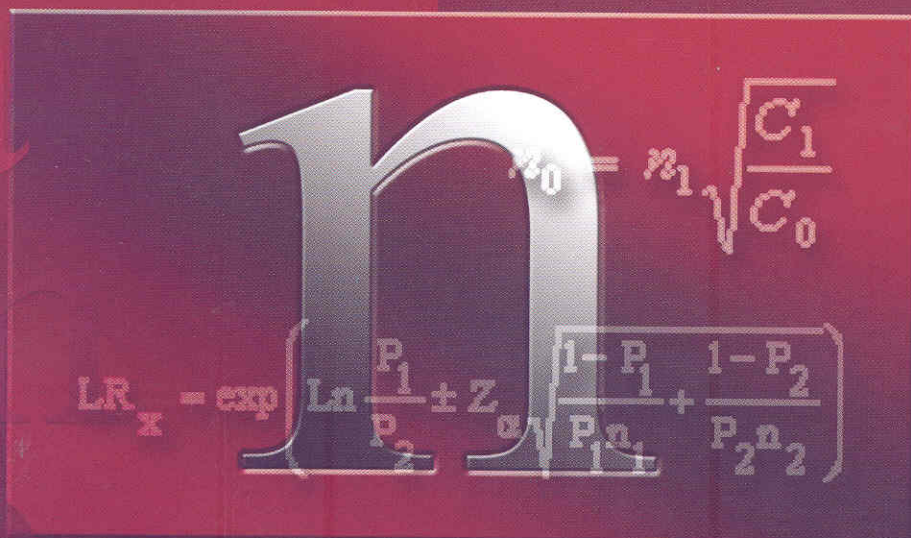


نمونه‌گیری و محاسبه حجم نمونه

در مطالعات علوم پزشکی



مؤلفین:

نوشین شاهقلی
نادر صدیق
دکتر مازیار مرادی
مریم هاشم‌نژاد

دکتر علی چهرئی
محسن صابری
هما محمدصادقی
مهدی منتظر

ویراستار:

دکتر علی چهرئی

کمیته پژوهشی دانشجویی پزشکی دانشگاه علوم پزشکی ایران

& Sample Size
Estimation

نمونه گیری و محاسبه حجم نمونه در مطالعات علوم پزشکی / نویسندگان علی چهرئی ... [و دیگران] ؛
[تهیه کننده] کمیته پژوهشی دانشجویی پزشکی
دانشگاه علوم پزشکی و خدمات بهداشتی درمانی
ایران؛ ویراستار نهایی علی چهرئی.- تهران:
سارا، ۱۳۸۱.
۱۰۰ ص.: مصور، جدول، نمودار.

ISBN 964-93897-5-8

فهرست نویسی بر اساس اطلاعات فیپا.
کتابنامه: ص. ۹۹ - ۱۰۰.

۱. پزشکی -- تحقیق -- روشهای آماری.
۲. آمارگیری نمونه ای. الف. چهرئی، علی، ۱۳۵۶ -
ب. دانشگاه علوم پزشکی و خدمات بهداشتی درمانی
ایران. کمیته پژوهشی دانشجویی.

۶۱۰ / ۷۲

۹ چ ۹ / ر ۸۵۳

۸۱ ۱۸۵۶۱-م

کتابخانه ملی ایران

ناشر: انتشارات سارا

تلفن: ۵-۸۸۹۴۲۵۳

تیراژ: ۱۰۰۰ عدد

نوبت چاپ: اول ۱۳۸۱

اجرا: ایما نقشبینه

به آنانکه دستانمان را گرفتند تا راه رفتن بیاموزیم

و راه زندگی را نشانمان دادند

به پدران و مادرانمان

و به آنانکه دستاوردهای امروز و فردا را

با تلاشهای بی شائبه شان پایه ریختند

به پیشکسوتان پژوهش دانشجویی کشور

۱	(۱) مفاهیم پایه نمونه‌گیری
۲	سرشماری
۳	نمونه‌گیری
۵	جامعه
۵	نمونه
۵	اهداف نمونه‌گیری
۷	فواید نمونه‌گیری
۷	قالب نمونه‌گیری
۸	میزان پاسخ‌دهی
۱۰	اعتبار و پایایی
۱۱	خطا
۱۲	طراحی نمونه‌گیری
۱۱	(۲) انواع نمونه‌گیری
۱۲	روشهای نمونه‌گیری احتمالی
۱۲	نمونه‌گیری تصادفی ساده
۱۵	نمونه‌گیری تصادفی منظم
۱۶	نمونه‌گیری تصادفی طبقه‌بندی شده
۱۹	نمونه‌گیری تصادفی خوشه‌ای
۲۱	روشهای نمونه‌گیری غیر احتمالی
۲۱	نمونه‌گیری آسان
۲۲	نمونه‌گیری هدفدار
۲۵	روش نمونه‌گیری چند مرحله‌ای
۲۷	(۳) تورش
۲۸	انواع تورش
۲۹	انواع تورش انتخاب
۳۷	(۴) توزیع احتمالات
۳۸	توزیع نرمال یا توزیع گاوسی
۴۰	روش استفاده از جدول سطح زیر منحنی‌های نرمال استاندارد
۴۶	تئوری حد مرکزی

مقدمه

Introduction

با توجه به اینکه در تحقیقات علوم پزشکی در اکثر موارد امکان سرشماری بر روی کلیه بیماران موجود نمی‌باشد لذا نمونه‌گیری از اهمیت به سزایی در این تحقیقات برخوردار است. یک نمونه مناسب اولاً به لحاظ شکل باید مشابه جامعه هدف باشد و ثانیاً به لحاظ اندازه نیز متناسب باشد. با توجه به این مفهوم در فصول اول تا سوم این کتاب به بررسی کلیات نمونه‌گیری، انواع آن و سوگیری‌هایی که در جریان انتخاب نمونه اتفاق می‌افتد پرداخته‌ایم تا خصوصیت اول ذکرشده برای یک نمونه مناسب شرح داده شود و در فصول چهارم و پنجم به بیان چگونگی انتخاب یک نمونه با اندازه مناسب پرداخته شده است. لازم به ذکر است که سعی نویسندگان این کتاب تکیه بر عناوین ذکر شده با دیدگاه یک محقق رشته‌های پزشکی می‌باشد و سعی شده است از بحث‌های آماری صرف پرهیز گردد و مخصوصاً در قسمت محاسبه حجم نمونه تأکید بر بیان کاربردی روابط بوده است. در پایان امیدواریم که خوانندگان محترم با مطالعه این کتاب بتوانند در انتخاب نمونه‌های مناسب در مطالعات علوم پزشکی، تحقیقات کاربردی را به انجام برسانند.

تیر ۱۳۸۱

شورای نویسندگان

خواص عمومی توزیع نمونه‌گیری	۴۸
انواع خطا در طرح‌های تحقیقاتی	۴۹
(۵) روابط محاسبه حجم نمونه	۵۱
برآورد	۵۲
برآورد یک میانگین	۵۲
برآورد یک نسبت	۵۴
مقایسه	۵۶
مقایسه یک میانگین با عدد ثابت	۵۶
مقایسه دو میانگین	۵۸
مقایسه یک نسبت با یک عدد ثابت	۵۹
مقایسه دو نسبت	۶۰
محاسبه حجم نمونه بر مبنای نوع مطالعه	۶۵
مطالعات مشاهده‌ای - تحلیلی	۶۵
مطالعات همبستگی	۶۹
مطالعات مقایسه تابع بقا	۷۲
مطالعات ارزیابی تست‌های تشخیصی	۷۳
مطالعات کارآزمایی بالینی (Altman's Nomogram)	۷۶
محاسبه حجم نمونه با استفاده از نرم‌افزار EPI-Info	۷۷
(۶) حل تمرین	۸۳
(۷) پیوستها	۹۳
منابع	۹۹

مفاهیم پایه نمونه گیری

Basic Concepts of Sampling

- ✓ سرشماری و نمونه گیری
- ✓ جامعه هدف، جامعه آماری و نمونه
- ✓ فواید نمونه گیری
- ✓ قالب نمونه گیری
- ✓ نمونه معرف
- ✓ اعتبار و پایایی

فصل

۱

هر تحقیق علمی، با سؤالی درباره واقعیت آغاز می‌شود بنابراین پس از تجزیه و تحلیل منطقی سؤال آغازین، محقق باید به صورت برنامه‌ریزی شده به واقعیت رجوع کند و با مشاهده سیستماتیک و کنترل شده، اطلاعات لازم را از آن برگرداند، تحلیل کند و به سؤالش جواب دهد. واقعیت مقتضی مشاهده در هر تحقیقی را، تا حدود زیادی سؤال تحقیق معین می‌کند. با وجود این پس از تشخیص و تعیین واقعیت مقتضی باید پرسید آیا کل واقعیت باید مشاهده شود یا مشاهده جزئی از آن کافی است؟ اگر مشاهده جزئی کافی است این جزء چگونه باید انتخاب شود تا معرف آن کل باشد؟

در پاسخ به این سؤال است که مسأله نمونه‌گیری مطرح می‌شود. بنابراین نمونه‌گیری مرحله‌ای از مراحل به هم پیوسته تحقیق علمی و یکی از عناصر اساسی روش‌شناسی (Methodology) علم جدید است. نمونه‌گیری به عنوان یک ابزار تحقیقی سبب تسهیل کار تحقیق می‌شود. نمونه‌گیری این امکان را برای محقق فراهم می‌کند تا با صرف امکانات کمتر به نتایج مورد نظرش دست یابد. در واقع سیر تکوینی و گسترش نمونه‌گیری، سیر تکاملی و تزايدی ایفای نقش کرده است؛ این نقش در آمارگیری‌های وسیع به خوبی مشهود است. این روزها نمونه‌گیری جای خود را در آمارگیری‌ها کاملاً باز کرده است، تا آنجا که می‌رود سرشماری را هم از درون تهی کند. در تحقیقات آکادمیک هم مطالعه نمونه‌ای دامنه وسیعی یافته است تا آنجا که می‌توان آن را در کنار اندازه‌گیری، یکی از ارکان روش‌شناسی تحقیق به شمار آورد. در واقع استفاده از نمونه‌گیری دیگر محدود به مطالعات پیمایشی (Survey) نیست و جای آن در انواع دیگر تحقیق هم باز شده است.

در ذیل روشهایی که می‌تواند جهت انتخاب نمونه به کار رود، بررسی می‌شود:

سرشماری (Census)

در این روش تمام افراد جمعیت مورد مطالعه تحت بررسی قرار می‌گیرند. این روش معمولاً زمانی انتخاب می‌شود که تعداد افراد جامعه مورد مطالعه کم بوده و از نظر وقت و هزینه مشکلی در پیش نباشد و با اعتبار علمی مطالعه با انتخاب نمونه کوچک تفاوت چندانی ندارد.

نکته: هر سرشماری، ضرورت ندارد که الزاماً موفقیت‌آمیز بوده و شامل همه عناصر جامعه باشد به عنوان مثال معلوم شده است که در سرشماری‌های مربوط به جوامع آمریکا و کانادا هر چند سعی زیادی برای در بر گرفتن همه افراد می‌شود ولی تعداد قابل توجهی از افراد جامعه وارد نمونه نمی‌شوند.

نمونه‌گیری (Sampling)

در بعضی از تحقیقات تعداد افراد، محلها یا مراکز مورد مطالعه کم هستند و می‌توان همه آنها را در مطالعه گنجاند. با این وجود، اغلب مسائل تحقیقاتی عده کثیری از افراد، محلها یا مراکز را در برمی‌گیرند و به این ترتیب گنجاندن فقط قسمتی از آنها در امر تحقیق الزامی می‌گردد بنابراین در این شرایط باید نمونه‌ای از کل جمعیت انتخاب کرد که معرف جمعیت مورد بررسی باشد یعنی کلیه خصوصیات مهم جامعه را در برداشته باشد. (به مضدق مثل معروف مشت نمونه خروار است)

در واقع نمونه‌گیری، یک پروسه یا تکنیک است که نمونه‌های مناسب را از یک جمعیت انتخاب می‌کند و آنگاه با بررسی این نمونه‌ها، نتیجه کلی جمعیت تعمیم داده می‌شود.

در اینجا لازم است که برخی از واژه‌ها که در ارتباط با نمونه‌گیری می‌باشند توضیح داده شود:

۱- جامعه (Population)

مجموعه‌ای از واحدهاست که در چیز یا چیزهایی مشترک باشند. این واحدها ممکن است افراد انسانی، واحدهای سازمانی مثل بیمارستانها، مراکز بهداشتی- درمانی، خانه‌های بهداشت و غیره باشند و یا نمونه‌هایی از مواد مثل نمونه‌های شیر خوراکی، آب آشامیدنی، مواد غذایی و غیره باشند و یا نمونه‌های آزمایشگاهی مانند نمونه‌های خون، مدفوع، و غیره را شامل شوند. با وارد کردن برخی از ویژگیها می‌توان جمعیت را به دو یا چند زیر جمعیت یا طبقه تقسیم کرد.

ویژگیهای جمعیت، برای مشخص کردن دو نوع جمعیت به کار می‌روند:

۱- ویژگیهای بالینی و دموگرافیک: جمعیت هدف را تعیین می‌کنند مثلاً تمام نوجوانان مبتلا به آسم.

۲- ویژگیهای جغرافیایی و زمانی: جمعیت در دسترس را تعیین می‌کنند مثلاً نوجوانان مبتلا به آسم در سال ۱۳۸۰ در شهر تهران.

دقت در انتخاب جمعیت مناسب نیل به هدفهای پروژه را تسهیل می‌کند. در انتخاب جمعیت بایستی نکات ذیل مد نظر باشد:

- آیا می‌توان از جمعیت مورد نظر اطلاعات مورد نیاز را کسب کرد؟
- آیا امکان همکاری این جمعیت وجود دارد یا احتمالاً به علت بیش از اندازه مورد پژوهش قرار گرفتن

در گذشته، مقاومت نشان می‌دهد؟

- اگر هدف یک پیگیری طولانی است، آیا جمعیت چنان متحرک است که برقراری تماس با موارد مطالعه

دچار اشکال احتمالی نشود؟

✓ نکته: مطالعه بر روی برخی از جمعیت‌ها از اعتبار تعمیم‌پذیری یافته‌ها می‌کاهد؛ مثلاً جمعیت‌های داوطلب،

جمعیت‌های بیمارستانی و جمعیت‌هایی که فرد در آن بیش از یکبار به حساب می‌آید از این نوع است.

در تحقیقات علوم پزشکی می‌توان جامعه مورد بررسی را از چند دیدگاه ذیل نگریست:

۱-۱) جمعیت آماری (Statistical or Sampled Population)

جمعیت آماری را گاهی جمعیت اصلی، جامعه آماری، جمعیت بزرگتر و جمعیت کل نیز می‌نامند. منظور از جمعیت آماری جامعه‌ای است که موضوع مورد تحقیق در ارتباط با آن انجام می‌شود. به عنوان مثال در بررسی مسائل دانشجویان کشور، جمعیت آماری مورد نظر کل دانشجویان کشور اعم از دولتی یا غیر دولتی می‌باشد. مجموعه واحدهایی که حداقل در یک صفت مشترکند، یک جامعه آماری را مشخص می‌کنند و این صفت مشترک در مثال فوق دانشجویان بودن می‌باشد. پس لازم است که محقق از قبل ویژگیها یا صفات مورد نظر خود را در جامعه آماری مشخص نماید.

گفته شد که جامعه آماری حداقل دارای یک صفت مشترک است. یک جامعه آماری ممکن است دهها ویژگی داشته باشد. به عنوان مثال در بررسی مسائل دانشجویان می‌توان صفات دیگری را هم در نظر گرفت و جامعه را به دانشجویان دختر غیر بومی محدود کرد. بنابراین صفات به ترتیب عبارت از دانشجوی بودن، از جنس مؤنث بودن و بالاخره غیر بومی بودن در محل تحصیل خود می‌باشند. پس تعیین ویژگیهای معیار برای جامعه آماری باعث محدودتر کردن آن از یک سو و سپس تسهیل امر نمونه‌گیری از دیگر

سو می‌باشد. همچنین در هر جمعیت آماری لازم است محقق واحد آماری خود را روشن سازد. منظور از واحد آماری، افراد یا اشیایی است که اطلاعات مورد نیاز از طریق آنها یا درباره آنها جمع‌آوری می‌شود. جمعیت آماری، آمار مورد نیاز محقق را در جهت تأمین هدفهای تحقیق فراهم می‌آورد و واحدهای آماری وسیله تأمین این نیازها هستند.

۱-۲) جامعه هدف (Target Population)

گروه بزرگی از مردم که نتایج مطالعه به آنها تعمیم داده می‌شود و به معنی مجموعه خاصی از معیارهای ورود است که تنظیم ویژگیهای بالینی و جمعیتی افرادی را که برای موضوع پژوهش کاملاً مناسب هستند مقدور می‌کند. تعریف جامعه هدف باید جامع و مانع باشد. یعنی این تعریف باید چنان باشد که از لحاظ زمانی و مکانی همه واحدهای مورد مطالعه را در بر بگیرد و در ضمن از شمول واحدهایی که نباید به مطالعه آنها پرداخته شود جلوگیری به عمل آید.

۱-۳) جمعیت در دسترس (Accessible Population)

معرف جمعیت هدف است و تعریف کاربردی از جامعه هدفی با معیارهای ورود بیشتر است تا مشخص کند افراد مورد پژوهش از کجا و در طی چه زمانی پیدا خواهند شد. به طور مثال آنهایی را که مطالعه روی ایشان غیر اخلاقی است یا مایل به شرکت در مطالعه نیستند و یا به نظر می‌رسد که اطلاعات بدی بدهند حذف کند. پژوهشگر برای انتخاب جمعیت در دسترس دو راه عمده دارد:

❖ انتخاب مبتنی بر بیمارستان

نمونه‌های مبتنی بر بیمارستان بالینی، ارزان بوده و جذب آنها در مطالعه آسان است ولی عوامل محدودش کننده‌ای که تعیین می‌کنند چه کسی به بیمارستان یا درمانگاه مراجعه کند اثر مهمی دارند. شکل رایج این مسأله درمانگاه تخصصی در سطح سوم مراقبت پزشکی است که مبتلایان به انواع خطرناک یا مشکل یک بیماری را که تصویری از پیش‌آگهی نامناسب و تظاهرات پیش یا افتاده را با هم ارائه می‌کنند یکجا جمع می‌نماید.

برای مثال مطالعات بروز اختلالات تشنجی در کودکانی که تشنجهای ناشی از تب داشته‌اند برآوردی به دست می‌دهد که دامنه آن از ۱/۵٪ در کل جمعیت تا ۵۸٪ در مراجعین درمانگاه تخصصی اطفال متغیر است.

❖ انتخاب مبتنی بر جمعیت

راه عمده دیگر انتخاب جمعیت در دسترس، انتخاب شرکت‌کنندگان در خاله خودشان است، که نمونه‌ای معرف از منطقه خاص می‌باشد. چنین نمونه‌های مبتنی بر جمعیت، بویژه برای استفاده در بهداشت عمومی و طبابت بالینی در کل جامعه سودمند هستند.

✓ نکته: از دیدگاهی دیگر، جوامع را به دو دسته متناهی و نامتناهی تقسیم می‌کنند.

الف) جامعه متناهی شامل تعدادی متناهی از عناصر است؛ به طور مثال مجموعه رأی‌دهندگان واجد شرایط در یک شهر.

ب) جامعه نامتناهی شامل تعداد بی‌نهایت عنصر است. به طور کلی جامعه نامتناهی اشاره به فرایندی دارد و عناصر آن، اگر فرایند تحت شرایط یکسان عمل کند، شامل تمام برآمدهای فرایند است. به طور مثال فرایند ساخت تراشه‌های حافظه که عناصرش تراشه‌هایی هستند که وقتی فرایند تحت شرایطی یکسان تا بی‌نهایت به کار ادامه دهد تولید می‌شوند. وقتی جامعه نامتناهی است

اطلاعات مربوط به آن را فقط می‌توان از یک نمونه به دست آورد. بنابراین انجام مشاهده برای کسب اطلاع درباره یک فرایند همیشه اطلاعات نمونه‌ای را فراهم می‌نماید و مهم نیست که تعداد این مشاهدات چقدر زیاد باشد.

نمونه‌ای از جامعه نامتناهی ممکن است برای هدف دیگری به عنوان جامعه متناهی در نظر گرفته شود. بنابراین تراشه‌های حافظه کامپیوتر که در طی یک هفته تولید می‌شوند نمونه‌ای از جامعه نامتناهی مربوط به فرایند تولید است ولی وقتی این تراشه‌ها برای یک سازنده کامپیوتر فرستاده می‌شود که مایل است با نمونه‌گیری از آن، قابل قبول بودن محموله را تعیین کند تراشه‌های این محموله برای این هدف تشکیل یک جامعه متناهی را می‌دهند.

۲- نمونه (Sample)

نمونه، بخشی منتخب از جامعه تحت بررسی است به قسمی که بتوان از آن استنباطهایی درباره جامعه استخراج کرد. هر یک از واحدهای جمعیت که به عنوان عضوی از نمونه در معرض انتخاب است، را یک واحد انتخاب (Selection Unit) یا واحد نمونه‌ای (Sampling Unit) می‌خوانند. این واحد می‌تواند فرد یا واحد نهایی یا مجتمعی از افراد یا اصطلاحاً خوشه (Cluster) باشد. در نمونه‌گیری یک مرحله‌ای از افراد، فرد هم واحد نمونه‌ای و هم واحد مشاهده است. به عبارت دیگر در این حالت واحد نمونه‌ای بر واحد مشاهده‌ای منطبق است.

۳- اهداف نمونه‌گیری (Sampling Goals)

در خیلی از مطالعات به علت مشکلات عدیده از جمله کمی وقت، مخارج زیاد، کمبود نیروی انسانی و در دسترس نبودن وسایل و تجهیزات به اندازه کافی و پراکندگی جمعیت مورد نظر نمی‌توان تمام جمعیت را مورد مطالعه قرار داد؛ در این صورت، تحقیق بر روی نمونه‌ای از جمعیت انجام می‌گیرد. از طرفی در مطالعه نمونه، به جای تحقیق کل جمعیت محقق می‌تواند دقت بیشتری در جمع‌آوری اطلاعات به کار گیرد و بررسی عمیق‌تری به عمل آورد.

✓ نکته: از آنجا که محقق نتایج حاصل از بررسی نمونه را به جامعه مورد بررسی تعمیم می‌دهد لذا برای معتبر بودن نمونه باید به نکاتی توجه نمود:

- نمونه بایستی به خوبی انتخاب شود تا نماینده جمعیت باشد.
- نمونه بایستی به اندازه کافی انتخاب شود.
- نمونه بایستی به اندازه کافی مطالعه شود.

۴- فواید نمونه‌گیری (Sampling Benefits)

حداقل ۸ دلیل منطقی برای نمونه‌گیری وجود دارد:

۱-۴ عامل زمان (Time Factor)

نمونه سریعتر از جمعیت هدف مورد مطالعه واقع می‌شود زیرا سرعت، عامل بسیار مهمی است و مخصوصاً در مواردی که پزشک، نیاز فوری به جواب سؤال با اهمیت دارد، مانند زمانیکه یک بیماری مسری و کشنده در یک منطقه شیوع پیدا می‌کند به طوریکه کسی در مورد این بیماری جدید اطلاعی ندارد؛ در این حالت در صورت بررسی تمام افراد مبتلا و به دست آوردن نتایج، تعداد زیادی از افراد در این زمان می‌میرند.

۴-۲) عامل صرفه جویی اقتصادی (Economic Factor)

مطالعه یک نمونه به مراتب ارزانتر از مطالعه تمام افراد یک جمعیت است زیرا زیرمجموعه‌ای از جمعیت هدف مورد اندازه‌گیری و آزمایش قرار می‌گیرند. این بررسی بالاخص در طراحی مطالعات بزرگ که نیاز به پیگیری دراز مدت (مثلاً چند نسل یا نسلهای متوالی) دارند از اهمیت بسیار زیادی برخوردار است.

۴-۳) مشاهدات مخرب (Destructive Nature of the Observation)

در بعضی مواقع مطالعه تمام افراد در یک جمعیت امکان ندارد زیرا در برخی حالات، مراحل اجرایی مطالعه از قبیل اندازه‌گیری موجب از بین رفتن شیء یا موجود تحت مطالعه می‌گردد به عنوان مثال به منظور بررسی بهبود غضروف در سگ از طریق مطالعه بافت‌شناسی بعد از ۶ هفته بی‌حرکتی اندام، حیوان کشته می‌شود و نمونه‌های لازم جهت مطالعه بافت‌شناسی از غضروف تعیین می‌شود. در این مطالعه بررسی کل جمعیت هدف، نیاز به کشتن تمام سگها دارد که امکان‌پذیر نمی‌باشد.

در موارد محدود دیگری هدف پژوهشگر تفسیر اتفاقات آینده مانند مطالعه تعیین پیش‌آگهی بیماران مبتلا به آنمی آپلاستیک شدید بدون هیچگونه اقدام درمانی می‌باشد؛ در این مطالعه از آنجا که تعدادی از بیماران به واسطه کم‌خونی شدید نیاز به تزریق خون پیدا می‌کنند که در غیر این صورت خطر مرگ وجود دارد، مطالعه کل جمعیت هدف امکان‌پذیر نمی‌باشد.

۴-۴) صحت نتایج (Accuracy)

نتایج نمونه نسبت به نتایج جمعیت هدف اغلب از صحت بیشتری برخوردار است. زیرا برای نمونه زمان و منابع بیشتری می‌توان صرف نمود؛ علاوه بر این از طراحی روشهای گرانتر نمونه‌برداری که درجه صحیح بودن نتایج را افزایش می‌دهند می‌توان استفاده کرد.

۴-۵) کنترل اشتباهات آماری (Statistical Mistake)

اگر نمونه‌ها به طور درستی انتخاب شوند برای برآورد اشتباهات نتایج آماری، از روشهای مناسب احتمالات می‌توان استفاده کرد. با مراعات این نکته در نمونه‌گیری است که پژوهشگر می‌تواند درباره وضعیت داده‌ها در مطالعه خود اظهار نظر یا تفسیر علمی (احتمالاتی) نماید.

۴-۶) جمعیت‌های بسیار بزرگ (Very Large Populations)

در این مورد چون امکان تماس با تک تک افراد نمی‌باشد باید از نمونه‌گیری استفاده کرد.

۴-۷) استفاده از جمعیت در دسترس (The Partly Accessible Population)

دسترسی به برخی از جمعیت‌ها بسیار مشکل است؛ مثل افرادی که در زندانها هستند یا تعداد هواپیماهای سقوط کرده در دریاها یا دسترسی به برخی از جمعیتها به دلیل محدودیت وقت و بودجه مشکل می‌باشد یا اتفاقاتی که خیلی به تدرت اتفاق می‌افتند مثل وقوع سیل. در این شرایط نیز نمونه‌گیری، منطقی‌تر از سرشماری به نظر می‌رسد.

۴-۸) همگن کردن افراد مورد بررسی (Homogeneity)

نمونه‌ها می‌توانند به ترتیبی انتخاب شوند که موجب کاهش ناهمگنی گردد. برای مثال در مورد بیماری لوپوس اریتماتوی سیستمیک (SLE) به واسطه تفاوت‌های کلینیکی متفاوت بیماران، ناهمگنی وجود دارد. در چنین شرایطی جهت مطالعه جنبه‌های

معینی از یک بیماری، نمونه‌ای از جمعیت که دارای خصوصیات مورد نظر یا ناهمگنی کمتر می‌باشد مناسبتر از کل جمعیت خواهد بود.

۵- قالب نمونه‌گیری یا چارچوب نمونه‌گیری (The Sampling Frame)

شامل لیست تمام افراد، اشیاء یا مراکز است که محقق می‌خواهد از آنها نمونه انتخاب کند. گاهی اوقات چارچوب موجود برای نمونه‌گیری چارچوبی کامل از جامعه مورد نظر نیست. به عنوان مثال برای وکلای شهرستان، چارچوب نمونه‌گیری فهرستی از تمام وکلایی است که تا ماه قبل عضو کانون وکلای شهرستان بوده‌اند. این فهرست نام وکلای را که بعد از ماه قبل به کانون پیوسته‌اند و نام وکلایی را که عضو کانون نیستند شامل نمی‌شود. در مواردی نظیر مثال فوق، بین جامعه هدف و جامعه مورد پژوهش تمایز وجود دارد. جامعه هدف، تمام وکلای شهرستان که تا ماه گذشته عضو کانون وکلا بوده‌اند، خواهند بود. وقتی چارچوب موجود برای نمونه‌گیری از نظر عناصر با جامعه هدف خیلی متفاوت باشد، نتایج نمونه ممکن است فقط به گونه‌ای محدود با جامعه هدف مربوط باشند؛ لذا در عمل برای به دست آوردن چارچوبی که با جامعه هدف دقیقاً مطابقت داشته باشد کوشش‌های زیادی صورت می‌گیرد مانند استفاده از جدول ارقام تصادفی. گاهی اوقات عناصر چارچوب از قبل شماره‌گذاری شده‌اند مانند صورتحسابهایی که دارای شماره سریال هستند یا دانشجویانی که به آنها شماره دانشجویی نسبت داده شده است. این شماره‌ها را می‌توان برای تشخیص عناصر در عمل نمونه‌گیری به کار برد. گاهی این اعداد تخصیص یافته دارای فاصله‌اند مانند وقتی که یک صورتحساب از درجه اعتبار ساقط شده است و یا وقتی که دانشجوی فارغ‌التحصیل شده است.

۶- میزان پاسخ‌دهی (Response Rate)

نسبتی از افراد انتخاب شده می‌باشد که موافق ورود به مطالعه هستند.

✓ نکته: میزان پاسخ‌دهی بر روی اعتبار این استنتاج که نمونه معرف جمعیت است تأثیر می‌گذارد؛ چرا که افرادی که دسترسی به آنها مشکل است و کسایکه از شرکت در مطالعه خودداری می‌کنند یا افرادی که پاسخ داده‌اند متفاوت می‌باشند؛ این موضوع در تحقیق ایجاد سوگیری می‌کند. میزان عدم پاسخ به طور جدی روی قابلیت تعمیم مطالعه اثر می‌گذارد.

۷- اعتبار و پایایی

۷-۱) اعتبار (Validity)

اعتبار به معنی آن است که مقدار اندازه‌گیری شده با مقدار واقعی همخوانی داشته باشد بنابراین اعتبار یعنی درست بودن اندازه‌گیری که سبب نتیجه‌گیری دقیق و قابل اعتماد می‌شود. اعتبار را بر این اساس که در کدامیک از قسمتهای یک مطالعه مطرح گردد در چهار نوع بررسی می‌کنند که در اینجا به شرح دو نوع اعتبار که در مباحث نمونه‌گیری بیشتر مطرحند پرداخته خواهد شد:

❖ اعتبار خارجی (External Validity)

عبارت از قدرت تعمیم‌پذیری (Generalizability) نتایج یک مطالعه به جامعه آماری و به تبع آن به جامعه هدف است و به عبارت دیگر نشان می‌دهد که نتیجه‌گیرهای محقق از این نمونه خاص تا چه حد می‌تواند در مورد افراد و مکانهای دیگر (دیگر اعضای جامعه هدف) صدق نماید. عواملی که بر اعتبار خارجی مؤثرند علاوه بر دقت در انجام پژوهش شامل موارد ذیل هستند:

- روش نمونه‌گیری و انتخاب نمونه مناسب (Representative Sample)
- حجم نمونه مناسب.

❖ اعتبار داخلی (Internal Validity)

عبارت از توان یک مطالعه در حذف اثر متغیرهای مخدوش‌کننده (Confounding Variables) در رابطه‌ای که بررسی شده است می‌باشد. در واقع بیانگر اعتبار نتایج بدست آمده در نمونه مورد مطالعه صرف‌نظر از قدرت تعمیم آن به جامعه است. یکی از اهداف مهم محاسبه حجم نمونه مناسب ایجاد تعادلی بین امکانات محقق برای انجام پژوهش و اعتبار خارجی و داخلی می‌باشد و این به معنی آن است که همواره حجم نمونه بیشتر نشان‌دهنده پژوهشی با کیفیت بالاتر نیست.

✓ نکته: هرچه حجم نمونه بیشتر شود اعتبار خارجی بیشتر و اعتبار داخلی کمتر می‌شود.

۲-۷ پایایی (Reliability)

به معنی آن است که اندازه‌گیری، تکرارپذیر باشد و با تکرار نمونه‌گیری نتیجه نخست بدست آید، چه نتیجه درست باشد و چه نباشد. بنابراین پایایی، تکرارپذیری اندازه‌گیری است که در هر بار اندازه‌گیری نتیجه یکسانی حاصل شود. دو نوع پایایی وجود دارد:

❖ پایایی خارجی (External Reliability)

تکرارپذیری نتایج حاصل از مطالعه در جامعه هدف می‌باشد یعنی اگر نمونه‌های N تری از جامعه انتخاب شوند و مطالعه روی آنها تکرار شود نتایج یکسانی حاصل می‌گردد.

❖ پایایی داخلی (Internal Validity)

تکرارپذیری نتایج حاصل از مطالعه در داخل نمونه انتخاب شده می‌باشد بدین معنی که در تکرار اندازه‌گیری بر روی همان افراد مورد پژوهش نتیجه یکسان بدست آید.

۸- خطا (Error)

در علم تحقیق دو دسته خطای تصادفی و منظم وجود دارد:

۸-۱ خطای تصادفی (Random Error)

اشتباهی است که موجب نتیجه‌گیری غیردقیق (Imprecise) می‌شود بطوریکه اگر این نتیجه‌گیری چند بار تکرار شود نتایج متفاوتی بدست می‌آید. مثال ساده آن استفاده از ترازوی غیردقیقی است که وزن یک شخص ۷۵ کیلوگرمی را در ۵ بار اندازه‌گیری ۷۲، ۷۵، ۷۸، ۷۴ و ۷۶ ثبت می‌کند (در این مثال پایایی وسیله اندازه‌گیری اندک می‌باشد). این نوع خطاها اگر در حجم وسیع ادامه پیدا کنند یکدیگر را جبران و حذف می‌نمایند. از این رو افزایش حجم نمونه با ایجاد فرصت برای حذف خطاهای تصادفی ناهمسو سبب کاهش خطا می‌شود. در همان مثال بالا اندازه‌گیریهای ترازو دقیق نیست اما میانگین نهایی وزن، برآورد مناسبی از مقدار واقعی وزن فرد را نشان خواهد داد.

