \frac{D_i}{\sum_i D_i} \text{ یا } G_i = G_i \frac{(A_i H_i)}{(\sum_i A_i H_i)}$$

شمای کلی گردش مدل هسن را می توان در نمودار زیر نشان داد:





قابلیت توسعه هر منطقه جمع شود تا پتانسیل کلی توسعه به دست آید



قابلیت توسعه هر منطقه بر پتانسیل کلی توسعه تقسیم شود تا پتانسیل نسبی توسعه هر منطقه به دست آید



پتانسیل نسبی توسعه در رشد کل جمعیت ضرب شود تا رشد جمعیت هر منطقه به دست آید

نمودار (۵-۱): نمای گردش مدل هنس

مثال: در یک منطقه شهری فرضی به چهار ناحیه تقسیم شده است و تعداد کل شاغلان برای هر ناحیه مشخص شده است به دنبال تخصیص بهینه جمعیت بین نواحی شهری از طریق مدل هنس هستیم و فرض بر این است که توان به دست آمده از عمل تنظیم برابر ۲ است (حکمت نیا و موسوی، ۱۳۸۵: ۱۴۹).

جدول شماره (۸): داده های جمعیت و اشتغال برای منطقه فرضی

ناحیه	شرح	کل شاغلان	مجموع جمعیت	امکان رشد
۱		۴۵۰۰	؟	۱۵۰
۲		۳۵۰۰	؟	۱۰۰
۳		۶۵۰۰	؟	۱۵۰
۴		۲۵۰۰	؟	۱۲۵
مجموع		۱۷۰۰	۳۸۰۰	۵۲۵



جدول شماره (9): ماتریس مساحت/ زمان سفر

از \ به	D=1	D=2	D=3	D=4
D=1	2	6	3	8
D=2	6	3	3	4
D=3	3	3	4	6
D=4	8	4	6	2

مرحله اول: محاسبه شاخص دسترسی برای هر ناحیه از طریق رابطه (۱)

جدول شماره (10): محاسبه شاخص دسترسی برای هر ناحیه

ناحیه	D=1	D=2	D=3	D=4	
D=1					۱۴۸۲
D=2					۱۳۹۲
D=3					۱۳۶۵
D=4					۱۰۱۶

مرحله دوم: محاسبه کل پتانسیل توسعه برای هر یک از نواحی از طریق رابطه (۳)

جدول (11): محاسبه پتانسیل توسعه

ناحیه			
۱	۱۹۸۴	۱۵۰	۲۹۷۶۰۰
۲	۱۳۹۲	۱۰۰	۱۳۹۲۰۰
۳	۱۳۶۵	۱۵۰	۲۰۴۷۵۰
۴	۱۰۱۶	۱۲۵	۱۲۷۰۰

$$\sum_i D_i = ۷۶۸۵۵۰$$



مرحله سوم: محاسبه پتانسیل نسبی توسعه برای هر ناحیه، یعنی جاذبه آن با جاذبه تمام مناطق دیگر.

جدول شماره (12): محاسبه پتانسیل نسبی توسعه

ناحیه		
۱	۲۹۷۶۰۰	۰/۳۸۲
۲	۱۳۹۲۰۰	۰/۱۸۲
۳	۲۰۴۷۵۰	۰/۲۶۶
۴	۱۲۷۰۰۰	۰/۱۶۵
جمع	۷۶۸۵۵۰	۱

مرحله چهارم: تخصیص جمعیت به هر یک از نواحی با استفاده از جمعیت پیش بینی شده از طریق رابطه

(۵)

جدول شماره (13): پیش بینی جمعیت و تخصیص آن در هر یک از نواحی

ناحیه		
۱	۰/۳۸۲	۱۴۷۰۶
۲	۰/۱۸۲	۶۹۱۶
۳	۰/۲۶۶	۱۰۱۰۸
۴	۰/۱۶۵	۶۲۷۰
جمع	۱/۰۰۰	۳۸۰۰

مدل آنتروپی شانون^۱

از این مدل برای تجزیه و تحلیل و تعیین مقدار پدیده رشد بی قواره شهری^۲ استفاده می گردد. ساختار

کلی مدل به شرح زیر است: (Sudhira, et.al, 2003, PP.299-311)

^۱ Shannon's Entropy Model
^۲ Urban Sprawl Phenomenon



$$H = -\sum_{i=1}^n P_i \times \ln(P_i)$$

در رابطه بالا:

□: مقدار آنتروپی شانون،

□□: نسبت مساحت ساخته شده (تراکم کلی مسکونی) منطقه □ به کل مساحت ساخته شده مجموع

مناطق،

□: مجموع مناطق.

ارزش مقدار آنتروپی شانون از صفر تا $\ln(n)$ است. مقدار صفر بیانگر توسعه فیزیکی خیلی متراکم (فشرده) شهر است. در حالی که مقدار $\ln(n)$ بیانگر توسعه فیزیکی پراکنده شهری است. زمانی که ارزش آنتروپی از مقدار $\ln(n)$ بیشتر باشد رشد بی قواره شهری (اسپرال) اتفاق افتاده است. برای روشن تر شدن مطلب به مثال زیر توجه نمایید:

در یک شهر فرضی سه منطقه A و B و C را داریم که در سال ۱۳۷۵، مساحت ساخته شده هر کدام از مناطق به ترتیب ۴، ۹، ۶ هکتار بوده است. در سال ۱۳۸۵ این ارقام به ترتیب ۸، ۱۲ و ۱۰ هکتار افزایش یافته است. حال می خواهیم با استفاده از مدل آنتروپی شانون، چگونگی رشد فیزیکی را مشخص نماییم.

جدول شماره (14): محاسبه آنتروپی شانون برای سال 1375 در شهر فرضی

منطقه	مساحت ساخته شده (هکتار)		$\ln(P_i)$	$P_i \times \ln(P_i)$
۱	۴		-۱/۵۵۸۲	-۰/۳۲۸۰



۲	۹		-۰/۷۴۷۳	-۰/۳۵۳۹
۳	۶		-۱/۱۵۲۳	-۰/۳۶۴۰
کل	۱۹		$\sum P_i \times \square(P_i)=$	-۱/۰۴۵۹

H=1/0459

جدول شماره (15): محاسبه آنتروپی شانون برای سال 1385 در شهر فرضی

منطقه	مساحت ساخته شده (هکتار)		$\square(P_i)$	$P_i \times \square(P_i)$
۱	۸		-۱/۳۲۲۰	-۰/۳۵۲۴
۲	۱۲		-۰/۹۱۶۳	-۰/۳۶۶۵
۳	۱۰		-۱/۰۹۸۴	-۰/۳۶۶۲
کل	۳۰		$\sum P_i \times \square(P_i)=$	-۱/۰۸۵۱

H=1/0851



جداول شماره (۴-۴) و (۴-۵) نشان می دهد مقدار آنتروپی در سال ۱۳۷۵ برابر ۱/۰۴۵۹ بوده است، در حالی که حداکثر ارزش $Ln(۳) = ۱/۰۹۸$ است. نزدیک بودن مقدار آنتروپی به مقدار حداکثر نشانگر رشد پراکنده توسعه فیزیکی شهری است. این مقدار در سال ۱۳۸۵، ۱/۰۸۵ بوده است که نشان می دهد طی ده سال توسعه فیزیک به صورت پراکنده و غیرمترکم بوده است (حکمت نیا و موسوی، ۱۳۸۵: ۱۲۹).

ضریب آنتروپی

این مدل، معیاری برای سنجش توزیع جمعیت شهری و توزیع تعداد شهرها در طبقات شهری یک منطقه است. با استفاده از این مدل، می توان به میزان تعادل فضایی استقرار جمعیت و تعداد شهرها در سطح شبکه شهری، استانی، منطقه ای و ملی پی برد. ساختار کلی مدل به شرح زیر است: (Wheeler & Muller 1986, PP.384-385)

$$H = -\sum P_i \ln P_i$$

$$G = \frac{H}{\ln K}$$

□: مجموع فراوانی در لگاریتم نپری فراوانی،

P_i : فراوانی،

$Ln P_i$: لگاریتم نپری فراوانی،

K : تعداد طبقات،

□: میزان آنتروپی.

اگر آنتروپی به سمت صفر میل کند حکایت از تمرکز بیشتر و یا افزایش تمرکز یا عدم تعادل در توزیع جمعیت بین شهرها دارد و حرکت به طرف یک و بالاتر از آن توزیع متعادل تری را در عرصه منطقه ای نشان می دهد.

مثال: توزیع فضایی تعداد شهرها در طبقات شهری استان آذربایجان غربی با استفاده از ضریب آنتروپی مورد محاسبه واقع گردید. **جدول شماره (۴-۶)** نشان می دهد که میزان آنتروپی در سال ۱۳۶۵، ۰/۸۷ بوده که در سال ۱۳۷۵ به ۰/۸ کاهش یافته است. یعنی توزیع فضایی تعداد شهرها در طبقات شهری استان نسبت به ۱۳۶۵ تا حدودی نامتعادل شده است (حکمت نیا و موسوی، ۱۳۸۵: ۱۸۹).



جدول شماره (16): محاسبات تغییرات ضریب آنتروپی در طبقات شهری استان آذربایجان غربی

1375			1365			سال
						طبقات شهری
-0/1395	-3/1	0/045	-0/271	-1/99	0/136	کمتر از 4999
-0/336	-1/48	0/227	-0/3094	-1/7	0/182	5000-9999
-0/3094	-1/7	0/182	-0/3657	-1/15	0/318	10000-24999
-0/336	-1/48	0/227	-0/1395	-3/1	0/045	25000-49999
-0/271	-1/99	0/136	-0/33596	-1/48	0/227	50000-99999
-0/367	-1/02	0/36	-0/1395	-3/1	0/045	100000-249999
-0/36	-0/799	0/45	-0/1364	-3/1	0/045	250000-499999
-2/12	-11/569	1	-1/697	-15/62	1	Σ

$$H_{1365} = -1/697$$

$$H_{1375} = -11/569$$

$$G_{1365} = 0/87$$

$$G = 0/80$$

$$H=7 \rightarrow Ln_7 = 1/95$$

مدل قانون رتبه - اندازه^{۱۱}

اولین تجزیه و تحلیل جغرافیایی توزیع اندازه شهرها در نظام های شهری به اوایل قرن بیستم برمی گردد. فلیکس اوئرباخ^{۱۲}، جغرافیدان آلمانی در سال ۱۹۱۳ قانون مرتبه اندازه شهری را ارائه داد که بین اندازه شهرها و رتبه آنها رابطه معکوس وجود دارد (Carter, 1981, P.70). بعدها توسط کسانی همچون لوتکا^{۱۳} (۱۹۲۴)، گودریچ^{۱۴} (۱۹۲۶) و سینگر^{۱۵} (۱۹۳۶) مورد استفاده قرار گرفت (Delgado & Godinho, 2004, P.1). بالاخره در سال ۱۹۴۹ این نوع بررسی در شهرها توسط جورج زیپف (G. Zipf K.) به صورت کامل فرمول بندی و مورد عمل و بررسی واقع گردید. زیپف شهر دوم حدود^{۱۶}

^{۱۱} The Rank- Size Rule
^{۱۲} Felix. Auerbach
^{۱۳} Lotka
^{۱۴} Goodrich
^{۱۵} Singer



جمعیت شهر اول، شهر درجه سوم حدود^{۱۳} - شهر نخست و بالاخره جمعیت شهر □ حدود $\frac{1}{n}$ جمعیت شهر اول خواهد بود. او معتقد بود وجود همبستگی خطی مطرح است. بنابراین هر اندازه سیستم شهری یک کشور توسعه پیدا کند به توزیع نرمال نزدیک تر است. (Clark, 2000, PP.25-28) رابطه ریاضی چنین مفهومی را می توان به شرح زیر عنوان کرد: (Guerin, 1995, PP. 551-562)

$$P_r = \frac{p_1}{R^b}$$

p_1 : جمعیت شهر نخست در منطقه مورد نظر،

P_r : جمعیت شهر در مرتبه مورد نظر یا جمعیت شهر Rام،

R : مرتبه شهر در منطقه،

b : شیب خط مرتبه- اندازه.

سلسله مراتب شهرهای استان آذربایجان غربی با استفاده از تئوری زیپف نشان داده شده است. جدول شماره (۵-۶) بیانگر آن است که جمعیت شهر ارومیه به عنوان اولین شهر ۳ برابر دومین شهر (خوی)، ۴ برابر سومین شهر آن (بوکان) و ۱۴۲ برابر آخرین شهر یا ۲۲ام شهر (قوشچی) است.

جدول شماره (17): اندازه واقعی و تئوری مرتبه- اندازه شهرهای آذربایجان غربی 1375

اسم شهر	مرتبه	تعداد جمعیت واقعی	تعداد جمعیت در ارتباط با تئوری مرتبه- اندازه
ارومیه	۱	۴۳۵۲۰۰	۴۳۵۲۰۰



۲۱۷۶۰۰	۱۴۸۹۴۴	۳	خوی
۱۴۵۰۶۷	۱۲۰۰۲۰	۴	بوکان
۱۰۸۸۰۰	۱۰۷۷۹۹	۵	مهاباد
۸۷۰۴۰	۹۰۱۴۱	۶	میاندوآب
۷۲۵۳۳	۶۵۴۱۶	۷	سلماس
۶۲۱۷۱	۶۴۸۰۷	۸	نقده
۵۴۴۰۰	۴۲۵۶۹	۱۰	تکاب
۴۸۳۵۶	۳۳۸۰۵	۱۳	پیرانشهر
۴۳۵۲۰	۳۳۴۰۶	۱۴	ماکو
۳۹۵۶۴	۳۰۹۰۴	۱۵	سردشت
۳۶۲۶۷	۲۹۰۲۰	۱۶	شاهین دژ
۳۳۴۷۷	۲۳۵۶۹	۱۸	اشنویه
۳۱۰۸۶	۲۰۲۶۶	۲۱	قره ضیالدين
۲۹۰۱۳	۱۷۴۸۲	۲۵	شوط
۲۷۲۰۰	۱۳۰۱۲	۳۳	سیه چشمه
۲۵۶۰۰	۸۰۵۰	۵۳	فیروزق
۲۴۱۷۸	۷۶۸۶	۵۷	پلدشت
۲۲۹۰۵	۷۴۶۶	۵۸	تازه شهر
۲۱۷۶۰	۶۷۹۷	۶۴	محمدیار
۲۰۷۲۳	۵۷۰۸	۷۸	نوشین
۱۹۷۸۱	۳۰۹۳	۱۴۲	قوشچی

جورج زیپف رابطه (۱) را به صورت رابطه لگاریتمی بیان می کند که در آن شکل توزیع اندازه شهری حالت خاصی از توزیع پارتو است. در واقع در این رابطه $b=1$ است. در آن صورت جمعیت شهر n برابر $\frac{1}{n}$ شهر نخست خواهد بود. اگر $b=0$ باشد تمام شهرها به یک اندازه خواهد بود و اگر $b=\infty$ باشد فقط یک شهر وجود خواهد داشت. (Alperovich, 1984, PP.232-239).

(۲)

$$P_r = \log P_1 - b \log R$$

(۳)



$$\square = \frac{\text{Log } P_1 - \text{Log } R}{\text{Log } P_r}$$

بهترین روش برای استفاده توزیع رتبه-اندازه شهری و به دست آوردن شیب خط (b) استفاده از مدل رگرسیونی، روش حداقل مربعات است که مقدار b هرچقدر به طرف ۱ میل کند توزیع اندازه شهری به طرف توزیع لگاریتمی نرمال سوق خواهد نمود. اگر مقدار $b < 1$ باشد نشان دهنده اهمیت نسبی شهرهای متوسط و میانی در نظام شهری و اگر $b > 1$ باشد حاکی از تسلط نخست شهری در نظام شهری است. (Nishiyama, et.al, 2005).

ساختار ریاضی آن به شکل زیر است:

(۴)

$$y = a + bx$$

□: شیب خط،

□: مقدار ثابت،

□: لگاریتم رتبه شهر،

□: لگاریتم اندازه (جمعیت شهر).

توزیع لگاریتمی رتبه-اندازه شهرهای استان آذربایجان غربی در ۱۳۷۵ با استفاده از مدل رگرسیونی به دست آمده است که نشان می دهد:

- ۱- همبستگی معکوس قوی بین لگاریتم مرتبه شهرها (□) و لگاریتمی اندازه شهرها (□) برقرار بوده است. در واقع هرچه لگاریتم رتبه ها افزوده می شود از میزان لگاریتم جمعیت آنها کاسته می شود.
- ۲- مقدار ضریب خط یا شیب خط رتبه-اندازه با خط تعادل برای سال مورد بررسی ۱/۴۳۷- بوده است. یعنی عدم تعادل در شیب خط رگرسیون در نظام و شبکه شهری استان آذربایجان غربی برقرار بوده است که اعداد به دست آمده نشان دهنده تسلط نخست شهر (ارومیه) بر شبکه شهری استان است.
- ۳- معادله خط برای شبکه شهری استان به شرح ذیل به دست آمده است:



$$y = -1/4376x + 5/8$$

جدول شماره (18): لگاریتم رتبه- اندازه شهرهای استان آذربایجان غربی در سال 1375

رتبه شهر	نام شهر	جمعیت شهر	LogR(x)	LogP(y)
۱	ارومیه	۴۳۵۲۰۰	۰	۵/۶۳۹
۲	خوی	۱۴۸۹۴۴	۰/۳۰۱	۵/۱۷۳
۳	بوکان	۱۲۰۰۲۰	۰/۴۷۷	۵/۰۷۹
۴	مهاباد	۱۰۷۷۹۹	۰/۶۰۲	۵/۰۳۳
۵	میاندوآب	۹۰۱۴۱	۰/۶۹۹	۴/۹۵۵
۶	سلماس	۶۵۴۱۶	۰/۷۷۸	۴/۸۱۶
۷	نقده	۶۴۸۰۷	۰/۸۴۵	۴/۸۱۲
۸	تکاب	۴۲۵۶۹	۰/۹۰۳	۴/۶۲۹
۹	پیرانشهر	۳۳۸۰۵	۰/۹۵۴	۴/۵۲۹
۱۰	ماکو	۳۳۴۰۶	۱	۴/۵۲۴
۱۱	سردشت	۳۰۹۰۴	۱/۰۴۱	۴/۴۹۰
۱۲	شاهین دژ	۲۹۰۲۰	۱/۰۷۹	۴/۴۶۳
۱۳	اشنویه	۲۳۵۶۹	۱/۱۱۴	۴/۳۷۲
۱۴	قره ضیالالدین	۲۰۲۶۶	۱/۱۴۶	۴/۳۰۶
۱۵	شوط	۱۷۴۸۲	۱/۱۷۶	۴/۲۴۳
۱۶	سیه چشمه	۱۳۰۱۲	۱/۲۰۴	۴/۱۱۴
۱۷	فیروزق	۸۰۵۰	۱/۲۳۰	۳/۹۰۶
۱۸	پلدشت	۷۶۸۶	۱/۲۵۵	۳/۹۳۹
۱۹	تازه شهر	۷۴۶۶	۱/۲۷۹	۳/۸۷۳
۲۰	محمدیار	۶۷۹۷	۱/۳۰۱	۳/۸۳۲
۲۱	نوشین	۵۷۰۸	۱/۳۲۲	۳/۷۵۶
۲۲	قوشچی	۳۰۹۳	۱/۳۴۲	۳/۴۹۰

رسل^۱ در تجزیه و تحلیل رتبه- اندازه شهری برای شهرهای انگلستان از فرمول زیر استفاده کرد. با استفاده از فرمول و استفاده از اندازه جمعیتی شهر نخست، جمعیت شهرهای رتبه های از ۲ تا ۱ را به



دست می آورد. به صورتی که به توزیع جمعیتی نرمال و رتبه-اندازه باشد. ساختار کلی مدل به شرح ذیل است: (Smith, 1967, PP.304-305)

$$P_r = P_1 \frac{\left(1 + \sqrt{r-1}\right)}{r}$$

P_r : جمعیت شهر در رتبه r در یک منطقه،

P_1 : جمعیت شهر نخست در یک منطقه،

r : رتبه اندازه شهر.

البته فرمول فوق در نظام های شهری مثل اکثر کشورهای در حال توسعه دارای الگوی نخست شهری^۱ هستند نمی تواند درست باشد. هیچ دلیلی وجود ندارد که شهرهای درجه دوم سوم تا n ام توزیع اندازه جمعیت شهری از شهر نخست پیروی کنند. در این گونه نظام شهری، تمرکز زدایی از نخست شهر با ارائه راهبردها و استراتژی های کلان می تواند راه حل بهینه باشد. در این زمینه خانم دکتر فاطمه بهفروز به ارائه فرمولی مبادرت ورزیده که در اکثر کشورهای جهان سوم و نظام های شهری آنها می تواند الگوی مناسب باشد. ساختار کلی مدل به شرح زیر است: (بهفروز، ۱۳۷۴، ص ۳۳)

$$P_{rth} = \frac{\sum P_{1-n} \div R_{rth}}{\sum \frac{1}{R_1} + \frac{1}{R_2} + \dots + \frac{1}{R_n}}$$

P_{rth} : جمعیت هر شهری که در مرتبه r قرار دارد،

$\sum P_{1-n}$: مجموع جمعیت واقعی شهرهای مورد مطالعه

R_{rth} : مرتبه شهر r ،

$\sum \frac{1}{R_1} + \frac{1}{R_2} + \dots + \frac{1}{R_n}$: مجموع نسبت های مرتبه ای تمام شهرهای مورد مطالعه.

فرمولی که استاد گرانقدر مبادرت به طراحی آن نموده است از مجموع جمعیتی شهرها استفاده نموده است که با استفاده از این فرمول می توان در نظام شهری دارای الگوی نخست شهر استفاده نمود. استفاده از این فرمول، جمعیت شهر نخست را کاهش می دهد و شهرهای رتبه های بعدی را به صورتی تنظیم می نماید که به الگوی رتبه-اندازه شهری نزدیک تر گردند.