۸-۲ خطای منظم یا تورش (Systematic Error or Bias)

اشتباهی است که نتایج را بطور منظم تغییر می‌دهد، بگونه‌ای که با تکرار آن آزمایش، همان اشتباه قبلی تکرار شود. برای مثال ترازویی که همواره وزن فرد ۷۵ کیلوگرمی را ۷۷ کیلوگرم و وزن فرد ۳۵ کیلوگرمی را ۳۷ کیلوگرم نشان می‌دهد (در هر اندازه‌گیری ۲ کیلوگرم از مقدار واقعی بیشتر گزارش می‌کند)، دارای سوگیری یا تورش است.

✓ نکته: نداشتن خطای تصادفی را دقت (Precision) یا پایایی (Reliability) و نداشتن خطای منظم را عدم سوگیری (UnBiasness) یا اعتبار (Validity) گویند.

✓ نکته: برآیند دو بردار پایایی و اعتبار را صحت (Accuracy) گویند و مطالعه‌ای صحت بیشتری دارد که سوگیری و خطای تصادفی کمتری داشته باشد (تصویر ۱-۱).



برای درک بهتر مفاهیم ذکر شده، به مثال تیراندازی به سیل که در ذیل می‌آید توجه کنید.

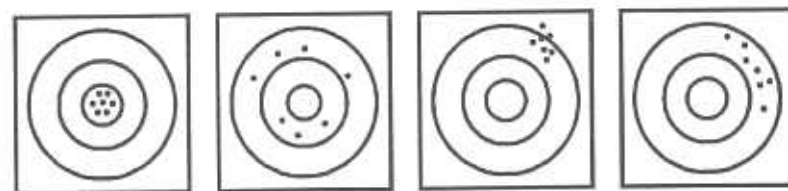
• اگر تفنگ سالمی را به تیرانداز ماهری دهند تا بطرف سیل شلیک کنند، همگی را به مرکز سیل می‌زنند.

به عبارتی در این روند هم اعتبار و هم پایایی وجود دارد (تصویر ۱-۲-۱).

• اگر همین تفنگ را به تیراندازی غیر ماهر دهند، پایایی از دست می‌رود چون تیراندازی او تکرارپذیر نیست ولی هنوز اعتبار دارد چراکه میانگین فاصله نقاط از مرکز سیل با احتمال بالایی در مرکز سیل خواهد بود (خطای تصادفی دارد ولی سوگیری ندارد) (تصویر ۱-۲-۲).

• اگر لوله تفنگ را کج کرده و آن را به تیرانداز ماهر دهند، اعتبار از دست می‌رود و تمامی نقاط به یک‌سوی سوگیری خواهند داشت ولی به دلیل مهارت تیرانداز، پایایی برقرار است و خطای تصادفی وجود ندارد (تصویر ۱-۲-۳).

• اگر لوله تفنگ را کج کرده و آن را به تیرانداز غیر ماهر دهند، اعتبار و پایایی هر دو از دست می‌روند زیرا به دلیل عدم مهارت، نقاط اصابت تکرار پذیر نیستند (خطای تصادفی) و به دلیل کج بودن لوله تفنگ، میانگین فاصله نقاط از مرکز سیل نیز منطبق بر مرکز سیل نمی‌باشد (سوگیری منظم) (تصویر ۱-۲-۴).



تصویر ۱-۲

۱

۲

۳

۴

✓ نکته: در مثال فوق کاملاً مشهود است که حجم نمونه مناسب می‌تواند اثر خطای تصادفی را از بین ببرد (حالت دوم) چراکه میانگین، درست و معتبر خواهد بود ولی افزایش حجم نمونه در کاهش خطای منظم اثر قابل توجهی نخواهد داشت.

✓ نکته: به طور خلاصه جهت تأمین اعتبار و پایایی تحقیق باید نکات ذیل مد نظر قرار گیرند:

- ۱- کنترل مداخله‌گرها از طریق حذف یا مشابه سازی
 - ۲- استفاده از ابزار دقیق
 - ۳- استفاده از روش مناسب جهت گردآوری اطلاعات نظیر پرسشنامه معتبر
 - ۴- استفاده از روش یکسان در مورد تمامی افراد مورد مطالعه
 - ۵- استفاده از روشهای تصادفی انتخاب نمونه‌ها تا حد امکان و انتخاب نمونه به تعداد کافی
 - ۶- استفاده از روشهای کور تا حد امکان
 - ۷- شناسایی و رفع سوگیریها
 - ۹- طراحی نمونه‌گیری (Sample Design)
- طراحی نمونه‌گیری شامل دو قسمت روش نمونه‌گیری و تخمین حجم نمونه می‌باشد.
- ۱-۹ روش نمونه‌گیری (Sampling Plan)

روش مورد استفاده برای انتخاب نمونه‌ها از جمعیت را گویند. برای این کار دو روش نمونه‌گیری احتمالی و غیر احتمالی وجود دارد. در نمونه‌گیری احتمالی (Probability Sampling) انتخاب افراد و واحدهای مطالعه به صورت تصادفی است. این نوع نمونه‌گیری خود شامل چند روش می‌باشد که در فصل دوم به آن پرداخته می‌شود.

در نمونه‌گیری غیر احتمالی (NonProbability Sampling) برای انتخاب نمونه، هیچ‌گونه روش تصادفی به کار گرفته نمی‌شود و نمونه‌گیری حالت غیر احتمالی به خود می‌گیرد که در مورد انواع آن به طور کامل در فصل آینده توضیح داده خواهد شد.

۲-۹) پروسه تخمین (Estimation Procedures)

این مرحله شامل محاسبه حجم کافی نمونه جهت تحقیق می‌باشد که با استفاده از روابط خاص، حداقل حجم نمونه لازم برای به دست آوردن یک نتیجه منطقی قابل تعمیم به جامعه، به دست می‌آید. درباره طریقه محاسبه حجم نمونه در فصل پنجم توضیح داده خواهد شد.

انواع نمونه‌گیری

Types of Sampling

✓ روشهای نمونه‌گیری احتمالی

✓ روشهای نمونه‌گیری غیر احتمالی

✓ روشهای نمونه‌گیری چند مرحله‌ای

فصل

۲

در انجام هر تحقیق می‌توان به دو صورت احتمالی (Probability) و غیر احتمالی (Nonprobability) نمونه‌گیری کرد. در نمونه‌گیری احتمالی شانس هر فرد برای ورود به نمونه معین و نامساوی با صفر می‌باشد و ذکر این نکته لازم است که شانس افراد لزوماً با هم برابر نمی‌باشد این در حالیست که در نمونه‌گیری غیر احتمالی شانس هر فرد برای ورود به نمونه معین نمی‌باشد.

استفاده از نمونه‌گیری غیر احتمالی چه در مرحله طراحی و چه در مرحله اجرا ساده‌تر از نمونه‌گیری احتمالی است و به زمان و سرمایه کمتری نیاز دارد. در حالیکه تنها نتایج به دست آمده از نمونه‌گیری احتمالی را می‌توان به جامعه آماری تعمیم داد و نمونه‌های غیر احتمالی چنین قابلیت را ندارند.

گاهی مطالعه را در ابتدا برای برآورد تقریبی وضعیت موجود و همچنین رفع نواقص اجرایی تحقیق با نمونه‌های غیر احتمالی انجام می‌دهند اصطلاحاً به آن مطالعه پیش‌آزمون (Pilot Study) می‌گویند و در گام بعدی، و پس از بررسی نتایج پیش‌آزمون و رفع نواقص موجود، تحقیق را با نمونه‌گیری احتمالی تکرار می‌کنند تا بتوان نتایج آن را به جامعه آماری تعمیم داد.

برای انتخاب یک نمونه احتمالی یا غیر احتمالی روشهای گوناگونی وجود دارد که در ادامه به معرفی و بررسی این روشها پرداخته خواهد شد.

روشهای نمونه‌گیری احتمالی (Probability Sampling)

نمونه‌گیری احتمالی را به روشهای ذیل انجام می‌دهند:

۱- نمونه‌گیری تصادفی ساده (Simple Random Sampling)

۲- نمونه‌گیری تصادفی منظم (Systematic Random Sampling)

۳- نمونه‌گیری تصادفی طبقه‌بندی شده (Stratified Random Sampling)

۴- نمونه‌گیری تصادفی خوشه‌ای (Cluster Random Sampling)

در ذیل به شرح هر یک از آنها می‌پردازیم:

نمونه‌گیری تصادفی ساده (Simple Random Sampling)

در این روش که از ساده‌ترین روشهای نمونه‌گیری احتمالی می‌باشد، شانس هر فرد برای ورود به نمونه با سایر افراد برابر است. برای انجام این روش در ابتدا باید لیستی از افراد جامعه مورد مطالعه که قالب نمونه‌گیری (Sampling Frame) نامیده می‌شود در اختیار داشت. سپس از این لیست به تعداد حجم نمونه، به صورت تصادفی افراد انتخاب می‌شوند و در نمونه قرار می‌گیرند. برای انتخاب تصادفی افراد می‌توان از قرعه کشی، جدول اعداد تصادفی و نرم‌افزارهای رایانه‌ای استفاده کرد که در ذیل به شرح این روشها می‌پردازیم:

۱- قرعه کشی: این روش بسیار ساده است. نام افراد روی کاغذهای جداگانه نوشته می‌شود، در ظرفی قرار گرفته و سپس از بین آنها انتخاب صورت می‌گیرد. در این حالت اگر کاغذها خوب مخلوط نگردند، احتمال افراد برای انتخاب شدن در جاهای مختلف ظرف متفاوت خواهد بود. برای رفع این مشکل می‌توان از کبره‌های گردان استفاده کرد که در این حالت، نام افراد را روی توپهای کوچک می‌نویسند و بعد در کبره‌ای توخالی قرار می‌دهند و با چرخاندن کبره مشکل فوق نیز برطرف می‌شود. قرعه کشی در جمعتهای بزرگ عملاً کارایی ندارد و بسیار وقت‌گیر می‌باشد.

۲- استفاده از جدول اعداد تصادفی (Table of Random Digits)

روش ساده و رایجی است و برای اجرای آن افراد موجود در قالب نمونه‌گیری، از شماره ۱ تا آخر شماره‌بندی می‌شوند؛ سپس با استفاده از جدول اعداد تصادفی عددی بین ۱ تا بالاترین شماره به تعداد حجم نمونه گزینش می‌گردند. در جدول اعداد تصادفی ارقام به طور کاملاً تصادفی کنار هم قرار گرفته‌اند. این جدول در انتهای تقریباً تمامی کتابهای آمار موجود است (ضمیمه ۱).

چگونه از جدول اعداد تصادفی استفاده می‌کنند؟

۱- محقق، جدول را رو بروی خود گذاشته و به صورت تصادفی عددی را از جدول انتخاب می‌کند (مثلاً با قرار دادن نوک مداد روی جدول با چشم بسته)

۲- محقق، تعداد رقم بزرگترین شماره در قالب نمونه‌گیری را در نظر می‌گیرد (k)؛ مثلاً اگر بزرگترین شماره ۴۰ باشد $K=2$ می‌شود یا اگر ۴۰۰ باشد $K=3$ خواهد بود سپس از عددی که پیش از این روی جدول اعداد تصادفی تعیین کرده است، شروع کرده و K رقم K رقم در جهتی که از قبل تعیین نموده، حرکت کرده و اعداد را می‌خواند. جهت حرکت می‌تواند به بالا، به پایین، چپ، راست و یا حتی به صورت مورب باشد. خواندن اعداد تا جایی ادامه می‌یابد که حجم نمونه تکمیل شود.

به عنوان مثال فرض کنید محقق قصد دارد از میان ۵۰۰ نفر، ۵۰ نفر را انتخاب کند. اسامی این افراد در قالب نمونه‌گیری موجود است و از ۱ تا ۵۰۰ شماره زده شده‌اند. بالاترین شماره در قالب نمونه‌گیری، ۵۰۰ است که ۳ رقمی می‌باشد. از نقطه شروع که تصادفی انتخاب شده، مثلاً در جهت رو به راست شروع به خواندن کرده و ۳ رقم ۳ رقم پیش می‌رود و تا آنجا ادامه می‌دهد که ۵۰ نفر انتخاب شوند. به این نکته توجه شود که ممکن است جدول، اعدادی مثل ۶۵۷ بدهد که فردی با این شماره در تحقیق موجود نمی‌باشد؛ در چنین شرایطی عدد بعدی خوانده می‌شود که یا ۰۰۳ به معنی فرد شماره ۳ است یا ۰۳۴، فرد شماره ۳۴ می‌باشد. توجه کنید که استفاده از جدول اعداد تصادفی نیز در جامعه‌های پر جمعیت و برای حجم نمونه‌های بالا کارایی ندارد و بسیار مشکل می‌باشد.

۳- استفاده از نرم افزارهای رایانه‌ای: برخی نرم افزارهای رایانه‌ای مثل EPI Info یا EXCEL می‌توانند اعداد تصادفی را در محدوده مورد نظر به محقق تحویل دهند و حتی در برنامه EXCEL با وارد کردن قالب نمونه‌گیری به برنامه می‌توان به جای عدد تصادفی، نام افراد انتخاب شده را از برنامه دریافت کرد. البته وارد کردن قالب نمونه‌گیری کار وقت‌گیری است؛ ولی اگر از قبل وارد رایانه شده باشد، مثل اسامی کارمندان یک شرکت، با یک COPY/PASTE ساده می‌توان آن را در EXCEL قرارداد.

تولید اعداد تصادفی با استفاده از نرم افزار EPI Info

پس از ورود به محیط نرم افزار EPI، منوی اصلی Program را انتخاب نموده و در این قسمت زیر منوی EPITABLE Calculator انتخاب می‌گردد. یکی از قسمتهای این زیر منو، زیر منوی Sample است که در این قسمت جعبه Random Number Table را فعال می‌نماییم. این پنجره زمانی کار برد خواهد داشت که محقق نیاز به تعدادی عدد تصادفی با تعداد رقمهای برابر داشته باشد. مثلاً به ۵۰ عدد سه رقمی نیاز داشته باشد که برای این کار تعداد اعداد مورد نیاز و تعداد رقم مورد نظر وارد می‌شوند. با انتخاب کلید Calculate، جدولی شامل آن تعداد عدد درخواستی با تعداد رقمهای وارد شده مشاهده می‌گردد (تصویر ۱-۲).



[] Random number table generator

How many random numbers 550
How many digits per number 5

Calculate

Reset

Quit

تصویر ۲-۱

همچنین در زیر منوی Sample، جعبه‌ای با نام Random Number List وجود دارد که جهت انتخاب اعداد تصادفی در یک دامنه مشخص و با تعداد مشخص کاربرد خواهد داشت. بعنوان مثال انتخاب ۱۰۰ عدد تصادفی از میان اعداد ۱ تا ۱۰۰۰۰، توجه کنید که کاربرد این قسمت با توجه به اینکه قابلیت انتخاب اعداد با تعداد رقمهای مختلف را دارا می‌باشد، بیش از قسمت قبلی است. این پنجره به ترتیب، تعداد عددهای مورد نیاز، کمینه (minimum) و بیشینه (maximum) دامنه اعداد وارد می‌گردند. همچنین در این پنجره کاربرد قدرت انتخاب یکی از دو حالت بدون جایگزینی یا با جایگزینی اعداد انتخاب شده را خواهد داشت. پس از انتخاب جعبه Calculate، لیستی از اعداد تصادفی در دامنه خواسته شده و به تعداد مورد نظر نشان داده خواهد شد (تصویر ۲-۲).

[] Random number List generator

How many random numbers 100
Minimum range of numbers 0
Maximum range of numbers 1000

[] Drawing with replacement

Calculate

Reset

Quit

تصویر ۲-۲

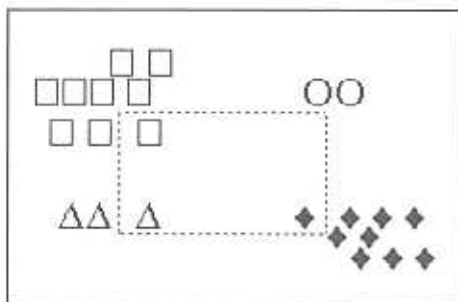
✓ نکته: در روش نمونه‌گیری تصادفی ساده حتما باید قالب نمونه‌گیری وجود داشته باشد.

✓ نکته: نمونه‌گیری تصادفی ساده، برای جامعه‌های بزرگ و حجم نمونه‌های بالا وقت گیر و گاهی غیر عملی است.



✓ نکته: نمونه‌گیری تصادفی ساده در جامعه‌هایی که از لحاظ صفت مورد نظر پراکندگی بالایی دارند، می‌تواند ایجاد تورش انتخاب (Selection Bias) نماید. به عنوان مثال فرض کنید محقق قصد دارد از یک کلاس ۱۰۰ نفره دانشجویان پزشکی، ۲۰ نفر را برای تعیین میانگین مقدار هموگلوبین خون انتخاب کند. به روش نمونه‌گیری تصادفی ساده، ۲۰ نفر را انتخاب می‌نماید که از قضا ۳۰ خانم و ۱۰ آقا انتخاب می‌شوند. حال اگر نسبت آقایان به خانمها در این کلاس بزرگتر از ۱ بر ۳ بوده باشد، از آنجایی که طبق بررسی متون، میانگین مقدار هموگلوبین خون خانمها از آقایان پایین‌تر است، محقق میانگینی پایین‌تر از میانگین واقعی کلاس به دست خواهد آورد.

با فرض کنید پراکندگی صفتی در جامعه به شکل تصویر ۲-۳ باشد. اگر نمونه‌ای که به روش تصادفی ساده انتخاب شده است، مستطیل نقطه چین باشد، در اینصورت همانگونه که در شکل مشخص است از موارد O حتی یک مورد در نمونه وجود ندارد، حال آنکه ممکن است روی نتیجه مطالعه تأثیر گذار باشد یا از موارد □ و ♦ یک نهم ولی از Δ یک سوم انتخاب شده است. در واقع در اینجا □ و ♦ نسبت به Δ شانس اثرگذاری کمتری بر نتیجه مطالعه دارند. پس در کل، نمونه‌گیری تصادفی ساده توجهی به پراکندگی صفت ندارد و آن را لزوما رعایت نمی‌کند.



تصویر ۲-۳

نمونه‌گیری تصادفی منظم (Systematic Random Sampling)

در این نمونه‌گیری نیز همچون نمونه‌گیری تصادفی ساده، شانس ورود هر فرد به نمونه با سایر افراد برابر است و به قالب نمونه‌گیری هم نیاز دارد ولی نسبت به نمونه‌گیری تصادفی ساده سریعتر است و پراکندگی را بهتر در نظر می‌گیرد.

در اینجا لازم است مفهوم فاصله نمونه‌گیری (Sampling Interval) بیان گردد.

اگر N تعداد افراد جامعه مورد مطالعه باشد و n تعداد حجم نمونه تحقیق باشد، به نسبت N/n فاصله نمونه‌گیری گویند و آن را با K نشان می‌دهند. برای انجام این روش عددی بین ۱ تا K به صورت تصادفی انتخاب می‌گردد. این عدد را با ۲ نشان می‌دهند. فرد با شماره ۲ انتخاب شده و بعد از k رقم k رقم جلو رفته و افرادی که شماره شان به دست می‌آید انتخاب می‌گردند و این کار آنقدر ادامه می‌یابد تا حجم نمونه کامل شود و بدین ترتیب افراد زیر انتخاب می‌شوند:

$$r, r+k, r+2k, \dots, r+(n-1)k$$

مثال: محقق قصد دارد از میان ۱۰۰ نفر دانش آموز، ۱۰ نفر را به روش نمونه گیری تصادفی منظم برای انجام طرحی با عنوان تعیین فراوانی آلودگی به شیش در دانش آموزان دبستانی شهر تهران در سال ۱۳۸۰ انتخاب کند:

فاصله نمونه گیری $k=10$ می باشد. حالا عددی بین ۱ تا ۱۰ به صورت تصادفی انتخاب می کند، مثلاً ۷ که به آن می گویند. افراد با شماره های زیر انتخاب خواهند شد:

۷، ۱۷، ۲۷، ۳۷، ۴۷، ۵۷، ۶۷، ۷۷، ۸۷، ۹۷

نکته ۱: همانگونه که در مثال فوق مشخص گردید در روش نمونه گیری تصادفی منظم، پراکندگی صفت بهتر در نظر گرفته می شود؛ چرا که محقق از تمام قسمتهای قالب نمونه گیری، فردی را انتخاب خواهد نمود، در حالی که اگر نمونه گیری تصادفی ساده انجام شود، به عنوان مثال شاید افراد ۶۰، ۶۱، ۱۷، ۲۹،، ۱۰ انتخاب شوند و هیچ نماینده ای از افراد با شماره های بالای ۶۰ وجود نداشته باشد.

نکته ۲: اگر افراد در قالب نمونه گیری به صورت تناوبی (Periodical) ردیف شده باشند، آنگاه ممکن است محقق دچار تورش انتخاب (Selection Bias) گردد.

مثال: فرض کنید که در قالب نمونه گیری، افراد به صورت زیر ردیف شده اند:

m = Male
f = female
m, m, f, f, m, m, f, f,
با $K=4$ و $r=2$ نمونه به صورت زیر خواهد شد:

m, m, m, m,

در این تحقیق خاتمه عمل حذف شده اند و این نمونه گیری دارای تورش انتخاب خواهد بود.

مثال: محقق می خواهد میانگین تعداد مراجعان به درمانگاه تخصصی روماتولوژی را در طول هفته به دست آورد. به طور تصادفی $k=7$ شده و ۲ هم روز شنبه انتخاب می شود، در این صورت، در این نمونه گیری همیشه تعداد مراجعان در روز شنبه ثبت می شود. با توجه به تحقیقات دیگر موجود در ایران، روزهای شنبه شلوغترین روز درمانگاهها می باشند؛ لذا میانگین تعداد مراجعان که در این تحقیق به دست می آید، بالاتر از میانگین اصلی آن خواهد بود.

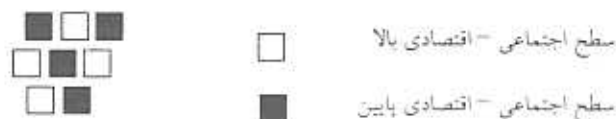
نمونه گیری تصادفی طبقه بندی شده (Stratified Random Sampling)

در نظر نگرفتن پراکندگی صفت، بزرگترین مشکلی است که شاید در روش نمونه گیری فوق پیش آید. برای رفع کامل این مشکل، محقق می تواند جامعه مورد مطالعه را بر اساس صفتی که پراکندگی آن اهمیت زیادی دارد، به چند طبقه (Stratum) تقسیم بندی کرده و سپس از هر طبقه به نسبت افراد موجود در آن به روش نمونه گیری تصادفی ساده انتخاب نماید. به عنوان مثال در تحقیقی با عنوان تعیین فراوانی ابتلا به شیش در دانش آموزان دبستانی شهر تهران در سال ۱۳۸۰ محققین می دانستند که سطح اجتماعی- اقتصادی خانواده ها روی ابتلا فرزندان به شیش تأثیر دارد و در ضمن، این سطح در شمال و جنوب شهر تهران متفاوت است. پس دبستانهای شمال شهر را یک طبقه (طبقه ۱) و دبستانهای جنوب شهر را طبقه دیگر (طبقه ۲) قرار دادند. این طبقه بندی بر اساس اختلاف سطح اجتماعی- اقتصادی خانواده های ساکن شمال و جنوب شهر می باشد. حال فرض کنید در طبقه اول ۲۰۰۰۰۰ دانش آموز و در طبقه دوم ۳۰۰۰۰۰ دانش آموز و حجم نمونه نیز ۵۰۰ نفر باشد. نسبت ۲ بر ۳ بین طبقات اول و

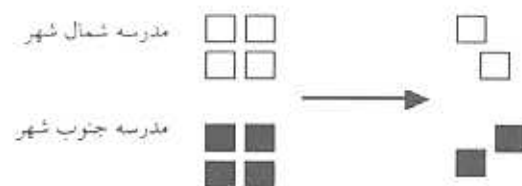
دوم باید در حجم نمونه نیز رعایت شود؛ یعنی ۲۰۰ نفر از دبستانهای شمال شهر و ۳۰۰ نفر از دبستانهای جنوب شهر به طور تصادفی ساده انتخاب گردند.

نکته: در نمونه گیری از طبقات می توان از نمونه گیری تصادفی منظم نیز استفاده کرد.

نکته: شرط یک طبقه بندی (Stratification) موفق این است که تا حد امکان پراکندگی صفت مورد نظر در داخل طبقه ها کم و در بین طبقات زیاد باشد. به تصاویر ۲-۴ و ۲-۵ توجه کنید.



تصویر ۲-۴

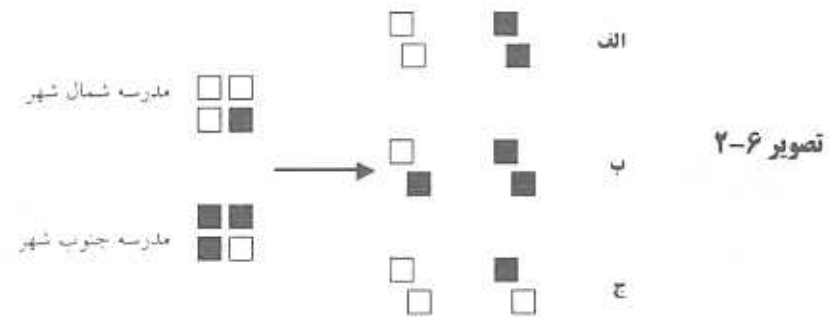


تصویر ۲-۵

تصویر ۲-۴ دانش آموزان دبستانی شهر تهران را نشان می دهد. در تصویر ۲-۵ مشخص می شود که در مدارس شمال شهر تمامی بچه ها از سطح اجتماعی- اقتصادی بالا و احتمال پایین تری برای ابتلا به شیش و در مدارس جنوب تمامی بچه ها از سطح اجتماعی- اقتصادی پایین و احتمال بالاتری برای ابتلا به شیش برخوردارند. یعنی پراکندگی درون طبقه ای خیلی کم (در اینجا صفر) و پراکندگی بین طبقه ای خیلی بالا است. این یک طبقه بندی معتبر است که در نهایت منجر به انتخاب یک نمونه معتبر و معرف (Representative) شده است.

در تصویر ۲-۶ حالتی ترسیم شده که در مدارس شمال بچه هایی با سطح اجتماعی- اقتصادی پایین و در مدارس جنوب هم بچه هایی با سطح اجتماعی- اقتصادی بالا وجود دارند. این به معنی پراکندگی بالای درون طبقه ای می باشد و البته مشخصاً پراکندگی صفت بین مدارس جنوب و شمال پایین آمده است. در نهایت ممکن است نمونه های متفاوتی برداشته شود. به عنوان مثال، نمونه الف معرف جامعه است ولی نمونه ب معرف جامعه نیست و سطح اجتماعی- اقتصادی پایین تر و در نتیجه فراوانی افراد مبتلا به شیش در آن بالاتر از حد واقعی جامعه می باشد. البته می توان نمونه های دیگری نیز داشت. نمونه ج را چگونه تفسیر می کنید؟ فراوانی افسردگی در چنین نمونه ای بالاتر از حد واقعی است یا پایین تر؟

✓ نکته: توجه کنید که برای طبقه‌بندی کردن درست و دقیق باید از جامعه مورد مطالعه و صفتهای موجود به درستی آگاه بود. در اینجا است که اهمیت بررسی متون (Literature Review) دقیق و کامل روشن می‌شود.



دو نوع نمونه‌گیری تصادفی طبقه‌بندی شده وجود دارد که در ذیل به آن می‌پردازیم.

۱- نمونه‌گیری تصادفی طبقه‌بندی شده متناسب (Proportional)

شانس تمامی افراد برای ورود به نمونه یکسان و برابر با یکدیگر می‌باشد. در این نوع، پس از طبقه‌بندی باید قالب نمونه‌گیری در هر طبقه وجود داشته باشد تا بتوان از هر طبقه به روش نمونه‌گیری تصادفی ساده نمونه گرفت. تمامی مثالهایی که تا اینجا ذکر شده‌اند از این نوع می‌باشند.

۲- نمونه‌گیری تصادفی طبقه‌بندی شده نامتناسب (Disproportional)

گاهی نمونه‌گیری از برخی طبقات بسیار پرهزینه تمام می‌شود. در این شرایط محقق ترجیح می‌دهد که از آن طبقات به نسبت کمتری نمونه‌گیری کند؛ یا گاهی پراکندگی صفت مورد پژوهش در برخی طبقات آنقدر بالا است که محقق، برای رسیدن به نمونه‌ای معرف، لازم می‌بیند که از این طبقات به نسبت بیشتری نمونه‌گیری شود.

✓ نکته: اینکه محقق از هر طبقه با توجه به هزینه و پراکندگی، به چه نسبتی نمونه‌گیری کند، قابل محاسبه است و مبتنی بر بررسی متون و یا انجام مطالعات پیش‌آزمون می‌باشد.

مثلا در همان مثال "تعیین فراوانی ابتلا به شیش در میان دانش‌آموزان دبستانی شهر تهران در سال ۱۳۸۰" محقق به دلیل هزینه بالای سفرهای درون شهری به مدارس جنوب شهر تصمیم به اجرای نمونه‌گیری طبقه‌بندی شده نامتناسب می‌گیرد و پس از محاسبه نسبتهای نمونه‌گیری مورد نیاز، براساس بودجه و بررسی متون قبلی، به ترتیب از شمال و جنوب تهران به نسبت ۳ به ۲ نمونه‌گیری می‌نماید. یعنی ۳۷۵ نفر از مدارس شمال شهر و ۱۲۵ نفر از مدارس جنوب شهر تهران انتخاب می‌کند.

✓ نکته: دقت این روش، اگر نسبت بین طبقات درست محاسبه نشده باشد، پایین‌تر از سایر روشهای نمونه‌گیری احتمالی است. برای بالا بردن دقت این روش، اغلب به وزن دادن داده‌های بدست آمده از نمونه‌ها در مرحله آنالیز اطلاعات می‌پردازند.

✓ نکته: با کمی دقت در می‌یابیم که با وجود نسبتهای متفاوت نمونه‌گیری، شانس هر فرد برای ورود به نمونه معین است و پس برخلاف سایر روشهای نمونه‌گیری احتمالی، شانس افراد برای ورود به نمونه با هم برابر نیست و میان طبقات مختلف فرق می‌کند.

نمونه‌گیری تصادفی خوشه‌ای (Cluster Random Sampling)

بیشترین کارسرد آن در برخورد با جمعیههای بزرگ و جمعیههایی است که نمونه‌گیری از آنها به روشهای گفته شده امکان ندارد و با بسیار مشکل و وقت‌گیر است. در این روش، جامعه را به چند خوشه (Cluster) تقسیم می‌کنند، سپس بطور تصادفی یک یا چند خوشه را جهت مطالعه انتخاب می‌نمایند.

مثال: در طرح "تعیین فراوانی ابتلا به شیش در دانش‌آموزان دبستانی شهر تهران در سال ۱۳۸۰"، محقق می‌داند که تهران مدارس زیادی دارد و نمونه‌گیری از تمام آنها به دلیل سفرهای درون شهری فراوان برایش مقدور نیست. در ضمن، طبق تحقیقات گذشته، فراوانی ابتلا به شیش در مدارس با هم تفاوت معنی‌داری ندارند، پس هر مدرسه را یک خوشه قرار داده و از میان مدارس به طور تصادفی چند مدرسه را انتخاب می‌کند.

✓ نکته: محقق در نمونه‌گیری تصادفی خوشه‌ای، عملاً تنها به بررسی قسمتی از جامعه می‌پردازد، چرا که چند خوشه را به جای تمام جامعه مورد بررسی قرار می‌دهد. لذا نباید خوشه‌های موجود، تفاوت زیادی داشته باشند؛ به عبارتی، برای یک نمونه‌گیری تصادفی خوشه‌ای معتبر باید پراکندگی صفت بین خوشه‌ها پایین باشد.

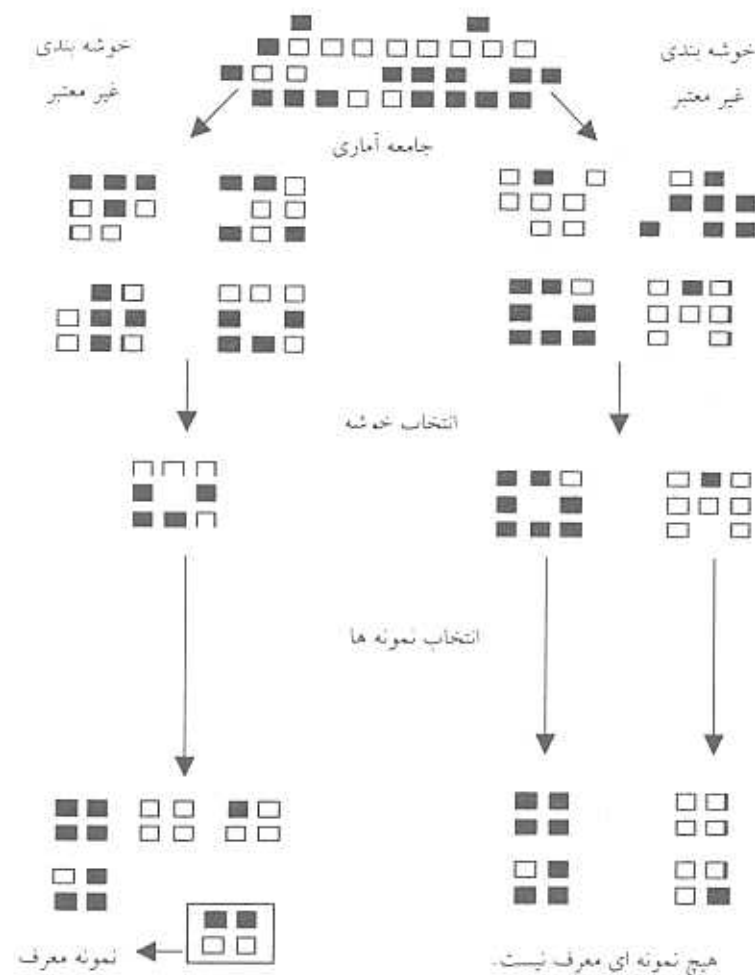
✓ نکته: خوشه‌ها باید معرف جامعه باشند؛ یعنی باید از هر دست فرد از نظر صفت مورد نظر، در یک خوشه وجود داشته باشد؛ به عبارتی برای یک نمونه‌گیری تصادفی خوشه‌ای معتبر، پراکندگی درون خوشه‌ای صفت، باید بالا باشد.