روش تعیین نخست شهر

مارک جفرسون^۱ جغرافیدان آمریکایی در سال ۱۹۳۹، در مقاله ای تخصصی و نوگرایانه «قانون نخست شهر»^۲ را ارائه نمود. بر اساس این قانون، نخست شهر در هر کشور همیشه به صورت یک شهر مستقل و بزرگ مورد توجه بوده و استثنأ بیان کننده توانایی و احساس ملی آن کشور می باشد (بهفروز، ۱۳۷۴، ۳۱۹). به عبارت دیگر نخست شهر عبارت است از تسلط جمعیتی، اقتصادی، اجتماعی و سیاسی یک شهر بر تمامی شهرهای دیگر در داخل یک نظام شهری (دراکاکیس، اسمیت، ۱۳۷۷، ص ۱۳). مارک جفرسون برای تعیین نخست شهر در چهل و چهار کشور پیشرفته جهان از «روش نسبی»^۳ استفاده نمود، که محاسبه بر اساس نسبت شهر نخست به شهر دوم انجام می گرفت $\left(\frac{P_1}{P_2}\right)$. کلارک پیشنهاد کرد که بهتر است به جای دو شهر، چهار شهر اول نظام شهری برای محاسبه انتخاب گردد. از همین رو الو اصطلاح «شاخص چهار شهر»^۴ را پیشنهاد کرد که در آن نسبت شهر نخست به سه شهر بعدی نظام شهری (مجموعاً چهار شهر) به کار گرفته می شد به صورت زیر:

$$\frac{P_1}{P_2 + P_3 + P_4}$$

که در آن P_1 جمعیت بزرگترین شهر و P_2 و P_3 و P_4 به ترتیب شهرهای بعدی به حساب می آمدند. سرانجام «مهتا»^۵ (۱۹۶۴) با اصلاح فرمول کلارک، بهترین روش برای تشخیص نخست شهری را نسبت اندازه شهر نخست به چهار شهر اول نظام شهری به صورت زیر پیشنهاد کرد:

$$\frac{P_1}{P_1 + P_2 + P_3 + P_4}$$

که در آن $\square$ با همان مفهوم در فرمول کلارک جمعیت شهر اول تا چهارم بود. بعدها ریچاردسون شاخص چهار شهر را با معیارهای قاعده رتبه- اندازه تطبیق داد. بدین صورت که اگر بر اساس قاعده رتبه- اندازه شهری، اندازه مطلوب شهرها در نظام شهری این گونه باشد که شهر او دو

^۱ Mark. Jefferson
^۲ The Law of Primate City
^۳ Proportion Technique
^۴ Four-City Index
^۵ Mehta



برابر شهر دوم، سه برابر شهر سوم و چهار برابر شهر چهارم باشد، بنابراین نسبت شهر اول به مجموع چهار شهر نخست نظام شهری باید برابر ۰/۴۸ باشد.

$$[(1 \div (1 + 0/5 + 0/33 + 0/25)) = 0/48]$$

که این توزیع بهترین و عادی ترین شکل برتری شهری خواهد بود. بر پایه چنین معیاری، درجه تسلط و برتری شهر اول بر نظام شهری بر اساس جدول زیر پیشنهاد شده است، که در آن دامنه تسلط و برتری مطلوب شهر نخست بین شاخص ۰/۴۱ تا ۰/۵۴ فرض شده است و برای فوق برتری، شاخص بین ۰/۶۵ تا ۱ پیشنهاد شده است (عظیمی، ۱۳۸۱: ۶۷-۵).

جدول شماره (19): درجه نخست شهری در نظام شهری بر پایه شاخص چهار شهر

نوع برتری شهری	شاخص چهار شهر
فوق برتری	۰/۶۵ تا ۱
برتری	۰/۵۴ تا ۰/۶۵
برتری مطلوب	۰/۴۱ تا ۰/۵۴
حداقل برتری	کمتر از ۰/۴۱

منبع: عظیمی، ۱۳۸۱: ۶۷.

مرور مختصری بر روش های ارزیابی در برنامه ریزی شهری و منطقه ای

برای تبیین مفهوم ارزیابی لازم است ابتدا بر فرایند برنامه ریزی مروری مختصر داشته باشیم. فرایند برنامه ریزی تلاش می کند تا چارچوبی مناسب را فراهم آورد که طی آن برنامه ریز بتواند برای رسیدن به راه حل بهینه اقدام کند (Lee, 1973: 2). این فرایند ابتدا به وسیله پتريک گدس در ۳ مرحله برداشت، تحلیل و طرح (برنامه) تبیین گردید. در سال ۱۹۴۷ با تصویب قانون برنامه ریزی شهر و روستا در انگلستان مورد تأیید و تأکید قرار گرفت. در دهه ۱۹۶۰، با پیدایش نگرش سیستمی، کوشش هایی برای تعریف مجدد فرایند برنامه ریزی صورت گرفت. برخی از محققین این فرایند را در ۷ مرحله و برخی دیگر در ۱۱ مرحله قابل انجام دانسته اند. در کلیه این فرایندها ارزیابی "به عنوان یکی از ارکان مهم فرایند برنامه ریزی مورد تأکید بوده است. به ترتیب که بعد از تعیین اهداف کلی و مقاصد برنامه ریزی و تهیه گزینه های مختلف، ارزیابی صورت می پذیرد تا با مقایسه گزینه های مختلف، بر اساس شایستگی نسبی



آنها گزینه یا آلترناتیو مطلوب انتخاب شود. در برنامه ریزی شهری و منطقه ای، روش های ارزیابی متعددی مورد استفاده قرار گرفته است. لیچفیلد و دیگران روش های ارزیابی مهمی که بیشتر مورد توجه و کاربرد بوده است را عنوان کرده اند که ما بیشتر آنچه در این پژوهش به آن خواهیم پرداخت را اشاره می کنیم :

- روش فهرست معیارها؛
- روش ماتریس دستیابی به اهداف ؛
- روش های ارزیابی بهینه یابی.

رابرتز روش های ارزیابی به کاربرده شده در زمینه برنامه ریزی را به دو گروه روش های ارزیابی جزئی و روش های ارزیابی جامع طبقه بندی می کند. معرفی روش های ارزیابی چند معیاری در سال های اخیر طبقه بندی های جدیدی را مطرح کرده اند. فالودی و وگد روش های ارزیابی به کار گرفته شده در برنامه ریزی شهری و منطقه ای را در ۳ گروه زیر طبقه بندی می کند:

الف) روش های ارزیابی پولی که در آنها چارچوب ارزیابی بر منبای مقادیر پولی صورت می پذیرد، مثل روش تحلیل تأثیر هزینه، روش هزینه - فایده و روش تحلیل آستانه ای.

ب) روش های ارزیابی جامع (کلی) که تنها پیامدهای مالی و پولی بلکه اثرات و پیامدهای غیر پولی گزینه ها نیز مورد تحلیل قرار می گیرند، مثل جدول ترازنامه برنامه ریزی و تحلیل تأثیر بر جامعه.

ج) روش های ارزیابی چند معیاری که در آنها امکان تحلیل و ارائه کلیه اطلاعات موجود در مورد گزینه ها بر اساس معیارهای متفاوت و چند بُعدی وجود دارد. این روش های ارزیابی ممکن است کاملاً کمی باشند (مثل روش ماتریس دستیابی به اهداف)، یا کلاً کیفی باشند

(مثل روش تحلیل نظام) و یا ترکیبی از اطلاعات کیفی و کمی (مثل روش های تحلیل اثرات زیست محیطی)

روش ارزیابی فرایند تحلیل سلسله مراتبی (AHP) جزو روش های ارزیابی چند معیاری است که در این مقاله به بررسی قابلیت کاربرد آن در برنامه ریزی شهری و منطقه ای خواهیم پرداخت .

چارچوب مفهومی فرایند تحلیل سلسله مراتبی

در ارزیابی هر موضوعی ما نیاز به معیار اندازه گیری با شاخص داریم، انتخاب شاخص مناسب به ما امکان میدهد که مقایسه درستی بین جایگزینیها یا آلترناتیوها به عمل آوریم. اما وقتی که چند یا



چندین شاخص برای ارزیابی در نظر گرفته میشود، کار ارزیابی پیچیده می کند و پیچیدگی کار زمانی بالا میگیرد که معیارهای چند یا چندینگانه باهم در فضا و از جنسهای مختلف باشند. در این هنگام کار ارزیابی و مقایسه از حالت ساده تحلیلی که ذهن قادر به انجام آن است خارج میشود و به یک ابزار تحلیل عملی قوی نیاز خواهد بود. یکی از ابزارهای توانمند برای چنین وضعیت هایی فرآیند تحلیل سلسله مراتبی است. (اسماعیلیان، ۱۳۸۹: ۱۲۴) این فرآیند یک سنتز ریاضی و یک شیوه جبری تصمیمگیری با مقیاس نسبی است. این روش با استفاده از یک شبکه سیستمی، شاخص های مختلف و ضوابط و معیارهای چندگانه با ساختارهای چند سطحی اولویتدار برای رتبهبندی یا تعیین اهمیت گزینههای مختلف یک فرآیند تصمیم گیری پیچیده مورد استفاده قرار میگیرد (معصومزاده و تراب زاده، ۱۳۸۳: ۷۰). روش AHP یکی از معروفترین فنون تصمیم گیری چندمنظوره است که در سال ۱۹۷۰ ابداع گردید. این روش هنگامی که عمل تصمیمگیری با چند گزینه رقیب و معیار تصمیم گیری روبروست میتواند مورد استفاده قرار گیرد. فرآیند AHP ترکیب معیارهای کیفی همراه با معیارهای کمی را به طور همزمان امکا نپذیر میسازد. اساس روش AHP مقایسههای زوجی یا دوبه دویی آلترناتیوها و معیارهای تصمیم گیری است (قدسی پور، ۱۳۷۹: ۶۷). برای چنین مقایسههای نیاز به جمع آوری اطلاعات از تصمیم گیرندگان است. این امر به تصمیم گیرنده این امکان را میدهد که فارغ از هرگونه نفوذ و مزاحمت خارجی تنها روی مقایسه دو معیار یا گزینه تمرکز کند. علاوه بر این مقایسه دو به دویی، به دلیل این که پاسخ دهنده فقط دو عامل را نسبت به هم میسنجد و به عوامل دیگر توجه ندارد، اطلاعات ارزشمندی را برای مسئله مورد بررسی فراهم میآورد و فرآیند تصمیم گیری را منطقی میسازد علاوه بر این، بر مبنای مقایسه زوجی بنا نهاده شده است که قضاوت و محاسبات را تسهیل میکند. همچنین میزان سازگاری و ناسازگاری تصمیم را نشان میدهد که از مزایای ممتاز این تکنیک در تصمیمگیری چندمعیاره است (باکویی، زنده روح؛ ۱۳۸۶: ۲۷).

فرآیند تحلیل سلسله مراتبی روشی است منعطف، قوی و ساده برای تصمیمگیری در شرایطی که معیارهای تصمیمگیری متضاد؛ انتخاب بین گزینهها را با مشکل مواجه میسازد (زبردست، ۱۳۸۰: ۱) و باید تصمیمگیری در یک فضای چند بعدی صورت پذیرد. در چنین شرایطی روشهای ارزیابی چند



معیاری؛ با توجه به این که در این روشها فرض بر این است که هر یک از معیارها محور یا جداگانهای هستند (توفیق، ۱۳۷۲: ۴۰). مورد استفاده قرار میگیرند.

این روش اکثراً برای تصمیمگیری در آنالیزهای عملیاتی و با ریسک بالا، جهت ارزیابی طرحهای جایگزین و در سطح کوچکتر جهت ارزیابی اثرات محیطی به کار برده میشود (Solnes، ۲۰۰۳: ۲۹۵). در این روش جدولی به نام جدول مقایسه دوبهدویی وجود دارد که امتیاز دهی براساس آن انجام میشود.

جدول شماره (20): مقایسه زوجی یا دویه دویی ال ساعتی

درجه اهمیت	امتیاز
به شدت مرجح	9
بسیار مرجح تا به شدت مرجح	8
بسیار مرجح	7
با ارجحیت قوی تا بسیار قوی	6
با ارجحیت قوی	5
با ارجحیت متوسط تا قوی	4
با ارجحیت متوسط	3
با ارجحیت متوسط تا مساوی	2
با ارجحیت مساوی	1

منبع: Al- Subhi Al-Harbi، 2001

مزیت این روش در این است که فاکتورهای عینی و ذهنی را در یک ساختار سیستماتیک شرکت داده و راه حلی ساده را در جهت حل مسائل تصمیمگیری جایگزین و یا ارائه می نماید (Al- Subhi Al-Harbi، ۱۹: ۲۰۰۱). روش AHP بر سه پایه شکل گرفته است، اول: ساختار مدل، دوم: مقایسه جایگزینها و سوم: نتیجه گیری از معیارها. امروزه استفاده از مدل AHP در جهت حل بسیاری از مسائل پیچیده در امر تصمیم گیری گسترش یافته است (Pakdin Amiri، ۱۹: ۲۰۱۰). به صورت دقیق تر می توان گفت که توسط AHP به مساله تصمیم گیری ساختار داده می شود و سپس گزینه های مختلف موجود بر اساس معیارهای مطرح در تصمیم گیری با هم مقایسه شده و اولویت انتخاب هر یک از آنها مشخص می شود (رهنما، ۱۳۸۸: ۴۲۳).

اساس روش AHP، نمایش سلسله مراتبی است که به حل مسائل پیچیده از طریق فرآیندهای ساده کمک می کند. این امر مستلزم این می باشد که معیارها به سطوح یکنواختی تجزیه شده باشد (Mepal



and others, ۲۰۱۰). کاربرد فضایی این مدل در قالب سیستم اطلاعات جغرافیایی توسط اوسوالد

مارینیوی در نرم افزار GIS Arc به کار گرفته شد (رهنما، ۱۳۸۸: ۴۲۳).

مراحل فرایند تحلیل سلسله مراتبی در تجزیه و تحلیل داده ها

برای توضیح مراحل فرایند تحلیل سلسله مراتبی از مثال زیر استفاده خواهد شد. فرض کنید از سه سایت A، B و C که به عنوان گزینه های مورد نظر برای مکان یابی مشخص شده اند. قرار است سایت مناسب برای اسکان براساس چهار معیار پستی و بلندی، دسترسی به آب، خطر زمین لرزه و دسترسی به شهر و راه انتخاب شود. معیار پستی و بلندی به دور زیر معیار ارتفاع از سطح دریا و شیب زمین؛ معیار خطر زمین لرزه به سه زیر معیار موقعیت سایت نسبت به گسله های بنیادی، موقعیت نسبت به رو مراکز زمین لرزه و شدت زلزله های تاریخی و دسنگاهی و معیار دسترسی به شهر و راه به دو زیر معیار

دسترسی به شهر و دسترسی به راه تقسیم شده اند

ساختن سلسله مراتبی برای تجزیه و تحلیل داده ها

در اولین اقدام، ساختار سلسله مراتبی مربوط به این موضوع را مشخص می کنیم (نمودار ۱). در این نمودار، ما با یک سلسله مراتب چهار سطحی شامل: هدف ها، معیارها، زیر معیارها و گزینه ها مواجه هستیم. تبدیل موضوع یا مسئله مورد بررسی به یک "ساختار سلسله مراتبی" مهم ترین قسمت فرایند تحلیل سلسله مراتبی محسوب می شود. زیرا در این قسمت با تجزیه مسائل مشکل و پیچیده، فرایند تحلیل سلسله مراتبی آنها را به شکلی ساده، که با ذهن و طبیعت انسان مطابقت داشته باشد، تبدیل می کند. به عبارت دیگر، فرایند تحلیل سلسله مراتبی مسائل پیچیده را از طریق تجزیه آن به عناصر جزئی که به صورت سلسله مراتبی به هم مرتبط بوده و ارتباط هدف اصلی مسئله با پایین ترین سطح سلسله مراتبی مشخص است، به شکل ساده تری در می آورد.

پستی و بلندی = D / دسترسی به آب = E / خطر زمین لرزه = F

دسترسی به شهر و راه = G / ارتفاع از سطح دریا = H / شیب زمین = I

موقعیت نسبت به گسل های بنیادی = J / موقعیت نسبت به رو مراکز زمین لرزه = K

شدت زلزله های تاریخی و دسنگاهی = L / دسترسی به راه = M / دسترسی به شهر = N



تبیین ضریب اهمیت معیارها و زیر معیارها

دارای اهمیت بیشتری است یا " دسترسی به منابع آب "؟ مبنای قضاوت در این امر مقایسه ای جدول ۹ کمیتی زیر (جدول ۱) است که براساس آن و با توجه به هدف بررسی ، شدت برتری معیار i نسبت به معیار j ، a_{ij} تعیین می شود . تمامی معیارها دو به دو مقایسه می شوند . مقایسه های دو به دو یک ماتریس $n \times n$ (در این حالت 4×4) ثبت می شوند و این ماتریس ، " ماتریس مقایسه دو دوئی معیارها " ، $A = [a_{ij}]_{n \times n}$ نامیده می شود . عناصر این ماتریس همگی مثبت بوده و با توجه به اصل " شروط معکوس " در فرایند تحلیل سلسله مراتبی (اگر اهمیت i نسبت به j برابر $\frac{1}{k}$ باشد ، اهمیت عنصر j نسبت به i برابر $\frac{1}{k}$ خواهد بود) . در هر مقایسه دو دوئی ، دو مقدار عددی a_{ij} و $\frac{1}{a_{ij}}$ را خواهیم داشت . در زیر ماتریس مقایسه دو دوئی معیارها برای مسئله مورد نظر ارائه شده است

		دسترسی خطر زلزله دسترسی پستی و بلندی			
		۱	۲	۳	۴
(۱) پستی و بلندی		1	$\frac{1}{9}$	$\frac{1}{7}$	$\frac{1}{5}$
(۲) دسترسی به آب		9	1	1	3
(۳) خطر زلزله		7	1	1	3
(۴) دسترسی به شهر و راه		5	$\frac{1}{3}$	$\frac{1}{3}$	1

$$= A$$

در این ماتریس ، مقدار عددی عنصر a_{21} (ردیف ۲ و ستون ۱) که ۹ می باشد ، نشان می دهد که معیار دسترسی به آب در مکان یابی در مقایسه با " پستی و بلندی " دارای اهمیت مطلق بوده و با توجه به شرط معکوس ، بنابراین مقدار عددی a_{12} برابر با $\frac{1}{9}$ خواهد بود . عناصر قطر این ماتریس ، با توجه به اهمیت برابر هر معیار نسبت به خود در دستیابی به هدف ، برابر با ۱ است .

برای محاسبه ضریب اهمیت معیارها ، چهار روش عمده زیر مطرح هستند :

۱. روش حداقل مربعات

۳. روش بردار ویژه

۲. روش حداقل مربعات لگاریتمی

۴. روش های تقریبی

از روش های فوق ، روش بردار ویژه بیشتر مورد استفاده قرار گرفته است . اما اگر ماتریس A دارای ابعاد بزرگتری باشد ، محاسبه مقادیر و بردارهای ویژه طولانی و وقت گیر خواهد بود ، مگر این که از نرم



افزارهای کامپیوتری برای حل آن کمک گرفته شود. به همین دلیل است که ساعتی چهار روش تقریبی زیر را ارائه کرده است ۱. مجموع سطری ، ۲. مجموع ستونی ، میانگین حسابی ، ۴. میانگین هندسی . (saaty ، ۱۹۹۰) در این بررسی ، روش میانگین هندسی به دلیل دقت بیشتر آن ، مورد استفاده قرار گرفته است . در این روش برای محاسبه ضریب اهمیت معیارها ، ابتدا میانگین هندسی ردیف های هندسی ردیف ماتریس A را به دست آورده و آنها را نرمالیزه می کنیم :

$$1: \left[(1)\left(\frac{1}{9}\right)\left(\frac{1}{7}\right)\left(\frac{1}{5}\right) \right]^{\frac{1}{4}} = 0.2374$$

۱: پستی و بلندی

$$2: \left[(9)(1)(1)(3) \right]^{\frac{1}{4}} = 2.2795$$

۲: دسترسی به آب

$$3: \left[(7)(1)(1)(3) \right]^{\frac{1}{4}} = 2.1497$$

۳: خطر زلزله

$$4: \left[(5)\left(\frac{1}{3}\right)\left(\frac{1}{3}\right)(1) \right]^{\frac{1}{4}} = \frac{0.8633}{5.5209}$$

۴: دسترسی به شهر و راه

ضریب اهمیت معیارها از نرمالیزه کردن این اعداد ، یعنی از تقسیم هر عدد به سر جمع آنها به دست

می آید .

$$w1 = \frac{0.2347}{5.5209} = 0.043$$

ضریب اهمیت پستی و بلندی

$$W2 = 0.4129$$

ضریب اهمیت دسترسی به آب

$$W3 = 0.3877$$

ضریب اهمیت خطر زلزله

$$W4 = 0.1564$$

ضریب اهمیت دسترسی به شهر و راه

همانطور که مشاهده می شود ، مجموع ضریب اهمیت معیارهای چهارگانه مزبور (سطح دوم سلسله

مراتبی) معادل یک است و این نشان دهنده نسبی بودن اهمیت معیارها است .

برای به دست آوردن ضرایب اهمیت زیر معیارها ، همان مراحل که در بالا به دست آوردن ضریب

اهمیت معیارها طی شده را انجام می دهیم . معیار پستی و بلندی از دو زیر معیار ، " ارتفاع از سطح دریا

" و " شیب زمین " تشکیل یافته است . بنابراین ، ماتریس مقایسه دودوئی معیارها را برای این دو زیر

معیار ، براساس همان جدول ۹ کمیتی ساعتی (جدول ۱) تشکیل می دهیم (ماتریس A۱) :



$$\square \text{ ارتفاع از سطح دریا} \quad \begin{bmatrix} 1 & \frac{1}{5} \\ 5 & 1 \end{bmatrix} = A_1$$

ضریب اهمیت این دو زیر معیار ، با استفاده از روش میانگین هندسی به دست می آید :

ضریب اهمیت ارتفاع از سطح دریا $W_H = 0.167$

ضریب اهمیت شیب زمین $W_I = 0.833$

به همین ترتیب ، ضرایب اهمیت زیر معیارهای دو معیار دیگر یعنی زلزله و دسترسی به شهر و راه را به دست می آوریم . مقایسه دو دوئی معیارها برای زیر معیارهای موقعیت نسبت به گسله های بنیادی ، موقعیت نسبت به رو مراکز زمین لرزه و شدت زلزله های تاریخی و دستگاهی به وسیله ماتریس A_2 و همین مقایسه برای زیر معیارهای دسترسی به شهر $\square \square \square$ به راه به وسیله ماتریس A_3 نشان داده شده است .

$$\begin{array}{l} \square \text{ موقعیت نسبت به گسله های بنیادی} \\ \square \text{ موقعیت نسبن به رو مراکز زمین لرزه} \\ \square \text{ شدت زلزله های تاریخی و دستگاهی} \end{array} \quad \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 2 \\ 1 & \frac{1}{2} & 1 \end{bmatrix} = A_2$$

ضریب اهمیت موقعیت نسبت به گسله های بنیادی $w_I = 0.327$

ضریب اهمیت موقعیت نسبت به رو مراکز زمین لرزه $w_k = 0.413$

ضریب اهمیت شدت زلزله های تاریخی و دستگاهی $w_l = 0.260$

$$\begin{array}{l} \square \text{ دسترسی به راه} \\ \square \text{ دسترسی به شهر} \end{array} \quad \begin{array}{cc} M & N \\ \begin{bmatrix} 1 & \frac{1}{7} \\ 7 & 1 \end{bmatrix} \end{array}$$

ضریب اهمیت دسترسی به راه $w_M = 0.125$

ضریب اهمیت دسترسی به شهر $w_N = 0.875$

تعیین ضریب اهمیت گزینه ها

بعد از تعیین ضرایب اهمیت معیارها و زیر معیارها ، ضریب اهمیت گزینه ها را باید تعیین کرد . در

این مرحله ، ارجحیت هر یک از گزینه ها در ارتباط با هر یک از زیر معیارها و اگر معیاری زیر معیار



نداشته باشد (مثل دسترسی به آب) مستقیماً با خود آن معیار ، مورد قضاوت و داوری قرار می گیرد .
 مبنای این قضاوت همان مقیاس ۹ کمیتی ساعتی است با این تفاوت که در مقیاس گزینه ها در ارتباط با
 هر یک از زیر معیارها (یا معیار ، حسب مورد) ، بحث کدام گزینه مهم تر است ؟ مطرح نیست، بلکه "
 کدام گزینه ارجح است ؟ و چقدر ؟ مطرح است . بنابراین ، مقیاس ۹ کمیتی ساعتی به شرح جدول ۲ ،
 مبنای قضاوت گزینه ها قرار خواهد گرفت :

جدول مقیاس ۹ کمیتی ساعتی برای مقایسه دو دوئی گزینه ها

امتیاز (شدت ارجحیت)	تعریف
۱	ترجیح یکسان (Equally preferred)
۳	کمی مرجح (Moderately preferred)
۵	ترجیح بیشتر (Strongly preferred)
۷	ترجیح خیلی بیشتر (very Strongly preferred)
۹	کاملاً مرجح (Extremely preferred)
۸,۶,۴,۲	ترجیحات بینابین (وقتی حالت های میانه وجود دارد)

فرایند به دست آوردن وزن (ضریب اهمیت) گزینه ها نسبت به هر یک از زیر معیارها شبیه تعیین
 ضریب اهمیت معیارها نسبت به هدف است . در هر دو حالت ، قضاوت ها بر مبنای مقایسه دو دوئی
 معیارها یا گزینه ها و براساس مقیاس ۹ کمیتی ساعتی صورت پذیرفته و نتیجه در ماتریس مقایسه دو
 دوئی معیارها یا گزینه ها ثبت شده و از طریق نرمالیزه کردن میانگین هندسی ردیف های این ماتریس ها
 ، ضرایب اهمیت مورد نظر به دست می آید . با این حال ، باید به یک تفاوت عمده در این مقایسه ها
 اشاره شود . مقایسه گزینه های مختلف نسبت به زیر معیارها و یا معیارها (اگر معیاری زیر معیار
 نداشته باشد) صورت می پذیرد . در صورتی که مقایسه با یکدیگر نسبت به " هدف " مطالعه صورت می
 پذیرفت . بنابراین ، به جای این که سؤال شود " معیار A ، در دستیابی به هدف چقدر از معیار B مهم تر



است؟" در مقایسه گزینه ها سؤال به این ترتیب مطرح می شود که "گزینه i در ارتباط با زیر معیار X چقدر برگزینه Z ارجحیت دارد؟"

در جدول ۳ که بیشتر به ماتریس ارزیابی معروف است، ارزش هر یک از گزینه ها در ارتباط با زیر معیارها و معیار دسترسی به آب (که زیر معیار ندارد) ارائه شده است. همان طور که ملاحظه می شود، زیر معیارها هم کمی هستند و هم کیفی، و این، نشان دهنده مزیت دیگر فرایند تحلیل سلسله مراتبی است که با ترکیبی از معیارهای کمی و کیفی سرو کار دارد.

جدول شماره (21). ماتریس ارزیابی برای مکانیابی مورد نظر

گزینه	ارتفاع از سطح دریا (متر)	شیب زمین (درصد)	دسترسی به آب	موقعیت نسبت به گسل های بنیادی	موقعیت نسبت به رو مراکز زمین لرزه	شدت زلزله های تفریحی و دستگاهی	دسترسی به راه	دسترسی به شهر
A	۵۶۹	۵	دسترسی خوب	مناسب	نامناسب	کم	دسترسی خوب	دسترسی نسبتاً خوب
B	۹۲۰	۱۰	دسترسی نسبتاً خوب	نامناسب	نسبتاً مناسب	متوسط	دسترسی نسبتاً خوب	دسترسی کم
C	۲۱۰۰	۳	دسترسی بسیار کم	نامناسب	مناسب	متوسط	دسترسی نسبتاً خوب	دسترسی نسبتاً کم

در زیر ماتریس های مقایسه دو دوئی گزینه ها در ارتباط با هر یک از زیر معیارها و نیز معیار دسترسی به آب ارائه شده است:

$$\begin{bmatrix} 1 & 3 & 8 \\ 1/3 & 1 & 4 \\ 1/8 & 1/4 & 1 \end{bmatrix}$$

ارتفاع از سطح دریا

$$\begin{bmatrix} 1 & 6 & 1/4 \\ 1/6 & 1 & 1/7 \\ 4 & 7 & 1 \end{bmatrix}$$

شیب زمین

□ □ □

□ □ □



$$\begin{bmatrix} 1 & 8 & 8 \\ 1/8 & 1 & 1 \\ 1/8 & 1 & 1 \end{bmatrix}$$

دسترسی به آب

$$\begin{bmatrix} 1 & 5 & 8 \\ 1/5 & 1 & 4 \\ 1/8 & 1/4 & 1 \end{bmatrix}$$

موقعیت نسبت به گسل های بنیادی

$$\begin{bmatrix} 1 & 1/5 & 1/7 \\ 5 & 1 & 1/5 \\ 7 & 5 & 1 \end{bmatrix}$$

موقعیت نسبت به رو مراکز زمین لرزه

$$\begin{bmatrix} 1 & 5 & 5 \\ 1/5 & 1 & 1 \\ 1/5 & 1 & 1 \end{bmatrix}$$

شدت زمین لرزه های تاریخی و دستگاهی

$$\begin{bmatrix} 1 & 6 & 4 \\ 1/6 & 1 & 1/3 \\ 1/4 & 3 & 1 \end{bmatrix}$$

دسترسی به راه

$$\begin{bmatrix} 1 & 5 & 5 \\ 1/5 & 1 & 1 \\ 1/5 & 1 & 1 \end{bmatrix}$$

دسترسی به شهر

ضریب اهمیت گزینه ها در ارتباط با زیر معیارها، از طریق نرمالیزه کردن میانگین هندسی ردیف های

ماتریس مقایسه دو دوئی و به شرح زیر تعیین می شود :

$$\text{ارتفاع از سطح دریا} \begin{cases} A: \sqrt[3]{24} = 2.885 \\ B: \sqrt[3]{4/3} = 1.101 \\ C: \sqrt[3]{1/32} = \frac{0.315}{4021} \end{cases} \Rightarrow \begin{cases} W_A = 0.685 \\ W_B = 0.240 \\ W_C = 0.075 \end{cases}$$

$$\text{شیب زمین} \begin{cases} A: \sqrt[3]{6/4} = 1.145 \\ B: \sqrt[3]{1/42} = 0.288 \\ C: \sqrt[3]{28} = \frac{3.037}{4.47} \end{cases} \Rightarrow \begin{cases} W_A = 0.256 \\ W_B = 0.064 \\ W_C = 0.0680 \end{cases}$$



$$\text{دسترسی به آب} \quad \left\{ \begin{array}{l} A: \sqrt[3]{40} = 3.420 \\ B: \sqrt[3]{4/5} = 0.928 \\ C: \sqrt[3]{1/32} = \frac{0.315}{4.663} \end{array} \right. \Rightarrow \left\{ \begin{array}{l} W_A = 0.733 \\ W_B = 0.064 \\ W_C = 0.068 \end{array} \right.$$

$$\text{موقعیت نسبت به گسل های بنیادی} \quad \left\{ \begin{array}{l} A: \sqrt[3]{64} = 4.000 \\ B: \sqrt[3]{1/8} = 0.500 \\ C: \sqrt[3]{1/8} = \frac{0.500}{5} \end{array} \right. \Rightarrow \left\{ \begin{array}{l} W_A = 0.800 \\ W_B = 0.10 \\ W_C = 0.10 \end{array} \right.$$

$$\text{موقعیت نسبت به رو مراکز زمین لرزه} \quad \left\{ \begin{array}{l} A: \sqrt[3]{1/35} = 0.306 \\ B: \sqrt[3]{1} = 1.000 \\ C: \sqrt[3]{35} = \frac{3.271}{43577} \end{array} \right. \Rightarrow \left\{ \begin{array}{l} W_A = 0.067 \\ W_B = 0.218 \\ W_C = 0.715 \end{array} \right.$$

$$\text{ت زمین لرزه های تاریخی و دستگاهی} \quad \left\{ \begin{array}{l} A: \sqrt[3]{25} = 2.924 \\ B: \sqrt[3]{1/5} = 0.585 \\ C: \sqrt[3]{1/5} = \frac{0.585}{4.054} \end{array} \right. \Rightarrow \left\{ \begin{array}{l} W_A = 0.714 \\ W_B = 0.143 \\ W_C = 0.143 \end{array} \right.$$

$$\text{دسترسی به راه} \quad \left\{ \begin{array}{l} A: \sqrt[3]{25} = 2.924 \\ B: \sqrt[3]{1/5} = 0.585 \\ C: \sqrt[3]{1/5} = \frac{0.585}{4.054} \end{array} \right. \Rightarrow \left\{ \begin{array}{l} W_A = 0.714 \\ W_B = 0.143 \\ W_C = 0.143 \end{array} \right.$$

$$\text{دسترسی به شهر} \quad \left\{ \begin{array}{l} A: \sqrt[3]{24} = 2.884 \\ B: \sqrt[3]{1/18} = 0.382 \\ C: \sqrt[3]{1/12} = \frac{0.437}{3.703} \end{array} \right. \Rightarrow \left\{ \begin{array}{l} W_A = 0.779 \\ W_B = 0.103 \\ W_C = 0.118 \end{array} \right.$$



تعیین امتیاز نهایی (اولویت) گزینه ها

تا این مرحله ، ضرایب اهمیت معیارها و زیر معیارها در ارتباط با هدف مطالعه و نیز ضرایب اهمیت (امتیاز) گزینه ها در ارتباط با هر یک از زیر معیارها و نیز معیار " دسترسی به آب " تعیین شده است . در این مرحله ، از تلفیق ضرایب اهمیت مزبور ، " امتیاز نهایی " هر یک از گزینه ها تعیین خواهد شد . برای این کار از " اصل ترکیب سلسله مراتبی " ساعتی که منجر به یک بردار اولویت با در نظر گرفتن همه قضاوت ها در تمامی سطوح سلسله مراتبی می شود ، استفاده خواهد شد

$$g_{ij} = \sum_{k=1}^n \sum_{i=1}^m W_k W_i \quad (\text{امتیاز نهایی (اولویت) گزینه } j)$$

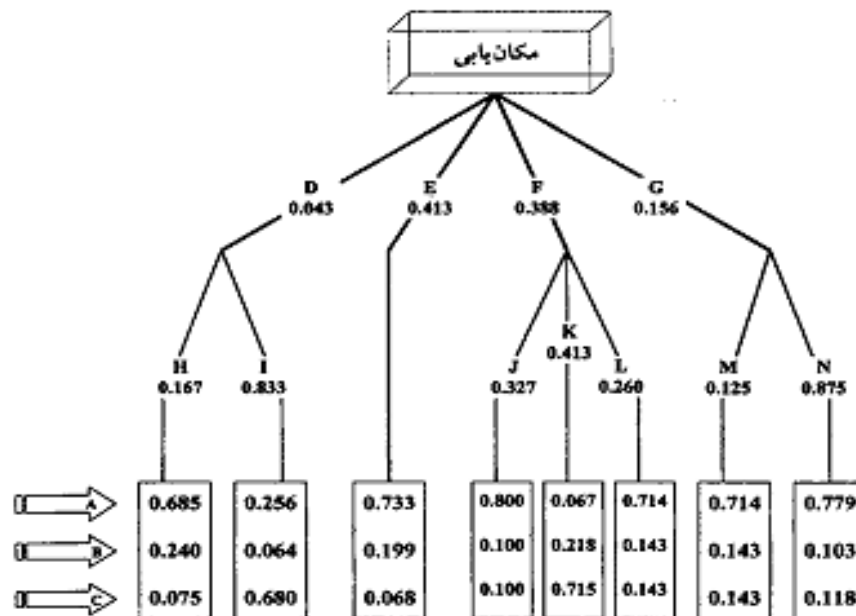
که در آن :

W_k ضریب اهمیت معیار K

W_i همضریب اهمیت زیر معیار i

g_{ij} امتیاز گزینه j در ارتباط با زیر معیار i

ضرایب اهمیت معیارها ، زیر معیارها و امتیاز گزینه ها در ارتباط با هر یک از زیر معیارها در نمودار ۲ ارائه شده است .



نمودار ۲. ضرایب اهمیت معیارها، زیر معیارها و گزینه ها در ساختار سلسله مراتبی

نحوه تعیین امتیاز نهایی گزینه ها، براساس اصل ترکیب سلسله مراتبی و با

استفاده از ضرایب اهمیت ارائه شده در نمودار ۲ در جدول ۳ ارائه شده است. امتیازات

نهایی گزینه ها نشان می دهد که گزینه A (سایت A) برای اهداف مکان یابی بهترین

گزینه ها هستند. گزینه های B و C نیز به ترتیب دوم و سوم بهترین گزینه ها هستند.



بررسی سازگاری در قضاوت ها

یکی از مزیت های فرایند تحلیل سلسله مراتبی امکان بررسی سازگاری در قضاوت های انجام شده برای تعیین ضریب اهمیت معیارها و زیر معیارها است . به عبارت دیگر در تشکیل ماتریس مقایسه دو دوئی معیارها (ماتریس A) ، چقدر سازگاری در قضاوت ها رعایت شده است ؟ وقتی اهمیت معیارها نسبت به یکدیگر برآورد می شود ، احتمال ناهماهنگی در قضاوت ها وجود دارد . یعنی اگر A_i از A_j مهم تر باشد و A_j از A_k مهم تر باشد . اما علیرغم همه کوشش ها ، رجحان ها و احساس های مردم غالباً ناهماهنگ و نامتعدی هستند . پس باید سنجه ای را یافت که میزان ناهماهنگی داوری ها را نمایان سازد (توفیق ، ۱۳۷۲ ، ص ۴۲).

مکانیزی که ساعتی (Saaty , ۱۹۸۸) برای بررسی ناسازگاری در قضاوت ها در نظر گرفته است ، محاسبه ضریبی به نام ضریب ناسازگاری $(I.R.)^{(۷۴)}$ است که از تقسیم شاخص ناسازگاری $(I.I.)^{(۸۴)}$ به شاخص تصادفی بودن $(R.I.)$ حاصل می شود . چنانچه این ضریب کوچک تر یا مساوی ۰/۱ باشد ، سازگاری در قضاوت ها مورد قبول است و گرنه باید در قضاوت ها تجدید نظر شود . به عبارت دیگر ماتریس مقایسه دو دوئی معیارها باید مجدداً تشکیل شود :

$$I.I. = \frac{\lambda_{max} - n}{n - 1} \quad \text{شاخص ناسازگاری}$$

شاخص تصادفی بودن با توجه به تعداد معیارها (n) از جدول زیر قابل استخراج است :

جدول شماره (۲۲): شاخص تصادفی بودن ($R.I.$)

15	14	13	12	11	10	9	8	7	6	5	4	3	2	n
1/59	1/57	1/56	1/48	1/51	1/49	1/45	1/41	1/32	1/24	1/12	0/9	0/58	0	R.I.

مأخذ: Bowen, 1993: 349



در روش میانگین هندسی که یک روش تقریبی است، به جای محاسبه مقدار ویژه ماکزیمم

($\lambda_{\max}$) از L به شرح زیر استفاده می شود :

$$L = \frac{1}{n} \left[\sum_{i=1}^n (AW_i / W_i) \right]$$

که در آن AW_i برداری است که از ضرب ماتریس مقایسه دو دوئی معیارها (ماتریس A) در بردار W_i (بردار وزن یا ضریب اهمیت معیارها حاکی از آن است که سازگاری قضاوت ها در تعیین ضرایب اهمیت معیارهای چهارگانه برای مکان یابی ارائه شده است :

بررسی سازگاری در قضاوت ها برای تعیین ضرایب اهمیت معیارها :

1. محاسبه بردار AW

$$\begin{matrix} 1 \\ 2 \\ 3 \end{matrix} \begin{bmatrix} 1 & \frac{1}{9} & \frac{1}{7} & \frac{1}{5} \\ 9 & 1 & 1 & 3 \\ 7 & 1 & 1 & 3 \\ 5 & \frac{1}{3} & \frac{1}{3} & 1 \end{bmatrix} \times \begin{bmatrix} 0.043 \\ 0.4129 \\ 0.3877 \\ 0.1564 \end{bmatrix} = \begin{bmatrix} 0.1755 \\ 1.6568 \\ 1.5708 \\ 1.6383 \end{bmatrix}$$

2. محاسبه L

$$L = \frac{1}{n} \left[\sum_{i=1}^n (AW_i / W_i) \right]$$

$$L = \frac{1}{4} \left[\frac{0.1755}{0.043} + \frac{1.6568}{0.4129} + \frac{1.5708}{0.3877} + \frac{1.6383}{0.1564} \right]$$

$$L=4.0567$$

3. محاسبه شاخص سازگاری CI

$$CI = \frac{L-n}{n-1}$$



$$CI = \frac{4.0567 - 4}{4 - 1} = 0.0189$$

4. محاسبه ضریب سازگاری CR

$$CR = \frac{CI}{RI} = \frac{0.0189}{0.9}$$

$$CR = 0.021 < 0.1 \quad \text{O.K.}$$

یعنی سازگاری در قضاوت ها رعایت شده است .

ریشه و منشا و تفسیر های اولیه از توسعه

پیشینه توسعه بسیار قوی و دارای قدمتی همزمان با تغییرات طبیعی ، اجتماعی در سطح کره زمین می باشد که به تدریج جایگزین مفاهیم دیگری نظیر ترقی ، تکامل و رشد شد . (مطیعی لنگرودی حسن ، ۱۳۸۲ ، ۵۸) در ابتدا از علوم طبیعی استخراج گردیده و در مورد فرآیند تغییر در جوامع بشری به کار گرفته شد . واژه توسعه در نخستین کاربردش به زبان فرانسه و انگلیسی در سال ۱۷۵۲ به معنای رسیدن به اهداف با ایدهایی طبق یک طرح یا برنامه بود . سپس این واژه بعنوان مراحل مشخصی در برنامه خود و بعد به مثابه توالی بیولوژیکی تغییر از یک دانه و تخم گیاه به گل به کار رفت . در معنای لغوی واژه توسعه نیز همین کاربرد مستتر است . به معنی خارج شدن از پوشش و لفاف می باشد . یا بروز و ظهور نمودن همه آنچه که بطور بالقوه در چیزی وجود دارد یا رشد و تکامل یک ارگانیسم از نوع و حالتی ساده به نوع و حالتی کامل تر پیچیده تر در حد بلوغ و کامل . از آن پس تحت تاثیر قالب فکری داروینیسیم ، توسعه برای توصیف تغییرات بیولوژیکی به کار گرفته شد و با واژه رشد مترادف گردید . (موثقی احمد ، ۱۳۸۳ : ۲۲۵)

به این ترتیب توسعه به معنای فرآیندی شد که از طریق آن استعداد های نهفته و توانایی های بالقوه یک شی یا ارگانیسم شکوفا می شود تا اینکه به شکل طبیعی و کامل به بلوغ نهایی اش برسد . با این استعاره امکان نشان دادن هدف و توسعه بعد برنامه آن فراهم شد . در دوره و فاصله بین ولف تا داروین (۱۷۵۹ تا ۱۸۵۹) توسعه از مفهوم دگرگونی به سوی شکلی مناسب از وجود به مفهوم دگرگونی به سمت شکل همیشه کامل متحول شد و با واژه تکامل هم خوانی پیدا کرد . این استعاره در ربع آخر قرن هجدهم میلادی در حوزه اجتماعی به کار گرفته شد و هربرت اسپنسر به مطالعه و بررسی در اجتماع پرداخته و تکامل و پیشرفت اجتماعی را مطرح کرد و نه تنها جامعه را با یک ارگانیسم بیولوژیکی مقایسه کرد بلکه



تاکید نمود که جامعه یک ارگانیسم است . کم کم تئوری طبیعت با فلسفه تاریخ ترکیب شد تا وحدتی سیستماتیک و منسجم را برقرار کرد . بالاخره کارل مارکس تکامل دورانهای تاریخی را مطرح نمود . البته مارکس توسعه را در تغییرات اقتصادی می دید که بیشتر کیفی هستند تا کمی . اقتصاددانان کلاسیک به توسعه به عنوان تحول فیزیکی زمین ، کار و سرمایه به اشکالی با قابلیت تولید و بهره وری بیشتر توجه داشتند و از توسعه اقتصادی سخن گفتند . بدین ترتیب توسعه مفهومی شد که بر تکامل نظامهای اجتماعی بشری از اشکال ساده تر به اشکال پیچیده تر ، بالاخره در حد بلوغ و کمال دلالت داشت . از سوی دیگر واژه های دیگری همچون ترقی مدرنیته نوسازی و حتی غربی سازی هم مطرح شد و به مفهوم توسعه پیوند خورد . (همان : ۲۲۷)

توسعه چیست

توسعه در لغت به معنای رشد تدریجی در جهت پیشرفته تر شدن ، قدرتمند تر شدن و حتی بزرگتر شدن است . بروکفیلد در تعریف توسعه می گوید « توسعه را باید بر حسب پیشرفت به سوی اهداف رفاهی نظیر کاهش فقر ، بیکاری و نابرابری تعریف کنیم » (صرافی ، ۱۳۷۷ : ۲۵)

دادلی سیزر « توسعه را جریانی چند بعدی می داند که تجدید سازمان و سمت گیری منافع کل نظام اقتصادی ، اجتماعی را به همراه دارد » به عقیده او توسعه علاوه بر بهبود میزان تولید و درآمد شامل دگرگونی اساسی در ساختمانهای نهادی اجتماعی ، اداری و همچنین وجهه نظرهای عمومی مردم است . توسعه در بسیاری از موارد حتی آداب و رسوم و عقاید مردم را نیز در بر می گیرد . (قدیری مسعود ، ۱۳۸۴ : ۳)

آقای مصطفی ازکیا « توسعه را به معنی کاهش فقر بیکاری ، نابرابری ، صنعتی شدن بیشتر ، ارتباطات بهتر ، ایجاد نظام اجتماعی مبتنی بر عدالت و افزایش مشارکت مردم در امور سیاسی جاری » تعریف می کنند .