تصویر ۷-۲ یک خوشه‌بندی معتبر را نشان می‌دهد. توضیح اینکه در اینجا پراکندگی بین خوشه‌ها بسیار کم است، چرا که کاملاً شبیه هستند و در هر کدام ۴ مربع سیاه و ۴ مربع سفید وجود دارد، ولی در داخل هر خوشه پراکندگی بالا است چون هم مربع سفید و هم مربع سیاه وجود دارد و همانگونه که در انتهای شکل دیده می‌شود احتمال به دست آوردن نمونه معرف با چنین خوشه‌بندی زیاد است.

شکل ۸-۲ یک خوشه‌بندی غیر معتبر را نشان می‌دهد. توضیح اینکه در اینجا پراکندگی بین خوشه‌ها بسیار بالا است، چرا که نسبت مربعهای سیاه و سفید در تمام آنها یکسان نیست، ولی در داخل هر خوشه پراکندگی پایین است چون در هر خوشه یکی از مربعهای سفید یا سیاه تعداد غالب را دارند. احتمال به دست آوردن نمونه معرف با چنین خوشه‌بندی پایین است.

✓ نکته: با کمی توجه در نکات فوق، می‌توان گفت که دقت نمونه‌گیری تصادفی خوشه‌ای از نمونه‌گیری تصادفی طبقه‌بندی شده کمتر است، پس به حجم نمونه بیشتری نیاز دارد.

✓ نکته: به طور نظری در یک نمونه‌گیری تصادفی خوشه‌ای شانس تمام خوشه‌ها برای انتخاب شدن برابر است و این احتمال وجود دارد که نمونه نهایی از هر کدام از خوشه‌ها باشد و این موضوع نیاز به قالب نمونه‌گیری را مشخص می‌سازد. به عنوان مثال در همان مثال قبلی محقق باید نام و آدرس تمام مدارس شهر تهران را داشته باشد چرا که هر کدام از این مدارس شانس انتخاب شدن دارند.



تصویر ۸-۲

تصویر ۷-۲

✓ نکته: گاهی در عمل، محقق برخی از خوشه‌ها را نمی‌شناسد؛ مثلاً نام و آدرس برخی از مدارس را ندارد و یا گاهی امکان نمونه‌گیری از خوشه خاصی را ندارد در این حالت برخی خوشه‌ها عملاً شانس انتخاب شدن ندارند و محقق در واقع به انتخاب نمونه‌ها از خوشه‌های در دسترس می‌پردازد که این موضوع نمونه‌گیری را از حالت تصادفی خارج کرده و از اعتبار آن می‌کاهد. مثلاً اگر هر بیمارستان دولتی شهر تهران یک خوشه باشد، دانشجوی دانشگاه علوم پزشکی ایران ترجیح می‌دهد از بیمارستانهای دانشگاه خودش نمونه بگیرد تا بیمارستانهای دانشگاههای دیگر، بخصوص که در اغلب موارد اجازه نمونه‌گیری از آن بیمارستانها را هم به او نمی‌دهند. دقت کنید که اگر پراکندگی درون خوشه‌ای بسیار بالا باشد و پراکندگی بین خوشه‌ای بسیار پایین باشد نکته ذکر شده، اشکال کمتری ایجاد می‌کند. چرا که هر خوشه حال از هر کجا که باشد تا حد مناسبی معرف جامعه است.

روشهای نمونه‌گیری غیر احتمالی (Nonprobability Sampling)

همانگونه که گفته شد شانس افراد جامعه آماری برای ورود به نمونه (Sample) در روشهای غیر احتمالی معین نمی‌باشد و به همین دلیل در این روش نمونه‌گیری، نتایج حاصله قابل تعمیم به جامعه هدف نیست. روشهای غیر احتمالی خود به دو دسته عمده آسان (Convenience) و هدفدار (Purposive) تقسیم می‌شوند که در ذیل به شرح آن می‌پردازیم.

آسان یا اتفاقی (Accidental or Convenience or Availability or Haphazard): محقق هیچگونه ایده‌ای را در انتخاب نمونه‌ها به کار نمی‌برد. فقط از میان افراد در دسترس به تعداد حجم نمونه انتخاب می‌کند. مثلاً در درمانگاه یا مطب از میان اولین مراجعاتی که شرایط ورود به نمونه را دارند (معیارهای ورود را پر می‌کنند) به تعداد حجم نمونه، نمونه‌گیری می‌نماید و وقتی حجم نمونه کامل شد به مراجعان بعدی که ممکن است قابلیت ورود به نمونه را نیز داشته باشند توجهی نمی‌کند. مثال: مطالعه‌ای با عنوان «بررسی افسردگی ۳ ماهه اول دوران بارداری در زنان مراجعه کننده به درمانگاه پره‌ناتال بیمارستان قاطعیه همدان در سال ۱۳۷۹» با حجم نمونه ۲۰۰ نفر با روش نمونه‌گیری غیر احتمالی آسان انجام شده است. بدین شکل که محققین پرسشنامه استاندارد افسردگی (Beck) را به اولین ۲۰۰ نفر مراجعه کننده به بیمارستان مزبور داده بودند و با افرادی که پس از پر شدن حجم نمونه‌شان مراجعه می‌کردند کاری نداشتند.

نمونه‌گیری متوالی (Consecutive Sampling)

حالتی از نمونه‌گیری آسان است و در این روش محقق سعی می‌کند تا آنجا که از لحاظ هزینه و نیروی انسانی در توان اوست از میان نمونه‌های در دسترس به انتخاب بپردازد و به حجم نمونه محاسبه شده توجهی ندارد به عنوان مثال در مثال فوق، حتی پس از پر شدن حجم نمونه باز هم از میان مراجعان به درمانگاه یا مطب تا زمانی که هزینه‌های تحقیق و نیروی انسانی و زمان اجازه دهد به نمونه‌گیری ادامه خواهد داد.

✓ نکته: توجه به این نکته لازم است که اگر نمونه‌گیری متوالی آنقدر ادامه پیدا کند که مقدار قابل توجهی از جامعه آماری را فرا بگیرد، خیلی به سرشماری (Census) نزدیک می‌شود و در یک نگاه کلی این روش از میان روشهای نمونه‌گیری غیر احتمالی دقیق‌ترین روش است و حتی می‌تواند نمونه‌ای معرف (Representative) جامعه هدف در اختیار محقق قرار دهد.



نمونه‌گیری هدفدار (Purposive Sampling):

محقق برای ورود واحدهای نمونه‌گیری به نمونه‌اش، ویژگی خاصی از آنها را در نظر گرفته و به پیاده کردن ایده‌ای می‌پردازد. روشهای نمونه‌گیری هدفدار فراوانند و از تنوع بالایی برخوردار می‌باشند. در اینجا به بررسی پرکاربردترین‌های آنها در طرحهای بهداشتی و تحقیقات علوم اجتماعی می‌پردازیم.

۱- نمونه‌گیری مورد شایع (Modal Instance, Typical Case Sampling)

نمونه‌گیری از افرادی صورت خواهد گرفت که از نظر صفتی خاص که در مطالعه اهمیت دارد، در نما جامعه قرار می‌گیرند. یعنی بیشترین فراوانی را دارند. مثلاً در مطالعه‌ای با عنوان "بررسی علل عدم گرایش زنان به استفاده از وسایل تنظیم خانواده" نمونه‌ها را از میان زنانی با خانواده‌های متوسط از نظر اجتماعی - اقتصادی انتخاب می‌کنند. این نوع نمونه‌گیری بخصوص در آگاهی‌سنجی‌ها کاربرد دارد و افرادی که از لحاظ سن، درآمد، سطح تحصیلات، ... در نما جامعه قرار دارند وارد نمونه می‌شوند.

۲- نمونه‌گیری ناهمگونی (Dimensional or Heterogeneity or Maximum Variation Sampling)

اغلب در حجم نمونه‌های پایین انجام می‌شود و اساس کار بر نمونه‌گیری از تمامی افرادی است که احتمال تأثیرگذاری آنها بر نتیجه مطالعه می‌رود. به عبارتی، از طیف وسیعی از افراد نمونه‌گیری می‌کنند. برای مثال در مطالعه فوق نمونه‌ها را از میان زنانی با خانواده‌هایی که از نظر سطح اجتماعی - اقتصادی خیلی بالا، بالا، متوسط، پایین و خیلی پایین می‌باشند انتخاب می‌نمایند و سعی می‌کنند که از هر دسته تعدادی از افراد را وارد مطالعه کرده و در نمونه داشته باشند یا مثال دیگر اینکه در یک نظرسنجی از دانشجویان یک دانشگاه بین‌المللی سعی می‌شود که از تمامی ملیتهای موجود در دانشگاه دست کم یک مورد وارد نمونه گردد.

۳- نمونه‌گیری همگونی (Homogenous Sampling)

محقق گاهی در اجرای روش قبل دچار مشکل می‌شود و ناچار به این روش روی می‌آورد به عنوان مثال در مثال بالا به جای نظر سنجی از تمامی ملیتهای موجود در دانشگاه تنها به بررسی یک ملیت می‌پردازد این روش اجرای نظرسنجی و آنالیز اطلاعات را آسان می‌کند ولی روشن است که نسبت به روش قبل دقت کمتری دارد و تنها مزیت آن سادگی‌اش است. ✓ نکته: نمونه‌گیری مورد شایع (Modal Instance Sampling) نوعی نمونه‌گیری همگونی است.

۴- نمونه‌گیری مورد انتهایی (Extreme and Deviant Case Sampling)

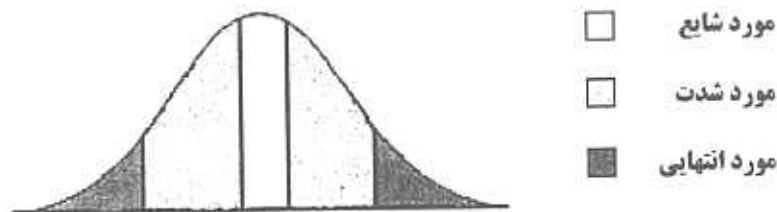
مطالعه افرادی است که از لحاظ ویژگی مورد نظر محقق بسیار غیر معمول و یا به عبارت دیگر در دو سر طیف می‌باشند. به عنوان مثال در مطالعه‌ای با عنوان "فاکتورهای دخیل در روند تحقیقات دانشجویی و تأثیر این عوامل بر نگرش دانشجویان پزشکی" افرادی که در تحقیقات دانشجویی به موفقیت خاص و کمیابی دست پیدا کرده‌اند یا افرادی که شکستهای کم نظیری داشته‌اند وارد نمونه می‌شوند. این نوع نمونه‌گیری به محقق کمک می‌کند که جامعه را بهتر بشناسد و محقق باید دقت کند که نمونه به دست آمده معرف جامعه نمی‌باشد.

۵- نمونه‌گیری شدت (Intensity Sampling)

مطالعه افرادی است که ویژگی مورد نظر را با شدت بالایی نشان می‌دهند مثل نمونه‌گیری توأم از دانش‌آموزان با سطح علمی بالای متوسط و پایین متوسط. ذکر این نکته لازم است که اصولاً افرادی ویژگی را با شدت بالا نشان می‌دهند که در دو طرف میانگین قرار می‌گیرند و گرچه در داشتن فا داشتن صفت مورد نظر شدت دارند اما در دو سر انتهایی طیف نمی‌باشند.



✓ نکته: برای افتراق بین سه روش نمونه‌گیری مورد شایع، مورد انتهایی، و شدت به تصویر ۹-۲ توجه کنید:



تصویر ۹-۲

۶- نمونه‌گیری مورد بحرانی (Critical Case Sampling)

نمونه‌گیری از موارد بحرانی است و اساس این نمونه‌گیری بر آن است که اگر مطالعه برای موارد بحرانی و حساس جواب دهد برای موارد ساده‌تر و معمولی نیز جواب خواهد داد. به عنوان مثال اگر دانش‌آموزی با ضریب هوشی (IQ) پایین بتواند آزمونی را بگذراند دانش‌آموزان با ضریب هوشی بالاتر نیز می‌توانند این آزمون را بگذرانند یا اگر یک روش درمانی روی بیماران با مراحل بالای بیماری مؤثر باشد روی بیماران با مراحل پایین‌تر نیز مؤثر خواهد بود. به عنوان مثال در مطالعه‌ای که اثر یک داروی جدید در کنترل سرطانی خاصی را بررسی می‌کنند، چون افرادی که مراحل اولیه بیماری را می‌گذرانند و امید زیادی به درمان با روشهای رایج دارند، حاضر به مصرف داروی جدید نمی‌شوند، ناچار دارو روی بیماران که مراحل پیشرفته بیماری را دارند و امید زیادی به بهبودی با روشهای قبلی ندارند، مورد مطالعه قرار خواهند گرفت. اگر دارو روی این بیماران جواب دهد، به احتمال قوی روی بیماران که در مراحل اولیه بیماری هستند نیز جواب خواهد داد.

۷- نمونه‌گیری مورد تأییدکننده و ردکننده (Confirming and Disconfirming Case Sampling)

محقق مطالعه‌ای را روی افراد با ویژگی خاصی انجام داده و نتیجه‌ای نیز گرفته است؛ حال جهت بیشتر کردن وسعت تعمیم پذیری مطالعه، از استثناهای و افرادی که سابقاً در مطالعه قرار نمی‌گرفتند نیز نمونه‌گیری کرده و به بررسی آنها می‌پردازد. در واقع نمونه‌گیری از افرادی است که نتیجه مطالعه انجام شده، در مورد آنها روشن نیست و محقق قصد مشخص کردن آن را دارد. برای مثال مطالعه‌ای با عنوان "بررسی علل عدم گرایش زنان به استفاده از وسایل تنظیم خانواده" به روش نمونه‌گیری مورد شایع انجام شده است؛ بدین شکل که از زنانی با سطح تحصیلات متوسط نمونه‌گیری کرده‌اند. حال محقق برای گسترش دایره تعمیم‌پذیری مطالعه، این طرح را تکرار می‌کند و این بار از زنانی که سطح تحصیلاتی بالاتر یا پایین‌تر از متوسط دارند نمونه می‌گیرد و نتیجه مطالعه قبلی را در این نمونه‌ها بررسی می‌کند.



۸- نمونه‌گیری فرصت‌طلب (Opportunistic Sampling)

استفاده از موقعیتهای پیش آمده در طول مطالعه است. بسیار پیش می‌آید که حین انجام یک مطالعه، ایده‌های جدید به ذهن محقق برسد یا فرصت برای انجام یک کار جانبی فراهم شود. در این موارد محقق دست به نمونه‌گیری فرصت‌طلب می‌زند؛ یعنی از شرایط ایجاد شده استفاده می‌کند و به یک کار موازی با تحقیق اصلی نیز می‌پردازد. برای مثال در مطالعه‌ای که با عنوان "مقایسه پماد جلدی پیروکسیکام و اسانس سالویا در بهبود استئوآرتریت زانو در کمیتة تحقیقات دانشجویی دانشگاه علوم پزشکی ایران انجام شد، افرادی وارد تحقیق می‌شدند که دارای استئوآرتریت خفیف (Mild) یا متوسط (Moderate) بودند. در تعیین شدت بیماری، ناچار افرادی که استئوآرتریت شدید (Severe) داشتند نیز شناسایی می‌شدند. محققین از فرصت استفاده کرده و فرم جمع‌آوری اطلاعات خود را برای این افراد که وارد نمونه طرح مورد نظر نیز نمی‌شدند پر کردند تا برآوردی از شدت استئوآرتریت در مراجعان به آن درمانگاه داشته باشند.

۹- نمونه‌گیری داوطلبی (Volunteer Sampling)

نمونه‌گرفتن افرادی است که خود داوطلب هستند. این روش نمونه‌گیری بیشتر در مطالعات کارآزمایی بالینی (Clinical Trial) به کار می‌رود و با اینکه بسیار خطا دارد ولی دقت بالای کارآزمایی‌های بالینی تا حد قابل قبولی این خطا را جبران می‌کند. مثال بارز آن افرادی هستند که داوطلب درمان با یک داروی جدید می‌شوند.

۱۰- ۱۱- نمونه‌گیری سهمیه‌ای (Quota Sampling) و نمونه‌گیری قضاوتی (Judgmental Sampling)

در این دو روش ابتدا محقق جامعه مورد مطالعه را بر اساس صفتی که اهمیت خاصی در مطالعه داشته باشد به چند گروه تقسیم می‌کند سپس از هر گروه به نسبت جمعیت آن وارد نمونه می‌نماید. تفاوت این دو روش در این است که در نمونه‌گیری سهمیه‌ای، نسبت جمعیت گروه‌ها ثابت شده است و در بررسی متون به آن رسیده‌اند ولی در نمونه‌گیری قضاوتی، این نسبتها صرفاً زائیده ذهن محقق هستند. مثلاً محقق در نمونه‌گیری سهمیه‌ای می‌داند که نسبت جمعیت گروه‌ها ۲:۲:۶ است، ولی در نمونه‌گیری قضاوتی این نسبت را حدس می‌زند و اطمینان ندارد.

✓ نکته: این دو روش و خصوصاً نمونه‌گیری سهمیه‌ای بسیار شبیه به روش نمونه‌گیری تصادفی طبقه‌بندی شده (Stratified Random Sampling) هستند که از روشهای نمونه‌گیری احتمالی می‌باشد، ولی تفاوت در اینجاست که در نمونه‌گیری تصادفی طبقه‌بندی شده از هر گروه به روش احتمالی مثلاً تصادفی ساده (Simple Random) نمونه انتخاب می‌گردد ولی در روشهای سهمیه‌ای و قضاوتی از هر گروه به روش غیر احتمالی و اغلب آسان (Convenience) نمونه‌گیری صورت می‌گیرد.

۱۲- نمونه‌گیری گلوله برفی (Snowball or Chain Sampling)

رایجترین روش نمونه‌گیری از جمعیت‌هایی است که جایگاه اجتماعی - فرهنگی خوبی ندارند و یا جمعیت‌هایی که به دلایلی حاضر به همکاری با محقق نمی‌باشند. مثلاً افرادی که دچار انحرافات اخلاقی هستند، افرادی که به تهیه و توزیع مواد مخدر می‌پردازند، معتادین، مجرمین و حتی بیمارانی که به دلیل درک نادرست جامعه گوشه‌گیر شده‌اند؛ مثل بیماران مبتلا به ایدز.



در این روش، محقق ابتدا یک یا چند تن از افراد جمعیت مورد نظر را پیدا کرده، آنها را توجیه می‌نماید و لزوم انجام طرح را برایشان روشن می‌کند و بعد، از آنها می‌خواهد که فرم جمع‌آوری اطلاعات یا پرسشنامه را به درون جمعیت برده و به افراد دیگر جمعیت مورد نظر که می‌شناسد برسانند و از آنها نیز بخواهند که این رنجیره را ادامه دهند. پس از مدتی فرمها را همان افراد اولیه جمع‌آوری کرده و به محقق بر می‌گردانند. در این روش محقق با نمونه‌هایش برخورد نزدیک ندارد.

✓ نکته: ایراد اصلی این روش نمونه‌گیری، نمونه‌های از دست رفته (Loss) فراوان آن است؛ چرا که بسیاری از افراد اینگونه جمعیتها حاضر به همکاری حتی با کسی مثل خودشان نیستند.

۱۳- نمونه‌گیری فرد ماهر (Expert Sampling)

نمونه‌گیری از افرادی است که در مورد خاصی تخصص دارند و اغلب به دو منظور انجام می‌شود:

۱- گاهی مطالعه‌ای انجام یافته است و جهت ارزیابی مطالعه و یا احتمالاً سند قرار دادن نظر متخصصین، به این نمونه‌گیری پرداخته می‌شود. مثلاً محقق به تأثیرات مثبت یک رژیم غذایی در درمان دیابت دست یافته است، حال از متخصصان نیز نظر می‌خواهد، شاید آنها این رژیم غذایی را برای بیماران خود تجویز می‌کرده‌اند. در مثال در طرحی که با عنوان "تعیین اعتبار و پایایی سؤالات امتحان زنان و زایمان دانشگاه علوم پزشکی ایران در مرداد ۱۳۸۰" انجام شد، محققین پس از انجام مطالعه و تعیین ارزش هر سؤال، نتایج حاصله را با نظر اساتید دانشگاه نیز مقایسه کرده و ایده آنان را نیز در ارزیابی مطالعه‌شان دخیل نمودند.

۲- گاهی محقق نمونه‌هایش را در اختیار ندارد و یا دسترسی به آنها برایش مشکل است، مثلاً معتادین نریقی. لذا به جای نمونه‌گیری از معتادین نریقی، به نمونه‌گیری از متخصصانی خواهد پرداخت که روی این افراد کار کرده‌اند؛ مثلاً پزشکان مستقر در درمانگاه‌های ترک اعتیاد.

روش نمونه‌گیری چند مرحله‌ای (Multi Stage)

استفاده از چند روش نمونه‌گیری نوامی می‌باشد که این روشها می‌توانند همگی احتمالی یا همگی غیر احتمالی و یا مخلوطی از روشهای احتمالی و غیر احتمالی با هم باشند. برای مثال در مطالعه‌ای با عنوان "تعیین میانگین قد دانش‌آموزان کلاس اول مدارس شهر تهران" می‌توان یک نمونه‌گیری چند مرحله‌ای را طراحی و اجرا کرد: بدین صورت که محقق ابتدا شهر تهران را براساس موقعیت اجتماعی - اقتصادی خانواده‌ها به چهار طبقه (Stratum) شمال، جنوب، شرق و غرب تقسیم کرده، سپس در هر منطقه، هر مدرسه را یک خوشه (Cluster) قرار داده، از میان خوشه‌ها به روش تصادفی ساده (Simple Random) یک یا چند خوشه را انتخاب می‌نماید. در مدارس انتخاب شده نیز مجدداً می‌توان هر کلاس را یک خوشه قرار داده باز به روش تصادفی ساده برخی خوشه‌ها را انتخاب کنند و در داخل هر کلاس با استفاده از دفتر حضور و غیاب آنها، که قالب نمونه‌گیری خواهد بود به روش تصادفی ساده یا منظم دست به انتخاب نمونه‌تهایی بزنند. توجه دارید که در این مثال تمام روشها احتمالی بوده‌اند.

تصادفی طبقه‌بندی شده ← خوشه‌ای ← تصادفی ساده ← خوشه‌ای ← تصادفی ساده ← تصادفی ساده یا منظم

✓ نکته: نمونه‌گیری در بیشتر طرحهایی که در سطح جمعیت‌های بزرگ انجام می‌شوند به صورت چند مرحله‌ای است.

✓ نکته: تقسیم حجم نمونه در مراحل مختلف باید رعایت شود.

✓ نکته: نمونه‌های غیر احتمالی لزوماً معروف جامعه نمی‌باشند ولی باید سعی شود با طراحی مناسب، تا حد امکان آنها را معروف و دقیقتر انتخاب کرد. یکی از این روشها، وارد کردن روشهای احتمالی در نمونه‌گیری چند مرحله‌ای است یعنی در سلسله مراتب نمونه‌گیری چند مرحله‌ای، یک یا چند مرحله تصادفی وارد می‌شود، اگر چه بیشتر مراحل غیر احتمالی باشند.

تورش انتخاب

Selection Bias

✓ مفهوم تورش
✓ تورش انتخاب

فصل

۳

تورش، اشتباه سیستماتیک است که نتایج را بطور منظم و در جهتی خاص تغییر داده، منجر به تخمین نادرستی از رابطه بین عامل مواجهه و پیامد می‌شود و حتی با تکرار مطالعه هم این تخمین نادرست جهت دار بدست می‌آید. ذکر این نکته لازم است که جهت دار بودن اشتباه در تورش، آن را از خطای تصادفی که در هر نمونه‌ای روی می‌دهد و در صورت تکرار مطالعه به همان شکل ایجاد نمی‌شود، متمایز می‌سازد. تورش را خطای اندازه‌گیری (Measurement Error) یا خطای سیستماتیک (Systematic Error) نیز می‌نامند تا بتوان تورشها را از خطای تصادفی (Random Error) با تغییرات تصادفی (Random Variation) تمیز داد.

امکان ایجاد تورش در هر یک از مراحل تحقیق وجود دارد که در ذیل به آن اشاره خواهد شد:

- مرحله مطالعه نظری: بعنوان مثال اطلاعات ناکامل، ناقص یا اشتباه از منابع استخراج شود.
- مرحله انتخاب نمونه: در این بخش، به تفصیل به بررسی این تورشها خواهیم پرداخت.
- مرحله انجام طرح آزمایشگاهی: اشتباه در مراحل انجام طرح نیز می‌تواند منجر به ایجاد تورش شود.
- مرحله اندازه‌گیری عامل مواجهه و پیامد: تعاریف نادرست یا متفاوت از عوامل مواجهه و نیز پیامد، همبطور اشتباه در اندازه‌گیری آنها، منجر به ایجاد تورش می‌گردد.
- آنالیز اطلاعات بدست آمده: به عنوان مثال سوگیری که آنالیز کننده یک تحقیق در اثبات ارتباط بین پیامد و عامل مواجهه تورش ایجاد می‌نماید.
- تفسیر داده‌ها: تورش در مرحله تفسیر داده‌ها نیز می‌تواند ایجاد شود که در برخی موارد، از روش کورکردن (Blinding) برای رفع آن استفاده می‌شود.
- انتشار نتایج: گاهی محقق قسمتی از نتایج را که به نفع فرضیه ذهنی از پیش تعیین شده است، منتشر می‌کند.

انواع تورش

تورشها را به سه دسته کلی انتخاب، اطلاعات و مخدوش‌کننده تقسیم می‌کنند که در ذیل به شرح آن می‌پردازیم.

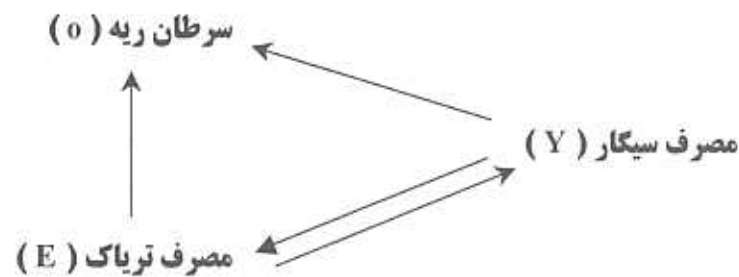
- تورش انتخاب (Selection Bias): در مرحله انتخاب افراد مورد پژوهش بوجود می‌آید و نتیجه آن در مطالعات تحلیلی تخمین نادرست از رابطه بین عامل مواجهه و پیامد و در مطالعات توصیفی بر آورد جهت دار غیر صحیح از متغیر مورد سنجش می‌باشد.
- تورش اطلاعات (Information Bias): این تورش به جمع‌آوری نادرست اطلاعات مربوط می‌شود و اسامی دیگری مثل تورش اندازه‌گیری (Measurement Bias)، تورش دسته‌بندی (Classification Bias) و تورش مشاهده (Observation Bias) نیز دارد و بطور کلاسیک به سه دسته طبقه‌بندی می‌گردد.

-تورش ناشی از فرد مورد مطالعه (Subject)

-تورش ناشی از مصاحبه‌گر (Interviewer)

-تورش ناشی از دستگاه اندازه‌گیری (Instrument)

- تورش مخدوش‌کنندگی (Confounding Bias): این نوع تورش وجود یک فاکتور خارجی که رابطه بین مواجهه و پیامد مورد نظر مطالعه را مخدوش می‌کند، مطرح می‌نماید و بدین ترتیب رابطه واقعی را مخدوش می‌سازد. به عبارت دیگر رابطه مشاهده شده بین عامل مواجهه و پیامد، در واقع مربوط به اثر سایر متغیرها خواهد بود و عدم مشاهده یک رابطه نیز ممکن است مربوط به عدم کنترل چنین متغیرهایی باشد. متغیری مخدوش‌کننده است که رابطه‌ای با عامل مواجهه داشته باشد، عامل خطر برای پیامد به شمار آید و در عین حال یک فاکتور واسطه‌ای نباشد (تصویر ۱-۳).



تصویر ۱-۳

در نمودار فوق، Y فاکتور مخدوش‌کننده محسوب می‌شود.

با توجه به اینکه تورش انتخاب معمولاً در ضمن نمونه‌گیری هر تحقیق اتفاق می‌افتد، در این بخش به بررسی انواع تورشهای انتخاب می‌پردازیم.

انواع تورش انتخاب

۱- تورش معروفیت (Popularity Bias)

این تورش زمانی بوجود می‌آید که یک موهبه یا روش درمانی خاص (اتریشی، جراحی و...) باعث ایجاد اشتیاق بیشتر در استفاده از آنها شود. بعنوان مثال، بیماران مثلاً به آلرژی، بعلمت معروفیت کلینیک آلرژی، بیشتر به این کلینیک مراجعه می‌کنند. اگر محقق در نمونه‌گیری، به این مسأله توجه نکرده و از چنین کلینیکی نمونه‌گیری انجام دهد نمونه بدست آمده، نمایانگر جامعه نخواهد بود و دچار تورش گردیده است.

۲- تورش متمایل به مرکز (Centripetal Bias)

این تورش زمانی رخ می‌دهد که به یک پزشک یا مرکز درمانی خاص، افرادی یا بیماری‌ها یا مواجهه معین مراجعه می‌نمایند؛ نمونه‌ای که از این مرکز گرفته می‌شود، نمایانگر جامعه نخواهد بود. به عنوان مثال اگر برای نمونه‌گیری در بررسی اپیدمیولوژی درد مزمن،



قرار باشد به کلینیک های درمانی مراجعه شود، مراجعه به کلینیک درد (که بیماران دچار درد مزمن به آن مراجعه می کنند)، طرح را دچار تورش خواهد کرد.

۳- تورش فیلتر ارجاع (Referral filter Bias)

وقتی که گروهی از بیماران اولیه به مراقبت های سطح دوم (مراقبت در مراکزی که بعنوان نخستین سطح ارجاع بکار می روند) و سطح سوم (مراقبت های بیمار تخصصی در مراکز ویژه و بیمارستانها) ارجاع می یابند، تعداد موارد نادر، تشخیص های چندگانه و موارد منجر به فوت، افزایش می یابد. در این حالت اگر نمونه از بین افرادی انتخاب شود که به سطح سوم مراقبت ها ارجاع شده اند (مثلاً در بیمارستان بستری شده اند) ممکن است شدت بیماری وخیم تر از واقعیت نشان داده شود و چنانچه این نوع نمونه گیری در مطالعات تحلیلی تنها در یکی از دو گروه مورد و شاهد انجام گیرد، احتمال برخورد با عامل مواجهه در این دو گروه متفاوت خواهد بود. بعنوان مثال در طرح بررسی ارتباط دیابت وابسته به انسولین و استرس های زندگی* اگر محقق بیماران دیابتی را از مراجعین به ICU بیمارستان انتخاب کند (به این معنا که این افراد مبتلا به بیماری شدیدتری بوده و استرس های زیادی تحمل کرده اند) و در مقابل، گروه شاهد را از افراد عادی انتخاب نماید، ممکن است به اشتباه میزان استرس را در افراد دیابتی که خیلی از آنها هم تا به حال در ICU بستری نشده اند، بیشتر از افراد جامعه گزارش کند.

۴- تورش تردید تشخیص (Diagnostic Suspicion Bias)

داشتن اطلاعاتی در مورد مواجهه نمونه با یک علت احتمالی (نژاد، دریافت یک داروی خاص، داشتن ناراحتی ثانویه، قرار گرفتن در یک اپیدمی) هم بر شدت و هم بر نتیجه فرآیند تشخیص، اثر می گذارد. این تورش بیشتر در مطالعات همگروهی رخ می دهد. به عنوان مثال، در تحقیق تعیین ارتباط مواجهه آریست و ایجاد پلاک های پلورال در عکس قفسه سینه* با توجه به این که تفسیر عکس ریه می تواند بر حسب داشتن یا نداشتن اطلاعات کلینیکی از بیمار تغییر کند لذا دالشن وضعیت مواجهه نمونه می تواند بر تخمین پیامد موثر باشد و پلاک های پلورال در افراد مواجهه یافته با آریست بیشتر تشخیص داده شوند. با بی اطلاع نگه داشتن محقق از شرایط مواجهه، می توان این خطا را به حداقل رسانید و یا می توان تست پاراکلینیکی را در دو حالت دانستن و ندانستن شرایط کلینیکی بیمار تفسیر و سپس آن دو را مقایسه کرد.

۵- تورش دسترسی به تشخیص (Diagnostic Access Bias)

این تورش زمانی رخ می دهد که افراد از نظر جغرافیایی، زمانی یا اقتصادی، دسترسی متفاوتی به روش های تشخیص داشته باشند. در چنین حالتی، دسترسی متفاوت افراد به امکانات تشخیص، باعث ایجاد تفاوت در تعیین فراوانی و شدت بیماری یا عامل مواجهه می گردد. به عنوان مثال چنانچه محقق قصد برآوردی از شیوع آترواسکلروز در دو شهر متفاوت را داشته باشد و در یک شهر امکان استفاده از آنژیوگرافی موجود باشد، واضح است که برآورد بیماری های عروقی کرونر در شهر واجد آنژیوگرافی بالاتر از شهر فاقد این دستگاه می باشد.

۶- تورش رواج تشخیص (Diagnostic Vague Bias)

این تورش زمانی ایجاد می شود که بیماری یکسانی در زمانها و مکانهای مختلف، متفاوت تشخیص داده شود. به عنوان مثال اگر بیماری ای که در انگلستان بروشیت تشخیص داده می شود در امریکای شمالی، آمفیزم تشخیص داده شود، نتایج به دست آمده در هریک از این دو کشور، قابل تعمیم به دیگری نیست. این تورش در بیماری هایی که تظاهرات شبیه به هم دارند و یا در مراکز درمانی مختلف با شاخص های متفاوتی تشخیص داده می شوند ایجاد می شود.

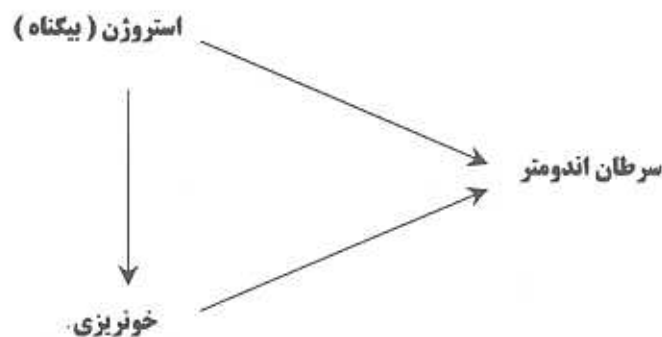


۷- تورش خالص بودن تشخیص (Diagnostic purity Bias)

زمانیکه گروه های خالص تشخیصی مانع از همراهی با مرگ و میر شوند، در این صورت نمایانگر جامعه هدف نخواهد بود. به عنوان مثال در بررسی ایجاد افسردگی و مدت بستری در بیمارستان، اگر نمونه ها از افرادی انتخاب شوند که بیشتر از دو هفته در بیمارستان بستری بوده اند، ممکن است تعدادی از افراد که قبل از دو هفته فوت شده اند (شاید به علت افسردگی)، نادیده گرفته شوند.

۸- تورش آشکار سازی (Unmasking or Detection Signal Bias)

یک مواجهه بی ضرر (Innocent Exposure)، اگر علاوه بر ایجاد بیماری، علامت یا نشانه ای ایجاد کند که منجر به تحقیقات و جستجو برای وجود بیماری شود، مورد ظن قرار می گیرد. این تورش بیشتر در مطالعات با جهت رو به جلو (همگروهی و کارآزمایی بالینی) مشاهده می شود. به عنوان مثال در تحقیق بررسی ارتباط مصرف استروژن بعد از یائسگی و سرطان اندومتر*، استروژن، مواجهه بی ضرر است که باعث خونریزی می شود و علامتی است که پزشک را به سرطان اندومتر رهنمون می نماید. در این مثال، گروه مورد، استروژن مصرف کرده و بعلاوه خونریزی ایجاد شده محقق بیشتر به دنبال یافتن سرطان اندومتر خواهد بود اما در گروه کنترل که استروژن مصرف نمی کنند و خونریزی هم ندارند، دقت لازم در بررسی سرطان اندومتر بعمل نمی آید (تصویر ۲-۳). ذکر این نکته لازم است که در همه مطالعات باید احتمال وجود چنین عاملی مدنظر باشد.



تصویر ۲-۳

۹- تورش تقلید (Mimicry Bias)

یک مواجهه بی ضرر (Innocent Exposure)، اگر علاوه بر ایجاد بیماری، ناراحتی ضعیفی هم ایجاد کند که با بیماری شباهت داشته باشد، مورد ظن قرار می گیرد. بعنوان مثال زایمان زودرس هم می تواند باعث ایجاد سندرم زجر تنفسی و هم عفونت استرپتوکوک شود و علائم این دو بیماری مشابه یکدیگرند (تصویر ۳-۳).

زایمان زودرس (E^+) بیگناه

علائم مشابه

سندرم زجر تنفسی

عقوت استریتوکوکی

تصویر ۳-۳

۱۰- تورش عقیده قبلی (Previous Opinion Bias)

اگر محقق از شیوه ها و نتایج تشخیص قبلی بیمار آگاه باشد، این آگاهی ممکن است بر روش و نتایج تشخیص بعدی همان بیمار تأثیر بگذارد. مثال بارز آن، این است که در گرفتن فشار خون افراد مسن، مرحله کاهش صدا را فشار دیاستولیک قرار می دهند و در افراد جوان یا ورزشکار نقطه قطع صدا را فشار دیاستولیک قرار می دهند.

۱۱- تورش اندازه نامناسب نمونه (Wrong Sample Size Bias)

در حجم نمونه بسیار کوچک هیچ چیز را نمی توان ثابت کرد در حالیکه در حجم نمونه بسیار بزرگ، همه چیز قابل اثبات است. لذا در حجمهای بسیار زیاد نمونه، وجود ارتباط می تواند به دلیل حساسیت بیش از حد مطالعه به کشف ارتباط هر چند ضعیف باشد. بنابراین تعیین حجم نمونه مناسب برای جلوگیری از این تورش، ضروری است.

۱۲- تورش میزان پذیرش (Admission Rate or Berkson Bias)

این تورش زمانی رخ می دهد که میزان پذیرش متفاوتی در مطالعه وجود داشته باشد. برای مثال، زمانی که بیماران (موارد)، دفعات بیشتری نسبت به افراد شاهد در بیمارستان پذیرفته شوند به عبارت دیگر، علت این تورش، میزانهای پذیرش متفاوت برای موارد بیماری و افراد شاهد است. این سوگیری، اولین بار در جریان ارزیابی یک مطالعه قبلی صورت گرفت که در آن نتیجه گرفته شده بود که سل می تواند اثر محافظتی بر سرطان ریه داشته باشد. این استنتاج از یک مطالعه مورد-شاهدی بدست آمده بود که در جریان آن ارتباط منفی بین سل و سرطان ریه مشخص شده بود. در این مطالعه، فراوانی سل در بیماران سرطانی بیشتری کمتر از فراوانی سل در افراد شاهد بستری بدون ابتلا به سرطان بود. علت وجود این نتایج، این بود که نسبت کوچکتری از بیماران مبتلا به هر دو بیماری سرطان و سل، بستری می شوند و بدین ترتیب، نسبت کوچکتری برای انتخاب موارد در این مطالعه موجود است، چرا که احتمال مرگ بیماران مبتلا به هر دو بیماری یا هم، بیشتر از بیماران مبتلا به سل یا سرطان ریه به تنهایی می باشد.

تورش میزان پذیرش، مسأله مهمی است که باید به آن توجه شود زیرا بسیاری از مطالعات مورد-شاهدی از بیماران بستری در بیمارستان بعنوان منبع انتخاب موارد بیماری و افراد شاهد استفاده می کنند. یک راه کنترل این تورشها، استفاده از یک گروه شاهد بدون تورش است؛ بهترین روش انجام این کار، انتخاب افراد شاهد از طیف وسیعی



از بیماریهای مختلف یا از میان جمعیت افراد سالم می باشد. همچنین می توان با انتخاب موردها از منابع غیر از بیماران بستری در بخش این سوگیری را رفع کرد.

۱۳- تورش میزان شیوع یا میزان بروز نیمن (Prevalence-Incidence or Neyman or Survival Bias)

تورش شیوع هنگامی روی می دهد که بیماری با مرگ زودرس (تعدادی از بیماران قبل از تشخیص قوت می کنند) یا موارد خاموش (مواردی که هنگام شروع بیماری دلیل مبتلی بر در معرض عامل خطر بودن آنها از بین می رود)، مشخص می شود. هنگامی که فاصله زمانی خالی بین تماس با عامل مواجهه و انتخاب افراد برای مطالعه وجود داشته باشد و بدترین موارد مرده باشند، تورش شیوع رخ می دهد. در یک مطالعه همگروهی که پیش از آغاز بیماری شروع می شود، موارد ابتلا به طور صحیح قابل کشف می باشند. با این حال مطالعه مورد شاهدهی که در زمان دیرتری شروع می شود، فقط موارد باقی مانده، یعنی بیمارانی را شامل می شود که فوت نکرده اند. در این مطالعات می توان از طریق محدود کردن معیارهای لازم برای مطالعه موارد تازه تشخیص داده شده یا میزان بروز، از ارتکاب این تورش جلوگیری کرد.

برای نشان دادن این تورش فرض کنید که دو گروه از مردم در اختیار محقق باشد، یک گروه با عامل خطر برای یک بیماری مشخص (به عنوان مثال مصرف مواد مخدر به عنوان عامل خطر زخم پپتیک سوراخ شده) و گروه دیگر فاقد عامل خطر. در این مطالعه هزار بیمار مصرف کننده مواد مخدر و هزار نفر بدون مصرف مواد مخدر برای مدت ۱۰ سال پیگیری شده اند. نتایج حاصله از این مطالعه در جدول ۱-۳ خلاصه شده است:

زنده ماندن بدون ابتلا به زخم پپتیک سوراخ شده	فوت به علت زخم پپتیک سوراخ شده	زنده ماندن و ابتلا به زخم پپتیک سوراخ شده	
۷۰۰	۲۵۰	۵۰	مصرف مواد مخدر (E^+)
۹۰۰	۲۰	۸۰	عدم مصرف مواد مخدر (E^-)

جدول ۱-۳

از یافته های فوق بر می آید که در بیماران مصرف کننده مواد مخدر احتمال ابتلا به زخم پپتیک سوراخ شده بیشتر از افرادی است که فشارخون طبیعی دارند و احتمال مرگ آنها در اثر زخم پپتیک نیز بیشتر است. حال اگر چنین مطالعاتی به صورت مورد شاهدهی انجام شود، نتایج مشابه جدول ۲-۳ بدست می آید.



عدم مصرف مواد مخدر مصرف مواد مخدر

مبتلا به زخم پتیک سوراخ شده	۵۰	۸۰
عدم ابتلا به زخم پتیک سوراخ شده	۷۰۰	۹۰۰

جدول ۲-۳

نسبت شانس (Odds Ratio) در این مطالعه برابر است با $0.8 \times 700 / 50 \times 900 = 0.56$ و به نظر می رسد که مصرف مواد مخدر، یک عامل محافظت کننده در برابر ابتلا به زخم پتیک سوراخ شده می باشد. لذا این تورش بیشتر در مطالعات با جهت روبه عقب همچون مطالعات مورد شهادتی، اتفاق می افتد.

۱۴- تورش انتخاب روش درمانی (Precedure Selection Bias)

روشهای درمانی مشخص ممکن است برای افرادی که در معرض خطر کمتری هستند، ترجیح داده شود. به عنوان مثال در انتخاب بیماران برای درمان دارویی در مقابل درمان جراحی، افراد با شدت بیماری کمتر ترجیحاً با دارو درمان خواهند شد.

۱۵- تورش از دست دادن داده های کلینیکی (Missing Clinical Data Bias)

از دست دادن داده های کلینیکی ممکن است به علت تفاوت در طرز نوشتن اطلاعات در پرونده بیماران باشد. به عنوان مثال اطلاعاتی ممکن است در پرونده نوشته نشوند چون طبیعی هستند یا اندازه گیری شده اند یا اندازه گیری ولی گزارش نشده اند و این در حالیکه داده های غیر طبیعی حتماً نوشته خواهند شد و این امر باعث ایجاد سوگیری سیستمیک در داده ها می شود.

۱۶- تورش زمان شروع (Starting Time Bias)

اگر امکان تعیین دقیق زمان مواجهه یا شروع بیماری میسر نباشد، احتمال این تورش وجود دارد. به عنوان مثال، در تعیین بقا بیماران مبتلا به نورویلاستوما امکان تعیین دقیق شروع بیماری وجود ندارد و عملاً زمان تشخیص بیماری (که احتمالاً مدتی بعد از زمان شروع بیماری است) به عنوان زمان شروع در نظر گرفته می شود و در نتیجه تعیین میزان بقا دچار تورش می گردد.

۱۷- تورش شاهد غیر هم عصر (Historical Control Bias or Non-Contemporaneous Bias)

تغییرات دوره ای در تعاریف عوامل مواجهه، تشخیص، بیماری و درمان، باعث می شود که شاهد های غیر هم عصر قابل مقایسه نباشند به عنوان مثال می توان به بیماری لوپوس اشاره کرد که در طی سالهای متفاوت تعاریف گوناگونی جهت تشخیص آن شده است و در صورت استفاده از بیماران دوره های گذشته باید به این تفاوت تعاریف توجه شود. در مطالعات مورد شهادتی، چنانچه یکی از دو گروه مورد و شاهد از دوره های گذشته انتخاب شود باید به تعاریف این بیماری در آن زمان توجه نمود.

۱۸- تورش بیماری غیر قابل پذیرش (Unacceptable Disease Bias)



زمانی که یک بیماری، از جهات اخلاقی در جامعه پذیرفته شده نباشد (بیماری ابتدز)، میزانهای آن همواره کمتر از مقدار واقعی گزارش می شوند.

۱۹- تورش مهاجرت (Migration Bias)

اگر مهاجران در طرح شرکت داشته باشند، ممکن است به طور سیستمیک با افراد بومی متلفه، متفاوت باشند. به عنوان مثال، اگر در بررسی ارتباط سل و ابتلا به سرطان ریه تعداد مهاجران افغانی در گروه مورد بیش از گروه شاهد باشد، به علت شیوع بیشتر سل در مهاجران افغانی، ممکن است محقق به اشتباه، به ارتباط بین دو مورد فوق دست یابد.

۲۰- تورش عضویت (Membership Bias)

وجود این تورش، به علت وجود دسته ها و گروه هایی است که از قبل، تشکیل شده اند. علت دیگر این تورش آن است که خصوصیتی که سبب تعلق افراد به گروه ها می شود، یا پیامد مورد نظر ارتباط داشته باشد. برای مثال، پژوهشگران به خاطر رعایت اخلاق پژوهش قادر به انجام کار آزمایی بالینی برای بررسی اثرات استعمال سیگار بر سرطان ریه نبوده اند. لذا تعدادی از پژوهشگران ادعا نموده اند که تنها سیگار نمی باشد که سبب سرطان ریه می گردد بلکه عامل های دیگری که در افراد سیگاری شایع تر است، مسؤول این ارتباط می باشند. اطلاع از تورش عضویت برای خواننده مقالات پزشکی بسیار اهمیت دارد. زیرا از آن نمی توان پیشگیری کرد و این نوع تورش موجبات مشکل شدن مطالعه ارتباط اثر عامل های خطر موجود با نحوه زندگی می گردد. مسأله ای مشابه تورش عضویت، بنام اثر کارگر سالم (Healthy Worker Effect) وجود دارد. اثر کارگر سالم در همه گیری شناسی مطرح است و زمانی به وجود آمد که مشخص شد کارگران شاغل در یک محیط مخاطره آمیز، یا سربازان و دریانوردان، میزان بقا بالاتر از عموم مردم دارند، که البته علت این امر، وجود پیش شرط سلامتی در حد مطلوب برای استخدام این افراد بود.

۲۱- تورش عدم پاسخ (Refusal or Non-response Bias)

این تورش به علت خودداری یا عدم توانایی نمونه ها برای شرکت در مطالعات، یا عدم توانایی گروه تحقیق در دستیابی به نمونه ها به وجود می آید و بدین ترتیب، میزان مواجهه در افرادی که به دلیلی در مطالعه شرکت نکرده اند با موردها و شاهد هایی که در مطالعه شرکت می کنند متفاوت خواهد شد. به عنوان مثال در طرح تعیین ارتباط زخم پتیک سوراخ شده با مصرف الکل، اگر افراد اطلاعات دقیقی از مصرف الکل در اختیار محقق نگذارند، محقق به اشتباه به عدم ارتباط این دو می رسد. در چنین مواردی که میزان عدم شرکت بالاست در مرحله اول می بایست مقدمات و اطلاعاتی که راجع به موضوع مطالعه به شرکت کنندگان ارائه شده است مورد تجدید نظر قرار گیرد و مصاحبه کنندگانی که بیشترین عدم پاسخ را داشته اند، مجدداً آموزش داده شوند تا بتوانند به ترغیب افراد بی میل به شرکت در مطالعه بپردازند و اگر هیچکدام از روش های فوق مؤثر نبود باید اثر احتمالی این مشکل را ارزیابی نمود که این امر با آنالیز بدترین حالت (Worst Case Analysis) امکان پذیر است. در این روش، آنالیز با فرض وقوع بدترین حالت ممکن انجام می گیرد و در آن تمامی موردهایی که در مطالعه شرکت نکرده اند، مواجهه یافته و تمامی شاهد هایی که از شرکت خودداری کرده اند مواجهه نیافته در نظر گرفته می شوند حال اگر نتیجه بدست آمده باز هم عدم وجود رابطه بین عامل مواجهه و پیامد را نشان دهد بدین معنا خواهد بود که عدم پاسخ نمونه ها، اثری بر روی نتایج نداشته است ولی در غیر این صورت، باید نسبت به اعتبار نتایج مطالعه، محتاطانه برخورد کرد.

۲۲- تورش داوطلب (Volunteer Bias)



معمولاً افرادی که داوطلبانه در مطالعه شرکت می‌کنند ممکن است به طور سیستماتیک تفاوت‌هایی با سایر افراد شرکت کننده داشته باشند. به عنوان مثال افراد سرطانی که وضعیت و خیمتری دارند و به داروهای مختلف مقاومت نشان داده‌اند، انگیزه بیشتری در استفاده از داروهای جدید خواهند داشت که این امر، اثر دارو را دچار تورش می‌نماید.

۲۴- تورش وضعیت اقتصادی-اجتماعی (Socioeconomic Bias)

در صورت انتخاب موردهایی مبتلا به بیماریهای خفیف یا خاصی که فقط گروه ویژه‌ای از افراد جامعه مثل طبقه ثروتمند را درگیر می‌کند، نباید شاهد از طبقه پایین یا متوسط انتخاب شود، زیرا احتمالاً در این طبقات، مراجعه جهت درمان چنین بیماریهایی صورت نمی‌گیرد، لذا در چنین شرایطی باید حداکثر تشابه از نظر اقتصادی و اجتماعی، بین گروه مورد و شاهد، برقرار نمود. به عنوان مثال، در یک طرح مورد-شاهدی که به بررسی عوامل موثر در ایجاد ملاسمای پوستی پرداخته است، اغلب گروه بیماران مراجعه کننده به بیمارستان از طبقه متوسط و بالای جامعه بوده‌اند. اگر گروه کنترل از کل جامعه انتخاب شده باشند، ممکن است عامل خطر تعداد فرزندان در ایجاد ملاسمای کمتر از حد واقعی (Underestimate) برآورد شود. زیرا در گروه مورد، تعداد فرزندان کمتر بوده ولی به دلیل اهمیت دادن بیشتر به ملاسمای پوستی مراجعات فراوان‌تری از گروه شاهد که تعداد فرزند بیشتری داشته‌اند و کمتر مراجعه نموده‌اند، وجود داشته است.

۲۵- تورش ناشی از طول مدت اقامت در بیمارستان (Length of Hospital Stay Bias)

این تورش زمانی رخ می‌دهد که موردها بجای اینکه از لیست پذیرش یا ترخیص (Admission or Discharge lags) انتخاب شوند، از پرونده بیمارانی که در زمان مطالعه هنوز بستری هستند انتخاب گردند. در این حالت بیمارانی که مدت طولانی‌تری بستری بوده‌اند، و احتمالاً به نوع شدیدتر بیماری نیز مبتلا هستند، نسبت به بیمارانی که بیمارشان خفیفتر بوده و زودتر ترخیص شده‌اند، شانس بیشتری برای ورود به مطالعه دارند و از طرف دیگر، افرادی که فوت کرده‌اند نیز شانس انتخاب شدن نخواهند داشت.

۲۶- تورش مراقبت (Surveillance Bias)

در یک مطالعه مورد-شاهدی، احتمال دارد که عامل مواجهه با دقت بیشتری در گروه مورد بررسی گردد، و یا در مطالعات همگروهی، مراقبت بیشتری از گروه مواجهه یافته جهت کشف بیماری صورت گیرد. به عنوان مثال در بررسی ارتباط شیوع دیابت وابسته به انسولین و استرس زندگی که به صورت مطالعه مورد-شاهدی انجام گرفته است احتمال این وجود دارد که محقق در مبتلایان به دیابت وابسته به انسولین بیش از گروه شاهد به دنبال یافتن استرسهای قبلی بگردد. یکی از راههایی که می‌توان وجود این تورش را ارزیابی نمود، آنالیز جداگانه نمونه‌ها براساس فراوانی معاینات و پیگیریها می‌باشد. این تورش در مطالعات مقطعی-تحلیلی هم اتفاق می‌افتد.

توزیع احتمالات

Probability Distribution

✓ متغیر تصادفی

✓ توزیع نرمال و توزیع Z

✓ تئوری حد مرکزی

✓ سطح اطمینان

✓ خطای آلفا، بتا و توان آزمون

فصل

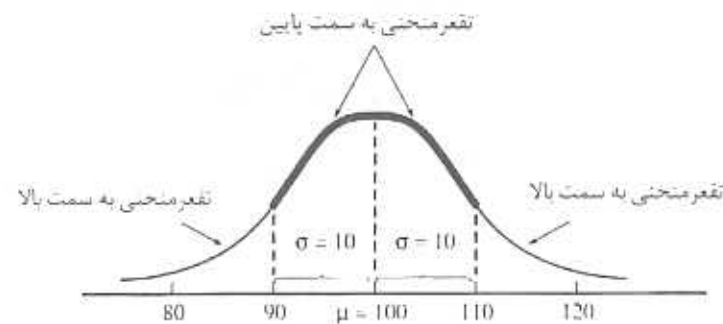
۴

به ویژگی یا خصوصیت مورد نظر در یک مطالعه که در افراد مورد پژوهش تغییر می‌کند متغیر (Variable) گویند. واژه متغیر یک واژه قابل درک می‌باشد زیرا مقدار یک صفت در افراد مختلف تغییر می‌کند و این تغییرات نتیجه تغییرات بیولوژیک موجود بین یکایک افراد یا اشتباهات ناشی از اندازه‌گیری یا ثبت آنها می‌باشد. یک متغیر تصادفی (Random Variable) شامل تغییری در مطالعه است که افراد آن به طور تصادفی انتخاب می‌شوند؛ مقادیر خصوصیت یا صفت مورد نظر می‌تواند به صورت توزیع فراوانی به سادگی رسم شود. بنابراین خلاصه کردن مقادیر یک متغیر تصادفی در یک توزیع فراوانی را توزیع احتمالات (Probability Distribution) گویند. به طور مثال اگر محقق بخواهد توزیع احتمالات فشار خون یک نمونه ۵۰ نفری را رسم نماید باید به ازای هر مقدار فشار خون احتمال آن را به دست آورد؛ فرضاً اگر به ازای فشار خون ۱۲۰، ۱۳۰، ۱۴۰ و ۱۵۰ وجود داشته باشد مقدار

$$\frac{4}{50} = 0.08 \text{ احتمال برابر است با:}$$

توزیع نرمال یا توزیع گاوسی (Normal Distribution)

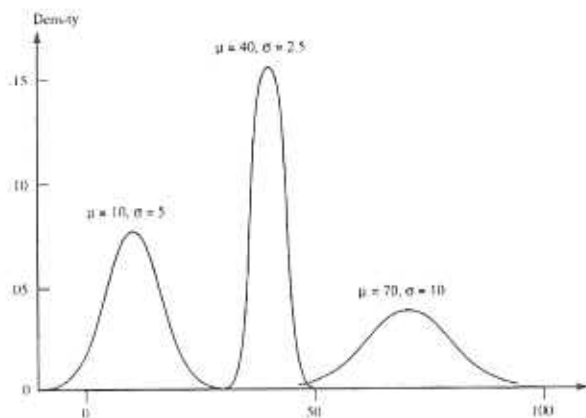
توزیع نرمال، توزیعی از نوع پیوسته است زیرا هر مقداری را اعم از صحیح و اعشاری می‌تواند قبول کند. منحنی این توزیع به صورت دنگوله‌ای یکدست قرینه نسبت به خط قائمی که از میانگین جمعیت (μ) می‌گذرد می‌باشد. انحراف معیار توزیع جمعیت با حرف σ معرفی شده است و برابر یا فاصله افقی میانگین از هریک از دو نقطه عطف منحنی است. (نقطه عطف نقطه‌ای است که تحدب منحنی به ثقل تبدیل شود) در توزیع نرمال مقادیر μ و σ مشخص‌کننده توزیع و یا دو پارامتر توزیع می‌باشند (تصویر ۴-۱).



تصویر ۴-۱

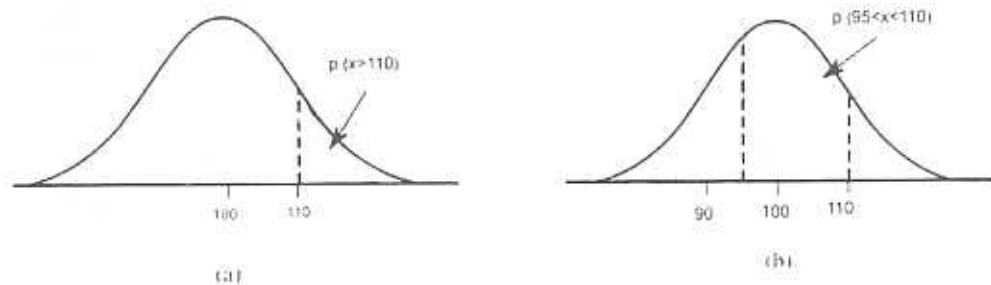
یعنی از میان بیشمار توزیعهای نرمال متفاوت، یک توزیع نرمال مشخص به وسیله σ و μ مشخص می‌شود. از آنجایی که توزیع نرمال یک توزیع احتمالات است بنابراین سطح زیر منحنی برابر با یک می‌باشد، از طرفی چون توزیعی قرینه است لذا نصف سطح منحنی در طرف چپ خط قائم گذرنده از μ می‌باشد و نیمی دیگر در طرف راست آن.

در تصویر ۴-۲ منحنی‌های نرمال متعددی نشان داده شده است؛ یا کمی دقت در تصویر مشخص می‌شود که هرچه انحراف معیار کمتر شده منحنی مربوط به آن باریک و بلندتر می‌شود. بنابراین وقتی انحراف معیار کوچک است سطح زیادی باید نزدیک مرکز منحنی وجود داشته باشد و شانس بیشتری برای مشاهده یک مقدار نزدیک میانگین موجود می‌باشد.



تصویر ۴-۲

حال در منحنی توزیع X (که X برابر فشار خون می‌باشد) اگر این مدل توصیف احتمال پراکندگی واقعی باشد می‌توان احتمال وقوع یک فشار خون بالای ۱۱۰ را با احتمال وقوع فشار خون بین ۹۵-۱۱۰ به سطوح زیر نمودار مقایسه نمود (تصویر ۴-۳).



تصویر ۴-۳

متأسفانه محاسبه مستقیم این چنین احتمالاتی (مساحت زیر نمودار) ساده نبوده و برای محاسبه ساده‌تر این مقادیر به جدول سطوح زیر نمودار برای یک نمودار مرجع که منحنی "توزیع نرمال استاندارد" نامیده شده مراجعه می‌شود.

■ **تعریف:** توزیع نرمال استاندارد یک توزیع نرمال با میانگین صفر و انحراف معیار یک می‌باشد. مرسوم است که صرف Z برای نشان دادن متغیری که توزیعش به وسیله منحنی نرمال استاندارد مشخص شده استفاده شود و غالباً واژه "منحنی Z " به جای منحنی نرمال استاندارد به کار می‌رود.

پدیده‌های نادری یافت می‌شوند که توزیع آنها شبیه به توزیع Z باشد ولی این توزیع باز هم برای محققین مهم می‌باشد زیرا در محاسبه سایر توزیع‌های نرمال به کار می‌رود. زمانی که در تحقیقی، احتمال بر پایه سایر منحنی‌های نرمال وجود دارد ابتدا آن احتمال به معادلی که شامل اطلاعات تحقیق در زیر یک منحنی نرمال استاندارد است تبدیل شده و سپس جدول منحنی Z برای یافتن سطح مورد نظر استفاده می‌شود.

روش استفاده از جدول سطح زیر منحنی‌های نرمال استاندارد (تصویر ۴-۴)

z^*	.00	.01	.02	.03	.04	.05
0.0	.5000	.5040	.5080	.5120	.5160	.5199
0.1	.5398	.5438	.5478	.5517	.5557	.5596
0.2	.5793	.5832	.5871	.5910	.5948	.5987
0.3	.6179	.6217	.6255	.6293	.6331	.6368
0.4	.6554	.6591	.6628	.6664	.6700	.6736
0.5	.6915	.6950	.6985	.7019	.7054	.7088
0.6	.7257	.7291	.7324	.7357	.7389	.7422
0.7	.7580	.7611	.7642	.7673	.7704	.7734
0.8	.7881	.7910	.7939	.7967	.7995	.8023
0.9	.8159	.8186	.8212	.8238	.8264	.8289
1.0	.8413	.8438	.8461	.8485	.8508	.8531
1.1	.8643	.8665	.8686	.8708	.8729	.8749
1.2	.8849	.8869	.8888	.8907	.8925	.8944
1.3	.9032	.9049	.9066	.9082	.9099	.9115
1.4	.9192	.9207	.9222	.9236	.9251	.9265
1.5	.9332	.9345	.9357	.9370	.9382	.9394
1.6	.9452	.9463	.9474	.9484	.9495	.9505
1.7	.9554	.9564	.9573	.9582	.9591	.9599
1.8	.9641	.9649	.9656	.9664	.9671	.9678

تصویر ۴-۴

$$P(Z < 1.42)$$

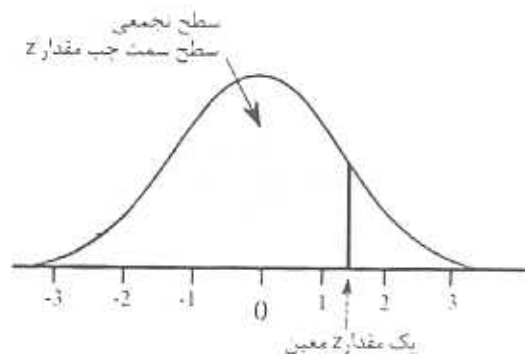
برای هر عدد Z^* بین $-3/89$ تا $+3/89$ و گرد شده تا دو رقم اعشار روابط زیر صادق است:

$$P(Z < Z^*) = P(Z \leq Z^*) = P(\text{سطح زیر منحنی } Z \text{ در سمت چپ } Z^*)$$

عدد Z نشان دهنده یک متغیر تصادفی بوده که از توزیع نرمال استاندارد پیروی می‌کند. حال برای یافتن این احتمال باید به روش زیر عمل نمود:

- ۱- باید سطری را که با Z^* مشخص شده و عدد دو طرف معیار (برای مثال $1/4$ اگر $Z=1/42$) را نشان می‌دهد پیدا نمود.
 - ۲- ستونی که با دومین رقم اعشار Z^* مشخص شده (مثلاً $0/02$ اگر $Z=1/42$) پیدا شود.
 - ۳- عدد محل تقاطع ستون و سطر یافته شده احتمال مورد نظر می‌باشد ($P(Z < Z^*)$)
- ✓ نکته: در کار با توزیع نرمال، یک محقق باید دو مهارت عمومی داشته باشد:
- ۱- محقق باید بتواند توزیع نرمال را برای محاسبه احتمالات بکار ببرد که عبارت از سطوح زیر یک منحنی نرمال و بالای یک فاصله مشخص می‌باشد.

- ۲- محقق باید بتواند مقادیر انتهایی (Extreme) در توزیع، همانند بزرگترین 5% ، کوچکترین 1% و بیشترین 5% انتهایی (که شامل بزرگترین $2/5\%$ و کوچکترین $2/5\%$ است) را محاسبه نماید.
- جدول ضمیمه ۲ نشان‌دهنده سطح تجمعی زیر منحنی، از نوع نشان داده شده در تصویر ۵-۵ برای مقادیر مختلف Z می‌باشد.



تصویر ۵-۴

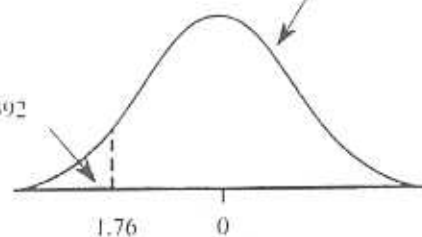
کمترین مقداری که برای سطح تجمعی داده شده است $-3/89$ است که در پائین‌ترین قسمت منحنی Z است. کمترین مقدار بعدی $-3/88$ و سپس $-3/87$ است که هر بار $0/01$ افزایش می‌یابد و نهایتاً با مقدار $+3/89$ خاتمه می‌یابد.

مثال ۱- احتمال $P(Z < -1.76)$ در محل تقاطع سطر $-1/7$ و ستون $0/6$ در جدول Z یافته می‌شود.

$$P(Z < -1.76) = 0.0392$$



منحنی z:

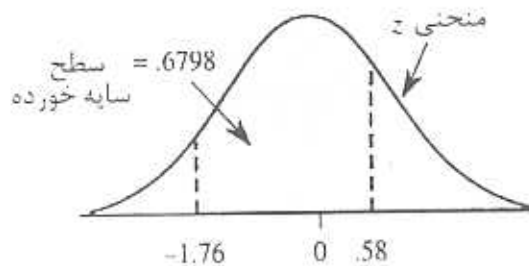
سطح
سایه خورده = 0.0392

تصویر ۴-۶

مثال ۲- احتمال $P(Z < -4.12)$: این احتمال در جدول ضمیمه ۲ وجود ندارد و ستون 4.1- موجود نیست. ولی در هر حال این مقدار باید کمتر از $P(Z < -3.89) = 0.0000$ (کمترین مقدار Z در جدول) باشد. به سبب اینکه $P(Z < -3.89) = 0.0000$ (صفر یا ۱ رقم اعشار معنی دار) مشخص می‌شود که $P(Z < -4.12) \approx 0$ و به همین ترتیب نتیجه گرفته می‌شود که $0 \approx P(Z > 4.12)$

مثال ۳- احتمال اینکه Z بین $-1/76$ و $0/58$ باشد:

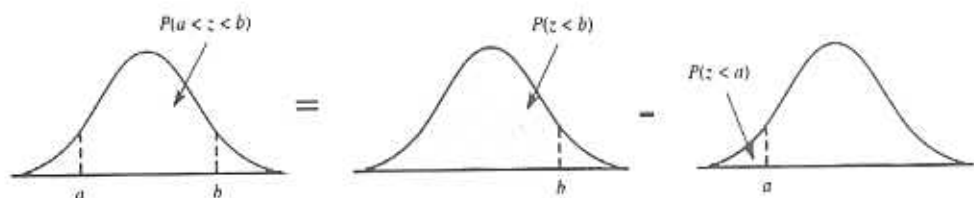
$$P(-1.76 < Z < 0.58) = P(Z < 0.58) - P(Z < -1.76) = 0.7190 - 0.0392 = 0.6798$$



تصویر ۴-۷



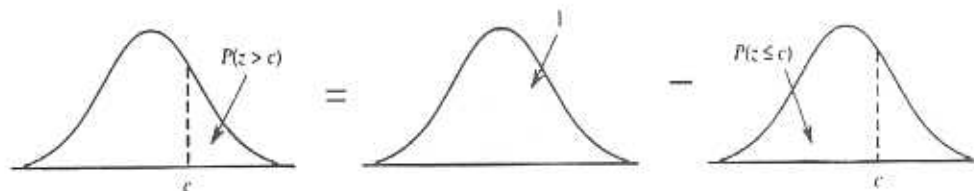
در نتیجه احتمال این که Z بین فاصله در حد تحتانی a و فوقانی b قرار گیرد به این صورت محاسبه می‌گردد:



تصویر ۴-۸

مثال ۱- احتمال اینکه مقدار Z از $1/96$ تجاوز نماید

$$P(Z > 1.96) = 1 - P(Z \leq 1.96) = 1 - 0.9750 = 0.0250$$



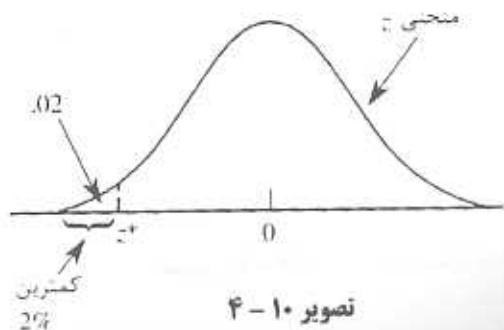
تصویر ۴-۹

مشخص نمودن مقادیر انتهایی:

محقق می‌خواهد مقادیری که کمترین ۲٪ توزیع نرمال استاندارد را می‌سازند بیابد و هدف پیدا نمودن این مقداری است که آنرا

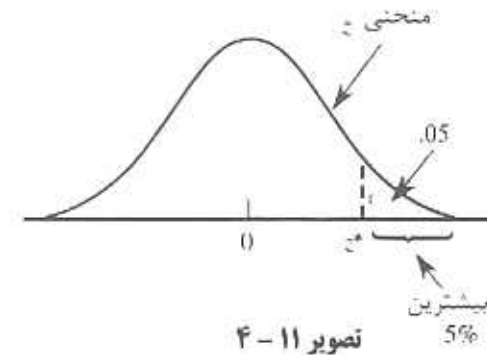
$$P(Z < Z^*) = 0.02$$

تصویر ۱۰-۱ نشان می‌دهد که سطح تجمعی برای Z^* برابر $0/02$ است. بنابراین محقق باید سطح تجمعی $0/02$ را در جدول ضمیمه ۲ بیابد که نزدیکترین سطح تجمعی در جدول $0/0202$ در سطر $2/0$ - و ستون $0/05$ است. بنابراین محقق $Z^* = -2/05$ را استفاده می‌نماید که نزدیکترین تقریب در جدول است و مقادیر زیر $2/05$ - کمترین ۲٪ پراکندگی نرمال را می‌سازند.