زنده یاد دکتر حسین عظیمی از مجموع نظرات علمای توسعه : توسعه را به معنای بازسازی جامعه براساس اندیشه ها و بصیرت های تازه تعبیر می نماید . این اندیشه ها و بصیرت های تازه در دوران مدرن شامل سه اندیشه علم باوری ، انسان باوری ، آینده باوری است . برای این منظور نیل به توسعه سه اقلام



اساسی درک و هضم اندیشه های جدید ، تشریح و تفضیل این اندیشه ها و ایجاد نهادهای جدید برای تحقق عملی این اندیشه ها صورت پذیرد.

تودارو توسعه را به معنی ارتقاء مستمر کل جامعه و نظام اجتماعی به سوی زندگی بهتر می داند . به هر تقدیر امروزه ما از مفهوم توسعه فرآیندی همه جانبه (نه فقط اقتصادی) که معطوف به بهبودی تمامی ابعاد زندگی مردم یک جامعه (به عنوان لازم و ملزوم) است . ابعاد مختلف توسعه عبارتند از توسعه اقتصادی ، اجتماعی ، توسعه سیاسی ، توسعه فرهنگی ، توسعه اجتماعی و توسعه امنیتی (موثقی ، ۱۳۸۳ : ۲۲۳).

ضرورت استفاده از برنامه ریزی منطقه ای

رشد و توسعه به عنوان یک مقوله اقتصادی - اجتماعی ابتدا از سوی اقتصاد دانان و سپس جامعه شناسان و پژوهشگران برخی از علوم چون جغرافیا مورد توجه و اساس برنامه ریزی قرار گرفت . از جمله ، دشواری های همیشگی در بررسی ادبیات توسعه اقتصادی و دگرگونی های اجتماعی ، مشخص کردن مفهوم رشد و توسعه است . توسعه را می توان فرآیندی سیاسی ، اجتماعی و اقتصادی دانست که منتج از استانداردهای زندگی بوده و باعث بهبود سطح زندگی بخش های در حال افزایش جمعیت شود . فرآیند توسعه به حدی اهمیت دارد که باید به موازات رشد جمعیت قابل مشاهده باشد و توصیفی اینچنین ، مفاهیم کیفی و منصفانه را به عنوان اهداف قطعی در هر فرآیند توسعه در بر می گیرد . به عبارت ساده می توان گفت با توجه به این که هدف اصلی توسعه حذف نابرابری هاست بهترین مفهوم توسعه ، رشد همراه با عدالت اجتماعی است (معصومی و حبیبی ، ۱۳۸۳ : ۱۴۸) . اهداف کلی برنامه ریزی منطقه ای و یا توسعه اقتصادی نیز برقراری عدالت اجتماعی و توزیع رفاه و ثروت در بین افراد جامعه است و برای دستیابی به اهداف فوق در هر جامعه ای نیاز به تهیه ، تدوین و در نهایت اجرای برنامه های متنوع و متعدد می باشد. لذا در برنامه ریزی منطقه ای بررسی و شناخت وضعیت نواحی ، قابلیت ها و تنگناهای آن از اهمیت بسزایی برخوردار است . امروزه آگاهی از نقاط ضعف نواحی نوعی ضرورت جهت ارایه طرح ها در برنامه ها محسوب می شود . بطوری که استفاده از شاخص های اقتصادی ، اجتماعی ، فرهنگی ، بهداشتی و ... می تواند معیار مناسب هم برای تعیین جایگاه آن نواحی و همچنین عاملی در جهت رفع مشکلات و نارسایی های مبتلا به خود برای نیل به رفاه اقتصادی و سلامت اجتماعی در جهت رسیدن به



توسعه باشد. (موسوی و حکمت نیا، ۱۳۸۴: ۵۶) زیرا یکی از مشکلات اساسی توسعه فضایی و ناحیه ای گسیختگی، سازمان فضایی و عدم سلسله مراتب مبتنی بر رابطه میان سکونت گاه ها است. در همین راستا تعیین و تشکیل سلسله مراتبی از سکونت گاه ها که بتواند چارچوب موثری برای توزیع جمعیت، فعالیت ها، خدمات و کارکرد ها در سطوح مختلف باشد ضروری است. بنابراین بکارگیری معیارها و روشهای کمی جهت سطح بندی سکونت گاه ها در سیستم فضایی مناطق نه تنها موجب شناخت تفاوت میان سکونت گاه ها می گردد بلکه این سطح بندی معیاری برای تعیین مرکزیت و همچنین تعیین انواع خدمات مورد نیاز و تعدیل نابرابری بین سکونت گاه ها است. امروزه با پیشرفت روش های آماری و رایانه ای در مطالعات جغرافیایی استفاده از شاخصهای مختلف در زمینه های گوناگون متداول تر می شود. سطح بندی سکونت گاه ها است. (حکمت نیا و موسوی، ۱۳۸۵: ۲۰۹). اما بایستی اذعان کرد که استفاده از این شاخص ها منوط به دسترسی داشتن اطلاعاتی جامع و کامل از نواحی است. در دسترس داشتن اطلاعات خام و بدون پردازش هرگز برنامه ریزان را در رسیدن به اهداف خود یاری نخواهند کرد. بنابراین دسترسی به اطلاعات آماری در زمینه ی شاخصهای مختلف در نواحی و پردازش آنها با استفاده از مدل های آماری مهمترین گام جهت نیل به اهداف مورد نظر می باشد. شرایط کشور ما نیز از لحاظ توسعه و زیر ساخته های آن در نواحی جغرافیایی کشور در اثر برنامه ریزی نامطلوب ملی و متمرکز، تفاوت شدیدی را در روند توسعه آشکار ساخته است و در زمینه اقتصادی و اجتماعی آن یکی دو منطقه و یا در نهایتا چند منطقه دارای مسئولیت اصلی را در زمینه ایجاد درآمد ملی و بر خوداری از خدمات عمومی و بالطبع شکوفایی اقتصادی، اجتماعی به قیمت عقب نگه داشتن مناطق دیگر بوده است. چنین وضعیتی در اکثر مناطق و استان های کشور صادق است.

فرآیند سنجش سطوح توسعه مناطق

برای سنجش سطوح توسعه مناطق باید مراحل را به شرح ذیل طی کرد:

الف - تعیین هدف مطالعه و تدوین چارچوب آن

در این مرحله لازم است محقق هدف یا اهداف مطالعه خود را به طور دقیق تعیین کند. او باید مشخص کند که چه جنبه ای از توسعه مد نظر است آیا هدف سنجش توسعه مناطق به مفهوم کلی است یا جنبه خاصی از توسعه (نظیر توسعه صنعتی، کشاورزی، اجتماعی، اقتصادی و...) را مد نظر دارد.



پس از مشخص شدن آن باید مفهوم مورد نظر را بطور دقیق و کامل تعریف عملیاتی شود. عبارت دیگر اگر هدف، بررسی و سنجش سطح توسعه اجتماعی است در این صورت باید مفهوم توسعه اجتماعی به طور جامع و مانع تعریف گردد. ارائه یک تعریف عملی در این مرحله محقق را از سردرگمی و پراکنده گوئی نجات داده و او را در انتخاب شاخص های توسعه کمک می کند.

ب - تعیین سطوح مطالعه

پس از تعیین هدف مطالعه و تعریف عملیاتی موضوع تحقیق باید سطوح مطالعه خود را تعیین نماید. مطالعات منطقه ای در سطوح مختلف نظیر روستا، دهستان، بخش، شهرستان، استان و یا مناطق کلان دیگر انجام می شود. تعیین سطح مطالعه راه را برای انجام مراحل بعدی هموار می کند. مهمترین عامل در تعیین سطح مطالعه، دسترسی به آمار و اطلاعات می باشد. از آنجا که در مطالعات توسعه منطقه ای عمدتاً از آمار ثانویه و جمع آوری شده توسط مراکز آمار و یا سایر سازمان ها و وزارتخانه ها استفاده می شود، در نتیجه سطح دسترسی به آمار و نوع آن در تعیین سطح مطالعه نقش تعیین کننده ای دارد.

ج - شناخت نوع آمار قابل دسترسی

پس از تعیین هدف مطالعه، تعریف عملیاتی موضوع، تدوین چارچوب نظری مناسب و مشخص کردن سطح مطالعه محقق باید آمار و اطلاعات قابل دسترس را مورد بررسی قرار داده و لیستی از متغیر هایی که در مورد آن آمار وجود دارد تهیه کند مشخص بودن نوع اطلاعات موجود محقق را در تعیین شاخصهای توسعه یاری می دهد. در ایران اینگونه آمارها در سطح استان، شهرستان، بخش، دهستان و حتی روستا در مرکز آمار ایران وجود دارد. اما هرچه قدر سطح مطالعه پایین تر باشد و به سطح روستا نزدیکتر گردد میزان آمار قابل قابل دسترس محدودتر می شود. (کلانتری، ۱۳۸۰: ۱۱۰)

د - انتخاب شاخص های توسعه

شناخت بهتر و دقیق تر از وضعیت مکان های جغرافیایی در زمینه های مختلف در سطوح متفاوت منوط به در دسترس داشتن اطلاعات کامل و پردازش شده از مکان های مورد نظر است. برای نیل به این مهم از یک سری شاخص های ترکیبی اقتصادی، اجتماعی، فرهنگی، آموزشی، بهداشتی و ... استفاده می شود. این شاخص های ترکیبی می تواند سطحی از آسایش و رفاه و رشد و توسعه مکان های جغرافیایی را براساس معیارهای انتخاب شده نشان دهند. (میر نجف و موسوی) تعیین این شاخص



مهمترین قدم در مطالعات توسعه منطقه ای است و در واقع بیان آماری پدیده های موجود در منطقه است . برای بیان اهمیت شاخص های توسعه و نقش آن در بیان آماری پدیده ها ضروری است تا مفاهیم مربوط به متغیر و شاخص به طور عمیق تر روشن شود . متغیر ارقامی هستند که نمی توانند سطح توسعه ی مثلا تعداد پزشکان ، تعداد شاغلان بخش صنعت و ... زیرا که بالا بودن شاغلان و پزشکان دلیلی بر توسعه مکان های مورد نظر نمی تواند باشد . چون امکان دارد مکانی جمعیت بیشتری داشته باشد که در این صورت طبیعیست که تعداد پزشکان و یا متغیر های دیگر نیز بالا خواهد بود . در حالی که شاخص ها ارقامی هستند که برای اندازه گیری و سنجش نوسان های عوامل متغیر در طول زمان به کار می رود (آسایش ، ۱۳۷۴: ۲۹) به عبارت دیگر با تبدیل متغیر به سرانه ها ، نسبت ها مختلف درصدها و ... می توان به شاخص سازی متغیر ها اقدام کرد . مثل تعداد پزشکان در ده هزار نفر جمعیت ، درصد شاغلان بخش صنعت نسبت به کل شاغلان . مساله بعدی در ارتباط با متغیر و شاخص ها مربوط به مفاهیم آنهاست بعضی از شاخص ها منفی بیانگر عدم توسعه یک مکان جغرافیای است بنابراین در موقع شاخص سازی بایستی از انتخاب این شاخص ها صرف نظر کرد یا اینکه این شاخص ها را معکوس و از مقادیر کم کرد (کلانتری ، ۱۳۸۰: ۱۱۴) .

تشریح روش تاکسونومی

آنالیز تاکسونومی یکی از روش های درجه بندی مناطق از لحاظ توسعه یافتگی است که برای طبقه بندی های مختلف در علوم به کار برده می شود . نوع خاص آن تاکسونومی عددی ، که بنا به تعریف ، ارزیابی عددی شباهت ها و نزدیکی ها بین واحد های تاکسونومیک و درجه بندی آن عناصر می باشد (زیاری ، ۱۳۷۸: ۱۳۷) . بطوری که بین عناصر تشکیل دهنده موضوعات و وقایع که در یک گروه قرار می گیرد حداکثر تشابه یا نزدیکی وجود دارد و نیز این گروه در مجموع با عناصر تشکیل دهنده گروه های دیگر دارای حد اکثر اختلاف است . بعبارت دیگر تفاوت بین عناصر درونی یک گروه در حداقل و بین گروه با عناصر دیگر در حد اکثر است از این رو گروه های معمولاً خوشه ای از موضوعات و وقایع هستند که بر حسب تشابه یا نزدیکی عناصرشان تعریف می شود . در این روش هدف همگونی موضوعات مختلف بر اساس فاصله آنها نسبت به همدیگر اندازه گیری شود و این بدان معناست که هر مورد در فضای



تاکسونومیک قرار گرفته و فواصل آنها محاسبه گردد. (دلیر، ۱۳۸۰: ۱۵۵) در این روش معمولاً یکی از نقاط به عنوان نقطه یا منطقه ایده آل انتخاب شده و نقاط یا مناطق دیگر را بر مبنای آن درجه بندی می کنند بدین ترتیب تفاوت یا فاصله هر منطقه از آن منطقه ایده آل معین می شود .

این روش ها اولین بار به وسیله آندرسون در سال ۱۷۶۳ میلادی پیشنهاد شد، و در سال ۱۹۶۸ نیز به وسیله پرفسور هلوینگ از مدرسه عالی اقتصاد یونسکو به عنوان وسیله ای برای طبقه بندی توسعه یافتگی مناطق مطرح شد. بر اساس این روش ، درجه توسعه یافتگی بین صفر و یک بوده و هر چه درجه به دست آمده برای منطقه ای به صفر نزدیک تر باشد ، نشان دهنده میزان توسعه یافتگی بالاتر است. با مرتب کردن مناطق به ترتیب درجه توسعه یافتگی، رتبه مناطق به دست می آید. ایراد اساسی روش تاکسونومی افزون بر حساسیت بسیار نتیجه نهایی نسبت به تعداد نما گر ها به ویژه نما گرهای آماری هم جهت این است که در صورتی که داده ها نسبت به همبستگی داشته باشند نتیجه گمراه کننده خواهد بود. به عبارت دیگر، در این روش برای حذف همبستگی بین نما گر ها هیچ تلاشی صورت نمی گیرد.

مراحل روش تاکسونومی

مرحله اول اجرای روش، تشکیل ماتریس اطلاعات است. از آنجا که هریک از شهرستان های موجود دارای مشخصات و شاخص هایی است، آنها را به صورت زیر مرتب می کنیم که P تعداد شهرستان ها و X_n شاخص های مورد بررسی می باشد :

$$P1 = X_1, X_2, \dots, X_n$$

$$P2 = X_1, X_2, \dots, X_n$$

$$Pm = X_1, X_2, \dots, X_n$$

این اطلاعات به صورت ماتریس $X=[X]_{m.n}$ نوشته می شود که برای انجام مراحل بعدی آماده باشد. در مرحله دوم، بدین دلیل که هر یک از شاخص ها دارای واحد اندازه گیری متفاوتی می باشند، ماتریس داده ها را با کسر از میانگین هر شاخص و تقسیم بر انحراف معیار همان شاخص، استاندارد می کنیم که از این طریق همگی شاخص ها به واحد مشترک تبدیل شده که دارای میانگین صفر و انحراف معیار یک هستند. در این مرحله ماتریس X تبدیل به ماتریس نرمال $Z=[Z]_{m.n}$ می شود.



در مرحله سوم، فاصله هر شهرستان را نسبت به دیگر شهرستان ها برای هر یک از شاخص ها پیدا کرده تا همگنی و سنخیت آنها را مورد ارزیابی قرار دهیم. اگر D را فاصله میان دو شهرستان بدانیم خواهیم داشت :

که a و b نشان دهنده شهرستان های مختلف، Z_{aj} مقدار شاخص Z_{bj} و a و b مقدار شاخص Z_{bj} شهرستان b است.

اگر در مورد همه شاخص ها این عمل انجام شود و فاصله تمامی شهرستان ها از یکدیگر محاسبه شود، ماتریس $D=[D]_{m,m}$ تشکیل می شود که متقارن بوده و عناصر قطری آن برابر صفر می باشد.

در مرحله چهارم جهت ارزیابی همگنی شهرستان ها ماتریس $d=[d]_{m,m}$ را به گونه ای تشکیل می دهیم که هر یک از عناصر آن حداقل سطرهای ماتریس D بدون در نظر گرفتن مقادیر صفر است. سپس برای تعیین محدوده همگنی، ابتدا میانگین و انحراف معیار ماتریس d را بدست می آوریم و آنگاه محدوده همگنی (L_1, L_2) را تعریف می کنیم. اگر $\bar{d}$ را میانگین و S_d را انحراف معیار ماتریس d بدانیم :

$$S_d = \sqrt{\frac{\sum_{i=1}^m (d_i - \bar{d})^2}{m}} \quad \bar{d} = \frac{\sum_{i=1}^m d_i}{m}$$

آنگاه محدوده همگنی را به صورت زیر تعریف می کنیم

$$L_1 = \bar{d} + 2 S_d \quad L_2 = \bar{d} - 2 S_d$$

که در اینجا مقادیر L_1 و L_2 به گونه ای است که $L_2 < b < L_1$ و b را محدوده همگنی یا فاصله همگنی می نامند. حال اگر حداقل فاصله هر شهرستان از سایر شهرستان ها که در ماتریس d محاسبه شده در محدوده b واقع شود، همگنی برقرار است و به مرحله بعدی می رویم ولی اگر یکی از عناصر ماتریس d در محدوده همگنی واقع نشود، آن شهرستان را از کلیه داده ها جدا نموده و محاسبات را از اول و بدون وجود شهرستان غیر همگن ادامه می دهیم.

مرحله پنجم به محاسبه درجه شهرستان های همگن می پردازد. در این قسمت بزرگترین عدد هر شاخص در ماتریس شهرستان های همگن را انتخاب کرده و آن را مقدار ایده آل می نامیم (Z_{aj}). در ادامه



فاصله هر شهرستان از مقدار ایده آل را برای هر یک از شاخص ها بدست آورده و برای هر فاصله از فرمول زیر استفاده می کنیم.

$$C_{io} = \sqrt{\sum_{j=1}^n (Z_{ij} - Z_{oj})^2}$$

را به صورت زیر تعریف Co در ادامه سرمشق توسعه می کنیم :

$$Co = \overline{C_{io}} + 2 S_{io}$$

که $\overline{C_{io}}$ میانگین و S_{io} انحراف معیار Co است و از طریق فرمول زیر بدست می آید :

$$\overline{C_{io}} = \frac{\sum_{i=1}^m C_{io}}{m} \quad S_{io} = \sqrt{\frac{\sum_{i=1}^m (C_{io} - \overline{C_{io}})^2}{m}}$$

حال اگر Fi را درجه توسعه یافتگی شهرستان ها بدانیم، بدین صورت بدست می آید :

$$Fi = \frac{C_{io}}{Co}$$

در واقع درجه توسعه یافتگی از تقسیم فاصله شهرستان ها از شهرستان ایده آل بر سرمشق توسعه بدست می آید. دامنه تغییرات Fi محدوده بین صفر و یک است که هر چه به صفر نزدیک تر باشد، شهرستان مزبور توسعه یافته تر و برخورداری و هر چه مقدار Fi به یک نزدیک تر باشد، شهرستان توسعه نیافته تر و محروم تر تلقی می شود. (عماد زاده و دیگران ، ۱۳۸۴ : ۵۵)

تشریح روش تحلیل عاملی^۱

مدل تحلیل عاملی روشی برای خلاصه کردن اطلاعات زیاد می باشد که اولین بار توسط اسپرمن^۲ در سال ۱۹۰۴ ابداع گردید . مدل پیشنهادی او که در واقع بر مبنای مدل خطی ساده استوار بود به خوبی می توانست روابط بین متغیرهای مختلف را توجیه کند (ترابی و جهانبخش، ۱۳۸۳ : ۱۵۶). این مدل یکی از شیوه های بررسی روابط بین مجموعه متغیرهاست و هدف اصلیش این است که یگانگی ها را در میان



متغیرهای متعدد کشف کند و تعداد زیادی متغیر را به محدودی متغیرهای زیربنایی یا عامل تقلیل دهد (حسینی و دیگران، ۱۳۸۶: ۷۰). تحلیل عاملی کاربرد مختلفی دارد اگر در تحلیل عاملی هدف خلاصه تعداد شاخص به عوامل معنی دار باشد، باید از تحلیل عاملی نوع R استفاده گردد؛ در صورتی که هدف ترکیب و تلخیص تعدادی از مکان ها با نواحی جغرافیایی در گروه های همگن در درون یک سرزمین باشد، از تحلیل نوع Q باید استفاده شود (کلانتری، ۱۳۸۲: ۲۸۱). در مطالعات جغرافیایی تحلیل عاملی نوع R بیشتر برای سطح بندی مناطق، شهر ها و روستاها به کار برده می شود.