تصویر ۴-۱۰

حال برای محقق این نیاز پیش آمده که بزرگترین ۵٪ مقادیر Z را محاسبه نماید $P(Z > Z^*) = 0.05$ (تصویر ۴-۱۱)



از آن رو که ضمیمه ۲ همیشه با سطوح تجمعی دار می‌دهد (سطح سمت چپ) قدم اول مشخص نمودن سطح سمت چپ Z^* است.

$$Z^* = 1 - 0.05 = 0.95$$

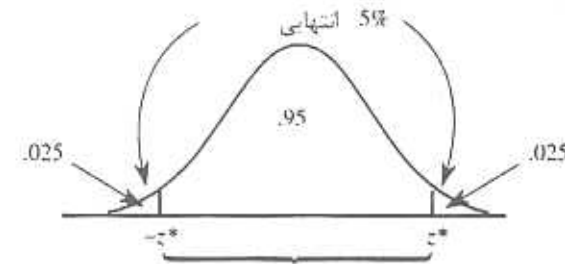
در جدول ضمیمه ۲ مشاهده شده که ۰/۹۵ دقیقاً در میانه ۰/۹۴۹۵ (با مقدار Z ۱/۶۵) و ۰/۹۵۰۵ (با مقدار Z ۱/۶۴) قرار دارد. از این رو:

$$Z^* = \frac{1.64 + 1.65}{2} = 1.645$$

✓ نکته: در مثال فوق اگر یکی به ۰/۹۵ نزدیکتر می‌بود باید Z مربوط به آن استفاده می‌شد.

۱/۶۴۵ بزرگترین ۵٪ توزیع نرمال استاندارد را می‌سازد، و براساس تقارن، ۱/۶۴۵ - جدا کننده کمترین ۵٪ مقادیر Z از بقیه است.

مثال: محاسبه بیشترین ۵٪ انتهایی (The Most Extreme)



۹۵٪ میانی

تصویر ۴-۱۲

یعنی در واقع باید ۹۵٪ میانی را از ۵٪ انتهایی جدا نمود.

با توجه به تقارن توزیع نرمال این ۵٪ به دو قسمت چپ و راست ۲/۵٪ تقسیم می‌شود و اگر Z^* ۲/۵٪ سمت راست را جدا نماید، $-Z^*$ ۲/۵٪ سمت چپ را جدا می‌نماید؛ پس برای محاسبه Z^* سطح تجمعی برای Z^* را محاسبه می‌گردد:

$$Z^* \text{ سطح سمت چپ } = 0.95 + 0.025 = 0.975$$

و Z^* برابر با ۱/۹۶ است پس مقادیر بزرگتر از ۱/۹۶ و کوچکتر از ۱/۹۶، ۵٪ بیشترین آنها می‌باشند.

تعیین احتمالات

برای محاسبه احتمال در هر توزیع نرمال ابتدا باید متغیرهای مربوط را استاندارد نمود و سپس از جدول سطح منحنی Z استفاده کرد، اگر X متغیری باشد که رفتار از منحنی نرمال پیروی کند که میانگین μ و انحراف معیار برابر σ داشته باشد (تصویر ۴-۱۳):

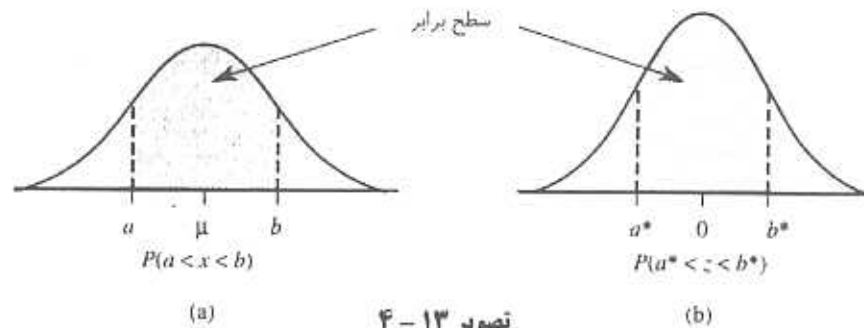
$$P(x < b) = P(z < b^*)$$

$$P(a < x) = P(a^* < z)$$

$$P(a < x < b) = P(a^* < z < b^*)$$

نکته: مقادیر a^* و b^* از روابط زیر بدست می‌آید:

$$a^* = \frac{a - \mu}{\sigma}, \quad b^* = \frac{b - \mu}{\sigma}$$

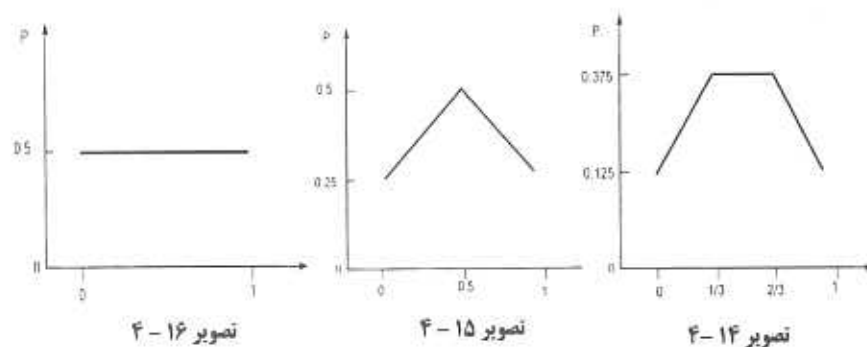


توجه: در فرمولهایی که در فصلهای بعد مشاهده خواهید کرد برای خلاصه تر نمودن فرمولها بجای (سطح زیر نمودار $P(Z > Z^*)$) از Z^* (سطح زیر نمودار Z) استفاده می‌شود.

همچنین در برخی فرمولها جدول سطح زیر منحنی های نرمال، سطح زیر نمودار را بطور دوطرفه محاسبه می‌نماید. یعنی مثلاً بجای $P(Z > 1.96) = 0.95$ در جدول یک دامنه، $P(Z > 1.96) = 0.45$ در جدول دو دامنه محاسبه میشود. یعنی به ازای Z^* بین ۰ تا ۳.۸۹ سطح زیر نمودار را بین ۰ تا ۰.۵۰ تغییر می‌نماید. از جدول یک دامنه برای اثبات یا رد فرضیه آلترناتیو یک دامنه و از جداول دو دامنه برای اثبات یا رد فرضیه آلترناتیو دو دامنه استفاده می‌گردد.

تئوری حد مرکزی (Central Limit Theory)

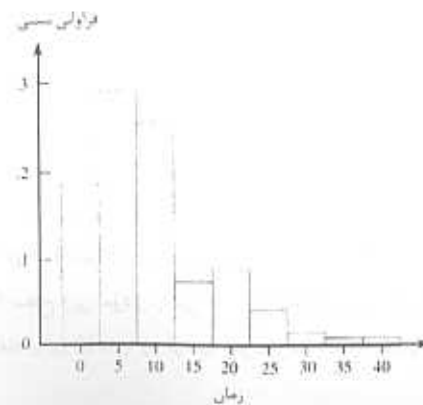
بیان تئوری: این تئوری یکی از تئوریهای پایه‌ای احتمالات می‌باشد. به طور خیلی ساده تئوری حد مرکزی بیان می‌دارد که پراکندگی مجموع تعداد بسیار زیاد متغیرهای غیر وابسته و با توزیع یکسان، بدون ارتباط به پراکندگی اولیه آنها تقریباً نرمال خواهد بود. تصور کنید که سکه‌ای وجود دارد که یکروی آن عدد ۰ و یکروی دیگر آن عدد ۱ حک شده باشد. در یکصد پرتاب این سکه نتایج بر روی نمودار تقریباً به صورت تصویر ۴-۱۱ می‌باشد.



اگر این بار دوسه‌بار یکصد پرتاب شوند و هر بار میانگین دو عدد ثبت شود نمودار آن به صورت تصویر ۴-۱۵ خواهد بود. اگر دوباره این عمل با ۳ سکه انجام شود نمودار حاصل بصورت تصویر ۴-۱۶ خواهد بود (توجه: نمودارهای فوق به راحتی توسط قانون احتمالات قابل محاسبه می‌باشد).

با دقت به نمودارهای فوق متوجه خواهید شد که نمودار از حالت مستطیل شکل (Rectangular) به سمت توزیع نرمال خواهد رفت و نکته قابل توجه دیگر اینکه میانگین توزیع در تمامی نمودارها همان میانگین اولیه و برابر عدد ۰/۵ است.

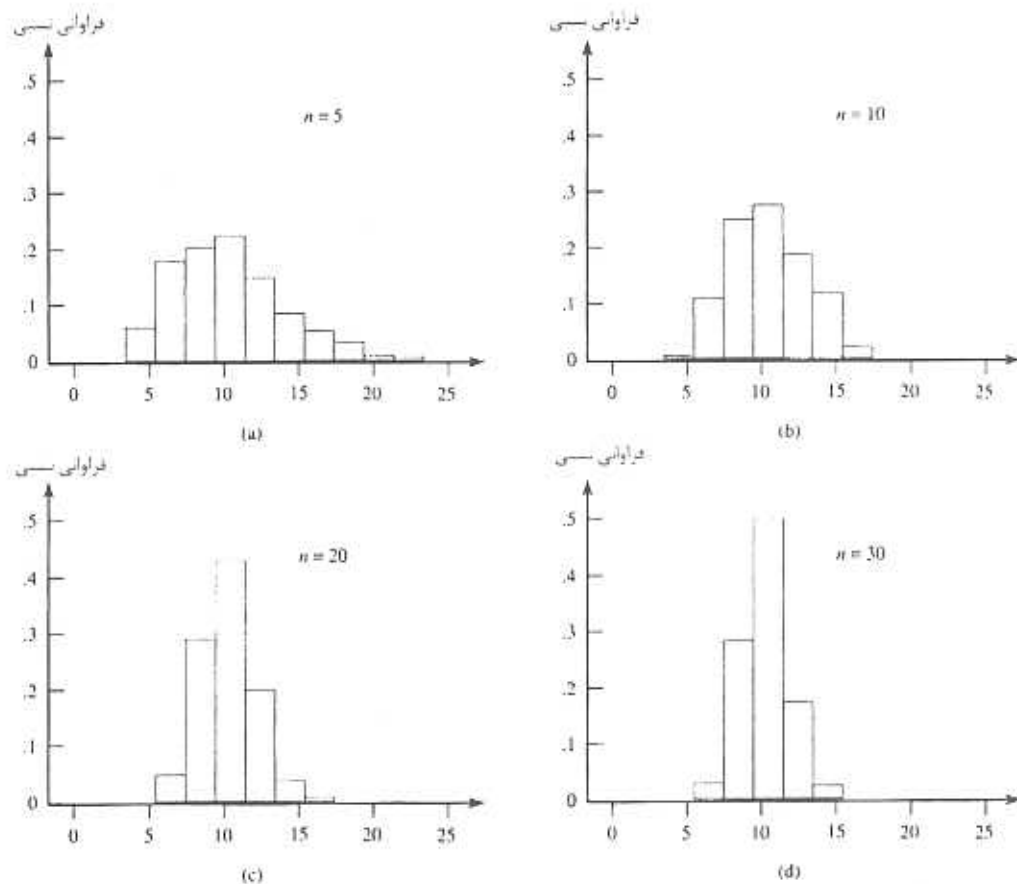
نمودار تصویر ۴-۱۷ نشان‌دهنده میزان تأخیر دانش آموزان یک کلاس ۲۵۱ نفره در کلاس آمار پزشکی می‌باشد (این نمودار



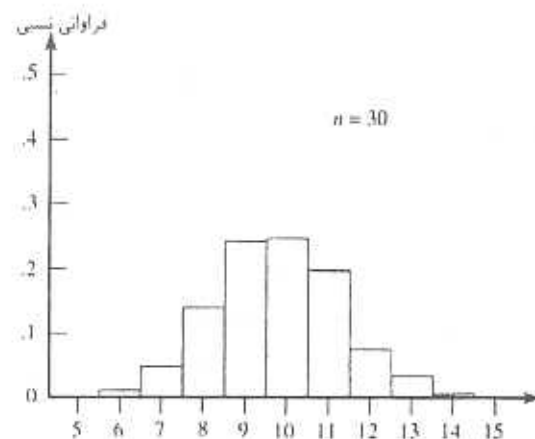
تصویر ۴-۱۷

نشان‌دهنده آن است که اکثر دانشجویان در دقایق اولیه کلاس تأخیر دارند ولی فراوانی کسانی که بیش از ۱۵ دقیقه تأخیر دارند بسیار کم است. نتیجتاً این نمودار یک نمودار دارای چولگی (Skewness) می‌باشد.

اگر محقق با تحلیل این نمودار مثلاً یک حجم نمونه تصادفی برابر ۵ نفر اختیار نماید و میانگین تأخیر آنها را محاسبه و رسم کند و این عمل را ۵۰۰ بار انجام دهد، نموداری مطابق اولین نمودار تصویر ۴-۱۸ به دست خواهد آمد که اگر حجم نمونه انتخاب شده افزایش داده شود ($n=10, n=20, n=30$) گسترش نمودار در طرفین مرکز آن کم می‌شود و نکته دیگر آنکه اگر آخرین نمودار (d) را با بزرگنمایی بیشتر مشاهده کنیم (تصویر ۴-۱۹) مشاهده می‌شود که منحنی با چولگی، کاملاً به منحنی نرمال تبدیل شده است.



تصویر ۴-۱۸



تصویر ۱۹-۴

خواص عمومی توزیع نمونه‌گیری $\bar{X}$

اگر X یک متغیر تصادفی با میانگین μ و انحراف معیار σ باشد و $X_1, X_2, \dots, X_n$ نمونه‌هایی تصادفی از توزیع X باشند آنگاه:

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

✓ نکته: بر توزیع $\bar{X}$ قوانین ذیل حاکم خواهد بود:

- قانون ۱: $\mu_{\bar{X}} = \mu$
- قانون ۲: $\sigma_{\bar{X}} = \frac{\sigma}{\sqrt{n}}$ (این قانون تا زمانی که کمتر از ۵٪ جمعیت در نمونه باشد صحیح است.)
- قانون ۳: زمانی که توزیع جامعه نرمال باشد توزیع $\bar{X}$ نیز برای هر مقدار n نرمال خواهد بود.
- قانون ۴: (تئوری حد مرکزی): زمانی که n به اندازه کافی بزرگ باشد توزیع نمونه $\bar{X}$ به شکل یک نمودار نرمال خواهد بود حتی اگر توزیع جامعه اصلی نرمال نباشد.

به عنوان مثال اگر محقق تعداد زیادی نمونه ۲۵ نفری تصادفی از جامعه انتخاب نماید توزیع میانگین قد افراد نمونه وی به چه صورت خواهد بود؟ (در آمارهای ملی میانگین و انحراف معیار قد افراد ایرانی به ترتیب برابر ۱۷۰ و ۱۵ سانتی‌متر است)

$$\mu_{\bar{X}} = 170 \text{ cm}, \sigma_{\bar{X}} = \frac{\sigma}{\sqrt{n}} = \frac{15}{\sqrt{25}} = 3$$



در مثالی دیگر فرض کنید یک شرکت داروسازی ادعا می‌کند که مصرف دوز ۱۵۰ mgr در روز یکی از داروهایش که بجز درمائی کوچک نیز دارد، حداکثر غلظت خون ۱۸ $\mu\text{gr/lit}$ یا $\sigma=1$ می‌دهد. محقق از ۳۶ فرد که این دارو را با دوز ۱۵۰ mgr مصرف نموده‌اند نمونه خونی گرفته و میانگین خونی برابر ۱۸/۴ $\mu\text{gr/lit}$ به دست آورده است. وی از این آزمایش چه نتیجه‌ای خواهد گرفت؟

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}} = \frac{1}{\sqrt{36}} = 0.1667$$

اگر ادعای شرکت دارو سازی درست باشد باید:

$$\mu_{\bar{x}} = \mu = 18 \mu\text{gr/Lit}$$

آیا ممکن است محقق فوق‌الذکر نتیجه بگیرد که دوز پیشنهاد شده توسط شرکت داروسازی زیاد است؟

احتمال به دست آوردن میانگین برابر و یا بیش از ۱۸/۴ برابر است با:

$$P(\bar{x} \geq 18.4) = P(Z \geq \frac{18.4 - 18}{0.1667}) = P(Z \geq 2.4) = 0.0082$$

در این مثال ابتدا عدد ۱۸/۴ به توزیع نرمال استاندارد تبدیل می‌شود و سپس با استفاده از جدول Z معین می‌گردد که سطح زیر منحنی Z بزرگتر از ۲/۴، برابر ۰/۰۰۸۲ می‌باشد. یعنی اگر قرار باشد به صورت تصادفی با نمونه‌گیری‌های ۳۶ تایی میانگینی برابر یا بیشتر از ۱۸/۴ $\mu\text{gr/lit}$ به دست آید شانس این محقق برابر ۰/۰۰۸۲ می‌باشد که عدد بسیار کوچکی است و احتمالاً ادعای این شرکت مبنی بر مناسب بودن دوز دارو درست می‌باشد.

انواع خطا در طرح‌های تحقیقاتی

در انجام هر طرح تحقیقاتی، یک فرضیه صفر وجود دارد که در عالم واقع این فرضیه درست یا غلط خواهد بود و پس از پایان طرح تحقیقاتی نیز فرضیه صفر پذیرفته یا رد می‌شود.

با توجه به موارد فوق بین صحیح یا غلط بودن فرضیه صفر و قبول یا رد این فرضیه چهار حالت ایجاد خواهد شد (جدول ۱-۴):
۱- چنانچه فرضیه صفر واقعاً درست باشد و در آنالیز نتایج تحقیق نیز قبول گردد، هیچگونه خطایی اتفاق نیفتاده است و احتمال این حالت با قدرت تعمیم‌پذیری فرضیه تحقیق به جامعه (Confidence Coefficient) برابر است.

۲- چنانچه فرضیه صفر درست باشد اما در تحقیق مورد نظر رد گردد، مطالعه دچار خطای نوع اول شده است (رد فرضیه صفر صحیح) که احتمال آنرا با علامت α نمایش می‌دهند و بطور معمول سطح آنرا معادل ۵٪ در نظر می‌گیرند.

۳- چنانچه فرضیه صفر واقعاً غلط باشد اما در تحقیق مورد نظر پذیرفته شود خطای نوع دوم اتفاق افتاده است (پذیرفتن فرضیه صفر غلط و یا رد فرضیه آلترناتیو صحیح) که احتمال آنرا با علامت β نشان می‌دهند و بطور معمول در طرح‌های تحقیقاتی حداکثر ۲۰٪ در نظر گرفته می‌شود.

۴- چنانچه فرضیه صفر غلط باشد و در تحقیق مورد نظر رد گردد محقق دچار هیچگونه خطایی نشده است و احتمال این حالت میزان قدرت مطالعه (Power) را نشان می‌دهد. به عنوان مثال چنانچه احتمال این حالت در مطالعه‌ای ۸۰٪ گردد بدین معنی است که این مطالعه تا ۸۰٪ قدرت دارد که رابطه بین یک عامل خطر و پیامد آن را (در صورت وجود ازبابط) نشان دهد و اگر در چنین حالتی

در یک تحقیق ارتباطی بین عامل خطر و پیامد بدست نیامد نمی‌توان با قطعیت منکر وجود این ارتباط شد زیرا ۲۰٪ احتمال دارد که با وجود ارتباط واقعی بین عامل خطر و پیامد مطالعه، محقق توان اثبات این ارتباط را نداشته باشد.

✓ نکته: مقادیر ضریب اطمینان و α مکمل یکدیگر بوده و حاصل جمع آنها برابر یک خواهد بود، لذا با کم کردن α از مقدار یک، می‌توان میزان ضریب اطمینان یک مطالعه را محاسبه نمود.

فرضیه صفر
↓
غلط
↓
درست

احتمال خطای نوع دوم (β)	ضریب اطمینان (CC)
قدرت (POWER)	احتمال خطای نوع اول (α)

جدول ۱-۴

قبول
↓
فرضیه
↓
رد

✓ نکته: مقادیر قدرت (Power) و β مکمل یکدیگر بوده و حاصل جمع آنها برابر یک می‌گردد. لذا مقدار قدرت یک مطالعه را در بدو انجام مطالعه می‌توان با کم کردن β از عدد یک تخمین زد. پس در یک تحقیق چنانچه فرضیه صفر رد گردد احتمال ایجاد خطای نوع اول وجود دارد لذا گزارش بیشترین احتمال رد فرضیه صفر در صورتی که درست باشد یا همان P value الزامی است. در چنین شرایطی احتمال ایجاد خطای نوع دوم متغی است.

چنانچه محقق نتواند فرضیه صفر را رد نماید (قبول فرضیه صفر) دیگر احتمال ایجاد خطای نوع اول وجود نخواهد داشت لذا گزارش مقدار P value الزامی نیست اما با توجه به اینکه در چنین حالتی احتمال خطای نوع دوم وجود دارد محقق باید معین نماید که این مطالعه جهت اثبات وجود ارتباط بین عامل خطر و پیامد چه میزان قدرت دارد لذا محاسبه قدرت (Power) واقعی مطالعه و گزارش آن در چنین حالتی الزامی خواهد بود.

قدرت (Power)، خطای نوع اول، حجم نمونه و کمترین اختلاف مورد نظر (Effect Size) سیستم بسته‌ای را تشکیل می‌دهند که با مشخص نمودن سه عامل، عامل چهارم قابل محاسبه خواهد بود.

قبل از انجام مطالعه با داشتن تخمینی از خطای نوع اول و قدرت مورد انتظار و کمترین اختلاف مورد نظر می‌توان حجم نمونه را محاسبه نمود. در واقع در ابتدای یک مطالعه تحلیلی، محقق به دنبال آن است که کمترین میزان حجم نمونه را با کمترین خطای نوع اول و دوم و در نتیجه بیشترین توان تعمیم‌پذیری (CC) و قدرت (Power) به ترتیب در صورت رد و قبول فرضیه صفر بدست آورد و در ابتدای یک مطالعه توصیفی محقق سعی دارد کمترین حجم نمونه را با کمترین خطای نوع اول و در نتیجه بیشترین توان تعمیم‌پذیری (CC) بدست آورد. لازم به ذکر است که با توجه به ماهیت مطالعات توصیفی امکان ایجاد خطای نوع دوم در آنها وجود ندارد و لذا محاسبه حجم نمونه این‌گونه مطالعات به خطای نوع دوم ارتباطی نخواهد داشت.

محاسبه حجم نمونه

Sample Size Estimation

✓ محاسبه حجم نمونه بر آورد

✓ محاسبه حجم نمونه در مطالعات مشاهده‌ای

تحلیلی

✓ محاسبه حجم نمونه در مطالعات کار آزمایی

بالینی

✓ محاسبه حجم نمونه در مطالعات ارزیابی

تستهای تشخیصی

✓ محاسبه حجم نمونه در مطالعات بقا

✓ محاسبه حجم نمونه همبستگی

✓ محاسبه حجم نمونه با استفاده از نرم افزار

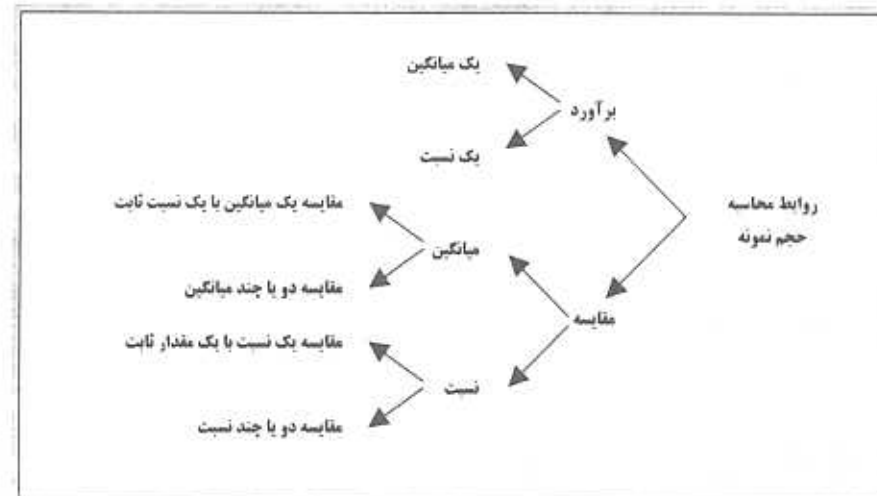
EPI-Info 6

✓ و ...

فصل

۵

برای تعیین نوع رابطه‌ای که جهت محاسبه حجم نمونه در یک طرح تحقیقاتی مورد استفاده است، محقق در ابتدا باید تعیین کند که هدف از انجام طرح، نهایتاً برآورد کردن یک مقدار یا مقایسه کردن چند مقدار می‌باشد و در مرحله بعد، محقق باید ماهیت این مقادیر را مشخص کند؛ بدین معنی که این مقدار، یک میانگین یا یک نسبت می‌باشد و در نهایت باید مشخص شود که هدف محقق مقایسه میانگین یا نسبت با یک مقدار ثابت است یا مقایسه دو یا چند میانگین یا دو یا چند نسبت می‌باشد. در هر یک از این موارد از رابطه‌ای خاص جهت محاسبه حجم نمونه استفاده می‌گردد که در ذیل به آن می‌پردازیم (تصویر ۵-۱).



تصویر ۵-۱

برآورد

برآورد یک میانگین

برای محاسبه حجم نمونه در طرحهایی که هدف آنها برآورد یک میانگین می‌باشد از رابطه (۵-۱) استفاده می‌شود:

$$n = \frac{(Z^2 \times \sigma^2)}{d^2} \quad \text{رابطه ۵-۱}$$

که در این رابطه:

d : تعداد نمونه‌های مورد نیاز برای انجام طرح

α : احتمال خطای نوع اول (رد فرضیه H_0 صحیح) و $Z_{1-\alpha/2}$ بیان کننده ضریب اطمینان (Confidence Coefficient)

طرح در تعمیم نتایج به جامعه است؛ از آنجا که هرچه دقت مطالعه بیشتر باشد نیازمند حجم نمونه بیشتری خواهد بود، $Z_{1-\alpha/2}$ با n رابطه مستقیم دارد.

σ : انحراف از معیار (پراکندگی) متغیر مورد نیاز در جامعه هدف است و از آنجاییکه هر چه پراکندگی صفت در افراد مورد پژوهش بیشتر باشد، جهت شباهت نمونه به جامعه هدف، باید تعداد نمونه‌های بیشتری انتخاب گردد، σ^2 یا n رابطه مستقیم دارد و با افزایش پراکندگی صفت مورد بررسی در جامعه، نیاز به حجم نمونه بالاتری ایجاد می‌شود.

d : فاصله اطمینانی (Confidence Interval) است که نهایتاً جهت گزارش نتایج طرح و تعمیم آنها به جامعه، مورد نیاز است. به عنوان مثال گزارش می‌شود که میانگین هموگلوبین در افراد مورد پژوهش 12 ± 3 بوده است. عدد ۳ در این مثال برآوردی از میزان دقت می‌باشد. واضح است که هرچه پهنای این فاصله کمتر باشد، دقت مطالعه بیشتر خواهد بود و نیاز به حجم نمونه بیشتری خواهد داشت. از این رو فاصله اطمینان با حجم نمونه رابطه معکوس دارد و جهت تغییرات آنها عکس یکدیگر می‌باشد.

مثال: در مطالعه "تعیین میانگین وزن نوزادان متولد شده در بیمارستان حضرت علی‌اصغر (ع) در سال ۱۳۸۰" هرگاه براساس نتایج مطالعات در سایر مراکز، $d=0.4g$ و $\sigma=3$ باشد و $\alpha=0.05$ فرض شود، حجم نمونه مورد نیاز برای انجام طرح بدین ترتیب محاسبه می‌گردد:

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$n = \frac{(1.96)^2 \times 3^2}{(0.4)^2} \approx 216$$

یعنی محقق باید براساس روش نمونه‌گیری خود، ۲۱۶ نوزاد که در سال ۱۳۸۰ در بیمارستان حضرت علی‌اصغر متولد شده‌اند را انتخاب کرده و میانگین وزن آنها را محاسبه کند.

✓ نکته: هرگاه محقق اطلاعاتی در مورد میزان فاصله اطمینان مطلوب نداشته باشد، می‌تواند d را عددی بین 0.1σ تا 0.32σ قرار دهد.

مثال: در مطالعه "تعیین میانگین سطح سرمی کلسترول در افراد ۴۰-۵۰ ساله بستری در بخش CCU بیمارستان فیروزگر در سال ۱۳۷۹" هرگاه براساس نتایج مطالعات قبلی $\sigma=50$ باشد، حجم نمونه مورد نیاز بدین ترتیب محاسبه می‌شود ($\alpha=0.05$): از آنجاییکه محقق از مقدار d مطلع نیست، آن را براساس تخمین بالینی، برابر با 0.2σ در نظر می‌گیرد:

$$d = 0.2\sigma = 10$$

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$n = \frac{(1.96)^2 \times 50^2}{10^2} \approx 96$$

یعنی محقق باید براساس روش نمونه‌گیری طرح، ۹۶ بیمار ۴۰-۵۰ ساله که در سال ۱۳۷۹ در بخش CCU بیمارستان فیروزگر بستری بوده‌اند انتخاب کرده و میانگین سطح کلسترول سرم آنها را محاسبه نماید.

✓ نکته: رابطه ۵-۱ صرفاً جهت محاسبه حجم نمونه در مطالعات مقطعی توصیفی کاربرد دارد.

برآورد یک نسبت

در مطالعاتی که هدف آنها برآورد نسبت یک صفت مشخص در جامعه است، برای محاسبه حجم نمونه مورد نیاز جهت انجام طرح، از رابطه ۵-۲ استفاده می‌گردد:

$$n = \frac{Z^2 \times P(1-P)}{d^2} \quad \text{رابطه ۵-۲}$$

که در این رابطه هر پارامتر به صورت زیر تعریف می‌شود:

n : تعداد نمونه‌های مورد نیاز برای انجام طرح

Z : ضریب اطمینان و دقت مطالعه

p : پیش‌فرضی از نسبت فراوانی وجود صفت مورد نظر در جامعه

d : فاصله اطمینان متغیر مورد نظر جهت گزارش نتیجه آن

مثال: در مطالعه "تعیین فراوانی بیماری دیابت در مراجعه‌کنندگان به درمانگاه حضرت رسول اکرم (ص) در سال ۱۳۸۰" هرگاه براساس بررسی متون انجام شده، $p=0.8$ و $d=0.08$ باشد، با در نظر گرفتن $\alpha=0.05$ حجم نمونه مورد نیاز بدین ترتیب محاسبه می‌شود:

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$n = \frac{(1.96)^2 \times 0.8 \times (1-0.8)}{(0.08)^2} \approx 97$$

یعنی محقق باید بنا به روش نمونه‌گیری طرح، ۹۷ بیمار که در سال ۱۳۸۰ به درمانگاه غدد بیمارستان حضرت رسول اکرم (ص) مراجعه کرده‌اند را انتخاب کرده و محاسبه کند که چه نسبتی از آنها مبتلا به دیابت می‌باشند.

✓ نکته: هرگاه محقق پیش‌فرضی در مورد میزان d مطلوب نداشته باشد، می‌تواند d را برابر با عددی بین $0.1p$ تا $0.2p$ انتخاب کند.

مثال: در مطالعه "تعیین فراوانی نوزادان نارس متولد شده در بیمارستان اکبرآبادی در سال ۱۳۸۰" هرگاه براساس بررسی متون، $p=0.1$ باشد، با در نظر گرفتن $\alpha=0.05$ ، حجم نمونه بدین ترتیب محاسبه می‌گردد:

از آنجاییکه محقق از مقدار d مطلع نیست، آنرا برابر با $0.1p$ در نظر می‌گیرد:

$$d = 0.2P = 0.02$$

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$n = \frac{(1.96)^2 \times 0.1 \times (1-0.1)}{(0.02)^2} \approx 864$$

یعنی محقق بنا به روش نمونه‌گیری طرح، ۸۶۴ نوزاد که در سال ۱۳۸۰ در بیمارستان اکبرآبادی متولد شده‌اند را انتخاب کرده و تعیین می‌کند که چه نسبتی از آنها نارس متولد شده‌اند.

✓ نکته: هرگاه محقق بخواهد فراوانی چند متغیر با تست‌های پیش‌فرض (p) متفاوت را در یک مطالعه تعیین کند، حجم نمونه بر مبنای تست پیش‌فرضی که به عدد 0.5 نزدیکتر است، محاسبه می‌گردد.

مثال: در مطالعه "تعیین فراوانی انواع گروه‌های خونی در بیماران مبتلا به تالاسمی بستری در بخش داخلی بیمارستان هفتم تیر"

هرگاه براساس بررسی متون، درصد فراوانی گروه خونی O ($P_1=0.43$)، درصد فراوانی گروه خونی B ($P_2=0.21$) و درصد فراوانی گروه خونی A ($P_3=0.24$) و درصد فراوانی گروه خونی AB ($P_4=0.10$) باشد، با فرض $\alpha=0.05$ ، حجم نمونه بدین ترتیب محاسبه می‌شود:

از آنجاییکه P_1 به عدد 0.5 نزدیکتر است، برای محاسبه حجم نمونه از آن استفاده می‌کنیم:

$$d = 0.2P = 0.2 \times 0.43 = 0.086$$

$$\alpha = 0.05 \Rightarrow Z_{\alpha} = 1.96$$

$$n = \frac{(1.96)^2 \times 0.43 \times (1-0.43)}{(0.086)^2} \approx 110$$

یعنی محقق باید براساس روش نمونه‌گیری طرح، ۱۱۰ بیمار مبتلا به تالاسمی که در بخش داخلی بیمارستان هفتم تیر بستری بوده‌اند را انتخاب کند و پس از مشخص شدن گروه خونی هر بیمار، فراوانی هر یک از گروه‌های خونی O, B, A, AB را در این بیماران گزارش نماید.

✓ نکته: چنانچه محقق هیچ برآوردی از مقدار p نداشته باشد برای بدست آوردن ماکسیمم حجم نمونه، کافی است مقدار p را برابر 0.5 فرض کند زیرا براساس یک قانون ریاضی هنگامیکه حاصل جمع دو عدد یک مقدار ثابت باشد، حاصل ضرب آن زمانی بیشینه است که آن دو عدد برابر در نظر گرفته شوند. در رابطه ۵-۲ حاصل جمع دو عدد p و $1-p$ همواره عدد ثابت ۱ خواهد بود در نتیجه حاصل ضرب این دو عدد زمانی که برابر 0.5×0.5 در نظر گرفته شوند، بیشینه می‌گردد.

✓ نکته: رابطه ۵-۲ صرفاً برای محاسبه حجم نمونه در مطالعات مقطعی توصیفی کاربرد دارد.

✓ نکته: همانگونه که از روابط محاسبه حجم نمونه برای برآورد یک میانگین و نسبت بر می‌آید، حجم جامعه هدف که نتایج مطالعه به آن تعمیم خواهد یافت در تعیین حجم نمونه تأثیری ندارد و آنچه اهمیت دارد پراکندگی صفت در جامعه می‌باشد. اما چنانچه حجم جامعه هدف کمتر از ۱۰۰۰۰ نفر باشد، آنگاه باید برای بدست آوردن حجم نمونه مناسب، حجم جامعه هدف را نیز دخیل نمود و جهت این منظور از رابطه ۵-۳ استفاده می‌گردد.



رابطه ۳-۵

$$n' = \frac{n}{1 + \frac{n}{N}}$$

در رابطه ۳-۵ هر یک از پارامترها به قرار زیر تفسیر می‌شوند:

n : حجم نمونه محاسبه شده اولیه از فرمولهای برآورد میانگین و نسبت

N : حجم جامعه هدف

n' : حجم نمونه اصلاح شده و مورد نیاز برای انجام مطالعه

✓ نکته: اثر نوع نمونه‌گیری (Design Effect) کلیه روابط محاسبه حجم نمونه جهت نمونه‌گیری به طریقه تصادفی ساده طراحی گردیده است. با کمی دقت در انواع روشهای نمونه‌گیری، مشخص خواهد شد که دقت نمونه‌گیری طبقه‌ای از نمونه‌گیری تصادفی ساده بیشتر می‌باشد زیرا در این نمونه‌گیری، جامعه به طبقاتی تقسیم و از تمام طبقات حتماً نمونه انتخاب می‌گردد؛ این درحالیست که در نمونه‌گیری تصادفی ساده ممکن است به طور تصادفی از بعضی افراد که وابسته به طبقه خاصی هستند نمونه‌ای انتخاب نگردد؛ لذا در صورتی که در تحقیق، نمونه‌گیری به صورت طبقه‌ای انجام گیرد به دلیل خطای کمتر نسبت به نمونه‌گیری تصادفی ساده به حجم نمونه کمتری نیاز می‌باشد.

در نمونه‌گیری خوشه‌ای به دلیل آنکه از تعدادی از خوشه‌ها نمونه‌گیری صورت نمی‌گیرد، دقت نمونه‌گیری نسبت به نمونه‌گیری تصادفی ساده کمتر می‌باشد لذا جهت انجام مطالعه با نمونه‌گیری خوشه‌ای به حجم نمونه بیشتری نیاز می‌باشد.

در محاسبه حجم نمونه، واژه‌ای به نام اثر نوع نمونه‌گیری (Design Effect) وجود دارد که حجم نمونه را برحسب نوع نمونه‌گیری اصلاح می‌نماید. اثر نوع نمونه‌گیری به لحاظ آماری برابر با نسبت واریانس با نمونه‌گیری غیر تصادفی ساده به واریانس با نمونه‌گیری تصادفی ساده می‌باشد. چنانچه نمونه‌گیری به صورت طبقه‌ای انجام گیرد با توجه به اینکه در این نوع طراحی نمونه‌گیری، پراکندگی صفت مورد نظر نسبت به نمونه‌گیری تصادفی ساده کمتر است، اثر نوع نمونه‌گیری عددی کمتر از یک خواهد شد و حجم نمونه کمتری محاسبه خواهد گشت. در حالیکه در نمونه‌گیری خوشه‌ای واریانس بزرگتری نسبت به نمونه‌گیری تصادفی ساده خواهیم داشت، لذا اثر نوع نمونه‌گیری عددی بزرگتر از یک خواهد شد و به حجم نمونه بیشتری نسبت به محاسبه اولیه نیاز خواهیم داشت. جهت محاسبه واریانس در دو حالت نمونه‌گیری تصادفی ساده و غیر از آن روابطی وجود دارد که از ذکر آن خودداری می‌گردد اما نرم‌افزار EPI Info، اثر نوع نمونه‌گیری را در قسمت EPITABLE Calculator زیر منو Describe محاسبه می‌نماید.

مقایسه

مقایسه یک میانگین با عدد ثابت

برای محاسبه حجم نمونه در مطالعاتی که هدف از انجام آنها مقایسه یک میانگین با یک مقدار ثابت، از رابطه ۴-۵ استفاده می‌شود:



رابطه ۴-۵

$$n = \frac{(Z_{1-\alpha/2} + Z_{1-\beta})^2 \sigma^2}{d^2}$$

که در این رابطه تعریف هر پارامتر به قرار زیر است:

n : تعداد نمونه‌های مورد نیاز برای انجام مطالعه

α : احتمال خطای نوع اول

β : احتمال خطای نوع دوم

σ^2 : واریانس یا پراکندگی متغیر مورد سنجش در جامعه هدف

d : حداقل اختلاف دارای ارزش هنگام مقایسه دو میانگین است که نام دیگر آن اندازه اثر (Effect Size) می‌باشد.

✓ نکته: اندازه اثر (Effect Size) یک شاخص آماری است و عبارت از حداقل اختلاف آماری معنی‌دار هنگام مقایسه دو گروه می‌باشد. معادل این شاخص در علوم بالینی، شاخصی تحت عنوان حداقل اختلاف بالینی مهم (Clinical Important Difference) است که عبارتست از حداقل اختلاف بین دو گروه که از نظر بالینی نیز دارای ارزش باشد. بعنوان مثال چنانچه مصرف سیگار در دو گروه مبتلا و غیر مبتلا به سرطان پروستات تنها یک درصد اختلاف داشته باشد، این میزان از نظر بالینی ارزش نخواهد داشت هرچند که از نظر آماری این اختلاف معنی‌دار باشد. چنانچه این اختلاف بیست درصد باشد، از نظر بالینی نیز ارزشمند خواهد بود؛ بنابراین در مطالعات بالینی در شرایط مطلوب باید این دو شاخص آماری و کلینیکی با هم برابر باشند.

اندازه اثر از جمله شاخصهای است که جهت محاسبه حجم نمونه برخی از مطالعات مورد نیاز است. از جمله مطالعاتی که هدف آنها مقایسه یک میانگین با یک عدد ثابت، مقایسه دو میانگین، مقایسه یک نسبت با یک عدد ثابت و مقایسه دو نسبت است. بنابراین اندازه اثر قبل از شروع مطالعه باید مشخص شده باشد.

اندازه اثر را به چند روش می‌توان بدست آورد:

- انجام بررسی متون و استفاده از اطلاعات مطالعات مشابه قبلی
- انجام مطالعه مقدماتی (Pilot Study)
- در مواردی که مطالعات مشابه قبلی وجود ندارد استفاده از روابط محاسبه اندازه اثر (جدول ۱-۵)
- انتخاب کیفی اندازه اثر به صورت کوچک، متوسط یا بزرگ براساس موقعیت و شرایط مطالعه و سپس تبدیل این متغیرهای کیفی به مقادیر کمی با استفاده از مقادیری که توسط آقای کوهن (Cohen) در سال ۱۹۷۷ ارائه شده است (جدول ۱-۵).

مثال: در مطالعه مقایسه میانگین قند خون ناشتای بیماران دیابتی با میانگین قند خون افراد عادی جامعه (۱۱۰ mg/dL) هرگاه

براساس بررسی متون، $\sigma = ۱۵$ و $d = ۳$ باشد، با فرض $\alpha = ۵\%$ و $\beta = ۲۰\%$ ، حجم نمونه بدین ترتیب محاسبه می‌شود:



$$\alpha = 0.05 \Rightarrow z_{1-\alpha/2} = 1.96$$

$$\beta = 0.2 \Rightarrow Z_{1-\beta} = 0.85$$

$$n = \frac{(1.96 + 0.85)^2 \times 15^2}{3^2} \approx 110$$

یعنی محقق براساس روش نمونه‌گیری طرح، ۱۱۰ بیمار مبتلا به دیابت را انتخاب کرده و میانگین قند خون ناشای آنها را محاسبه می‌کند و سپس میانگین بدست آمده را با عدد ثابت ۱۱۰ mg/dL که مربوط به افراد جامعه است، مقایسه می‌نماید.

نوع تست	فرمول محاسبه E.S	مقادیر E.S		
		کوچک	متوسط	بزرگ
T test for 2 group means	$d = \frac{ \bar{X}_1 - \bar{X}_2 }{\sigma}$	0.2	0.5	0.8
F test for independent means	$F = \frac{\sigma_{m_1}}{\sigma}$	0.1	0.25	0.4
Regression ($r \neq 0$)	R	0.1	0.3	0.5
Difference between 2 proportions	$h = \phi_1 - \phi_2 $ $\phi_1 = \text{Zarscine} \sqrt{P_1}$	0.2	0.5	0.8
Multiple regression ($R^2 \neq 0$)	$f^2 = \frac{R^2}{1 - R^2}$	0.2	0.15	0.35

جدول ۱-۵

مقایسه دو میانگین

برای محاسبه حجم نمونه در مطالعاتی که هدف از انجام آنها مقایسه دو میانگین می‌باشد از رابطه ۵-۵ استفاده می‌گردد.

$$n = \frac{2(Z_{1-\alpha/2} + Z_{1-\beta})^2 \sigma^2}{d^2} \quad \text{رابطه ۵-۵}$$

که در آن مفهوم متغیر ها مشابه رابطه ۵-۴ میباشد.

مثال: در مطالعه "مقایسه میانگین زمان پروترومبین پس از مصرف هپارین با وزن مولکولی کم و هپارین معمولی در بیماران مبتلا به آنژین غیر ثابت قلبی" هرگاه براساس بررسی متون، $\sigma = 3$ و $d = 1$ ثانیه باشد، با فرض $\alpha = 0.05$ و $\beta = 0.2$ ، حجم نمونه بدین ترتیب محاسبه می‌شود:



$$\alpha = 0.05 \Rightarrow z_{1-\alpha/2} = 1.96$$

$$\beta = 0.2 \Rightarrow Z_{1-\beta} = 0.85$$

$$n = \frac{2(1.96 + 0.85)^2 \times 3^2}{1^2} \approx 140$$

یعنی محقق براساس روش نمونه‌گیری مطالعه، ۱۴۰ بیمار که تحت درمان با هپارین با وزن مولکولی کم هستند و ۱۴۰ بیمار که تحت درمان با هپارین معمولی هستند را انتخاب می‌کند و میانگین زمان پروترومبین در هر گروه را محاسبه کرده و این دو گروه را با هم مقایسه می‌نماید.

✓ نکته: رابطه ۵-۵ برای محاسبه حجم نمونه در مطالعات مقطعی تحلیلی، مورد شاهدهی، همگروهی و کارآزمایی بالینی مورد استفاده قرار می‌گیرد.

مقایسه یک نسبت با یک عدد ثابت

برای محاسبه حجم نمونه در مطالعاتی که هدف از انجام آنها مقایسه یک نسبت با یک عدد ثابت است، از رابطه ۵-۶ استفاده می‌شود:

$$n = \frac{[Z_{1-\alpha/2} \sqrt{P_0(1-P_0)} + Z_{1-\beta} \sqrt{P_1(1-P_1)}]^2}{d^2} \quad \text{رابطه ۶-۵}$$

که در این رابطه مفهوم هر پارامتر به قرار زیر است:

n : تعداد نمونه‌های مورد نیاز برای انجام مطالعه

α : احتمال خطای نوع اول

β : احتمال خطای نوع دوم

P_0 : مقدار نسبت ثابت

P_1 : احتمال مواجه یافتن با عامل خطر در مطالعات با جهت رو به عقب و احتمال ایجاد بیماری در مطالعات با جهت رو به جلو است.

d : حداقل اختلاف دارای ارزش هنگام مقایسه دو گروه

✓ نکته: در مطالعات با جهت رو به عقب، بعنوان مثال مطالعه مورد شاهدهی، جهت حرکت مطالعه از بیماری (پیامد) به سمت عامل مواجهه است. در حالیکه در مطالعات با جهت رو به جلو بعنوان مثال مطالعه همگروهی، جهت مطالعه از عامل مواجهه به سمت بیماری (پیامد) است.

مثال: در مطالعه "مقایسه درصد فراوانی بیماری افسردگی در دانشجویان پزشکی با جامعه معمولی (۲۰٪)" هرگاه براساس بررسی

متون، $d = 0.1$ و $P_1 = 0.4$ باشد، با فرض $\alpha = 0.05$ و $\beta = 0.2$ ، حجم نمونه بدین ترتیب محاسبه می‌شود:

یعنی محقق باید براساس روش نمونه‌گیری مطالعه، ۱۴۴ دانشجوی پزشکی را انتخاب کرده و مشخص کند که چه نسبتی از این دانشجویان مبتلا به بیماری افسردگی‌اند و سپس این نسبت را با عدد ثابت ۰.۲ که مربوط به جامعه می‌باشد، مقایسه نماید.

✓ نکته: رابطه ۵-۵ برای محاسبه حجم نمونه در مطالعات مقطعی تحلیلی مورد استفاده قرار می‌گیرد.

$$\alpha = 0.05 \Rightarrow z_{1-\alpha/2} = 1.96$$

$$\beta = 0.2 \Rightarrow Z_{1-\beta} = 0.85$$

$$p_0 = 0.2$$

$$n = \frac{[1.96\sqrt{0.2(1-0.2)} + 0.85\sqrt{0.4(1-0.4)}]^2}{(0.1)^2} \approx 144$$

مقایسه دو نسبت

هرگاه هدف از انجام مطالعه‌ای، مقایسه دو نسبت باشد، برای محاسبه حجم نمونه از رابطه ۵-۷ استفاده می‌گردد:

$$n = \frac{[Z_{1-\alpha/2}\sqrt{2\bar{P}(1-\bar{P})} + Z_{1-\beta}\sqrt{P_0(1-P_0) + P_1(1-P_1)}]^2}{(P_1 - P_0)^2} \quad \text{رابطه ۵-۷}$$

که در این رابطه تعریف هر پارامتر به قرار زیر است:

1. تعداد نمونه‌های مورد نیاز برای انجام مطالعه

α : احتمال خطای نوع اول

β : احتمال خطای نوع دوم

$\bar{P}$: میانگین حسابی دو نسبت $(\frac{P_0 + P_1}{2})$

P_0 : در مطالعات با جهت رو به عقب، بیانگر احتمال مواجهه یافتن با عامل خطر در گروه کنترل و در مطالعات با جهت رو به جلو، بیانگر احتمال ایجاد بیماری در گروه مواجهه نیافته از جامعه هدف است.

P_1 : در مطالعات با جهت رو به عقب، بیانگر احتمال مواجهه یافتن در گروه مورد و در مطالعات با جهت رو به جلو، بیانگر احتمال ایجاد بیماری در گروه مواجهه یافته از جامعه هدف است.

✓ نکته: چنانچه محقق برآوردی از احتمال مواجهه یافتن (یا ایجاد بیماری) در گروه مورد را نداشته باشد، می‌توان P_1 را

با استفاده از P_0 و Odds Ratio (OR) در مطالعات مورد شهادی یا Relative Risk (RR) در مطالعات همگروهی و

کارآزمایی بالینی، بدست آورد:

$$P_1 = \frac{P_0 \times OR}{1 + P_0(OR - 1)} \quad \text{یا} \quad P_1 = \frac{P_0 \times RR}{1 + P_0(RR - 1)}$$

✓ نکته: OR و RR شاخصهایی‌اند که شدت ارتباط بین پیامد و عامل مواجهه را نشان می‌دهند.

Odd عبارتست از احتمال اینکه یک پیامد اتفاق افتد، تقسیم بر احتمال اینکه آن پیامد اتفاق نیفتد.

$$O = \frac{P}{1-P} \quad \begin{matrix} \text{(اتفاق یک پیامد)} \\ \text{(عدم اتفاق آن پیامد)} \end{matrix} \quad 0 \leq P \leq 1$$

OR در یک مطالعه مورد شهادی از نسبت بین Oddهای مواجهه با عامل خطر (E) در دو گروه مورد و شاهد محاسبه می‌شود:

	E^+	E^-	
D^+	a	b	$a+b$
D^-	c	d	$c+d$
	$a+c$	$b+d$	

$$Odd = O(E^+, D^+) = \frac{P(E^+, D^+)}{P(E^-, D^+)} = \frac{\frac{a}{a+b}}{\frac{b}{a+b}} = \frac{a}{b}$$

$$Odd = O(E^+, D^-) = \frac{P(E^+, D^-)}{P(E^-, D^-)} = \frac{\frac{c}{c+d}}{\frac{d}{c+d}} = \frac{c}{d}$$

$$OR = \frac{O(E^+, D^+)}{O(E^+, D^-)} = \frac{\frac{a}{b}}{\frac{c}{d}} = \frac{ad}{bc}$$



چنانچه OR برابر یک گردد بدان معناست که احتمال وجود عامل مواجهه در دو گروه مورد و شاهد برابر بوده و لذا ارتباطی بین پیامد و عامل مواجهه موجود نیست. چنانچه OR بزرگتر از یک گردد بدان معناست که احتمال وجود عامل مورد نظر در گروه مورد بیش از گروه شاهد می باشد و لذا آن عامل، یک عامل خطر برای ایجاد بیماری است و در صورتیکه OR کوچکتر از یک گردد بدان معناست که احتمال وجود عامل مورد نظر در گروه مورد کمتر از گروه شاهد است و لذا آن عامل، یک عامل پیشگیرنده یا محافظتی در ایجاد بیماری است.

RR با خطر نسبی در مطالعه همگروهی از نسبت بین احتمال بروز بیماری (پیامد) در گروه مواجهه یافته بدست می آید.

چنانچه RR برابر با یک گردد بدان معناست که احتمال بروز پیامد در دو گروه مواجهه یافته و مواجهه نیافته برابر است و در این شرایط هیچ ارتباطی بین پیامد و عامل خطر وجود ندارد. چنانچه RR بزرگتر از یک گردد بدان معناست که احتمال بروز پیامد در گروه مواجهه یافته بیش از گروه مواجهه نیافته است و در واقع عامل مواجهه، خطر بروز پیامد را بالا می برد و بعنوان یک عامل خطر محسوب می شود و چنانچه RR کوچکتر از یک گردد بدان معناست که احتمال بروز پیامد در گروه مواجهه یافته کمتر از گروه مواجهه نیافته است و در واقع عامل مواجهه احتمال بروز پیامد را کم می کند یعنی یک عامل محافظتی است.

	D^+	D^-	
E^+	a	b	$a+b$
E^-	c	d	$c+d$
	$a+c$	$b+d$	

$$P(D^+, E^+) = \frac{a}{a+b}$$

$$P(D^+, E^-) = \frac{c}{c+d}$$

$$RR = \frac{P(D^+, E^+)}{P(D^+, E^-)} = \frac{\frac{a}{a+b}}{\frac{c}{c+d}}$$



مثال: در مطالعه مورد شاهدهی "تعیین ارتباط زخم پیتیک سوراخ شده با مصرف سیگار" هرگاه براساس بررسی متون، $OR=3$ و

$P_0=0.3$ باشد، با فرض $\alpha=0.05$ و $\beta=0.1$ ، حجم نمونه بدین ترتیب محاسبه می شود:

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$\beta = Z_{1-\beta} = 1.28$$

$$P_1 = \frac{0.3 \times 3}{1 + 0.3(3-1)} = 0.56$$

$$\bar{P} = \frac{0.3 + 0.56}{2} = 0.43$$

$$n = \frac{[1.96\sqrt{2 \times 0.43(1-0.43)} + 1.28\sqrt{0.56(1-0.56)} + 0.3(1-0.3)]^2}{(0.56-0.3)^2} = 73$$

یعنی محقق براساس روش نمونه گیری مطالعه، ۷۳ بیمار مبتلا به زخم پیتیک سوراخ شده (گروه مورد) و ۷۳ نفر غیر مبتلا به زخم پیتیک سوراخ شده (گروه شاهد) را انتخاب کرده و نسبت افراد سیگاری و غیر سیگاری در هر گروه را با مراجعه به سابقه افراد مشخص می کند و در نهایت نسبت افراد سیگاری در دو گروه مورد و شاهد را با هم مقایسه می نماید.

مثال: در مطالعه همگروهی "بررسی ارتباط بین سطح بالای لاکتات سرم و میزان مرگ و میر بیماران بستری در بخش مراقبتهای ویژه" هرگاه براساس بررسی متون، $RR=4$ و $P_0=0.2$ باشد، با فرض $\alpha=0.05$ و $\beta=0.2$ ، حجم نمونه بدین ترتیب محاسبه می شود:

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$\beta = Z_{1-\beta} = 0.85$$

$$P_1 = \frac{0.2 \times 4}{1 + 0.2(4-1)} = 0.5$$

$$\bar{P} = \frac{0.2 + 0.5}{2} = 0.35$$

$$n = \frac{[1.96\sqrt{2 \times 0.35(1-0.35)} + 0.85\sqrt{0.2(1-0.2)} + 0.5(1-0.5)]^2}{(0.5-0.2)^2} = 39$$

یعنی محقق براساس روش نمونه گیری مطالعه، ۳۹ بیمار بستری در بخش مراقبتهای ویژه که سطح لاکتات سرم آنها بالا است (گروه مواجهه یافته) و ۳۹ بیمار بستری در بخش مراقبتهای ویژه که سطح لاکتات سرم آنها طبیعی است (گروه مواجهه نیافته) را انتخاب کرده و پس از گذشت زمان پیگیری، مشخص می کند که در هر گروه، چه تعداد از بیماران فوت نموده اند و چه تعداد زنده مانده اند و در نهایت نسبت بیمارانی که فوت نموده اند را در دو گروه با سطح لاکتات بالا و سطح لاکتات طبیعی مقایسه می کند.



✓ نکته: رابطه ۵-۷ برای محاسبه حجم نمونه در مطالعات مقطعی - تحلیلی، مورد - شاهدهی، همگروهی و کارآزمایی بالینی مورد استفاده قرار می‌گیرد.

✓ نکته: در اکثر تحقیقات خصوصاً تحقیقاتی که نیاز به پیگیری بیماران دارد (همچون همگروهی و کارآزمایی بالینی)، تعدادی از افراد مورد پژوهش پس از مدتی به دلایل معلوم و حتی گاهی نا معلوم از ادامه حضور در پژوهش سر باز می‌زنند؛ لذا جهت جلوگیری از اتلاف زمان در تحقیق و همچنین برآورد هزینه در ابتدای تحقیق باید حجم نمونه نهایی را برحسب پیش‌بینی از داده‌های از دست رفته تصحیح نمود:

$$n' = n \times \frac{1}{1 - P} \quad \text{رابطه ۵-۷}$$

که در آن هریک از پارامترها به قرار زیر تعریف می‌گردد:

n' : حجم نمونه تصحیح شده

n : حجم نمونه اولیه

P : پیش‌بینی درصد داده‌های از دست رفته

به عنوان مثال در طرح تحقیقاتی "تعیین اثر روزه‌داری بر حجمهای ریوی" که در کمینه پژوهشی دانشجویی دانشگاه علوم پزشکی ایران انجام یافت، جهت تعیین این اثر تعدادی از افراد قبل از ماه رمضان و دو بار در ماه رمضان و بعد از ماه رمضان مورد ارزیابی حجمهای ریوی توسط دستگاه اسپرومتری قرار گرفتند. حجم نمونه ابتدایی محاسبه شده در این تحقیق با استفاده از رابطه مقایسه میانگینها ۱۱۰ نفر محاسبه گردید اما برآورد محققین از درصد افرادی که در نوبت اول مراجعه کرده و در نویتهای بعدی به دلایل گوناگون از جمله عدم روزه‌داری، نداشتن وقت و غیره مراجعه نمی‌کردند ۲۵٪ بود لذا براساس رابطه ۵-۸ حجم نمونه در ابتدای تحقیق ۱۴۶ نفر پیشنهاد گردید تا در انتهای طرح و در مرحله آنالیز حداقل ۱۱۰ نفر جهت تجزیه و تحلیل آماری وجود داشته باشند.

✓ نکته: چنانچه گروههای مورد مقایسه در یک تحقیق بیش از دو گروه باشند، حجم نمونه نهایی باید تصحیح گردد زیرا روابط ۵-۵ و ۵-۶ صرفاً جهت مقایسه میانگین و نسبت در دو گروه می‌باشند. برای تصحیح حجم نمونه محاسبه شده از رابطه

$$n' = \sqrt{K} n \quad \text{رابطه ۵-۹}$$

۵-۸ استفاده می‌گردد:

که در آن هریک از پارامترها به ترتیب زیر تفسیر می‌شوند:

n' : حجم نمونه تصحیح شده

n : حجم نمونه اولیه

k : تعداد مقایسه هر گروه با گروههای دیگر یا به عبارتی تعداد کل گروهها منهای یک می‌باشد.

به عنوان مثال در مطالعه "مقایسه اثر درمانی اسپری سالبوتامول، آترونت، بکلومتازون و سالمترون در درمان آسم" چنانچه براساس رابطه ۵-۷ حجم نمونه ۵۰ نفر محاسبه گردد، با توجه به اینکه در این تحقیق ۴ گروه وجود داشته و در واقع هر گروه با سه گروه دیگر مقایسه می‌گردد، حجم نمونه نهایی به قرار زیر محاسبه می‌شود:

$$n' = \sqrt{3} \times 50 \approx 86$$

و بدین ترتیب محقق در هریک از گروههای فوق ملزم به تحقیق بر روی ۸۶ بیمار مبتلا به آسم می‌باشد.



محاسبه حجم نمونه بر مبنای نوع مطالعه

علاوه بر تقسیم‌بندی ذکر شده، روابط محاسبه حجم نمونه بر مبنای نوع مطالعه نیز، تقسیم‌بندی می‌شوند و روابط اختصاصی برای محاسبه حجم نمونه، در برخی از انواع مطالعات ارائه شده است که در ذیل به توضیح آن می‌پردازیم.

مطالعات مشاهده‌ای تحلیلی

این مطالعات خود به سه دسته مطالعات مقطعی - تحلیلی، مورد - شاهدهی و همگروهی تقسیم می‌شوند.

۱- مطالعه مقطعی - تحلیلی

برای محاسبه حجم نمونه در مطالعات مقطعی - تحلیلی هرگاه هدف محقق مقایسه یک میانگین با یک عدد ثابت و یا مقایسه دو میانگین باشد، از رابطه ۵-۴ و ۵-۵ هرگاه هدف مقایسه یک نسبت با یک عدد ثابت باشد از رابطه ۵-۶ و هرگاه هدف مقایسه دو نسبت باشد، از رابطه ۵-۶ استفاده می‌شود.

۲- مطالعه مورد - شاهدهی

از آنجاییکه هدف محقق در این نوع مطالعه، اغلب مقایسه نسبت یک متغیر در دو گروه شاهد و مورد است، برای محاسبه حجم نمونه مورد نیاز، از رابطه ۵-۷ استفاده می‌شود. ولی از آنجاییکه در اغلب موارد، بعلت نادر بودن بیماری مورد بحث در مطالعه مورد - شاهدهی، تعداد کمی بیمار جهت نمونه‌گیری در اختیار محقق است و با حضور این تعداد بیمار به همراه همین تعداد شاهد، آنالیز نتایج مطالعه به علت کم بودن حجم نمونه امکان پذیر نیست، می‌توان به بهای افزایش حجم نمونه کل، حجم نمونه گروه مورد را کاهش داد و به ازای هر مورد بیش از یک شاهد (حداکثر ۹ شاهد) انتخاب کرد، برای محاسبه حجم نمونه در شرایطی که به ازای هر مورد بیش از یک شاهد انتخاب شده است، از رابطه ۵-۱۰ استفاده می‌شود:

$$n = \frac{\left[Z_{1-\alpha/2} \sqrt{\left(1 + \frac{1}{c}\right) \bar{P}(1-\bar{P})} + Z_{1-\beta} \sqrt{P_1(1-P_1) + \frac{P_0(1-P_0)}{c}} \right]^2}{(P_1 - P_0)^2} \quad \text{رابطه ۵-۱۰}$$

که در این رابطه هریک از پارامترها به قرار زیر تعریف می‌شوند:

n : تعداد نمونه‌های مورد نیاز برای انجام طرح

α : احتمال خطای نوع اول

β : احتمال خطای نوع دوم

P_0 : احتمال مواجهه در گروه کنترل (شاهد)

P_1 : احتمال مواجهه در گروه مورد

$C = \text{Control/Case}$

C : نسبت شاهد به مورد

$$\bar{P} = \frac{P_1 + CP_0}{1 + C} \quad \bar{P}: \text{میانگین حسابی } P_1, P_0$$

✓ نکته: حجم نمونه بدست آمده، مربوط به گروه مورد است و حجم نمونه گروه شاهد از حاصل ضرب این مقدار در مقدار C بدست می‌آید.

مثال: در مطالعه "تعیین عوامل مؤثر بر مسمومیت کبدی ناشی از درمان ضد سل ریوی" هرگاه براساس اطلاعات بدست آمده در بررسی متون، $OR=3$ و $P_0=20\%$ باشد و نسبت شاهد به مورد برابر ۳ انتخاب شود با فرض $\alpha=5\%$ و $\beta=10\%$ ، حجم نمونه بدین صورت محاسبه می‌شود:

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$\beta = Z_{1-\beta} = 1.28$$

$$P_1 = \frac{0.2 \times 3}{1 + 0.2(3-1)} = 0.42$$

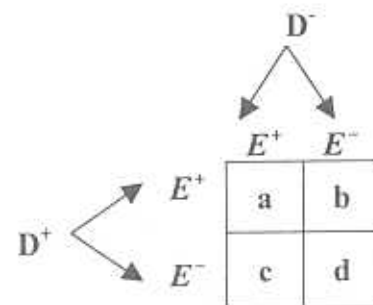
$$\bar{P} = \frac{0.3 \times 0.42}{1 + 3} = 0.25$$

$$n = \frac{\left[1.96 \sqrt{\left(1 + \frac{1}{3}\right) 0.25(1-0.25)} + 1.28 \sqrt{0.42(1-0.42) + \frac{0.2(1-0.2)}{3}} \right]^2}{(0.42 - 0.2)^2} = 57$$

یعنی محقق براساس روش نمونه‌گیری مطالعه، ۵۷ بیمار تحت درمان یا داروی ضد سل که دچار مسمومیت کبدی شده‌اند و $57 \times 3 = 171$ بیمار تحت درمان با داروی ضد سل که دچار مسمومیت کبدی نشده‌اند، را انتخاب کرده و متغیرهای مستقل مختلف را در این دو گروه مقایسه می‌نماید.

۳- مطالعه مورد-شاهدی همانند شده:

هرگاه در مرحله نمونه‌گیری مطالعه مورد-شاهدی، همانندسازی جفتی انجام شده باشد، یعنی به ازای هر مورد، یک شاهد که از تمام جهات مهم بجز عامل مواجهه با آن همانند شده است، انتخاب شود، حجم نمونه مورد نیاز طرح به روشی متفاوت محاسبه می‌شود. در ابتدا نتایج مطالعه در جدولی همانند جدول ۵-۲ جمع‌آوری می‌شود:



جدول ۵-۲

که در این جدول:

a: تعداد جفتهایی که عضو گروه بیمار و سالم آن هر دو مواجهه یافته‌اند (+/+).

b: تعداد جفتهایی که عضو گروه بیمار آن مواجهه یافته و شاهد همانند با آن، مواجهه نیافته است (+/-).

c: تعداد جفتهایی که عضو بیمار آن مواجهه نیافته و عضو سالم همانند با آن، مواجهه یافته است (-/+).

d: تعداد جفتهایی که اعضاء مورد و شاهد، هر دو مواجهه نیافته‌اند (-/-).