کاربرد اصلی این روش، کاهش تعداد متغیرها به عواملی می باشد که ممکن است در ظاهر وجود نداشته باشند؛ ولی در نهایت به شکل غیر وابسته باعث ایجاد اختلافات مکانی میشوند. از طرف دیگر می توان با استفاده از این روش تعیین کرد که هر یک از عاملها چه اندازه در ایجاد این اختلاف، نقش دارند. بر عکس روش تجزیه و تحلیل خوشه ای^۳ که متغیرها دارای ارزش مساوی هستند، در این روش ابتدا یک ماتریس همبستگی برای هر جفت از مناطق تشکیل می شود و با در نظر گرفتن اهمیت نقش هر یک در ایجاد عامل ها به ارزش آنها بار ویژه ای داده می شود و در نهایت متغیرهایی که ویژگی بر جسته تری نسبت به بقیه دارند، انتخاب می شوند و در دسته بندی مناطق مورد استفاده قرار می گیرند. البته باید توجه داشت که چگونگی انتخاب و کیفیت اطلاعاتی که به کار گرفته می شود، تأثیر قاطعی در نتیجه این روش خواهد داشت؛ اما با توجه به توسعه سریع عوامل نرم افزار کامپیوتری و روشهای آماری مانند تحلیلهای چند متغیره، روش یادشده کاربرد مؤثرتری در برنامه ریزی ها دارد (میاندهشتی، ۱۳۸۳: ۲۴)؛ درکل تحلیل عاملی می تواند به چهار پرسش عمده پاسخ دهد:

۱. برای تبیین الگوی روابط بین متغیرها به چند عامل مختلف نیاز است؟
 ۲. ماهیت این عوامل چیست؟
 ۳. عامل های نظری چگونه می توانند داده های مشاهده شده را تبیین کنند؟
 ۴. چه مقدار از واریانس هر متغیر مشاهده شده اساساً تصادفی یا یگانه است؟
- بنابراین هدف تحلیل عاملی کشف ساده ترین الگو از میان الگوهای مربوط به روابط میان متغیرهاست. این روش به دنبال درک این مطلب است که آیا متغیرهای مشاهده شده بر پایه تعداد کمتری متغیر)



عامل (به گونه وسیع و اساسی تبیین کرد (هومن و عسکری، ۱۳۸۴ : ۳). مراحل تحلیل عاملی را می توان به صورت زیر خلاصه کرد :

- تهیه ماتریس استاندارد .
 - محاسبه ماتریس ضرایب همبستگی .
 - استخراج عاملی .
 - چرخش عوامل .
 - محاسبه نمرات
- در تحلیل عاملی چند اصطلاح عمده وجود دارد که در زیر به آنها اشاره می شود :
- مقدار خالص^۴ : میزان واریانس تبیین شده به وسیله ی هر عامل را بین می کند .
 - عامل^۵ : عبارت است ترکیب خطی متغیرهای اصلی ، که نشان دهنده جنبه های خلاصه شده از متغیرهای مشاهده شده است .
 - بار عاملی^۶ : عبارت است از همبستگی بین متغیرهای اصلی و عوامل .
 - ماتریس عاملی^۷ : جدولی است که بار های عاملی کلیه متغیرها را در هر عامل نشان می دهد .
 - چرخش عاملی^۸ : فرآیندی است برای تعدیل محور عاملی به منظور دستیابی به عامل های معنی دار و ساده
 - وزن عاملی^۹ : وزن هایی هستند که متغیرها داده می شوند تا در تعیین امتیاز عوامل مشکل ایجاد نشود .
 - امتیاز عاملی^{۱۰} : وزن دهی عددی است که هر یک از نواحی پس از ضرب وزن عاملی از طریق معادل Z-S استاندارد به دست می آورد (طالبی و زنگی آبادی، ۱۳۸۰ : ۲۱) .

انواع مدل تحلیل عاملی

موارد استفاده تحلیل عاملی را به دو دسته کلی می توان تقسیم کرد

(الف) مقاصد اکتشافی ، (ب) مقاصد تاییدی

موارد استفاده اکتشافی نیز به دو رویکرد کلی تقسیم می شود:



مواردی که هدف آن پیدا کردن متغیرهای مکنون یا سازه های یک مجموعه متغیر اندازه گیری شده است. برای نیل به این هدف از روش تحلیل عامل مشترک (یا تحلیل عامل اصلی) و با استفاده از ماتریس همبستگی یا کواریانس متغیرهای اندازه گیری شده (نمره سوالات یک آزمون یا ریز نمرات آزمون ها) استفاده می شود. از لحاظ نظری متغیرهای مکنون یا سازه ها علل زیربنایی متغیرهای اندازه گیری شده است. رگرسیون متغیرهای اندازه گیری شده روی متغیرهای مکنون وزن هایی فراهم می آورد که بارهای عاملی نامیده می شود. تحلیل عامل مشترک، واریانس هر متغیر اندازه گیری شده را به دو واریانس مشترک و واریانس اختصاصی افزاز می کند. واریانس مشترک، تغییرات مشترک متغیرهای اندازه گیری شده را با متغیرهای مکنون نمایان می کند.

در موارد اکتشافی که هدف تلخیص مجموعه ای از داده ها باشد، از تحلیل مولفه های اصلی استفاده می شود.

در تحلیل مولفه های اصلی، واریانس کل متغیرهای مشاهده شده تحلیل می گردد. ماتریس همبستگی متغیرهای اندازه گیری شده دارای قطر اصلی ۱ است. در حالی که در تحلیل عامل مشترک در قطر اصلی ماتریس همبستگی میزان اشتراک (واریانس مشترک متغیر اندازه گیری شده و متغیرهای مکنون) قرار می گیرد. وقتی میزان اشتراک به عدد یک نزدیک باشد نتایج تمام روش های اکتشافی با نتایج مولفه های اصلی مشابه خواهد بود.

در تحلیل مولفه های اصلی، بر عکس تحلیل عامل مشترک، مولفه ها طوری برآورد می شود تا واریانس متغیرهای مشاهده شده را در کمترین ابعاد نشان دهد و مولفه های اصلی در واقع مجموع موزون متغیرهای مشاهده شده است. به عبارت دیگر در تحلیل مولفه های اصلی، متغیرهای مشاهده شده علل متغیرهای ترکیبی (مولفه ها) می باشد.

در تحلیل های عاملی تاییدی، که هدف پژوهشگر تایید ساختار عاملی ویژه ای می باشد، درباره تعداد عامل ها به طور آشکار فرضیه های بیان می شود و برازش ساختار عاملی مورد نظر در فرضیه با ساختار کواریانس متغیرهای اندازه گیری شده مورد آزمون قرار می گیرد (همان: ۲۳).

تحلیل عاملی را نیز بر حسب نمونه یا جامعه بودن آزمودنی ها و متغیرها به دو دسته ی توصیفی و استنباطی تقسیم می کنند.



جدول شماره (23): انواع تکنیک های استخراج عامل ها را بر حسب اکتشافی- تاییدی و توصیفی- استنباطی

نوع تحلیل	توصیفی	استنباطی
اکتشافی	- مولفه های اصلی - عامل مشترک (عامل اصلی) - تحلیل تصویر - تحلیل حداقل مانده	- تحلیل عاملی متعارف - حداکثر درست نمایی - تحلیل عاملی آلفا
تاییدی	- چند گروهی - Linear Structural Relationships - LISREL یا	- حداکثر درست نمایی تاییدی - LISREL

هدف از تحلیل عاملی

هدف اصلی مدل تحلیل عاملی این است که یگانگی ها را در میان متغیرهای متعدد کشف کند و تعداد زیادی متغیر را به معدودی متغیرهای زیربنایی یا عامل تقلیل دهد.

هدف از این کاهش را می توان حذف اطلاعات تکراری و مشترک بین متغیرها، ساده و آسان تر کردن تحلیل، صرفه جویی، معنی دار کردن و خلاصه سازی اطلاعات دانست.

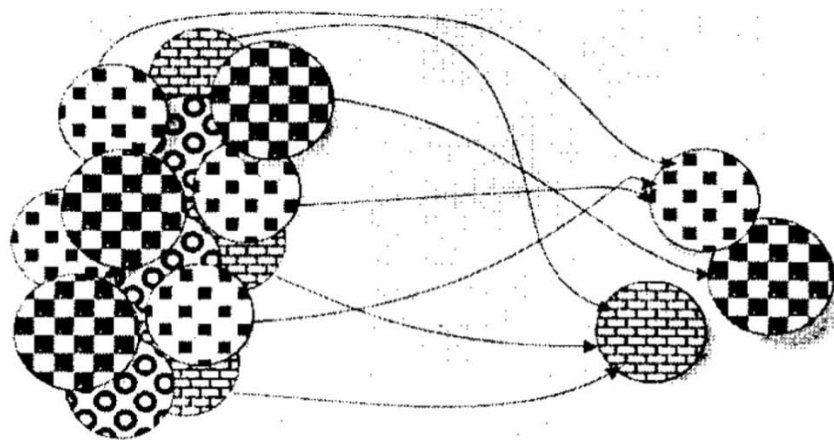
مجموعه کاهش یافته متغیرهای می تواند :

✓ زیر مجموعه ای از متغیرهای اولیه

✓ معرفی متغیرهای جدیدی باشد که از ترکیب متغیرهای اولیه به دست آمده

در شکل زیر هدف از تحلیل عاملی به نمایش گذاشته شده است. همان طوری که مشاهده می شود

جرم چندین دایره که دارای اشتراکاتی است در دوایر کمتری که اشتراکات کمتر یا اصولاً جرم مشترکی ندارند، ریخته شده است (رمضانی، ۱۳۷۶: ۵۶).



متغیرهای فراوان اولیه

عوامل ترکیبی از متغیرها

دارای اشتراکاتی کمتر، ولی معنی دار با حفظ اطلاعات هستند با مشترکات فراوانی که بیان آنها را مشکل می سازد

پیشینه مدل تحلیل عاملی

- نخستین کار درباره تحلیل عاملی ← چارلز اسپیرمن در سال ۱۹۰۰ (پدراین روش).
- پیشنهاد روش محورهای اصلی ← توسط کارل اسپیرمن در ۱۹۰۱ و توسعه آن توسط هتلینگ در ۱۹۳۰.
- ابداع و توسعه روش سانترئید ← ترستون در ۱۹۳۰ (برجسته ترین فرد در تحلیل عاملی نوین) ← ابداع روش ساختار ساده.
- توسعه کامپیوترهای پرسرعت ← حرکت از تئوری گراوی به سمت تحلیل عاملی اکتشافی (بخصوص تقریبا از تئوری عامل مشترک ترستون و فرمول بندی های هتلینگ)
- ابداع روشهای تحلیل عاملی در این دوره ← روش تحلیل تصویر (گاتمن)، روش تحلیل عاملی بنیادی (رائو، هریس)، تحلیل عاملی الفا (کیسر و کافری)، روش کمترین پس ماند (هامن، جونز)
- کاربرد روش تحلیل عاملی به عنوان یک تکنیک آماری ← انتشار مقاله لاولی درباره روش پیشینه احتمال در ۱۹۴۰ (همان: ۵۸)

مراحل اجرای تحلیل عاملی

برای اجرای یک تحلیل عاملی چهار گام اساسی ضرورت دارد:

۱- تهیه یک ماتریس همبستگی از تمام متغیرهای مورد استفاده در تحلیل و برآورد اشتراک



۲- استخراج عامل ها

۳- انتخاب و چرخش عامل ها برای ساده تر ساختن و قابل فهم تر کردن ساختار عاملی

۴- تفسیر نتایج

تهیه ماتریس همبستگی

تهیه ماتریس همبستگی از تمام متغیر های مورد مطالعه، اولی ن گام تحلیل عاملی است. در تهیه ماتریس همبستگی محقق باید تصمیم بگیرد که در قطر اصلی این ماتریس عدد ۱ یا عدد دیگری بگذارد. این عدد که اشتراک نامیده می شود، نشانگر نسبت واریانس مشترک بین هر متغیر و عامل هاست. مقدار اشتراک بین صفر و ۱ تغییر می کند. اشتراک صفر حاکی از این است که عامل های مشترک هیچ تغییری را در متغیر خاصی تبیین نمی کند، و اشتراک ۱ حاکی از این است که تمام تغییرات متغیر خاص توسط عامل های مشترک تبیین می شود. به عبارت دیگر اشتراک مساوی ۱ حاکی از این است که کل واریانس متغیر های مشاهده شده تحلیل عاملی می شود، در حالی که اگر واریانس مشترک متغیر های مشاهده شده و متغیر های مکنون (عامل ها) تحلیل عاملی شود، برآورد اولیه ای از اشتراک باید در قطر اصلی ماتریس همبستگی قرار گیرد. یکی از روش های معمول برای برآورد این اشتراک محاسبه مجذور همبستگی چندگانه (R^2) هر متغیر مستقل از روی سایر متغیر های مستقل است. این R^2 حد پایین برآورد اشتراک را فراهم می آورد. نخست این برآورد در قطر اصلی ماتریس همبستگی قرار می گیرد و ماتریس تحلیل عاملی می شود. از بارهای عاملی به دست آمده مجدداً اشتراک های جدید محاسبه می شود. چنان چه تفاوت این اشتراک ها از اشتراک های اولیه از مقدار ملاک (مثال ۰.۰۰۱) بیشتر باشد عمل محاسبه عامل ها و بار عاملی آن ها با قرار دادن اشتراک های جدید در قطر اصلی ماتریس تکرار $\square$ (Iteration) می گردد. اشتراک ها معمولاً در دو یا سه تکرار به اشتراک ملاک می رسد (شاهسونی و همکاران، ۱۳۸۵: ۱۲۷).

استخراج عامل ها

هدف مرحله استخراج عامل ها، به دست آوردن سازه های زیر بنایی است که تغییرات متغیر های مورد مشاهده را موجب شده است. SPSS نخست ترکیب هایی از متغیر ها را که همبستگی های آن ها بالاترین میزان از واریانس کل مشاهده شده را نشان می دهد، انتخاب می کند. این مجموعه عامل ۱ را



می‌سازد. عامل ۲، مجموعه متغیرهایی است که بالاترین سهم را در تبیین واریانس باقیمانده دارد. این شیوه برای عامل سوم، چهارم و عامل های بعدی ادامه پیدا می‌کند تا تعداد عامل های استخراج شده برابر با تعداد متغیر ها گردد.

همبستگی هر متغیر با هر عامل بار عاملی [□] نامیده می‌شود و مقدار آن بین ۱- و ۱+ تغییر می‌کند. واریانس تبیین شده توسط هر عامل برابر است با مجذور بار های عاملی آن. این واریانس مقدار ویژه [□] نامیده می‌شود. اولین مقدار ویژه همواره بیشترین بوده و از ۱ بزرگ‌تر می‌باشد. مقدار ویژه برای عامل های بعدی کوچک‌تر می‌باشد.

چرخش عامل ها

تمام عامل های استخراج شده مورد علاقه محقق نیست. هدف تحلیل عاملی تبیین پدیده های مورد نظر با تعداد کمتری از متغیر های اولیه است. در وهله اول هدف تعیین تعداد عامل هایی است که در تحلیل نگه داشته می‌شود. بنابراین عامل هایی باید نگه داشته شود که اعتبار صوری یا نظری داشته باشد. منتها قبل از فرایند چرخش نمی‌توان به معنی هر عامل به خوبی پی برد، بنابراین معمولاً از ملاک های ریاضی مانند ملاک کایزر یا آزمون اسکری کتل برای نگه داشتن عامل ها استفاده می‌شود.

براساس ملاک کایزر فقط عامل هایی نگه داشته می‌شوند که مجموع مجذور بارهای عاملی آن‌ها (مقدار ویژه) یک یا بیشتر باشد. این ملاک برای تحلیل عاملی آلفا مناسب است و برای سایر روش های تحلیل عاملی کران پایینی فراهم می‌آورد. در روش اسکری کتل نمودار مقدار ویژه برای هر عامل ترسیم می‌شود. در نقطه ای که شکل منحنی برای مقادیر ویژه به صورت افقی درآید، آن نقطه اسکری نامیده شده و عامل هایی که سمت چپ آن قرار دارد عامل های واقعی و آن هایی که در سمت راست آن قرار دارند عامل های خطا قلمداد می‌شود. در تفسیر نتایج آزمون اسکری ممکن است میان نظرات پژوهشگران درباره تعداد عامل های واقعی اختلاف نظر پدید آید. همچنین امکان دارد که بیش از یک اسکری موجود باشد. لذا لازم است علاوه بر آزمون اسکری آزمون های دیگری از جمله آزمون کایزر صورت گیرد.

پس از انتخاب عامل ها چرخش آن‌ها ضرورت دارد. هدف از چرخش عامل ها رسیدن به یک ساختار عاملی ساده است. در تحلیل عاملی، ساختار های عاملی متعددی برای یک ماتریس همبستگی وجود دارد.

□ Factor Loading
□ Eigen Value



اولین عامل غالباً یک عامل کلی است که تمام یا اکثر متغیرها بار عاملی بالایی روی این عامل دارد. عامل هایی بعدی معمولاً دو قطبی است و بارهای عاملی مثبت و منفی داشته و قابل تفسیر نمی باشد با چرخش ساختار عاملی روشن تر می شود.

مشهورترین ملاک برای خوبی یک ساختار عاملی، ملاک مشهور ساختار ساده ترستون است. طبق این ملاک هر متغیر باید حداقل یک بار عاملی غیر صفر داشته باشد. هر عامل باید فقط با چند متغیر همبستگی بالا داشته باشد. (منظور از همبستگی همان بار عاملی متغیر روی عامل است) و بار عاملی بقیه متغیرها روی این عامل باید اساساً صفر باشد. هر متغیر باید روی یک عامل بار عاملی بالا داشته باشد. اغلب شیوهای چرخش با توجه به این ملاک ها طرح ریزی شده است.

چرخش عامل ها به دو صورت متعامد (ناهمبسته) و مایل (همبسته) صورت می گیرد. در چرخش متعامد عامل های به دست آمده با هم همبستگی ندارند، در حالی که در چرخش مایل عامل ها با هم همبستگی دارند. روش های متعددی برای چرخش متعامد و مایل وجود دارد. از جمله چرخش های متعامد که غالباً مورد استفاده قرار می گیرد چرخش واریماکس است. از روش های چرخش مایل روش اوبلیمین را می توان نام برد.

بدیهی است که به کمک نرم افزارهای کامپیوتری از جمله SPSS می توان به سهولت تمام محاسبات لازم برای تحلیل عاملی را انجام داد. اما مهم ترین مرحله تحلیل عاملی تفسیر نتایج به دست آمده است. تفسیر

در یک ساختار عاملی آرمانی هر یک از متغیرها بار عاملی بالا (بزرگتر از ۰.۵) روی یکی از عامل ها و بار عاملی پایین (کمتر از ۰.۲) روی سایر عامل ها دارد. علاوه بر این، عامل هایی که بار عاملی بالا دارد، و اعتبار صوری آنها نیز مطلوب است و به نظر می رسد که خصیصه مکنونی را اندازه گیری می کند. چنین ساختار عاملی در واقع به ندرت اتفاق می افتد. غالباً یک متغیر روی چند عامل بار عاملی دارد و دو یا چند متغیر روی عامل نامناسبی بار عاملی دارد - محقق باید درک کافی از داده هایش داشته باشد و محاسبات تحلیل عاملی به تنهایی نمی تواند نتایج روشن فراهم آورد (همان: ۱۲۹).



معرفی نرم افزار SPSS

با پیشرفت علوم و گسترش تکنولوژی ، اهمیت استفاده از روشهای آماری در علوم مختلف بیش از پیش مورد توجه قرار گرفته است و آموختن آمار کاربردی در هر رشته جزء ملزومات گردیده است . فرآیند آنالیز آماری کمک می کند تا پژوهشگر بتواند از داده های اولیه، اطلاعات مورد نیاز خود را استخراج کند و در صورت لزوم نتایج را تعمیم دهد. اگر حجم داده ها بزرگ باشد، استفاده از روشهای مختلف آنالیز آماری بسیار خسته کننده و مشکل خواهد بود . امروزه انواع نرم افزارهای مختلف آماری موجود، قادرند انواع آنالیزهای آماری را انجام دهند . نرم افزار **SPSS** یکی از قدیمی ترین ، برنامه های کاربردی در زمینه تجزیه و تحلیل های آماری است ؛ نرم افزاری آماری با قابلیت های انجام توصیفی زیبا و گویا از اطلاعات، شامل رسم نمودارها و چارت های گوناگون و محاسبات مربوط به میانگین، انحراف معیار واریانس، میانه و غیره.

کلمه **SPSS مخفف Statistical package for social science** (نرم افزار آماری برای علوم اجتماعی) می باشد. این نرم افزار که یکی از نرم افزارهای تخصصی آمار است، بیشتر به بحث های آماری در حیطه علوم اجتماعی، روانشناسی و علوم رفتاری و ... می پردازد.