واضح است که برای بدست آوردن Odds Ratio در چنین حالتی، نسبت b به c بسیار مهم است زیرا این نسبت بیانگر نسبت مواجهه یافته‌ها در گروه بیمار که عضو مشابه آنها در گروه شاهد مواجهه نیافته‌اند، به مواجهه یافته‌ها در گروه شاهد که عضو مشابه آنها در گروه مورد مواجهه نیافته‌اند، می‌باشد.

$$OR = \frac{b}{c}$$

و در نهایت حجم نمونه بدین ترتیب محاسبه می‌شود:

$$m = \frac{\left[\frac{Z_{\alpha}}{2} + Z_{1-\beta} \sqrt{P(1-P)} \right]^2}{\left(P - \frac{1}{2} \right)^2} \quad \text{رابطه ۵-۱۱}$$

که در این رابطه:

m: تعداد جفتهای مورد نیاز برای انجام طرح می‌باشد که از لحاظ عامل مواجهه با یکدیگر متفاوتند (Discordant Pairs) یعنی جفتهایی که یکی از دو حالت (+/+) یا (+/-) هستند.

α : احتمال خطای نوع اول

β : احتمال خطای نوع دوم

$$\frac{OR}{1+OR} : P$$

و برای محاسبه حداقل تعداد جفتهای نهایی مورد نیاز از رابطه ۵-۱۱ استفاده می‌شود:

$$M = \frac{m}{[P_0(1-P_1) + P_1(1-P_0)]} \quad \text{رابطه ۵-۱۲}$$

که در این رابطه:

m: تعداد جفتهای مورد نیاز غیر همان

M: تعداد کل جفتهای مورد نیاز

P_0 : احتمال مواجهه یافتن در گروه کنترل

P_1 : احتمال مواجهه یافتن در گروه مورد



در واقع جهت دستیابی به m جفت غیر همسان از لحاظ عامل مواجهه، به M جفت همانند شده اولیه نیاز است. مثال: در مطالعه تعیین ارتباط استفاده از قرصهای ضد بارداری و ناهنجاریهای قلبی چنین "هرگاه براساس بررسی متون،

$P_0 = 0.3$ و $OR = 2$ باشد، با فرض $\alpha = 0.05$ و $\beta = 0.1$ حجم نمونه بدین ترتیب محاسبه می‌شود:

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$\beta = 0.1 \Rightarrow Z_{1-\beta} = 1.28$$

$$P = \frac{2}{1+2} = \frac{2}{3}$$

$$P_1 = \frac{P_0 \times OR}{1 + P_0(OR - 1)} = 0.46$$

$$m = \frac{\left[\frac{1.96}{2} + 1.28 \sqrt{\frac{2}{3} \times \frac{1}{3}} \right]^2}{\left(\frac{2}{3} - \frac{1}{2} \right)^2} = 90$$

$$M = \frac{90}{0.3(1 - 0.46) + 0.46(1 - 0.3)} \approx 186$$

یعنی محقق باید حداقل ۱۸۶ جفت مورد و شاهد همانند شده انتخاب نماید تا ۹۰ جفت گوناگون از نظر عامل مواجهه، بین آنها وجود داشته باشد.

۴- اصلاح حجم نمونه براساس منابع مالی در مطالعات مشاهده‌ای-تحلیلی

در برخی مطالعات مورد-شاهدی بین هزینه پایه برای هر فرد در گروه مورد و شاهد، تفاوت وجود دارد در این زمان با توجه به اینکه هزینه کل طرح (T) ثابت و مشخص می‌باشد نسبت مورد به شاهد را می‌توان طوری انتخاب کرد که قدرت (Power) مطالعه بیشینه (Maximum) گردد.

برای محاسبه حجم نمونه براساس منابع مالی از رابطه‌های ۵-۱۳ و ۵-۱۴ استفاده می‌گردد:

$$n_0 = \frac{T(\sqrt{C_0 C_1} - C_0)}{C_0(C_1 - C_0)} \quad \text{رابطه ۵-۱۳}$$

رابطه ۵-۱۳

که در این رابطه هر پارامتر به قرار زیر تعریف می‌گردد:

n_0 : حجم نمونه گروه شاهد

T: کل هزینه طرح



C_0 : هزینه لازم برای هر شاهد

C_1 : هزینه لازم برای هر مورد

$$n_1 = \frac{T(C_1 - \sqrt{C_0 C_1})}{C_1(C_1 - C_0)} \quad \text{رابطه ۵-۱۴}$$

که در این رابطه:

n_1 : حجم نمونه گروه مورد و سایر پارامترها مانند رابطه ۵-۱۲ می‌باشند.

به عبارت دیگر بین n_1 و n_0 باید چنین نسیبی برقرار باشد:

$$\frac{n_0}{n_1} = \sqrt{\frac{C_1}{C_0}}$$

مثال: در مطالعه مورد-شاهدی که محقق فقط اجازه هزینه کردن به میزان دویست و پنجاه هزار تومان دارد و هزینه مورد نیاز برای هر مورد ۱۷۰ تومان و هر شاهد ۱۴۰ تومان می‌باشد، مطالعه زمانی بیشترین قدرت را خواهد داشت که ۷۷۱ مورد و ۸۵۰ شاهد موجود باشد. یعنی نسبت شاهد به مورد برابر با ۱/۱ است.

$$n_0 = \frac{250000(\sqrt{170 \times 140} - 140)}{140(170 - 140)} = 850$$

$$n_1 = \frac{250000(170 - \sqrt{170 \times 140})}{170(170 - 140)} = 771$$

$$\frac{n_0}{n_1} = \sqrt{\frac{170}{140}} = \frac{850}{771} = 1.1$$

۵- مطالعات همگروهی

چنانچه هدف محقق در این نوع مطالعه مقایسه نسبت یک متغیر در دو گروه مراجعه یافته و مواجهه نیافته باشد، برای محاسبه حجم نمونه این مطالعات از رابطه ۵-۷ استفاده می‌شود، چنانچه هدف نهایی تحقیق مقایسه میانگینها در دو گروه باشد از رابطه ۵-۵ استفاده می‌گردد.

تعیین حجم نمونه در مطالعات همبستگی

مطالعات همبستگی به منظور بررسی وجود ارتباط بین دو متغیر کمی انجام می‌شوند و در نهایت ضریبی تحت عنوان ضریب همبستگی (r) محاسبه می‌گردد که نشان‌دهنده شدت ارتباط خطی بین دو متغیر است. دامنه تغییرات r بین -۱ و +۱ است. مقادیر



منفی ۲ نشان‌دهنده تغییرات دو متغیر در خلاف جهت هم می‌باشد که با افزایش یک متغیر، متغیر دیگر کاهش می‌یابد. بعنوان مثال با افزایش مقدار هورمون تیروکسین (T_4)، میزان هورمون محرک تیروئید (TSH) در سرم کاهش می‌یابد. مقادیر مثبت ۲ نشان‌دهنده تغییرات در جهت هم می‌باشد بدین معنی که با افزایش یک متغیر، متغیر دیگر هم افزایش می‌یابد. بعنوان مثال با افزایش میزان کلسترول سرم، احتمال ایجاد بیماری‌های عروق قلب افزایش می‌یابد. هرچه مقدار ۲ به عدد ۱ یا -۱ نزدیکتر شود، ارتباط بین دو متغیر قوی‌تر و هرچه به عدد صفر نزدیکتر شود، ارتباط بین دو متغیر ضعیفتر است. به عنوان مثال بین دو متغیر قد و وزن در بالغین برنجی جوامع، همبستگی بسیار بالایی وجود دارد ($r \approx 0.9$). اگرچه چنین مقادیر بالایی نامعمول هستند و بسیاری از روابط بیولوژیک، ضریب همبستگی بسیار کمتری دارند.

برای تعیین حجم نمونه یک مطالعه که در آنالیز آن از ضریب همبستگی استفاده می‌شود، محقق باید:

- فرضیه H_0 را بیان کند و مشخص نماید که فرضیه آلترناتیو (H_1) یک دامنه است یا دو دامنه.
- اندازه اثر (Effect Size) که معادل مقدار مطلق کمترین ضریب همبستگی (r) است که محقق برآورد می‌کند را مشخص نماید.
- α, β مورد نظر را تعیین کند.
- از رابطه ۵-۱۵ یا جدول ضمیمه ۳ استفاده کرده و مقدار حجم نمونه را مشخص کند.

$$n = \left[\frac{Z_\alpha + Z_\beta}{C} \right]^2 + 3 \quad \text{رابطه ۵-۱۵}$$

که در این رابطه هر پارامتر به قرار زیر تعریف می‌گردد:

n: تعداد کل نمونه‌های لازم

r: کمترین ضریب همبستگی مورد انتظار

C: $0.5 \times \ln[(1+r)/(1-r)]$ (مقادیر محاسبه شده بر حسب مقادیر مختلف در جدول پیوست ۳ خلاصه گردیده است).

مثال: در مطالعه "بررسی ارتباط بین سطح کورتیزین (Cortinine) ادرار (متغیری که شدت مصرف سیگار فعلی را نشان می‌دهد) و تراکم استخوان در افراد سیگاری" اگر براساس نتایج مطالعات قبلی، کمترین ضریب همبستگی بین مصرف سیگار (به صورت نخ سیگار در روز) و تراکم استخوان، $r = -0.3$ باشد و $\alpha = 5\%$ دو دامنه و $\beta = 10\%$ فرض شود، محاسبه حجم نمونه بدین صورت است:



$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$\beta = 0.1 \Rightarrow Z_\beta = 1.28$$

$$r = 0.3$$

$$C = 0.5 \times \ln \left[\frac{1+0.3}{1-0.3} \right] \approx 0.3$$

$$N = \left[\frac{1.96 + 1.28}{0.3} \right]^2 + 3 = 113$$

یعنی محقق باید بنا به روش نمونه‌گیری مطالعه، ۱۱۳ فرد سیگاری را انتخاب کرده و سطح کورتیزین ادرار و تراکم استخوان در هر فرد را تعیین نموده و سپس آنالیز همبستگی بین این دو متغیر را انجام دهد.

✓ نکته: نحوه استفاده از جدول پیوست ۱ که برای محاسبه حجم نمونه در مطالعاتی که از همبستگی استفاده می‌شود نیز به این قرار است که اولین ستون سمت چپ جدول نشان‌دهنده مقدار r است، با حرکت افقی به سمت راست و با توجه به مقادیر α, β ، حجم نمونه مورد نیاز، مشخص می‌شود.

✓ نکته: رابطه ۵-۱۵ حجم نمونه لازم برای رد فرضیه H_0 (همبستگی بین دو متغیر وجود ندارد) را نشان می‌دهد. اگر محقق بخواهد تفاوت یک ضریب همبستگی با یک مقدار غیر از صفر (بعنوان مثال $r = 0.2$) را نشان دهد برای محاسبه حجم نمونه باید از رابطه ۵-۱۶ استفاده کند.

$$n = \left[\frac{Z_\alpha + Z_\beta}{C_1 - C_2} \right]^2 + 3 \quad \text{رابطه ۵-۱۶}$$

که در این رابطه:

$$C_1 = 0.5 \times \ln \left[\frac{1+r_1}{1-r_1} \right] \quad C_2 = 0.5 \times \ln \left[\frac{1+r_2}{1-r_2} \right]$$

مثال: برای مقایسه شدت ارتباط بین میزان قند خون ناشتا و میزان هموگلوبین A1C با شدت ارتباط بین میزان قند خون ۲ ساعت بعد از غذا و میزان هموگلوبین A1C ($r = 0.6$) در بیماران دیابتی، هرگاه براساس بررسی متون انجام شده، کمترین ضریب همبستگی بین قند خون ناشتا و HbA1C، 0.4 باشد، با فرض $\alpha = 0.05$ و $\beta = 0.1$ حجم نمونه مورد نیاز بدین ترتیب محاسبه می‌شود:

$$C_1 = 0.5 \times \ln \left[\frac{1+0.6}{1-0.6} \right] = 0.693 \quad C_2 = 0.5 \times \ln \left[\frac{1+0.4}{1-0.4} \right] = 0.422$$

$$\alpha = 0.05 \Rightarrow Z_{1-\alpha/2} = 1.96$$

$$\beta = 0.1 \Rightarrow Z_{\beta} = 1.28$$

$$n = \left[\frac{1.96 + 0.85}{0.693 - 0.422} \right]^2 + 3 = 110$$

یعنی محقق باید براساس روش نمونه‌گیری مطالعه، ۱۱۰ بیمار دیابتی انتخاب کرده و میزان قند خون ناشتا و میزان هموگلوبین A1C را در هر بیمار تعیین نموده و پس از بررسی همبستگی بین این دو متغیر و محاسبه ضریب همبستگی، آن را با عدد ۰/۶ مقایسه کند.

✓ نکته: هرچه مقدار ضریب همبستگی مورد انتظار بیشتر باشد، تعداد حجم نمونه کمتر می‌شود؛ این امر را چنین می‌توان توجیه کرد که اگر رابطه‌ای با شدت بالا وجود داشته باشد، در حجم نمونه کم نیز باید به آن رابطه دست یافت. همچنین با توجه به رابطه ۵-۱۲ مشخص می‌گردد که افزایش ضریب همبستگی سبب افزایش n و در نتیجه کاهش حجم نمونه می‌شود.

محاسبه حجم نمونه جهت مقایسه تابع بقا

در بسیاری از طرحهای تحقیقاتی محقق در صدد یافتن اثر یک داروی جدید بر میزان بقا بیماران در مقایسه با داروهای فعلی می‌باشد. در چنین شرایطی برای محاسبه حجم نمونه از رابطه ۵-۱۷ استفاده می‌گردد.

$$n = \frac{4 \times (Z\alpha + Z\beta)^2}{o^2} \quad \text{رابطه ۵-۱۷}$$

که در این رابطه هر یک از پارامترها به صورت زیر تعریف می‌شوند:

n : تعداد وقوع پیامد که باید تا رسیدن به آن تعداد، مطالعه را ادامه داد

α : احتمال خطای نوع اول

β : احتمال خطای نوع دوم

o : لگاریتم نسبت خطر می‌باشد که از رابطه ۵-۱۸ بدست می‌آید (Log Hazard Ratio)

$$o = \log \Psi \quad \Psi = \frac{\log S_N(t)}{\log S_S(t)} \quad \text{رابطه ۵-۱۸}$$

به عبارت دیگر Ψ نشان می‌دهد که در زمان t خطر مرگ در گروه درمان شده با داروی جدید چند برابر خطر مرگ در گروه درمان شده با داروی استاندارد است.

مثال: بیمارانی که از هپاتیت مزمن رنج می‌برند بدلیل از کار افتادگی کبد دچار مرگ زودرس می‌شوند. در کارآزمایی بالینی طراحی شده جهت بررسی تأثیر یک داروی جدید در به تعویق انداختن مرگ بیماران، دو گروه از بیماران با درمان جدید و درمان

استاندارد مقایسه خواهند شد. اگر محقق بداند که با درمان استاندارد پس از ۵ سال احتمال بقای بیماران درمان شده با روش استاندارد ۰/۳۵ است و برآورد نماید که این احتمال برای بیماران درمان شده با روش جدید ۰/۵۵ می‌باشد، برای رد فرضیه برابری دو احتمال بقا با خطای $\alpha = 0.05$ و توان $1 - \beta = 0.90$ تعداد نمونه مورد نیاز به صورت زیر به دست می‌آید:

$$\Psi = \frac{\log S_N(t)}{\log S_S(t)} = \frac{\log(0.55)}{\log(0.35)} = 0.57$$

$$o = \log \Psi = \log(0.57) = -0.563$$

$$n = \frac{4 \times (1.96 + 1.28)^2}{(-0.583)^2} = 167 \approx 170$$

این بدان معنی است که الحاق بیماران به دو گروه را تا آنجا باید ادامه داد که در هر گروه حدوداً ۸۵ مرگ مشاهده شود؛ لذا عدد بدست آمده در این رابطه حجم نمونه کل و جهت هر دو گروه می‌باشد.

محاسبه حجم نمونه برای تحقیقات ارزیابی تستهای تشخیصی

پیش از آنکه به چگونگی محاسبه حجم نمونه برای مطالعات ارزیابی تستهای تشخیصی پرداخته شود لازم است که به طور خلاصه به برخی مفاهیم پایه که در ارزیابی تستهای تشخیصی کاربرد دارند و در محاسبه حجم نمونه نیز مورد استفاده قرار می‌گیرند اشاره گردد. مطالعه بیشتر در مورد این مفاهیم به خواننده واگذار می‌شود. در ساده‌ترین حالت، نتایج حاصله از ارزیابی یک تست تشخیصی را می‌توان در یک جدول 2×2 به صورت زیر خلاصه کرد:

TP: مثبت واقعی (True Positive): تعداد افرادی که تست آنها مثبت شده است و براساس استاندارد طلایی (Gold Standard) نیز بیمارند.

Gold Standard

Test Result		Dis ⁺	Dis ⁻
		TP	FP
	+		
	-	FN	TN
		n_1	n_2

FN: منفی کاذب (False Negative): تعداد افرادی که تست آنها منفی شده است اما براساس استاندارد طلایی بیمارند.

FP: مثبت کاذب (False Positive): تعداد افرادی که تست آنها مثبت شده است اما براساس استاندارد طلایی سالمند.

TN: منفی واقعی (True Negative): تعداد افرادی که تست آنها منفی است و براساس استاندارد طلایی نیز سالمند.



حساسیت (Sensitivity): توانایی یک تست در شناسایی افراد بیمار است که به صورت زیر تعریف می‌شود:

$$SEN = \frac{TP}{TP + FN}$$

ویژگی (Specificity): توانایی یک تست در شناسایی افراد سالم است که به صورت زیر تعریف می‌شود:

$$SPC = \frac{TN}{TN + FP}$$

✓ نکته: دو شاخص فوق را اغلب به صورت درصد گزارش می‌کنند.

نسبت درست‌نمایی مثبت (PLR: Positive Likelihood Ratio) این شاخص به تنهایی برای یک تست خاص مفهومی ندارد و در مقایسه با سایر تستها ارزش می‌یابد، بدین مفهوم که هرچه PLR یک تست بیشتر باشد ارزش آن تست در اثبات (Rule In) بیماری مورد نظر بیشتر است.

$$PLR = \frac{SEN}{1 - SPC}$$

نسبت درست‌نمایی منفی (NLR: Negative Likelihood Ratio) این شاخص نیز مانند PLR برای یک تست خاص به تنهایی مفهومی ندارد و در مقایسه با سایر تستها ارزشمند می‌شود، بدین مفهوم که هرچه NLR یک تست کمتر باشد ارزش آن تست در رد (Rule Out) بیماری مورد نظر بیشتر خواهد بود.

$$NLR = \frac{1 - SEN}{SPC}$$

از میان شاخصهای ذکر شده، PLR و NLR در تصمیم‌گیریهای بالینی ارزشمند می‌باشند. روشی که برای محاسبه حجم نمونه تستهای تشخیصی بیان می‌شود نیز با استفاده از نسبتهای درست‌نمایی (Likelihood Ratios) و مفهوم سطح اطمینان (Confidence Interval) می‌باشد. سطح اطمینان برای LR از رابطه کلی ۱۹-۵ به دست می‌آید:

$$LR = \exp \left(Ln \frac{P_1}{P_2} \pm Z_{\alpha} \sqrt{\frac{1-P_1}{P_1 n_1} + \frac{1-P_2}{P_2 n_2}} \right) \quad \text{رابطه ۱۹-۵}$$

که در آن هر یک از پارامترها به قرار زیر تعریف می‌گردد:

α : احتمال خطای آماری نوع یک

n_1 : تعداد افراد بیمار است (TP+FN) و n_2 : تعداد افراد سالم است (FP+TN)

• برای محاسبه PLR: $p_1 = SEN$, $p_2 = 1 - SPC$, $p_1 n_1 = TP$, $p_2 n_2 = FP$

• برای محاسبه NLR: $p_1 = 1 - SEN$, $p_2 = SPC$, $p_1 n_1 = FN$, $p_2 n_2 = TN$

✓ نکته: نمونه‌گیری در ارزیابی تستهای تشخیصی بدین صورت است که مطالعه را تا جایی ادامه می‌دهند که استاندارد طلایی (Gold Standard) دست کم n_1 نفر را بیمار و n_2 نفر را سالم تشخیص دهد.

✓ نکته: مجموع n_1 و n_2 حجم نمونه مورد نیاز برای ارزیابی تست تشخیصی مورد نظر است.



✓ نکته: محقق حساسیت و ویژگی را با توجه به آنچه در بررسی متون بدان رسیده و یا با استفاده از نتایج حاصله از مطالعات پیش‌آزمون تخمین زده و در رابطه فوق قرار می‌دهد. نسبت درست‌نمایی را نیز محقق براساس مورد استفاده و کارایی تست مورد نظر تعیین می‌نماید. با داشتن نسبت درست‌نمایی، حساسیت و ویژگی تخمینی می‌توان معادله را بر حسب n حل کرد و حجم نمونه مورد نیاز را به دست آورد.

مثال: محقق برای تعیین حجم نمونه تست تشخیصی که PLR آن ارزش زیادی دارد، مطالعه پیش‌آزمون انجام داده و به نتایج زیر رسیده است:

$$Sen = 0.8 \quad Spc = 0.73 \quad PLR = Sen / (1 - Spc) = 2.96$$

محقق اعتقاد دارد که این تست در صورتی ارزشمند خواهد بود که $PLR \geq 2$ باشد. برای تعیین حجم نمونه عدد ۲ را به عنوان حد پایین در فرمول سطح اطمینان قرار می‌دهد. حال با فرض اینکه $n_2 = n_1$ باشد، حجم نمونه را محاسبه می‌نماید:

$$2 = \exp \left(Ln \frac{0.8}{0.27} \pm 1.96 \sqrt{\left(\frac{1}{n} \right) \left[\left(\frac{0.2}{0.8} \right) + \left(\frac{0.73}{0.27} \right) \right]} \right)$$

با حل معادله فوق $n = 73/4$ می‌شود یعنی $n_1 = 74$ و $n_2 = 74$. حجم نمونه کل ۱۴۸ نفر خواهد بود. با این حجم نمونه می‌توان با اطمینان ۹۵٪ گفت که PLR تست مورد نظر بالاتر از ۲ خواهد بود.

در مثالی دیگر، تست تشخیصی (الف) تنها زمانی مفید است که $NLR \leq 0.4$ باشد؛ محقق در بررسی متون به $Sen = 90\%$ و $NLR = 0.7$ رسیده است. حال با فرض اینکه $n_1 = n_2$ باشد حجم نمونه مورد نیاز را محاسبه می‌کند. ابتدا باید SPC را با توجه به تعریف NLR به دست آورد:

$$NLR = 1 - Sen / Spc$$

$$0.2 \times Spc = 1 - 0.9 \Rightarrow Spc = 0.5$$

حال محقق Sen و SPC را داشته و ۰/۴ را به عنوان حد بالا در فرمول سطح اطمینان قرار می‌دهد.

$$0.4 = \exp \left(Ln \frac{0.1}{0.5} \pm 1.96 \sqrt{\left(\frac{1}{n} \right) \left[\left(\frac{0.9}{0.1} \right) + \left(\frac{0.5}{0.5} \right) \right]} \right)$$

با حل معادله، $n = 79/9$ به دست می‌آید. یعنی $n_1 = 80$ و $n_2 = 80$ و حجم نمونه کل ۱۶۰ نفر خواهد بود.

✓ نکته: بسیار پیش می‌آید که محققین بالینی در پیدا کردن بیمار برای ورود به مطالعه مشکل داشته باشند. مثلاً در مورد بیماریهایی با شیوع کم یا بیمارانی با شرایط خاص و نادر و یا حتی زمانی که تعداد زیادی از بیماران حاضر به همکاری نمی‌باشند. در اینگونه موارد محقق به ازای هر بیمار، دو یا بیشتر فرد سالم انتخاب می‌کند. به عبارت دیگر n_1 یا n_2 برابر نمی‌باشد بلکه n_2 چند برابر n_1 است. برای محاسبه حجم نمونه از همان معادله سطح اطمینان نسبت درست‌نمایی استفاده می‌شود با این تفاوت که اگر $n_2/n_1 = r$ باشد به جای n_1 ، $r \times n_2$ قرار داده و معادله را بر حسب n_2 حل می‌کنند. حجم نمونه کل برابر با $r \times n_2 + n_2$ است.

فرض کنید در مثال ۱ محقق پس از محاسبه n_1 دریافت که نمی‌تواند این تعداد بیمار را پیدا کند که حاضر به همکاری نیز باشند. ناچار به ازای هر بیمار ۵ فرد سالم انتخاب نمود. با توجه به داده‌های مثال ۱ می‌توان حجم نمونه جدید را به دست آورد:

$$r = n_1/n_2 = 0.2$$

$$2 = \exp \left[Ln \frac{0.8}{0.27} \pm 1.96 \sqrt{\frac{1-0.8}{(0.8)(0.2)n_2} + \frac{1-0.27}{(0.27)n_2}} \right]$$

پس از حل معادله برحسب n_2 ، $n_2 = 98/3$ بدست می‌آید. یعنی $n_2 = 99$ نفر و $n_1 = 0.2 \times 99 = 20$ نفر که حجم نمونه کل $n_1 + n_2 = 119$ نفر می‌باشد.

محاسبه حجم نمونه کارآزمایی‌های بالینی با استفاده از نئوموگرام آلتمن (Altman's Nomogram) نئوموگرامها، نمودارهایی هستند که براساس روابط آماری رسم می‌شوند و استفاده از این روابط را آسان می‌سازند. در اینجا به بررسی نئوموگرام آلتمن که برای محاسبه حجم نمونه در کارآزمایی‌های بالینی (Clinical Trials) و در سه تست آماری پرکاربرد (Paired t-Test, Chi-Squared Test, Unpaired t-Test) می‌باشد پرداخته می‌شود. لازم به ذکر است که تمامی نئوموگرامها برای محاسبه حجم نمونه نیستند بلکه بسیاری از آنها در موارد دیگر علوم آماری کاربرد دارند.

نئوموگرام آلتمن (Altman's Nomogram) دارای دو خط عمودی یکی برای توان (power) و دیگری برای اختلاف استاندارد شده (d: Standardised Difference) می‌باشد که در میان این دو، چند خط مورب و اگر قرار دارند که از صفر خط مربوط به اختلاف استاندارد شده، شروع می‌شوند. این خطوط مورب هر کدام معرف یک آلفا (احتمال خطای آماری نوع یک) هستند و روی آنها اعداد حجم نمونه قرار دارند.

✓ نکته: اختلاف (Difference)، حداکثر اختلاف میانگین به دست آمده، با میانگین واقعی است که از نظر بالینی قابل قبول باشد. برای استاندارد کردن آن و به دست آوردن اختلاف استاندارد شده، اختلاف را بر یک شاخص پراکندگی تقسیم می‌کنند. برای محاسبه اختلاف استاندارد شده در نئوموگرام آلتمن به تکنیک تستهای آماری مربوطه از جدول ۵-۳ استفاده می‌شود.

محاسبه حجم نمونه با نئوموگرام ساده است. مقدار در نظر گرفته شده برای توان آزمون را روی خط عمودی مربوط به آن مشخص می‌کنند، برای اختلاف استاندارد شده نیز به همینگونه و سپس این دو نقطه را به هم وصل می‌نمایند. خط رسم شده، هر خط مورب را در یک نقطه قطع می‌کند. هر کدام از این نقاط نشان‌دهنده یک حجم نمونه در یک α مشخص است.

✓ نکته: حجم نمونه‌ای که از نئوموگرام آلتمن بدست می‌آید (N)، برای تستهای آماری مختلف مفاهیم متفاوتی دارد. به عنوان مثال در مقایسه دو میانگین در دو گروه مستقل (Unpaired t-Test)، N حجم نمونه کل است که نیمی مورد و نیمی شاهد خواهد بود. مفهوم N برای این سه تست در جدول ۵-۳ آمده است.

مثال: اگر توان آزمون و اختلاف استاندارد شده ۸۰٪ در نظر گرفته شوند. خط رسم شده، خط مورب $\alpha = 5\%$ را در عدد ۴۹ و خط مورب $\alpha = 1\%$ را در عدد ۷۱ قطع می‌کند. بدین معنی که برای انجام مطالعه کارآزمایی بالینی که با تست آماری Paired t-Test

آنالیز می‌شود، با توان و اختلاف استاندارد شده ۸۰٪ با $\alpha = 5\%$ به ۹۸ نفر و با $\alpha = 1\%$ به ۱۴۲ نفر نیاز است. اگر تست آماری مورد نظر Chi-Squared Test باشد، با $\alpha = 5\%$ به ۴۹ نفر و با $\alpha = 1\%$ به ۷۱ نفر در کل نیاز است.

مفهوم N	اختلاف استاندارد شده	آزمون فرضیه
حجم نمونه کل	$\frac{\delta}{\sigma}$	Unpaired t-TEST (مقایسه دو میانگین در دو گروه مستقل)
حجم نمونه در هر گروه (مورد و شاهد)	$\frac{2\delta}{\sigma_d}$	Paired t Test (مقایسه دو میانگین در دو گروه وابسته)
حجم نمونه کل	$\frac{P_1 - P_2}{\sqrt{P(1-P)}}$	Chi Square Test (مقایسه دو نسبت در دو گروه مستقل)

جدول ۵-۳

پارامترهای جدول ۵-۳ به قرار زیرند:

δ : اختلاف

σ : انحراف معیار

$P_1 - P_2$: همان مفهوم d را برای متغیرهای کیفی دارد.

$\bar{P}$: میانگین حسابی P_1, P_2

σ_p : انحراف معیار اختلافها که اغلب از مطالعات پیش آزمون برآورد می‌شود.

✓ نکته: نئوموگرام آلتمن حجم نمونه کمتری را نسبت به رابطه‌ها می‌دهد ولی دقت کارآزمایی‌های بالینی این مشکل را حل می‌کند.

محاسبه حجم نمونه با استفاده از نرم‌افزار EPI-Info 6

جهت محاسبه حجم نمونه در تحقیقات علوم پزشکی نرم‌افزارهای گوناگونی طراحی گردیده است. بسیاری از این نرم‌افزارها که در بانک اطلاعات جهانی اینترنت نیز موجود می‌باشند صرفاً جهت محاسبه حجم نمونه هستند اما محاسبه حجم نمونه در قالب قسمتی از بعضی نرم‌افزارهای آنالیز در علوم پزشکی نیز گنجانده شده‌اند. از جمله این نرم‌افزارها، نرم‌افزار EPI Info (Epidemiological Information) می‌باشد که توسط سازمان بهداشت جهانی (WHO) طراحی گردیده است. در این نرم‌افزار امکاناتی فراهم شده که محقق توانایی محاسبه حجم نمونه در تحقیقات علوم پزشکی را می‌یابد.

برای وارد شدن به برنامه تحت سیستم DOS دستورات زیر تایپ می گردد (برنامه تحت ویندوز نیز قابل اجرا می باشد).

C:\>cd epi6

C:\EPI6>epi6

پس از اجرای برنامه بصورت فوق، پنجره اصلی این نرم افزار باز می شود که دارای منوهای متفاوتی است.

جهت محاسبه حجم نمونه، در منوی Program وارد زیرمنوی EPITABLE Calculator شوید. این زیرمنو خود دارای قسمتهای مختلفی است که یکی از مهمترین آنها زیرمنوی Sample می باشد. این زیرمنو دارای قسمتهایی به قرار زیر است.

Sample Size
Power Calculation
Random Number Table
Random Number List

جهت محاسبه حجم نمونه، جعبه رنگی را بر روی Sample Size قرار داده، و enter کنید. نرم افزار EPI، ۴ حالت مختلف برای محاسبه حجم نمونه در اختیار کاربر قرار می دهد که در ذیل به آنها اشاره خواهد شد.

برآورد یک نسبت

پس از انتخاب گزینه Single Proportion پنجره ای مشابه تصویر ۵-۲ باز می شود. در این پنجره اطلاعات شامل اندازه جامعه هدف (Size of the Population)، میزان دقت (Desired Precision)، شیوع قابل پیش بینی (Expected Prevalence) و اثر طرح (Design Effect) و احتمال خطای α از کاربر پرسیده می شود. پس از ورود این اطلاعات با فعال کردن جعبه Calculate، حجم نمونه محاسبه شده در قسمت پایین پنجره مشاهده می شود.

[] Sample size, Single proportion

Size of the population	999999	α risk	() 10%	Calculate Reset Edit Print Quit
Desired precision (%)	0.00	(.) 5%		
Expected prevalence (%)	0.00	() 1%		
Design effect	1.00	() 0.1%		
		() 0.01%		

تصویر ۵-۲

جعبه Reset اعداد نوشته شده را پاک کرده و جهت محاسبه حجم نمونه های دیگر پنجره را آماده می نماید. جعبه Edit توانایی تغییر اعداد داده شده را در اختیار کاربر قرار می دهد. جعبه Print از نتیجه حجم نمونه پرینت می گیرد و بوسیله جعبه Quit می توان از این قسمت نرم افزار خارج شد. در تصویر ۵-۳ مثالی از این بخش ذکر شده است.

[] Sample size, Single proportion

Size of the population	999	α risk	() 10%	Calculate Reset Edit Print Quit
Desired precision (%)	5	(.) 5%		
Expected prevalence (%)	50	() 1%		
Design effect	0.8	() 0.1%		
		() 0.01%		
Desired precision (%)	:	5.0		
Expected prevalence (%)	:	50.0		
Design effect	:	0.8		
Confidence level	:	95%		
Sample size	:	223		

تصویر ۵-۳

مقایسه دو نسبت

پس از انتخاب گزینه Two Proportions پنجره ای مشابه تصویر ۵-۴ باز می شود. در این پنجره اطلاعات شامل نسبت پیش فرض گروهها به یکدیگر (Ratio group1/group2)، درصد صفت مورد نظر در گروه اول (Percentage of group1)، درصد صفت مورد نظر در گروه دوم (Percentage of group2)، احتمال خطای α و توان آزمون سوال می گردد. با استفاده از جعبه Calculate حجم نمونه محاسبه شده نمایان می شود. مثالی از این مورد در تصویر ۵-۴ ذکر شد.