قابلیت های نرم افزار SPSS به شرح زیر است:

- تهیه خلاصه های آماری مانند گراف ها، جداول، آماره ها و ...
- انواع توابع ریاضی مانند قدر مطلق، تابع علامت، لگاریتم، توابع مثلثاتی و...
- تهیه انواع جداول سفارشی مانند جداول فراوانی، فراوانی تجمعی، درصد فراوانی و...
- انواع توزیع های آماری شامل توزیع های گسسته و پیوسته
- تهیه انواع طرح های آماری
- انجام آنالیز واریانس یکطرفه، دوطرفه، چندطرفه و آنالیز کوواریانس
- تکنیک های تجزیه و تحلیل سری های زمانی
- ایجاد داده های تصادفی و پیوسته
- محاسبه انواع آماره های توصیفی
- انواع آزمون های مرتبط با مقایسه میانگین بین دو یا چند جامعه مستقل و وابسته



➤ قابلیت مبادله اطلاعات با نرم افزارهای دیگر

➤ پیازش انواع مختلف رگرسیون

نوارهای SPSS

مانند اغلب نرم افزارها در ویندوز محیط اصلی SPSS دارای نوارهای متعددی است. نوارهای SPSS

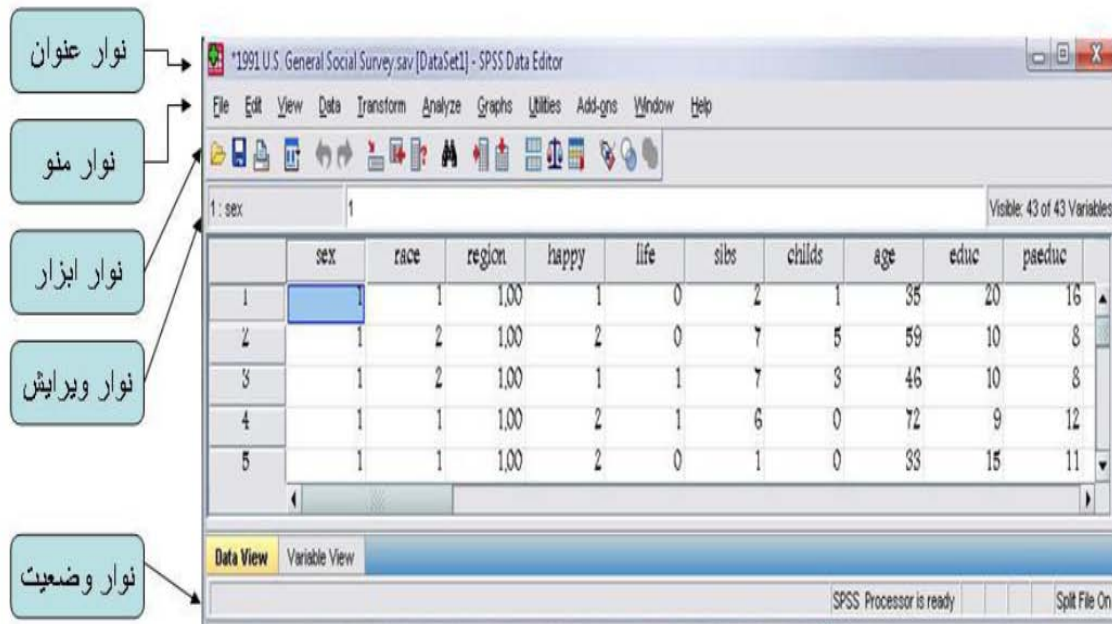
که در شکل زیر نشان داده شده اند، به صورت زیر اند :

۱- نوار عنوان : در این نوار اسم فایل جاری و مشخصات آن نشان داده می شود.

۲- نوار منو : اصلی ترین نوار SPSS نوار منو است و تقریباً کلیه فعالیت های مربوط به باز و بسته کردن و ذخیره فایل ها، ویرایش، تجزیه و تحلیل داده ها و تغییرات در روند اجرای نرم افزار، در گزینه های این نوار قرار دارند.

۳- نوار ابزار : گاهی استفاده از یک ابزار مناسب صرفه جویی در وقت است. ابزارهایی که بیشتر مورد نیاز هستند در نوار ابزار قرار دارند برای دسترسی سریعتر به این ابزارها روی آیکن های موجود در نوار ابزار کلیک کنید.

۴- نوار ویرایش داده ها : این نوار شامل سه قسمت است. در سمت چپ می توانید موقعیت هر سلول





را مشاهده کنید. در بخش میانی می توانید مقادیر سلول فعال را مشاهده نموده و ویرایش کنید. در قسمت سمت راست تعداد متغیرهای فایل داده ها نمایش داده می شود.

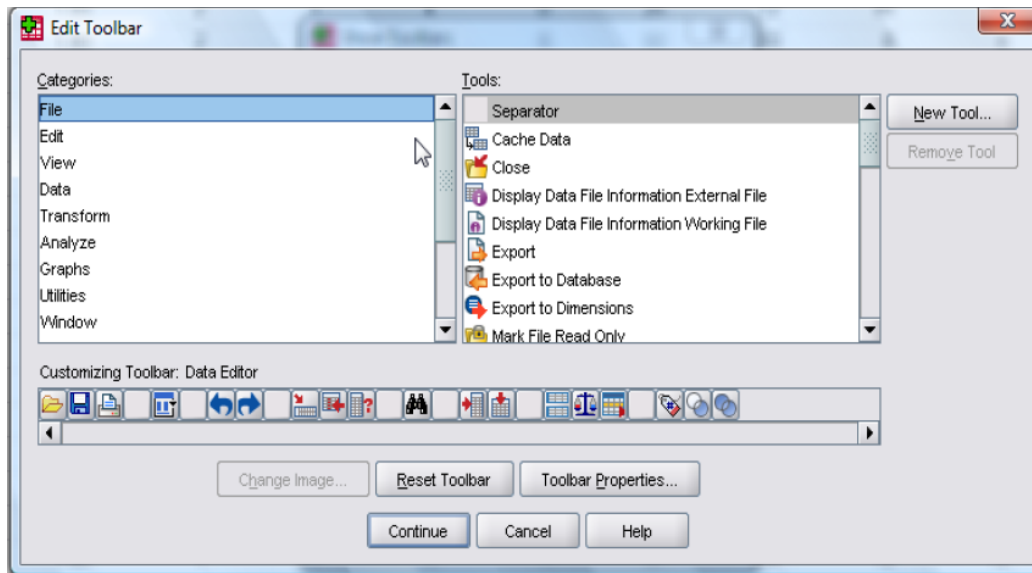
۵- نوار وضع یت : در این نوار فعالیت های در حال اجرای SPSS نمایش داده می شود. مثلاً اگر روی داده ها تغییری مانند Split یا Filter و مانند آن انجام داده اید، می توانید در سمت راست این نوار، پیغام آن را مشاهده کنید.

تهیه و تنظیم نوار ابزار دلخواه

اگر می خواهید نوار ابزار را به سلیقه خودتان تنظیم کنید، از فرمان

View/Toolbars/Customize

به کادر محاوره Show Toolbar وارد شوید. در این کادر محاوره می توانید ابزارها را اضافه و کم کنید. اندازه آیکون ها را کوچک و بزرگ کنید. همچنین قادر خواهید بود نوار ابزار موجود را به دلخواه تغییر داده یا نوار ابزار جدیدی بسته به نیاز خودتان ایجاد کنید. (به شکل زیر دقت کنید)

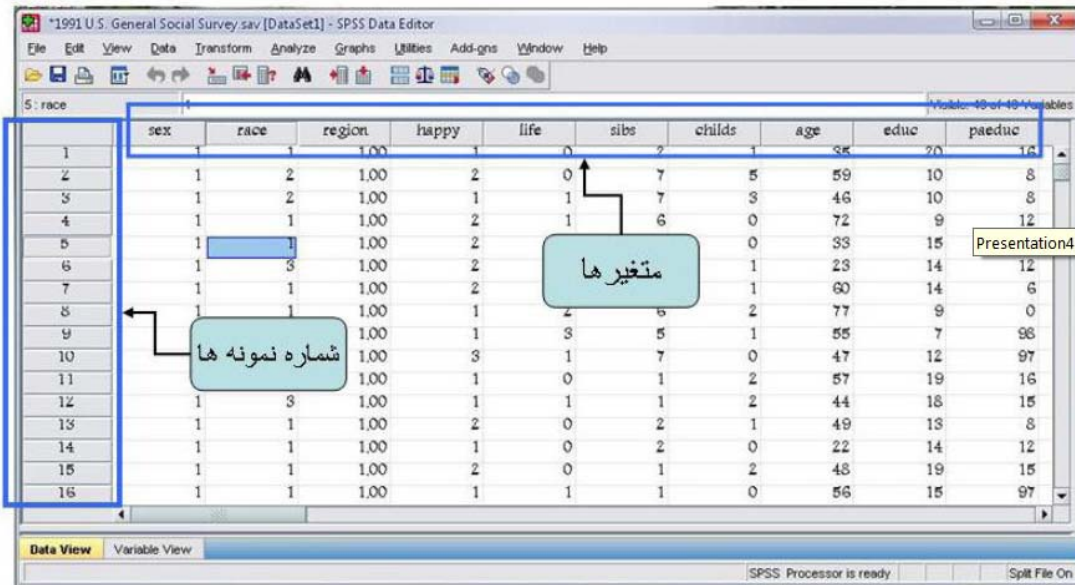


پنجره ویرایش گر داده ها (Data Editor)

در ابتدای شروع کار با SPSS که نرم افزار را فرا خوانی می کنید، اصلی ترین پنجره SPSS یعنی پنجره ویرایش گر داده ها را مشاهده خواهید کرد. این پنجره مانند سایر پنجره ها در ویندوز، از طریق کلیک کردن و کشیدن نوار عنوان جا به جا می شود و می توانید از طریق کلیک کردن و کشیدن گوشه ها



و حاشیه ها اندازه آن را تغییر دهید. پنجره ویرایش داده ها شامل ۲ پنجره فرعی است. یکی برای نمایش و ویرایش متغیرها به نام Variable View که در شکل زیر آن را مشاهده می کنید و دیگری برای نمایش و ویرایش داده ها به نام Data View است.



برای رفتن به هریک از پنجره های فرعی باید در پایین و سمت چپ این پنجره روی یکی از تب آن کلیک کنید. تب فعال به رنگ زرد است. همانگونه که در شکل مشاهده می کنید بخش Data View در این پنجره محیطی برای وارد کردن داده ها به نرم افزار و انجام ویرایش بر روی داده ها است. کاربرد در ابتدا باید داده ها را برای انجام تجزیه و تحلیل در این محیط وارد کند. برای این کار او باید متغیرهای مناسب را در بخش Variable View در همین پنجره وارد کرده و با استفاده از منوهای مختلف این پنجره، اصلاحات مناسب را روی داده ها اعمال کند و آنها را برای تجزیه و تحلیل های آماری آماده نماید. Data View به صورت یک ماتریس $n \times n$ است:



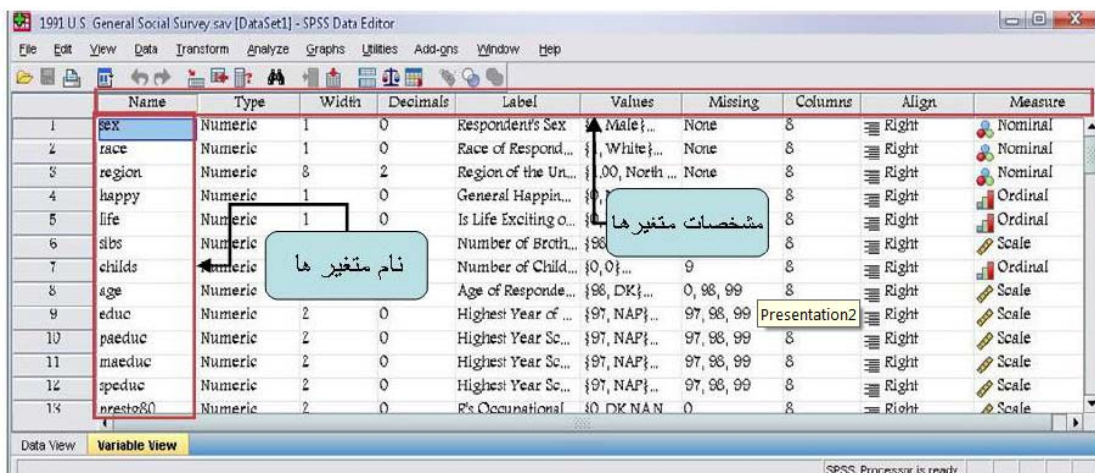
	gender	age	accident	var
1	Male	23	2	
2	Male	35	1	
3	Male	26	1	
4	Male	25	0	
5	Male	28	2	
6	Male	31	1	
7	Male	23	1	
8	Male	31	3	
9	Male	31		Snap1.bmp
10	Male	27	1	
11	Male	26	3	
12	Male	33	3	
13	Male	27	3	
14	Male	30	5	
15	Male	34	0	
16	Male	29	2	
17	Male	24	1	
18	Male	26	2	
19				

الف- سطرها : نشان دهنده ی تعداد نمونه هایی است که از آنها اطلاعات گرفته شده است.

ب- ستون ها: نشان دهنده متغیرهایی است که از هر نمونه پرسیده می شود.

ج- سلول ها: از تقاطع سطر ها و ستون ها سلول به وجود می آید که حاوی اطلاعات مربوط به نمونه ها هستند.

این بخش از پنجره Data Editor برای نمایش و ویرایش متغیرها است. هر کاربر قبل از وارد کردن داده ها باید ابتدا متغیرها را بطور مناسبی در این بخش تعریف کند. همان طور که در شکل پیدا است، در Variable View می توانید مشخصات متغیرها را مشاهده و ویرایش کنید.



جزئیات هر یک از مشخصه ها را در زیر توضیح داده ایم.



Name: در این مشخصه باید یک اسم برای متغیر یا سوالی که در پرسشنامه پرسیده‌اید، در نظر

بگیرید. انتخاب نام متغیرها در SPSS تابع قوانین زیر است :

✓ برخلاف نسخه‌های قبلی

spss که نام متغیر نمی توانست بیشتر از ۸ کاراکتر باشد، در نسخه‌های جدید مجاز هستیم تا ۶۴ کاراکتر برای نام اختصاص دهیم.

✓ نام متغیر می تواند شامل

حروف کوچک یا بزرگ، عدد یا یکی از کاراکترهای @ و # و _ و \$ باشد.

✓ از گذاشتن فاصله در نام یک

متغیر خودداری کنید.

✓ از گذاشتن کاراکتر های # و

\$ در ابتدای نام یک متغیر اجتناب کنید.

✓ نام متغیر می تواند با @

شروع شود.

✓ نام متغیر نباید با کاراکتر

های . یا _ تمام شود.

✓ نام متغیر نباید تکراری

باشد.

✓ نام متغیر نباید یکی از

کلمات کلیدی مانند ALL-AND-BY-EQ-GE-LE-LT-NE-NOT-OR-TO-

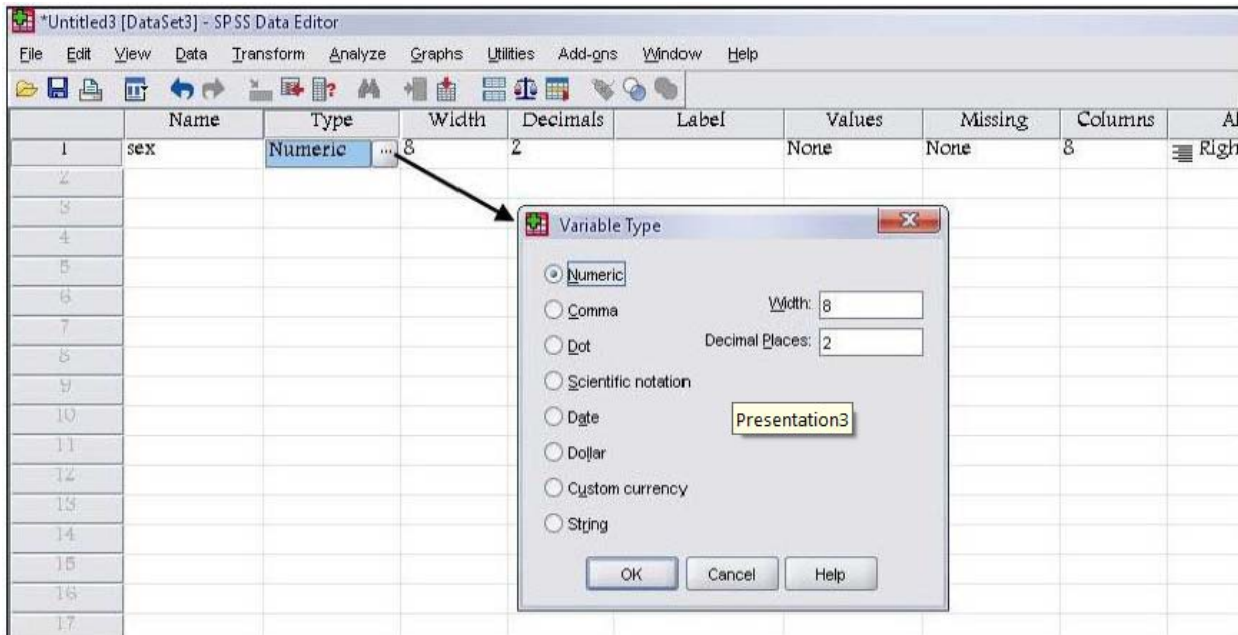
WITH که spss به عنوان عبارت محاسباتی از آنها استفاده می کند، باشد.

TYPE: در این قسمت که در شکل زیر مشاهده می کنید، باید نوع متغیر را که ممکن است عدد،

تاریخ، زمان، پول، نماد علمی و رشته باشد را مشخص کرد. برای این کار زیر عنوان TYPE کلیک کنید

تا سه نقطه داخل سلول فعال شود. با کلیک مجدد بر روی این سه نقطه پنجره ی Variable Type

مانند شکل زیر باز می شود.



گزینه های مربوط به این قسمت به صورت زیر است:

Numeric (عددی) به مقادیر عددی اختصاص دارد. اگر این گزینه را برای متغیر انتخاب می کنید مفهوم آن این است که مقادیر متغیر عددی خواهند بود. اگر داده ها اعشار هم داشته باشد در قسمت Width تعداد ارقام اعشار را مشخص کنید.

Comma (ویرگول) این گزینه هم به مقادیر عددی اختصاص دارد و در آن هر عدد سه رقم سه رقم از سمت راست با یک ویرگول جدا می شوند. مانند ۵،۶۵۱،۸۲۳

Dot (نقطه) این گزینه هم به مقادیر عددی اختصاص دارد و در آن هر سه رقم از سمت راست با یک نقطه جدا می شوند. مانند: ۵۴۸.۳۶۵.۲۵

Scientific notation (نماد علمی) اگر این گزینه را انتخاب کنید اعداد با نماد علمی نمایش داده می شوند.

Dollar (دلار) برای متغیرهایی است که حاوی مبالغی بر حسب دلار هستند و در آن برای نمایش مبالغ قالب های مختلفی پیش بینی شده است.



Custom currency (دلخواه) برای متغیرهایی است که حاوی مبالغی بر حسب یک واحد پولی غیر از دلار باشند و در آن برای نمایش مبالغ قالب های مختلفی پیشنهاد شده است.

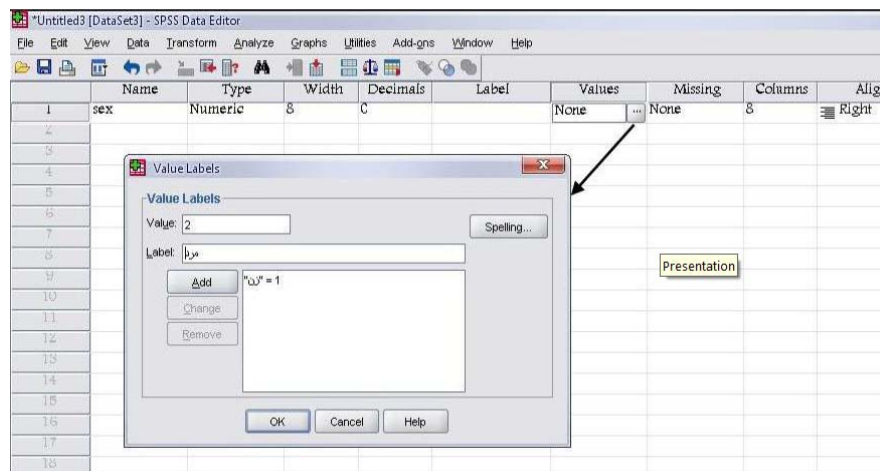
String (رشته ای) برای متغیرهایی است که حاوی رشته ای از کاراکتر ها هستند.

Width: طول رشته یا تعداد ارقام صحیح در نظر گرفته شده برای متغیر را می توان افزایش یا کاهش داد. علاوه بر اینکه در پنجره **Type Variable** می توانید طول رشته را مشخص کنید، این امکان در این قسمت نیز وجود دارد. کافی است در درون سلول زیر این مشخصه را کلیک کنید.

Decimal: در این قسمت تعداد رقم های اعشار یک عدد را می توان کاهش یا افزایش داد. علاوه بر اینکه در پنجره **Variable Type** نیز می توان تعداد رقم های اعشار را مشخص کنید، این امکان در این قسمت هم وجود دارد.

Label: در این قسمت می توانید توضیحات اضافه ای به صورت مشروح درباره ی متغیر وارد کنید تا در هنگام استفاده از متغیرها بتوانید آنها را دقیقاً از یکدیگر تمیز دهید.

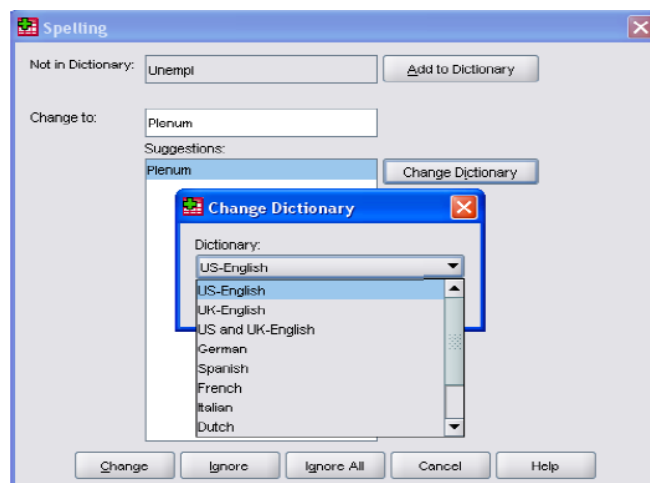
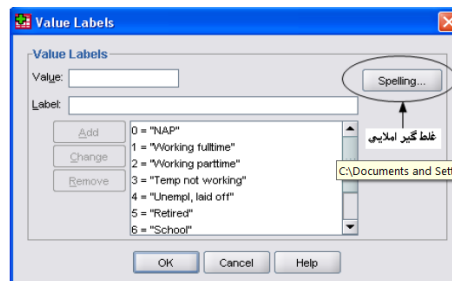
Values: در این قسمت می توانید سطوح یک متغیر کیفی (Ordinal; Nominal) را ارزش گذاری کنید. برای این کار در سلول زیر عنوان **Value** دوبار کلیک کنید تا سه نقطه داخل سلول فعال شود. با کلیک بر روی این سه نقطه پنجره **Value Label** باز می شود و می توانید در آن به هر یک از سطوح متغیر کیفی، (کد) یا عددی را نسبت دهید.





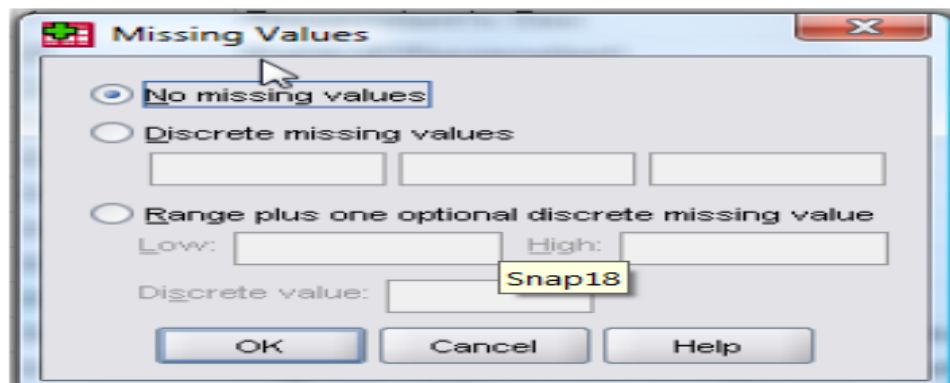
مثلا اگر می خواهید در متغیر جنسیت برای زن مقدار "۱" و برای مرد مقدار "۲" را در نظر بگیرید، در کادر Value عدد ۱ و در کادر Value Label کلمه "مرد" را تایپ کنید و سپس کلید Add را کلیک کنید. به همین صورت می توانید برای وارد کردن کد ۲ برای زن نیز عمل کنید و در انتهای کار OK را کلیک کنید.

از کلید های Change و Remove برای تغییر و پاک کردن کدهایی که قصد تغییر یا حذف آنها را دارید، استفاده کنید. یکی از امکانات نسخه های جدید نرم افزار SPSS قرار دادن غلط گیر املائی است. به همین منظور گزینه Spelling در این کادر محاوره گنجانده شده است.



Missing Value

اگر می خواهید مانند شکل زیر برای مقادیر گمشده یک متغیر، برچسب تعریف کنید ابتدا روی عنوان Missing در پنجره Variable View کلیک کنید تا سه نقطه درون سلول زیر آن فعال شود. سپس روی سه نقطه کلیک کنید به پنجره Missing Value باز شده و بتوانید مقادیر گمشده را کدگذاری کنید.



در این پنجره گزینه های زیر وجود دارد:

No missing value

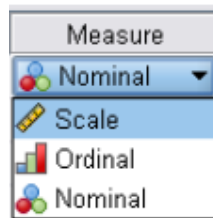
اگر در متغیر داده ی گم شده ای وجود ندارد، آن را تغییر ندهید (این گزینه از قبل علامت دار شده است)
Discrete missing values: با استفاده از این گزینه می توانید تا ۳ عدد را به عنوان داده ی فراموش شده در نظر بگیرید.

Range plus one optional discrete missing value: در این گزینه می توان فاصله ای از اعداد را به عنوان مقادیر فراموش شده در نظر گرفت.

Columns: تعداد ستون ها را برای هر متغیر می توان افزایش و کاهش داد. نرم افزار به طور پیش فرض برای هر ستون ۸ کاراکتر در نظر می گیرد. علاوه بر اینکه در پنجره Variable Type می توانید طول رشته را مشخص کنید، این امکان در این بخش نیز وجود دارد.

Align: در این مشخصه، تراز داده ها در سلول ها به صورت راست چین، چپ چین یا وسط چین قابل تغییر است.

Measure: تعیین اینکه مقیاس اندازه گیری متغیر چیست، موضوع بسیار مهمی است که نوع تحلیل های آماری بر اساس آن تغییر می کند. در این قسمت کاربر باید یکی از مقیاس های Scale (کمی)، Ordinal (ترتیبی)، Numerical (اسمی) را به درستس برای متغیر مورد نظر خود انتخاب نماید.



خروجی های SPSS (Outputs)

هنگامی که شما برنامه SPSS را باز می کنید، به طور همزمان ۲ پنجره باز می شود. یک پنجره Data Editor که پنجره اصلی SPSS است و پنجره ی دیگر خروجی ها Viewer است. پنجره خروجی SPSS مخصوص پیغام ها، هشدارها، نتایج حاصل از تحلیل مانند جداول آماری، نمودارها، متن برنامه ها، و غیره است. این پنجره مانند یک رابط بین نرم افزار و کاربرد است. شما وقتی از نرم افزار تقاضائی دارید حاصل عملیاتی که انجام می شود در یک فایل خروجی (Output) به شما گزارش می شود. از ابتدا که با نرم افزار شروع بکار می کنید تا زمانی که قصد خروج دارید، نتایج تقاضاهای شما را در این فایل ثبت می شود. وقتی قصد خروج از برنامه را دارید علاوه بر پیغامی مبنی بر بستن نرم افزار پیغام دیگری برای نگهداریه فایل خروجی از شما پرسیده می شود.

پنجره Viewer

همان طور که گفتیم این پنجره مکانی است که نتیجه کارتان را در آن مشاهده می کنید و شامل ۲ بخش است. بخش Outline Pane که نمای کلی اجزای خروجی یا سر فصل تمام نتایج موجود از Viewer را به نمایش می گذارد و بخش دوم در سمت راست به نام Display Pane است که در آن اشیای خروجی به نمایش گذاشته شده است.



SPSS Processor is ready

Output

- Log
- Summarize
- Title
- Notes
- Active Dataset
- Case Processing Summary
- Case Summaries

Summarize

[DataSet2] C:\Documents and Settings\Digaran\My Documents\My Pictures\P:

Case Processing Summary^a

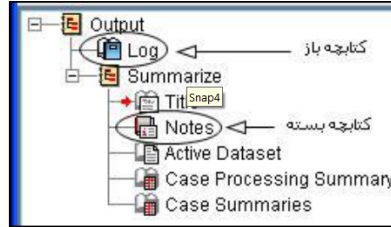
	Cases				Total	
	Included		Excluded		N	Percent
	N	Percent	N	Percent		
Sen	48	96.0%	2	4.0%	50	100.0%
Tedade Baradar Va Khahar	50	100.0%	0	0%	50	100.0%

a. Limited to first 50 cases.

Case Summaries^a

	Sen	Tedade Baradar Va Khahar
1	61	1
2	32	2
3	35	2

- با قرار دادن موس روی خطی که این دو بخش را از یکدیگر جدا کرده است می توانین پهنای هر یک از آنها را تغییر دهید. کافی است موس را روی مرز دو قسمت نگه داشته و کلیک موس را نگه دارید سپس آن را به چپ یا راست بکشید.
- در صورتی که خروجی ها بیشتر از یک صفحه باشند می توانید از دستگیره های کشویی برای مشاهده ی بقیه خروجی استفاده کنید.
- هر بخش از خروجی مانند یک نمودار با یک آیکون در سمت چپ Outline Pane در ارتباط است و هر آیکون نماینده قسمت خاصی از خرجی است.
- علاوه بر این Outline Pane کتابچه هایی را مشاهده می کنید. با کلیک بر روی هر یک از کتابچه ها در سمت چپ شیئی مربوط به آن در Display Pane قابل رویت است. یک کتاب بسته نشاندهنده ی آن است که شی مربوط به آن در خروجی فعلا قابل رویت نمی باشد



اگر می خواهید کتابچه ای را که در سمت چپ بسته است، باز کنید، روی آن دوباره کلیک کنید تا شی مربوط به آن را در سمت راست Viewer مشاهده کنید. برای مخفی کردن یک شی در خروجی، بر روی آیکن کتابچه باز آن در سمت چپ آن دو باره کلیک کنید. این عمل کتابچه را به حالت بسته و خروجی مربوط به آن را نیز مخفی می کند.

اگر می خواهید حاصل

عملیات مربوط به روند را مخفی کنید بر روی علامت منهای کنار نام آن روند کلیک کنید. کل آن بخش از خروجی مخفی شده و علامت منها به مثبت تبدیل می شود و بر عکس نمایش مجدد روند، بر روی علامت مثبت آن کلیک کنید.

ویرایش خروجی ها

ویرایش نتایج به دست آمده از تحلیل ها از این جهت ضروری است که ممکن است کاربران بخواهند از این نتایج در یک فایل متنی استفاده کنند. یکی از ویژگی های خروجی SPSS سادگی امکان ویرایش مطالب جدول تغییر مقادیر متغیراندازه و رنگ خطوط متغیر، قرار گرفتن سطر ها و ستون ها در جداول Pivot یا حذف بعضی از موارد که کاربر نمی خواهد در گزارش خود آنها را داشته باشد، وجود دارد. اگر روی هر بخش از خروجی مانند نمودار یا جدول دوبار کلیک کنید، آن بخش توسط چهارگوش خط چین شده ای احاطه شده و بلافاصله ویرایشگر آن فعال می شود. در این حالت بعضی از ویرایش های کلی را می توان روی هر بخش اعمال کرد.

اگر یک جدول خروجی را

برای ویرایش انتخاب کرده اید، برای پهن کردن یا باریک کردن یک ستون کافی است موس



را روی خطوط نگه دارید، مشکل یک پیکان دو طرفه را مشاهده کنید. سپس موس را کلیک کنید و آن را به چپ یا راست بکشید.

• در حالی خط چین اطراف

جدول وجود دارد روی هر مقدار یا عبارت دوبار کلیک کنید.

• حتی می توانید آن مقدار یا

عبارت را پاک کنید. برای این کار از دکمه Delete استفاده کنید.

• اگر قسمتی از جدول را

یکبار کلیک کنید به رنگ سیاه های لایت شده و سپس از راست کلیک می توانید آن قسمت

را Cut, Copy, Paste, Clear (Delete) کنید.

The screenshot shows the SPSS Statistics window with a table titled 'Statistics'. A context menu is open over the 'Mean' row, with the 'Copy' option selected. The table contains the following data:

		Respondent's Sex	Number of Brothers and Sisters
N	Valid	1517	1505
	Missing	0	12
Mean		1.58	
Median		2.00	
Std. Deviation		.494	
Variance		.244	
Range		1	
Percentiles	25	1.00	
	50	2.00	
	75	2.00	

Below the Statistics table is the 'Frequency' table for 'Respondent's Sex':

		Frequency	Percent
Valid	Male	636	41.9
	Female	881	58.1
	Total	1517	100.0

• گزینه Selected Table

جدول را به صورت کامل هایلایت می کند.

• در SPSS نسخه جدید شما

می توانید برای یک ستون از جدولی که انتخاب می کنید یک نمودار ستونی نقطه



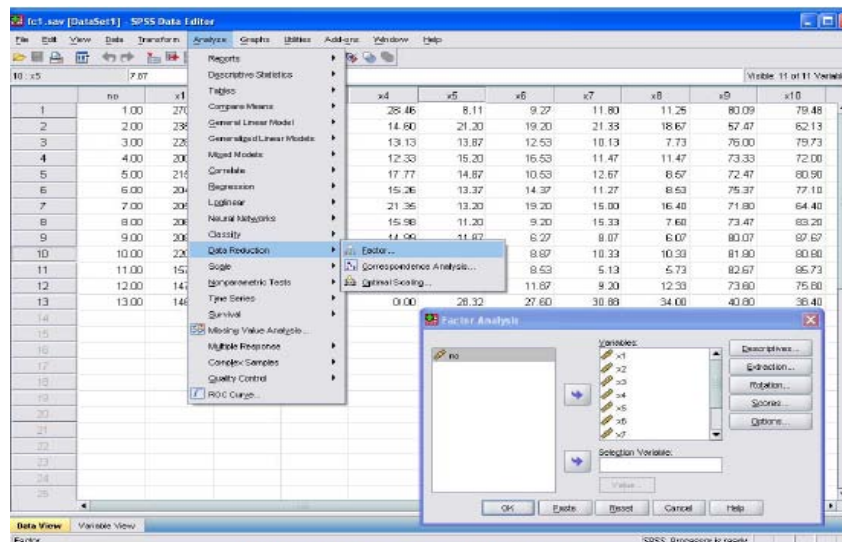
ای، خطی، ناحیه ای یا دایره ای رسم کنید. برای اینکار از گزینه Create Graph استفاده می کنیم.

اجرای تحلیل عاملی در نرم افزار SPSS

از روند زیر استفاده می شود:

Analyze>Data reduction>Factor analysis

بعد از انجام فرمان بالا، پنجره Factor Analysis ظاهر می شود (شکل ۱). متغیرها را در جعبه Variables وارد می کنند.



شکل ۱- روند اجرای فرمان Factor Analysis

محاسبه آماره های توصیفی

با فعال کردن کلید، Descriptive، امکان محاسبه آماره های زیر فراهم می شود (شکل ۲):

-در قسمت بالای این کادر، در بخش Statistics با انتخاب گزینه Univariate Descriptive آماره های تک متغیره از قبیل میانگین، انحراف معیار و تعداد مشاهده های مورد استفاده محاسبه می شود. با انتخاب گزینه Initial Solution برآوردهای اولیه ای از عامل ها محاسبه می شود. خروجی این فرمان برآورد اولیه ای از میزان اشتراک ها، مقادیر ویژه ماتریس همبستگی متغیرها، درصد کل واریانس توضیح داده شده به وسیله عامل های مشترک و نیز درصد تجمعی واریانس عام لها را ارائه می دهد.



در قسمت پایین کادر، مستطیل دیگری تحت عنوان Correlation Matrix

- با انتخاب گزینه Coefficient ماتریس ضرایب همبستگی متغیرهای انتخابی به صورت یک ماتریس مثلثی شکل محاسبه می شود که عناصر قطر اصلی آن یک است. هر چه عناصر خارج از قطر اصلی این ماتریس به یک نزدیکتر باشد، مطلوبیت روش تحلیل عاملی بیشتر خواهد بود.

- با انتخاب گزینه Significance Levels سطح معن یداری متناظر با ماتریس همبستگی به دست آمده محاسبه می شود.

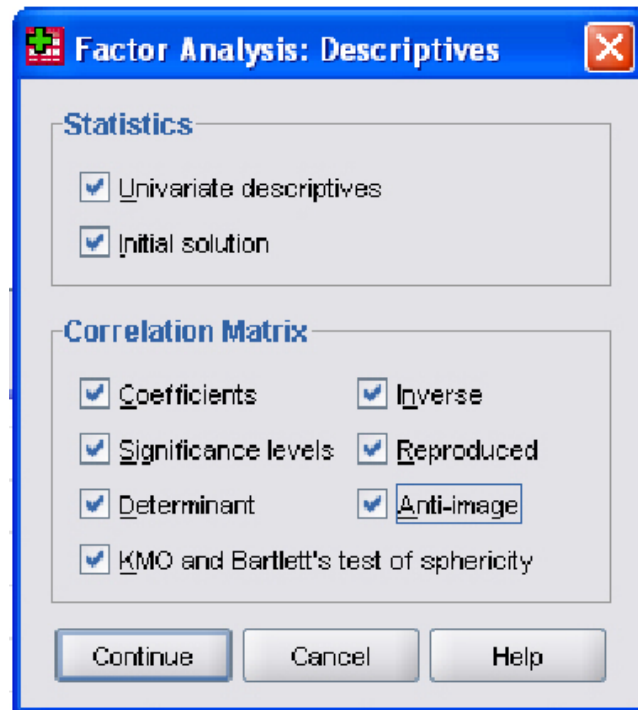
- با انتخاب گزینه Determinant دترمینان ماتریس ضرایب همبستگی محاسبه می شود. هر قدر مقدار به دست آمده کمتر باشد، انجام تحلیل عاملی معتبرتر خواهد بود.

- با انتخاب گزینه KMO and Bartlett's test of sphericity کفایت اندازه نمونه با استفاده از آزمون های KMO بارتلت تعیین می شود. به طور کلی، این گزینه شاخصی برای مقایسه مقادیر ضرایب همبستگی ساده و جزیی بر روی همه متغیرهاست. مقادیر بزرگ KMO بر رضایت بخش بودن تحلیل عاملی دلالت می کند و آزمون بارتلت نیز فرض یکه بودن ماتریس ضرایب همبستگی را آزمون می کند، به طوریکه اگر آزمون بارتلت معنادار نباشد (احتمال مربوطه بزرگتر از ۰/۵ باشد)، این امکان برای ماتریس همبستگی وجود دارد که یک ماتریس یکه باشد. این امر به معنای آن است که ماتریس مذکور برای تحلیل های بعدی مناسب نیست.

- با انتخاب گزینه Inverse معکوس ماتریس ضرایب همبستگی متغیرها محاسبه می شود.

- با انتخاب گزینه Reproduced ماتریسی محاسبه می شود که عناصر پایین قطر اصلی آن ضرایب همبستگی تبدیل یافته بین متغیرها، عناصر قطر اصلی میزان اشتراکات عامل ها و عناصر بالای قطر اصلی باقیمانده های روش تحلیل عاملی هستند. در مجموع روش تحلیل عاملی زمانی مفید است که مانده ها در این ماتریس کوچک و نزدیک به صفر باشند.

- با انتخاب گزینه Anti-image ماتریسی محاسبه می شود که عناصر آن ضرایب همبستگی جزیی با علامت مخالف و عناصر قطر اصلی این ماتریس بیانگر دقت نمون هگیری هستند. این مقادیر ورود متغیرها را به مدل تأیید می کنند



شکل ۲- کادر Descriptives فرمان Factor analysis

استخراج عوامل

برای تعیین روش استخراج عوامل روی کلید Extraction در کادر اصلی کلیک کرده تا کادری به نام Extraction Factor Analysis ظاهر شود (شکل ۳). در سمت چپ این کادر که با عنوان Method مشخص شده است، روشهای مختلف استخراج عامل ها عبارتند از:

- مؤلفه های اصلی
- حداکثر درستی
- عامل یابی محور اصلی
- حداقل مربعات غیروزنی
- عامل یابی آلفا
- عامل یابی تصویری

برای انتخاب هر کدام از روش های بالا باید بر روی علامت فلش که در سمت راست جعبه قرار دارد، کلیک کرد تا فهرست آنها نمایان شود. سپس بر روی روش مورد نظر کلیک تا به هنگام اجرای روش تحلیل عاملی از آن روش بهره گرفته شود. در این مثال روش مؤلفه های اصلی انتخاب می شود.

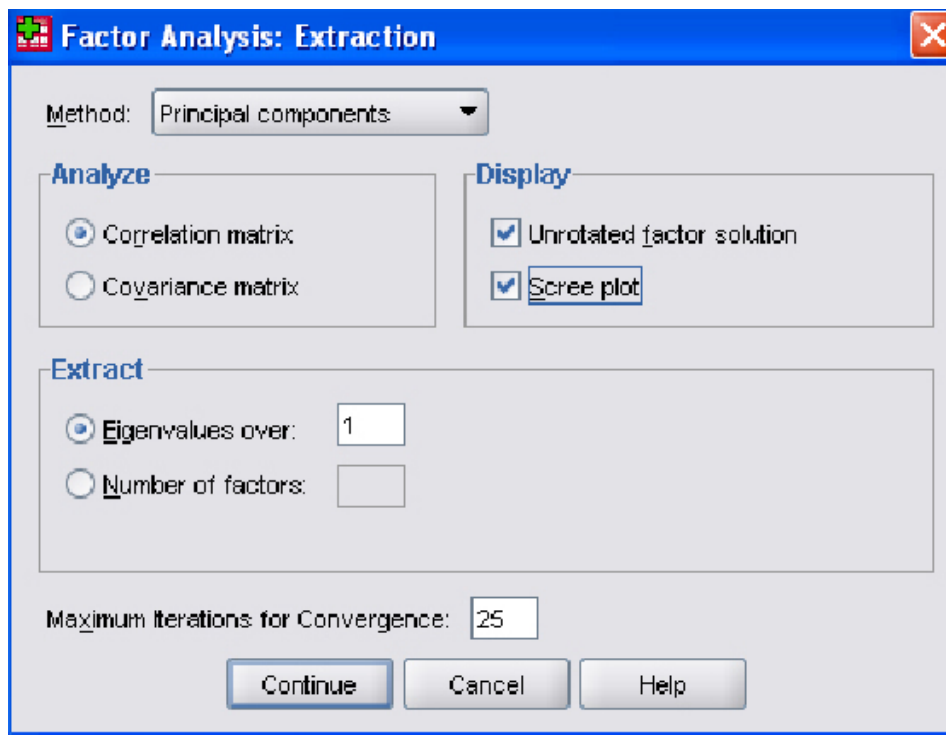


نکته: در صورتیکه هدف پژوهشگر خلاص هکردن متغیرها و دستیابی به تعداد محدودی عامل باشد، از روش مؤلفه های اصلی استفاده می شود.

در بخش Extract این کادر گزینه هایی وجود دارد که در آنها چگونگی انتخاب تعداد عامل ها مشخص می شود. با انتخاب گزینه Eigenvalue over تعداد عامل هایی که مقادیر ویژه متناظر آنها از عدد تعیین شده در جعبه مقابل آن بیشتر است، استخراج می شوند. همچنین اگر گزینه Number of Factors انتخاب شود، تعداد عامل ها را می توان به دلخواه در کادر مقابل این گزینه تعیین کرد.

در مستطیل سمت راست این کادر که با نام Display مشخص شده است، می توان ماتریس ضرایب عامل های غیردورانی را محاسبه و نموداری از مقادیر ویژه را برای عامل های مختلف رسم کرد. اگر گزینه Unrotated factor solution انتخاب شود، ضرایب عامل ها قبل از دوران محاسبه می شود. با انتخاب گزینه Scree plot نموداری رسم می شود که مقادیر ویژه عوامل انتخابی از بزرگترین تا کوچکترین مقدار را شامل می شود.

آخرین گزینه این کادر Maximum Iterations for Convergence نام دارد که با استفاده از آن می توان حداکثر تعداد دفعات تکرار برای همگرایی مدل را تعیین کرد.



شکل ۳- کادر Extraction فرمان Factor analysis

دوران عامل ها

در این مرحله به منظور بهبود روابط بین متغیرها و عامل های اولیه و اعمال تبدیلات خاص بر روی عامل ها، عمل دوران انجام می شود. با فعال کردن کلید Rotation در کادر اصلی، کادر Factor Analysis: Rotation ظاهر میشود که شامل بخش های زیر است (شکل ۴):

بخش Method روش های مختلف دوران را شامل می شود. به طور کلی روش های دوران به دو دسته متعامد و غیرمتعامد (همبسته) تقسیم می شوند. اگر پژوهشگر بخواهد تعداد زیادی متغیر مورد بررسی را به یک مجموعه کوچکتر از متغیرهای غیرمرتبط با هم تقلیل دهد، روش متعامد مناسب تر است. اما اگر هدف اصلی تحلیل عاملی ب هدست آوردن چند عامل باشد که از نظر تئوریکی معنی دار باشد، روش غیرمتعامد مناسب تر است.

با انتخاب فرمان None هیچ نوع دورانی بر روی ماتریس ضرایب اعمال نمی شود.

فرمان Varimax از جمله متداول ترین روش های دوران متعامد است که استقلال میان عامل های استخراجی را حفظ می کند. این روش متغیرهای دارای بار عاملی بزرگتر را به کمترین تعداد تقلیل می



دهد. این روش، جمع واریانس بارها در ماتریس عاملی را بیشترین مقدار می کند، به همین دلیل آن را واریماکس گویند. هنگامی از این روش استفاده می شود که هدف به دست آوردن عامل هایی است که دارای بار زیادی بر روی برخی از متغیرها و بار کم بر روی متغیرهای دیگر باشد.

در این روش تأکید بر ساده کردن ستون های ماتریس عاملی است، یعنی حداکثر امکان ساده کردن تا آنجایی حاصل می شود که بر روی یک ستون خاص ماتریس، فقط مقادیر (بارهای عاملی) صفر و یک قرار بگیرد. از این رو، مجموع تغییرات ایجاد شده در بارهای عاملی به حداکثر می رسد. در این حالت تفسیر عامل ها ساده می شود.

فرمان Quartimax هنگامی مورد استفاده قرار می گیرد که مدل تحلیل عاملی به تعداد عوامل کمتری وابسته باشد. این روش نیز از جمله روش های دوران متعامد است. هدف اصلی روش Quartimax ساده کردن سطرهای ماتریس عاملی است. یعنی اگر بار عاملی متغیر در یک عامل زیاد باشد، تا آنجاییکه ممکن است از تعلق متغیر به عامل های دیگر کاسته می شود. به عبارت دیگر بار عاملی اش در دیگر عامل ها پایین می آید. مشکل عمده این روش این است که در چرخش عامل ها تمایل به ایجاد یک عامل کلی و بزرگ وجود دارد.

فرمان Equamax نیز یکی از روش های متعامد است که در برگیرنده اهداف هر دو روش فوق است. بدین معنی که هم سطرها و هم ستون ها ساده می شود.

در مجموع، روش های متعامد فوق مقادیر نسبتاً بزرگ) از نظر قدر مطلق (یا صفر به ستون های ماتریس ضرایب عامل ها اختصاص می دهند. در چنین شرایطی عوامل به دست آمده یا با متغیرهای انتخابی وابستگی زیادی داشته یا از آنها کاملاً مستقل خواهند بود.

فرمان Direct oblmin از جمله روش های دوران غیرمتعامد یا مورب است. ویژگی این روش ساده تر بودن تفسیر عامل هاست، در حالیکه عامل های به دست آمده از آن در این حالت مستقل نخواهند بود. با انتخاب این روش، مقدار دلتا در جعبه Delta مشخص می شود، به طوریکه در حالت دلتا برابر با صفر وابستگی بین عام لهای تبدیل یافته به حداکثر مقدار خود خواهد رسید، در حالیکه به ازای مقادیر منفی با قدر مطلق بزرگ از وابستگی عامل ها کاسته شده و نتایج به دست آمده به نتایج روش های متعامد نزدیک می شود.



برای دستیابی به یک ساختار ساده، اگر راه حل متعامد انتخاب شود، روش Varimax بهترین روش است. اما در صورت انتخاب راه حل متمایل، چرخش Direct oblmin بهترین روش است. در این مثال روش Varimax انتخاب شد.

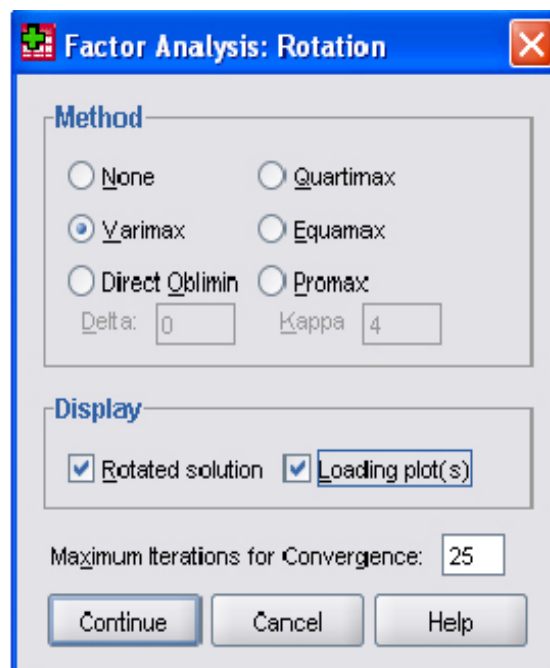
نکته: فیلد (۲۰۰۰) معتقد است که انتخاب نوع چرخش عاملی به این بستگی دارد که آیا دلیل نظری از یک طرف بر فرض همبستگی یا استقلال عاملها از همدیگر و از طرف دیگر برای چگونگی خوشه بندی متغیرها بر روی عاملها قبل از چرخش وجود دارد یا خیر؟ یک روش قابل قبول آن است که ابتدا هر دو روش اجرا شود. سپس نوع مناسب آن انتخاب شود.

در بخش Display گزینه های زیر وجود دارد:

- در صورت انتخاب یکی از روش های دوران، گزینه Rotated Solution قابل انتخاب بوده و با انتخاب آن ماتریس ضرایب عامل های دوران یافته، ماتریس تبدیل عوامل، ماتریس الگو و ماتریس همبستگی بین برآوردهای دورا نیافته عامل های مشترک محاسبه خواهد شد.

- با انتخاب گزینه Loading plot(s) نمودار سه بعدی از ضرایب متغیرهای موجود در ماتریس الگو رسم خواهد شد.

در پایین این کادر گزینه Maximum Iterations for Convergence وجود دارد که حداکثر تعداد دفعات تکرار فرآیند رابرای همگرایی روش تحلیل عاملی مشخص می کند.



شکل ۴- کادر Rotation فرمان Factor analysis



محاسبه نمره های عامل ها

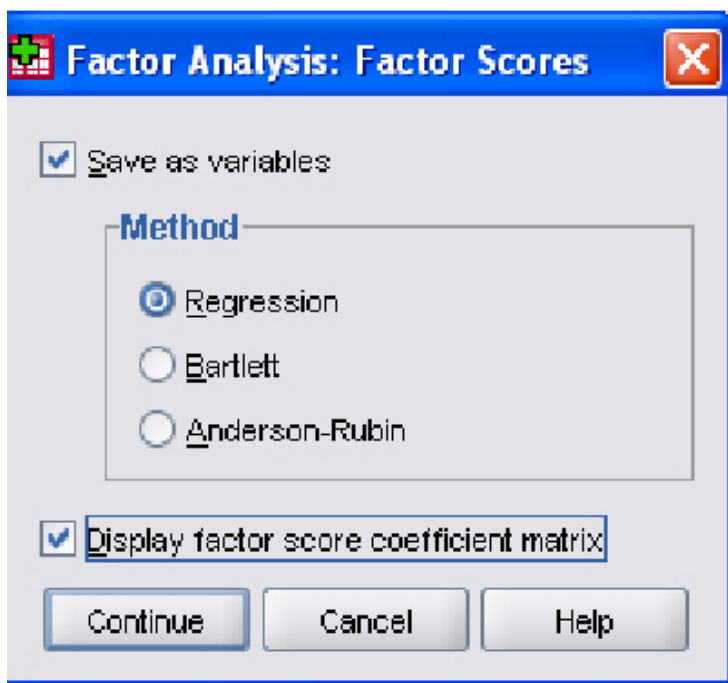
برای محاسبه نمره های عاملی بر روی کلید Scores در کادر اصلی کلیک تا کادری به نام Factor Scores Analysis: Factor ظاهر شود (شکل ۵). اولین گزینه این کادر Save as Variables است که با انتخاب آن می توان نمره های عاملی به دست آمده برای مشاهده ها را به صورت متغیرهای جدید در کنار داده های اولیه ذخیره کرد.

روش برآورد نمره های عاملی در بخش Method تعیین می شود. در برنامه SPSS سه روش منظور شده است. در روش رگرسیون بدون توجه به نوع دوران انتخابی، تعداد عامل های وابسته برآورد می شوند. در این روش، حتی زمانی که عامل ها مستقل فرض می شوند، می توانند با همدیگر همبستگی داشته باشند. شایان ذکر است که اگر روش خاصی برای برآورد انتخاب نشود، روش رگرسیونی بصورت خودکار اعمال می شود.

در روش Bartlett برای برآورد نمره های عاملی از روش حداقل مربعات وزنی استفاده می شود. روش Anderson-Rubin برای هر عامل نمره هایی با انحراف معیار یک محاسبه می کند و عامل ها مستقل از یکدیگرند. در این روش، به منظور برآورد ضرایب، روش حداقل مربعات معمولی به کار گرفته میشود. لازم به ذکر است که در صورت اعمال روش مؤلفه های اصلی، برای استخراج عوامل، نمره های عاملی حاصل از انتخاب هر سه روش یکسان خواهند بود.

آخرین گزینه موجود در این کادر Display Factor Score Coefficient Matrix است که انتخاب آن برآورد ماتریس ضرایب عامل های به دست آمده را نشان می دهد.

در ادامه پس از انتخاب گزینه های مورد نیاز بر روی کلید Continue کلیک تا بار دیگر کادر اصلی ظاهر شود.



شکل ۵- کادر Factor Scores فرمان Factor analysis

با کلیک کردن بر روی گزینه Options کادر اصلی شکل (۵-۶) ظاهر می شود که در بخش Missing Values آن چگونگی برخورد با داده های گمشده مشخص می شود. در این بخش سه گزینه زیر وجود دارد:

-گزینه Exclude cases listwise برای حذف مشاهده هایی به کار می رود که در یکی از متغیرهای خود دارای داده گمشده هستند.

-گزینه Exclude cases pairwise برای حذف مشاهده هایی به کار می رود که یک یا هر دو آنها دارای داده گمشده اند. این فرمان بر روی عملیاتی قابل اجراست که از دو متغیر به طور همزمان استفاده می شود.

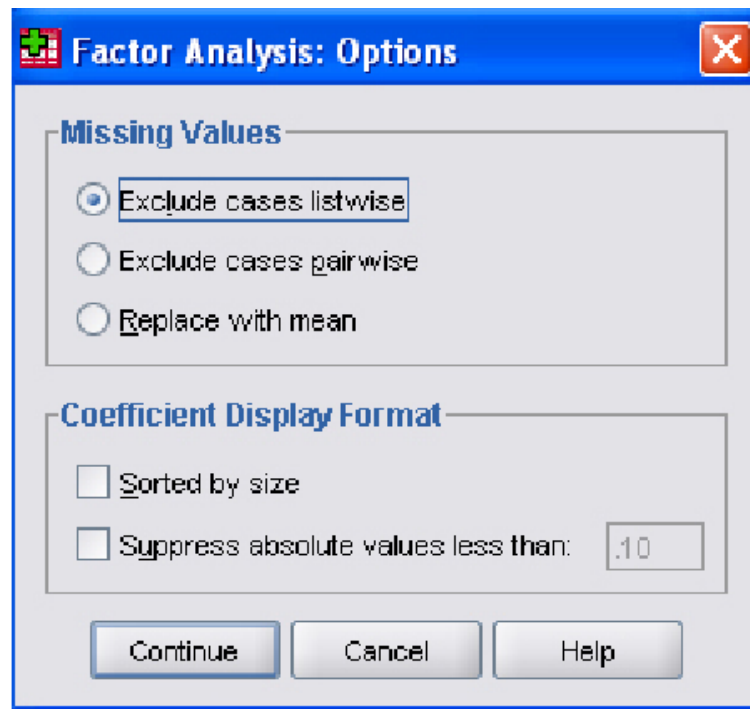
-گزینه Replace with Mean نه تنها مشاهده های دارای داده گمشده را حذف نم یکنند، بلکه داده های گمشده را با میانگین دیگر مشاهده ها جایگزین می کند.

در بخش Coefficient Display Format دو گزینه زیر وجود دارد:



-گزینه **Sorted by size** متغیرهایی که ضرایب عامل یا همبستگی بالایی دارند، در یک گروه جای می دهد و آنها را به ترتیب از بزرگ به کوچک مرتب می کند.

-گزینه **Suppress absolute values less than** ضرایبی را که قدر مطلق آنها از مقدار تعیین شده در جعبه مقابل آن کوچکتر است، از ماتریس ضرایب حذف می کند. با توجه به آنکه ضرایب عاملی بین صفر و یک هستند، عدد مندرج در جعبه مذکور باید در همین دامنه قرار داشته باشد.



شکل ۶- کادر Options فرمان Factor analysis



منابع و مأخذ:

۱. اسماعیلیان، زهرا (۱۳۸۹)، نقش مدیریت واحد در بحرانهای طبیعی شهری مطالعه موردی شهر اصفهان، رساله دکترا، دانشگاه تبریز.
۲. - رهنما، محمد رحیم (۱۳۸۸)، برنامه ریزی مناطق مرکزی شهرها (اصول، مبانی، تئوریه‌ها، تجربیات و تکنیکها)، چاپ اول انتشارات دانشگاه فردوسی، مشهد.
۳. - قدسیپور، سید حسن (۱۳۷۹)، فرآیند تحلیل سلسله مراتبی (AHP)، مرکز نشر دانشگاه صنعتی امیر کبیر (پلی تکنیک تهران).
۴. - معصومزاده، سید محسن، ترايزاده (۱۳۸۳)، «رتبه‌بندی تولیدات کشور»، فصلنامه پژوهشنامه بازرگانی، سال هشتم، شماره ۳۰.
۵. بهفروز، فاطمه (۱۳۷۴) زمینه های غالب در جغرافیای انسانی، انتشارات دانشگاه تهران.
۶. پرهیزگار، اکبر (۱۳۷۷) مدل های جاذبه و دسترسی در برنامه ریزی شهری و ناحیه ای، مجله مدرس، شماره هشتم، دانشگاه تربیت مدرس، تهران.
۷. پورمحمدی، محمد رضا (۱۳۸۲)، برنامه ریزی مسکن، چاپ دوم، تهران: انتشارات سمت.
۸. ترابی سیما، سعید جهانبخش (۱۳۸۳). تعیین متغیر های زمینه ای در طبقه بندی اقلیمی ایران معرفی و کاربری روش تحلیل عاملی و تجزیه مولفه های اصلی در تحلیل مطالعات جغرافیایی و اقلیم شناسی، فصلنامه تحقیقات جغرافیایی، شماره ۷۲.
۹. توفیق، فیروز (۱۳۶۹). مجموعه مباحث و روش های شهرسازی، مسکن، تهران: انتشارات مرکز مطالعات و تحقیقات مسکن و شهرسازی، وزات مسکن و شهرسازی.
۱۰. جوادی اشکلک، ا. (۱۳۷۶)، بررسی امکان پذیری توسعه فضایی مسکن در ناحیه مرکزی شهرها نمونه موردی شهر تبریز، پایان نامه کارشناسی ارشد رشته شهرسازی، دانشگاه شهید بهشتی.
۱۱. حسین زاده دلیر، کریم (۱۳۸۰)، برنامه ریزی ناحیه ای، تهران، انتشارات سمت، چاپ اول.
۱۲. حکمت نیا حسن، میر نجف موسوی (۱۳۸۵). کاربرد مدل در جغرافیا با تاکید بر برنامه ریزی شهری و ناحیه ای، یزد، انتشارات علم نوین، چاپ اول.
۱۳. حمزه مصطفوی، غلامحسین (۱۳۷۴). طراحی مسکن براساس نیاز، مجموعه مقالات دومین سمینار سیاست های توسعه مسکن در ایران، جلد اول، وزارت مسکن و شهرسازی.
۱۴. خلیل عراقی، منصور (۱۳۶۷). شناخت عوامل مؤثر در گسترش بی رویه ی شهر تهران، دانشگاه تهران.
۱۵. رفیعی، مینو، سنجش توسعه صنعتی مناطق کشور، مرکز تحقیقات شهرسازی و معماری ایران، ۱۳۷۰.



۱۶. رضانی ، محمد ابراهیم (۱۳۷۶). سنجش توسعه شمال خراسان بر پایه مدلهای برنامه ریزی ناحیه ای ، پایان نامه کارشناسی ارشد جغرافیای شهری ، دانشگاه تهران .
۱۷. زبردست، اسفندیار (۱۳۸۲). ارزیابی روش های تعیین سلسله مراتب و سطح بندی سکونتگاه ها در رویکرد عملکردهای شهری در توسعه روستایی، نشریه هنرهای زیبای دانشگاه تهران، شماره ۱۳.
۱۸. زیاری ، کرامت الله، مهدی دهقان (۱۳۸۲). بررسی وضعیت مسکن مسکن و برنامه ریزی آن در شهر یزد، نشریه صفه، شماره ۳۶، دانشکده معماری و شهرسازی دانشگاه شهید بهشتی، تهران.
۱۹. زیاری کرامت الله (۱۳۷۸). اصول و روش های برنامه ریزی منطقه ای ، یزد ، انتشارات دانشگاه یزد، چاپ سوم .
۲۰. سازمان مدیریت و برنامه ریزی استان فارس (۱۳۷۰) اوضاع اقتصادی و اجتماعی استان فارس مسکن مرکز انفرماتیک و مطالعات توسعه جنوب، نشریه ۷-۱۱. شیراز.
۲۱. شاهنوشی ، ناصر ، زهرا گلریز صائی، حمید رضا باقری (۱۳۸۵). تعیین سطح توسعه یافتگی شهر مشهد « مجموعه مقالات کنفرانس برنامه ریزی و مدیریت شهری مشهد » .
۲۲. صرافی مظفر (۱۳۷۷) مبانی برنامه ریزی توسعه منطقه ای ، انتشارات سازمان برنامه و بودجه ، چاپ اول، تهران
۲۳. طالبی هوشنگ ، زنگی آبادی علی (۱۳۸۰) . تحلیل شاخص ها و تعیین عوامل موثر در توسعه انسانی شهرهای بزرگ کشور ، فصلنامه تحقیقات جغرافیایی.
۲۴. عظیمی، ناصر، پویش شهرنشینی و مبانی نظام شهری، نشر نیکا، تهران، ۱۳۸۱.
۲۵. عماد زاده مصطفی، دلالی اصفهانی رحیم، صابر داریوش ، رتبه بندی شهرستان های استان.
۲۶. قدیر معصوم مجتبی ، کیومرث حبیبی (۱۳۸۳). سنجش و تحلیل توسعه یافتگی شهر ها و شهرستان های استان گلستان ، نامه های علوم اجتماعی ، شماره ۲۳ .
۲۷. قدیری معصوم (۱۳۷۹) سیری در مفاهیم و ابعاد مختلف توسعه ، مجله دانشکده ادبیات و علوم انسانی دانشگاه تهران ، زمستان . موثقی احمد (۱۳۸۳). توسعه ؛ سیر تحول مفهومی و نظری ، مجله دانشکده حقوق و علوم سیاسی ، شماره ۶۳ .
۲۸. کلانتری خلیل (۱۳۸۰). برنامه ریزی و توسعه منطقه ای ، تهران ، انتشارات خوشبین و انوار دانش ، چاپ اول .
۲۹. لی، کولین، مدلهای در برنامه ریزی شهری، مقدمه ای بر کاربرد مدلهای کمی در برنامه ریزی، ترجمه مصطفی عباس زادگان، انتشارات دانشگاه علم و صنعت، تهران، ۱۳۶۶.
۳۰. مطیعی لنگرودی حسن (۱۳۸۲). برنامه ریزی روستایی با تاکید بر ایران ، چاپ اول ، زمستان ۱۳۸۲.
۳۱. موسوی میر نجف ، حسن حکمت نیا (۱۳۸۴) . تحلیل عاملی و تلفیق شاخص ها در تعیین عوامل موثر بر توسعه انسانی نواحی ایران ، جغرافیا و توسعه .



۳۲. میاندشتی مهدی جدیدی (۱۳۸۳). توزیع متعادل منابع مالی به روش سطح بندی توسعه مناطق، فصلنامه پژوهش های اقتصادی، شماره ۱۱ و ۱۲.
۳۳. هومن حیدر علی (۱۳۸۴). تحلیل عاملی دشواری ها و تگناهای آن، مجله روان شناسی و علوم تربیتی، سال سی پنجم شماره ۲.

۳۴. Alperovich, G.A, The Size Distribution of Cities: On the Empirical Validity of the Rank-Size Rule, Journal of Urban Economics, Vol ۱۶, ۱۹۸۴.
۳۵. Carter, Harold, The study of Urban Geography, Routledge, NewYork, ۱۹۸۱.
۳۶. Chisholm, Michael, Human Geography: Evolution or Revolution, Richard Clay, Ltd, Bunyay, Suffolk, ۱۹۷۹.
۳۷. Clark, David, Urban World/Global City, Routledge, London, ۲۰۰۰
۳۸. Delgado, A.P, Godinho, I.M, The Evolution of City Size Distribution in Portugal: ۱۸۶۴-۲۰۰۱, Paper Presentd at the World Conference of International Regional Science Association, April ۲۰۰۴, Port Elizabet, South Africa, ۲۰۰۴.
۳۹. Guerin-Pace, F, Rank-Size Distributions and the Process of Urban Growth, Journal of Urban Studies, Vol ۳۲, No ۳۰, ۱۹۹۵
۴۰. Haggett, Petter, Locational Analysis vin Human Geography, London, ۱۹۶۸.
۴۱. Nishiyama. Y,et.at, Estimation and Testing for Rank Size Rule Regression under Pareto Distribution, School of Economics, Kyoto University, ۲۰۰۵.
۴۲. Wheeler, James O, Muller, Peter O, Economic Geography, John Wiley & Sons, Inc, Canada, ۱۹۸۶.
- 43.- Al-Sushi Al-Harbi، Kamal M.; (2001), Application of the AHP in project management, International Journal of Project Management, vol



۴۴. Solnes , Julius; (2003),
Environmental quality indexing of large industrial development alternatives using
AHP , Environmental Impact Assessment Review , vol 23.
۴۵. Pakdin Amiri , Morteza;
(2010), Project selection for oil-fields development by using the AHP and fuzzy
TOPSIS methods , Expert Systems with Applications , vol 37.