[] Sample size, Two proportions

Ratio group 1/group 2	1.000	Calculate Reset Edit Print Quit
Percentage group 1	0.00	
Percentage group 2	0.00	
α risk		
() 10% () 0.1%		
(.) 5% () 0.01%		
() 1%		
Power	() 60% () 90%	
	() 70% () 95%	
	(.) 80% () 99%	

تصویر ۵-۴

[] Sample size, Cohort study

Ratio of non exposed per exposed	3
Relative risk worth detecting	3
Attack rate among non exposed (%)	20

α risk	Power			Calculate
() 10%	() 0.1%	() 60%	() 90%	Reset
(.) 5%	() 0.01%	() 70%	() 95%	Edit
() 1%	(.) 80%	() 99%		Print

Power	80%	Quit
Confidence level	95%	
Number of Exposed	19	
Number of Non exposed	57	
Total #	76	

تصویر ۷-۵

مطالعه مورد شهادی

پس از انتخاب گزینه Cohort Study پنجره‌ای مشابه تصویر ۵-۸ باز می‌گردد که در آن به ترتیب نسبت تعداد شاهد به مورد، Odds Ratio، درصد مورد پیش‌بینی مواجهه یافتن در گروه کنترل و احتمال خطای α و توان آزمون وارد می‌شود. با فعال کردن جعبه Calculate حجم نمونه محاسبه شده مشاهده می‌شود (تصویر ۵-۹).

[] Sample size, Case-control study

Ratio of controls per cases	1.000
Odds ratio worth detecting	0.000
% of exposure among controls (%)	0.00

α risk	Power			Calculate
() 10%	() 0.1%	() 60%	() 90%	Reset
(.) 5%	() 0.01%	() 70%	() 95%	Edit
() 1%	(.) 80%	() 99%		Print

Quit

تصویر ۸-۵

[] Sample size, Two proportions

Ratio group 1/group 2	3
Percentage group 1	40
Percentage group 2	60

α risk	Power			Calculate
() 10%	() 0.1%	() 60%	() 90%	Reset
(.) 5%	() 0.01%	() 70%	() 95%	Edit
() 1%	(.) 80%	() 99%		Print

Power	80%	Quit
Confidence level	95%	
Sample required in group 1	71	
Sample required in group 2	213	
Total #	284	

تصویر ۵-۵

مطالعه همگروهی

پس از انتخاب گزینه Cohort Study پنجره‌ای مشابه تصویر ۵-۶ باز می‌گردد که در آن به ترتیب نسبت افراد مواجهه یافته به مواجهه یافته (Ratio of Nonexposed Per Exposed)، ریسک نسبی (Relative Risk Worth Detecting) و میزان بروز مورد پیش‌بینی بیماری (به درصد) در گروه مواجهه یافته (Attack Rate Among Nonexposed)، احتمال خطای α و توان آزمون پرسیده می‌گردد. با فعال کردن جعبه Calculate حجم نمونه محاسبه شده نمایان می‌شود (تصویر ۵-۷).

[] Sample size, Cohort study

Ratio of non exposed per exposed	1.000
Relative risk worth detecting	0.000
Attack rate among non exposed (%)	0.00

α risk	Power			Calculate
() 10%	() 0.1%	() 60%	() 90%	Reset
(.) 5%	() 0.01%	() 70%	() 95%	Edit
() 1%	(.) 80%	() 99%		Print

Quit

تصویر ۶-۵

[] Sample size, Case-control study

Ratio of controls per cases 3
 Odds ratio worth detecting 3
 % of exposure among controls (%) 20

Calculate

α risk () 10% () 0.1% () 60% () 90%
 () 5% () 0.01% () 70% () 95%
 () 1% () 80% () 99%

Power

Reset

Edit

Print

Quit

Confidence level : 95%
 Number of Cases : 50
 Number of Controls : 150
 Total # : 200

تصویر ۹-۵

✓ نکته: در نرم افزار EPI قسمت دیگری نیز جهت محاسبه حجم نمونه وجود دارد. این قسمت در زیرمَنوی اصلی Program, قسمت Statcalc Calculator می باشد. پس از انتخاب این گزینه پنجره ای باز می شود که یکی از قسمت های آن گزینه Sample Size & Power است. با انتخاب این گزینه پنجره جدیدی باز می شود که در آن امکان محاسبه حجم نمونه مشابه حالت های ذکر شده در قسمت EPITABLE Calculator است. تنها تفاوت موجود در نتایج این دو قسمت به شرح زیر است:

در قسمت Statcalc Calculator حجم نمونه برای مقادیر مختلف خطای نوع اول، توان آزمون و نسبت های مختلف تعداد دو گروه به یکدیگر، محاسبه می گردد. این در حالیست که در منوی EPITABLE Calculator حجم نمونه تنها براساس خطای نوع اول و توان آزمون وارد شده محاسبه می گردد.

حل تمرین

Solving Problems

✓ سوالات

✓ پاسخ ها

فصل

۶

سؤالات

۱- محقق در نظر دارد به منظور بررسی کفایت میزان ید در غذای مردم، متوسط غلظت ید در ادرار کودکان زیر ۱۰ سال منطقه را تعیین کند. در جامعه نرمال، دامنه اطمینان ۹۵ درصد غلظت ید در ادرار $(\mu \pm 2\sigma)$ بین ۶۰ تا ۱۲۰ میکروگرم در لیتر است؛ اگر در منطقه مورد بررسی در مجموع ۲۰۰۰ کودک زیر ۱۰ سال باشد، نمونه مورد نظر وی چند نفر را شامل شود تا میانگین را با دقت ۲ میکروگرم در لیتر در سطح معنی داری ۹۵ درصد برآورد نماید؟

۲- سرپرست مرکز ثبت داده‌های بیماران سرطانی (Cancer Registry)، داده‌های تمامی فوت‌شدگان سرطان مری در طول ۱۰ سال گذشته را در اختیار دارد. اگر وی مایل باشد طول عمر بیماران از زمان تشخیص تا مرگ را بداند، در هریک از موارد زیر چه نمونه‌ای مورد نیاز است؟ ($\alpha = 0.05$)

الف) تعیین متوسط طول عمر با حداکثر خطای ۱ ماه، در شرایطی که انحراف معیار طول عمر ۲ ماه باشد.

ب) تعیین نسبت افرادی که بیش از یک سال عمر می‌کنند یا حداکثر خطای ۱۰ درصد، در شرایطی که بقا یکساله در مطالعات دیگر ۳۰ درصد باشد.

۳- در نظر است مطالعه‌ای در بیماران سوخته به منظور تعیین سن و نسبت مردان در آنها انجام گیرد. اگر حداکثر خطا برای سن، معادل ۰/۱۵ انحراف معیار و در مورد جنس معادل ۰/۲ نسبت مردان در نظر گرفته شود، حجم نمونه مطلوب مطالعه چقدر است؟ ($\alpha = 0.01$)

۴- یک کلینیک تخصصی بیماریهای پستان ادعا دارد که طول عمر ۵ ساله بیماران مبتلا به سرطان مرحله III و IV پستان در آن مرکز، ۱۰ درصد از نسبت ارائه شده در کتب مرجع (که ۲۵ درصد ذکر شده) بیشتر است. برای آزمون این ادعا، چه تعداد بیمار مراجعه‌کننده به کلینیک فوق بایستی بررسی گردند؟ توان آزمون را ۹۰ درصد و احتمال خطای نوع اول را ۰/۰۵ در نظر بگیرید.

۵- داروی جدیدی برای درمان وریدی فشار خون شدید معرفی شده است. در نظر است اثر آن در کاستن فشار خون با نیتروپروساید مقایسه گردد. اگر اوربانیس فشار خون بیماران ۱۹۶ باشد، چه حجمی از نمونه لازم است تا با احتمال ۸۰ درصد اختلافی معادل ۱۰ میلی‌متر جیوه را بین دو دارو نشان دهد؟ ($\alpha = 0.05$) چنانچه بخواهیم اختلافی معادل ۵ میلی‌متر جیوه را شناسایی نماییم، چه حجم نمونه‌ای لازم است؟ در حالت اخیر اگر توان آزمون را ۹۰ درصد در نظر بگیریم، حجم نمونه چند نفر خواهد بود.

۶- محقق در نظر دارد رابطه بین زخم پتیک سوراخ شده و سابقه مصرف سیگار را در یک مطالعه مورد-شاهدی بررسی نماید. براساس یافته‌های قبلی، نسبت مصرف سیگار در افرادی که دچار عارضه سوراخ‌شدگی زخم پتیک نشده‌اند، ۲۰ درصد است؛ چنانچه حداقل OR ارزشمند برای محقق ۲ باشد، نمونه را چند نفر انتخاب کند تا با احتمال ۹۰ درصد اختلاف را در سطح معنی داری ۹۵ درصد نشان دهد؟ اگر به علت محدودیت موارد، به ازای هر نفر مورد، ۴ نفر شاهد را وارد مطالعه نماید، حجم نمونه به چه شکلی تغییر می‌کند؟ در این حالت چنانچه حداقل OR قابل قبول ۳ باشد حجم نمونه چند نفر خواهد شد؟

۷- در نظر است که در یک مطالعه مورد-شاهدی، به ازای هر بیمار مبتلا به سرطان مغزی یک شاهد که از نظر سن و جنس با وی مشابه است، انتخاب گردد و سابقه آنان از نظر تماس با پرتوهای دارای طول موج کوتاه بررسی شود. اگر حداقل OR قابل

قبول برابر ۳ باشد، نمونه را چند نفر انتخاب کنیم تا با احتمال ۹۰ درصد ارتباط بین بیماری و مواجهه معنی دار گردد؟ (احتمال مواجهه در گروه شاهد ۳۰ درصد و $\alpha = 0.05$ است.)

۸- اگر در حالت اول از مسأله ۷، کل بودجه‌ای که در اختیار محقق قرار دارد، ۲ میلیون ریال و بودجه لازم برای بررسی هر بیمار و شاهد به ترتیب ۸ و ۴ هزار ریال باشد، حجم نمونه را چقدر انتخاب کند تا توان آزمون به حداکثر مقدار برسد؟

۹- محقق در صدد است وجود همبستگی بین مدت بستری در بخش مراقبت‌های ویژه و غلظت سرمی منیزیم بیماران را بررسی نماید؛ اگر مایل باشد وجود یک همبستگی لااقل ۰/۴ (بدون توجه به جهت آن) را در سطح معنی داری ۹۹ درصد بررسی نماید، برای اینکار باید چه تعداد بیمار بستری در ICU را بررسی نماید؟ ($\beta = 0.1$)

۱۰- محقق یک سری معیار جدید برای تشخیص انفارکتوس حاد میوکارد (براساس علائم بالینی و یافته‌های نوار قلبی) ارائه کرده است و قصد دارد آنرا با افزایش آنزیم‌های قلبی (به عنوان معیار تشخیص) مقایسه نماید. براساس مطالعات قبلی، شیوع انفارکتوس حاد میوکارد در مراجعینی که دارای شرایط ورود به چنین مطالعه‌ای هستند، ۱۰/۲ است و محقق مایل است به همین نسبت بیمار در مطالعه وی حاضر باشد. اگر براساس یافته‌های مطالعه پیش‌آزمون، حساسیت و ویژگی روش پیشنهادی وی به ترتیب ۸۰ و ۷۵ درصد باشد و آزمون در صورتی ارزشمند باشد که نسبت درست‌نمایی مثبت (PLR) آن حداقل ۲/۵ باشد، چه حجم نمونه‌ای را پیشنهاد می‌کنید؟

۱۱- محقق قصد دارد کثرت استرس خانمهای باردار را در سه ماهه اول بارداری براساس یک پرسشنامه تعیین و آنها را به دو دسته کم استرس و پر استرس تقسیم کرده، بروز کم‌روزی را در نوزادان آنها بررسی نماید؛ اگر بروز تجمعی کم‌روزی در نوزادان مادران کم استرس ۱۰ درصد و حداقل خطر نسبی (RR) قابل قبول برای وی ۲/۵ باشد، چه تعداد نمونه را در هر گروه باید مورد بررسی قرار دهد؟ (در نظر بگیرید که در طول مطالعه ۲۰ درصد مادران از نظارت وی خارج می‌شوند و زایمان خود را در جای دیگری انجام می‌دهند.) ($\beta = 0.1, \alpha = 0.05$)

پاسخها

پاسخ سؤال ۱

دامنه اطمینان ۹۵ درصد در یک جامعه با توزیع نرمال، تقریباً ۴ برابر انحراف معیار را شامل می‌شود لذا انحراف معیار براساس

$$\text{داده‌های ارائه شده برابر است با } 20 = \frac{140 - 60}{4} \text{ سایر داده‌ها عبارتند از:}$$

$$d = 2$$

$$\alpha = 1 - 0.95 = 0.05$$

حجم نمونه براساس رابطه ۱-۵ به ترتیب زیر محاسبه می‌گردد:

$$n = \frac{Z^2 \times \sigma^2}{d^2} = \frac{3.84 \times 20^2}{2^2} = 384$$

با توجه به حجم جامعه هدف (۲۰۰۰ نفر)، حجم نمونه نهایی از رابطه ۳-۵، ۳۲۳ نفر به دست می‌آید:

$$n' = \frac{n}{1 + \frac{n}{N}} = \frac{384}{1 + \frac{384}{2000}} \cong 323$$

پاسخ سؤال ۲

الف) مطابق داده‌های ارائه شده، انحراف معیار برابر ۵، $d=1$ و $\alpha=0/05$ می‌باشد. با استفاده از رابطه ۱-۵، حجم نمونه برابر است با:

$$n = \frac{1/96^2 \times 5^2}{1^2} \cong 97$$

ب) داده‌های ارائه شده عبارتند از: $p=0/3$ ، $d=0/1$ و $\alpha=0/05$ ، حجم نمونه براساس رابطه ۲-۵ برابر است با:

$$n = \frac{1/96^2 \times 0/3(1-0/3)}{0/1^2} \cong 81$$

(لازم به ذکر است که در مورد حجم نمونه‌هایی که از طریق رابطه‌ها به دست می‌آیند، گرد کردن اعداد به سمت عدد طبیعی بالاتر صورت می‌گیرد.)

پاسخ سؤال ۳

در مواردی که مطالعه بیش از یک هدف را دنبال می‌کند، بایستی برای هر هدف حجم نمونه تعیین گردد و بیشترین آنها به عنوان حجم نمونه مطلوب در نظر گرفته شود. حجم نمونه لازم برای برآورد میانگین سن عبارتست از:

$$n = \frac{Z^2 \times \sigma^2}{d^2} = \frac{(2/57)^2 \times \sigma^2}{(0/15\sigma)^2} \cong 294$$

در مورد جنس، چون در مورد نسبت مورد انتظار در مردان، داده‌ای ارائه نشده است $P_{max}=0/5$ در نظر گرفته می‌شود و براین اساس d معادل ۰/۱ به دست می‌آید.

$$n = \frac{Z^2 \times p(1-p)}{d^2} = \frac{(2/57)^2 \times 0/5 \times 0/5}{(0/1)^2} \cong 166$$

با توجه به اینکه حجم نمونه محاسبه شده برای برآورد سن بیشتر است، حجم نمونه نهایی ۲۹۴ نفر خواهد بود.

پاسخ سؤال ۴

داده‌های ارائه شده در صورت مسأله عبارتند از:

$$\alpha = 0/05$$

$$\beta = 1 - \text{power} = 0/1$$

$$p_0 = 0/25$$

$$p_1 = 0/25 + 0/1 = 0/35$$

$$d = 0/1$$

از آنجایی که هدف محققین اثبات بالاتر بودن نسبت طول عمر ۵ ساله در مراجعین کلینیک نسبت به سایر بیماران است، فرضیه متقابل (آلترناتیو) به صورت یکطرفه در نظر گرفته می‌شود و در رابطه ۵-۵، بجای $Z_{1-\alpha/2}$ از $Z_{1-\alpha}$ استفاده می‌شود.

$$n = \frac{[Z_{1-\alpha} \sqrt{p_0(1-p_0)} + Z_{1-\beta} \sqrt{p_1(1-p_1)}]^2}{d^2}$$

$$= \frac{[1/64 \sqrt{0/25 \times 0/75} + 1/28 \sqrt{0/35 \times 0/65}]^2}{(0/1)^2} \cong 175$$

پاسخ سؤال ۵

در حالت اول، حجم نمونه برابر است با:

$$\text{Variance} = 196$$

$$d = 10 \text{ mmHg}$$

$$\alpha = 0/05$$

$$\beta = 1 - 0/8 = 0/2$$

$$n = \frac{2(Z_{1-\alpha/2} + Z_{1-\beta})^2 \sigma^2}{d^2} = \frac{2(1/96 + 0/84)^2 \times 196}{10^2} \cong 31$$

در صورتی که $d=5 \text{ mmHg}$ در نظر گرفته شود، حجم نمونه در هر گروه برابر خواهد بود با:

$$n = \frac{2(1/96 + 0/84)^2 \times 196}{5^2} \cong 123$$

در حالتی که توان آزمون ۹۰ درصد باشد ($\beta=0/1$)، حجم نمونه به صورت زیر محاسبه می‌گردد:

$$n = \frac{2(1/96 + 1/28)^2 \times 196}{5^2} \cong 165$$

پاسخ سؤال ۶

در این حالت، $p_0=0/25$ و $OR=2$ ارائه شده است؛ لذا p_1 برابر است با:

$$P_1 = \frac{P_0 \times OR}{1 + P_0(OR - 1)} = \frac{0/25 \times 2}{1 + 0/25(2 - 1)} = 0/4$$

سایر داده‌های مسئله به شرح زیر است:

$$\alpha = 0/05$$

$$Power = 0/9 \Rightarrow \beta = 1 - 0/9 = 0/1$$

$$\bar{P} = \frac{P_0 + P_1}{2} = \frac{0/25 + 0/4}{2} = 0/325$$

براساس رابطه ۵-۷، حجم نمونه در هر گروه در حالت اول برابر است با:

$$n = \frac{\left[Z_{1-\alpha/2} \sqrt{2\bar{P}(1-\bar{P})} + Z_{1-\beta} \sqrt{P_0(1-P_0) + P_1(1-P_1)} \right]^2}{(P_1 - P_0)^2}$$

$$= \frac{[1/96 \times 0/66 + 1/28 \times 0/65]^2}{(0/4 - 0/25)^2} \cong 201$$

در حالت دوم که به ازای هر مورد، ۴ نفر شاهد انتخاب شده حجم نمونه براساس رابطه ۱۰-۵ به طریق زیر محاسبه می‌گردد:

$$\bar{P} = \frac{0/4 + 4 \times 0/25}{1 + 4} = 0/28$$

$$n = \frac{\left[Z_{0/975} \sqrt{\left(1 + \frac{1}{4}\right) \times 0/28 \times 0/72} + Z_{0/9} \sqrt{0/4 \times 0/6 + \frac{0/25 \times 0/75}{4}} \right]^2}{(0/4 - 0/25)^2} \cong 124$$

تعداد نمونه در گروه شاهد، ۴ برابر گروه مورد (۴۹۶ = ۱۲۴ × ۴ نفر) خواهد بود. در حالت سوم، حجم نمونه براساس روابط زیر

محاسبه می‌گردد:

$$P_1 = \frac{0/25 \times 3}{1 + (0/25 \times 2)} = 0/5$$

$$\bar{P} = \frac{0/5 + 4 \times 0/25}{1 + 4} = 0/3$$

$$n = \frac{\left[1/96 \sqrt{\frac{5}{4} \times 0/3 \times 0/7} + 1/28 \sqrt{0/5 \times 0/5 + \frac{0/25 \times 0/75}{4}} \right]^2}{(0/5 - 0/25)^2} \cong 47$$

۴۷ نفر در گروه مورد که حجم نمونه در گروه شاهد، ۴ برابر گروه مورد (۱۸۸ نفر) می‌باشد.

پاسخ سؤال ۷

در این مطالعه که هر مورد با یک نفر شاهد جور (match) شده، مقایسه به صورت زوجی (Paired) انجام می‌گیرد. در این حالت تعداد جفت‌هایی که از نظر عامل مواجهه با یکدیگر متفاوتند (m) به ترتیب زیر محاسبه می‌گردد:

$$\alpha = 0/05$$

$$\beta = 1 - Power = 0/1$$

$$OR = 3$$

$$P = \frac{OR}{1 + OR} = \frac{3}{1 + 3} = 0/75$$

$$m = \frac{\left[\frac{Z_\alpha}{2} + Z_{1-\beta} \sqrt{P(1-P)} \right]^2}{\left(P - \frac{1}{2}\right)^2} = \frac{\left[\frac{1/96}{2} + 1/28 \sqrt{0/75 \times 0/25} \right]^2}{(0/75 - 0/5)^2} = 38$$

حجم نمونه نهایی از طریق روابط زیر محاسبه می‌گردد:

$$P_0 = 0/3$$

$$P_1 = \frac{P_0 \times OR}{1 + P_0(OR - 1)} = \frac{0/3 \times 3}{1 + (0/3 \times 2)} = 0/56$$

$$M = \frac{m}{[P_0(1-P_1) + P_1(1-P_0)]} = \frac{38}{[0/3(1-0/56) + 0/56(1-0/3)]} \cong 72$$

به عبارت دیگر، بایستی ۷۲ نفر مورد و ۷۲ نفر شاهد که دویّه دو نظر سن و جنس جور شده‌اند، انتخاب گردد.

پاسخ سؤال ۸

با توجه به داده‌های ارائه شده، حجم نمونه در گروه مورد به ترتیب زیر محاسبه می‌گردد:

$$T = 2 \times 10^6$$

$$C_1 = 8 \times 10^3$$

$$C_2 = 4 \times 10^3$$

$$n_1 = \frac{T(C_1 - \sqrt{C_1 C_0})}{C_1(C_1 - C_0)} = \frac{2 \times 10^6 (8 \times 10^3 - \sqrt{8 \times 4 \times 10^6})}{8 \times 10^3 (8 \times 10^3 - 4 \times 10^3)} = 146$$

در این حالت حجم نمونه در گروه شاهد برابر است با:

$$n_0 = n_i \sqrt{\frac{C_1}{C_0}} = 146 \times \sqrt{\frac{8 \times 10^3}{4 \times 10^3}} \cong 207$$

پاسخ سؤال ۹

داده‌های ارائه شده در این مورد عبارتند از:

$$r = 0/9, \alpha = 0/09, \beta = 0/1$$

حجم نمونه از طریق روابط زیر محاسبه می‌گردد:

$$C = 0/5 \times \ln \frac{1+r}{1-r} = 0/5 \times \ln 2/33 \cong 0/42$$

$$n = \left[\left(Z_{1-\alpha/2} + Z_{1-\beta} \right) \div C \right]^2 + 3$$

$$= \left[(2/56 + 1/28) \div 0/42 \right]^2 + 3 = 86$$

پاسخ سؤال ۱۰

داده‌های ارائه شده در صورت مسئله عبارتند از:

$$Sen = 80\% \Rightarrow P_1 = 0/8$$

$$Spc = 75\% \Rightarrow P_2 = 1 - Spc = 0/25$$

$$PLR = 2/5$$

$$\frac{n_1}{n_2} = \frac{2}{8} \Rightarrow n_1 = 0/25 n_2$$

$$LR = \exp \left(\ln \frac{P_1}{P_2} \pm Z_{1-\alpha/2} \sqrt{\frac{1-P_1}{P_1 n_1} + \frac{1-P_2}{P_2 n_2}} \right)$$

$$2/5 = \exp \left(\ln \frac{0/8}{0/25} - 1/96 \sqrt{\frac{1-0/8}{0/8 \times (0/25 n_2)} + \frac{1-0/25}{0/25 n_2}} \right)$$

$$2/5 = \exp \left(1/16 - \frac{3/92}{\sqrt{n_2}} \right)$$

از حل معادله فوق، $n_2 = 266$ نفر بدست می‌آید. حجم نمونه کلی برابر است با:

$$n = n_2 + n_1 = n_2 + 0/25 n_2 \cong 333$$

پاسخ سؤال ۱۱

در این مطالعه که به صورت همگروهی آینده‌نگر طراحی شده، نسبت عدم پیگیری مطالعه توسط شرکت‌کنندگان (Loss to follow up)، ۲۰ درصد پیش‌بینی شده است. سایر داده‌های مسئله به ترتیب زیر است:

$$\alpha = 0/05, \beta = 0/1, P_0 = 0/1, RR = 2/5$$

میزان بروز جمعیتی کم وزنی نوزاد در گروه مادران پر استرس عبارت است از:

$$P_1 = \frac{P_0 \times RR}{1 + P_0(RR - 1)} = \frac{0/1 \times 2/5}{1 + 0/1(2/5 - 1)} \cong 0/22$$

$$\bar{P} = \frac{P_0 + P_1}{2} = 0/16$$

براساس رابطه ۵-۷:

$$n = \frac{\left[Z_{1-\alpha/2} \sqrt{2\bar{P}(1-\bar{P})} + Z_{1-\beta} \sqrt{P_0(1-P_0) + P_1(1-P_1)} \right]^2}{(P_1 - P_0)^2}$$

$$= \frac{(1/96 \times \sqrt{2 \times 0/16 \times 0/84} + 1/28 \sqrt{0/1 \times 0/9 + 0/22 \times 0/78})^2}{(0/22 - 0/1)^2}$$

$$= 194$$

حجم نمونه نهایی در هر گروه براساس رابطه ۵-۸ برابر است با:

$$n' = \frac{n}{1-f} = \frac{194}{1-0/2} = 243$$

پیوست ها

Appendixes

✓ جدول ارقام تصادفی

✓ جدول توزیع Z

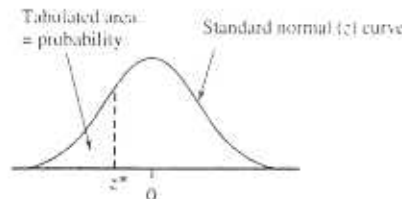
✓ نوموگرام آلتمن

✓ و ...

پیوست ۱. جدول ارقام تصادفی

54814	68020	28998	51687	40088	35458	24708	01815	53776
90106	50899	07394	91071	22411	61643	64435	62552	64316
47185	31782	48894	68790	51852	36918	05737	90653	61123
81354	57296	39329	52263	43194	51624	42429	61367	41207
85467	85622	95778	05347	00445	51334	29445	99176	30091
27924	34167	57060	57535	32278	16949	04960	04116	91467
58319	88164	94130	07743	16917	15681	93572	90753	49117
40762	66702	72425	99117	49298	87265	14195	81392	19794
60504	26749	68743	39139	44495	11944	12970	56523	62411
50074	97517	97450	54251	51777	21073	03909	26519	39578
81147	57508	93479	87826	28965	74474	97468	80149	17834
74689	28937	59819	03105	61325	83145	44684	72958	91824
14802	25982	48024	15461	37570	44685	47386	09504	77831
68501	34194	85355	58411	46559	41694	99678	88268	86674
48734	92671	85252	34228	34228	91289	56331	14683	36493
84102	81699	97352	54509	93196	51204	43351	11818	41179
28432	32875	83834	09862	12720	64569	42218	20726	80866
91458	82524	75523	01276	19591	47473	90251	99103	72947
45475	30389	69732	81962	30243	96199	33546	39672	83760
23557	78437	44957	98728	65674	34701	83398	54102	65845
50395	91850	52004	04844	28848	19728	96571	13317	70859
69991	12755	97916	57639	43445	90463	85556	35469	19749
52980	43608	20592	72527	63583	46443	51329	87219	55198
50776	37035	53765	55196	68659	71429	25225	91942	51132
73714	79868	23880	92254	72984	07792	81306	24277	82366
61547	16575	68520	59869	67299	73565	77316	96682	18031
87737	01058	76012	76247	75616	51335	70364	78942	40564
98669	08334	40520	78389	56495	48122	48122	48599	67579
81575	46690	92814	44456	29227	48122	70532	13852	48436
05975	47110	32733	46929	98261	52193	83215	53192	83109

پیوست ۲. جدول توزیع z یکطرفه



z	.00	.01	.02	.03	.04	.05	.06	.07	.08	.09
-5.8	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001
-5.7	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001
-5.6	.0002	.0002	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001
-5.5	.0002	.0002	.0002	.0002	.0002	.0002	.0002	.0002	.0002	.0002
-5.4	.0003	.0003	.0003	.0003	.0003	.0003	.0003	.0003	.0003	.0002
-5.3	.0005	.0005	.0005	.0004	.0004	.0004	.0004	.0004	.0004	.0003
-5.2	.0007	.0007	.0006	.0006	.0006	.0006	.0006	.0005	.0005	.0005
-5.1	.0010	.0009	.0009	.0009	.0008	.0008	.0008	.0008	.0007	.0007
-5.0	.0013	.0013	.0013	.0012	.0012	.0011	.0011	.0011	.0010	.0010
-4.9	.0019	.0018	.0018	.0017	.0016	.0016	.0015	.0015	.0014	.0014
-4.8	.0026	.0025	.0024	.0023	.0023	.0022	.0021	.0021	.0020	.0019
-4.7	.0035	.0034	.0033	.0032	.0031	.0030	.0029	.0028	.0027	.0026
-4.6	.0047	.0045	.0044	.0043	.0041	.0040	.0039	.0038	.0037	.0036
-4.5	.0062	.0060	.0059	.0057	.0055	.0054	.0052	.0051	.0049	.0048
-4.4	.0082	.0080	.0078	.0075	.0073	.0071	.0069	.0068	.0066	.0064
-4.3	.0107	.0104	.0102	.0099	.0096	.0094	.0091	.0089	.0087	.0084
-4.2	.0139	.0136	.0132	.0129	.0125	.0122	.0119	.0116	.0113	.0110
-4.1	.0179	.0174	.0170	.0166	.0162	.0158	.0154	.0150	.0146	.0143
-4.0	.0228	.0222	.0217	.0212	.0207	.0202	.0197	.0192	.0188	.0185
-3.9	.0287	.0281	.0274	.0268	.0262	.0256	.0250	.0244	.0239	.0233
-3.8	.0359	.0351	.0344	.0336	.0329	.0322	.0314	.0307	.0301	.0294
-3.7	.0446	.0436	.0427	.0418	.0409	.0401	.0392	.0384	.0375	.0367
-3.6	.0548	.0537	.0526	.0516	.0505	.0495	.0485	.0475	.0465	.0455
-3.5	.0668	.0655	.0643	.0630	.0618	.0606	.0594	.0582	.0571	.0559
-3.4	.0808	.0793	.0778	.0764	.0749	.0735	.0721	.0708	.0694	.0681
-3.3	.0968	.0951	.0934	.0918	.0901	.0885	.0869	.0853	.0838	.0823
-3.2	.1151	.1131	.1112	.1093	.1075	.1056	.1038	.1020	.1003	.0985
-3.1	.1357	.1335	.1314	.1292	.1271	.1251	.1230	.1210	.1190	.1170
-3.0	.1587	.1562	.1539	.1515	.1492	.1469	.1446	.1423	.1401	.1379
-2.9	.1841	.1814	.1788	.1762	.1736	.1711	.1685	.1660	.1635	.1611
-2.8	.2119	.2090	.2061	.2033	.2005	.1977	.1949	.1922	.1894	.1867
-2.7	.2420	.2389	.2358	.2327	.2296	.2266	.2236	.2206	.2177	.2148
-2.6	.2743	.2709	.2676	.2643	.2611	.2578	.2546	.2514	.2483	.2451
-2.5	.3085	.3050	.3015	.2981	.2946	.2912	.2877	.2843	.2810	.2776
-2.4	.3446	.3409	.3372	.3336	.3300	.3264	.3228	.3192	.3156	.3121
-2.3	.3821	.3783	.3745	.3707	.3669	.3632	.3594	.3557	.3520	.3483
-2.2	.4207	.4168	.4129	.4090	.4052	.4013	.3974	.3936	.3897	.3859
-2.1	.4602	.4562	.4522	.4483	.4443	.4404	.4364	.4325	.4286	.4247
-2.0	.5000	.4960	.4920	.4880	.4840	.4801	.4761	.4721	.4681	.4641

یہوست ۳. جدول مقادیر

$$C=0.5[\ln(1+r)/(1-r)]$$

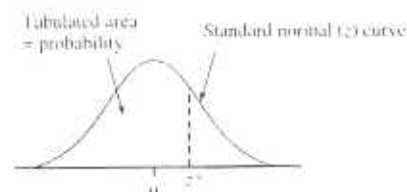
<i>r</i>	.00	.01	.02	.03	.04	.05	.06	.07	.08	.09
.0	.00000	.01000	.02000	.03001	.04002	.05004	.06007	.07012	.08017	.09024
.1	.10034	.11045	.12058	.13074	.14093	.15114	.16139	.17167	.18198	.19234
.2	.20273	.21317	.22366	.23419	.24477	.25541	.26611	.27686	.28768	.29857
.3	.30952	.32055	.33165	.34283	.35409	.36544	.37689	.38842	.40006	.41180
.4	.42365	.43561	.44769	.45990	.47223	.48470	.49731	.51007	.52298	.53606
.5	.54931	.56273	.57634	.59014	.60415	.61838	.63283	.64752	.66246	.67767
.6	.69315	.70892	.72500	.74142	.75817	.77530	.79281	.81074	.82911	.84795
.7	.86730	.88718	.90764	.92873	.95048	.97295	.99621	1.02033	1.04537	1.07143
.8	1.09861	1.12703	1.15682	1.18813	1.22117	1.25615	1.29334	1.33308	1.37577	1.42192
.9	1.47222	1.52752	1.58902	1.65839	1.73805	1.83178	1.94591	2.09229	2.29756	2.64665

پیوست ۴. حجم نمونه برای مطالعات همبستگی

(Correlation)

One-sided $\alpha =$ Two-sided $\alpha =$ $\beta =$ r^*	0.005 0.01			0.025 0.05			0.05 0.10		
	0.05	0.10	0.20	0.05	0.10	0.20	0.05	0.10	0.20
0.05	7.118	5.947	4.663	5.193	4.200	3.134	4.325	3.424	2.469
0.10	1.773	1.481	1.162	1.294	1.047	782	1.078	854	616
0.15	783	655	514	572	463	346	477	378	273
0.20	436	365	287	319	259	194	266	211	153
0.25	276	231	182	202	164	123	169	134	98
0.30	189	158	125	139	113	85	116	92	67
0.35	136	114	90	100	82	62	84	67	49
0.40	102	86	68	75	62	47	63	51	37
0.45	79	66	53	58	48	36	49	39	29
0.50	62	52	42	46	38	29	39	31	23
0.60	40	34	27	30	25	19	26	21	16
0.70	27	23	19	20	17	13	17	14	11
0.80	18	15	13	14	12	9	12	10	8

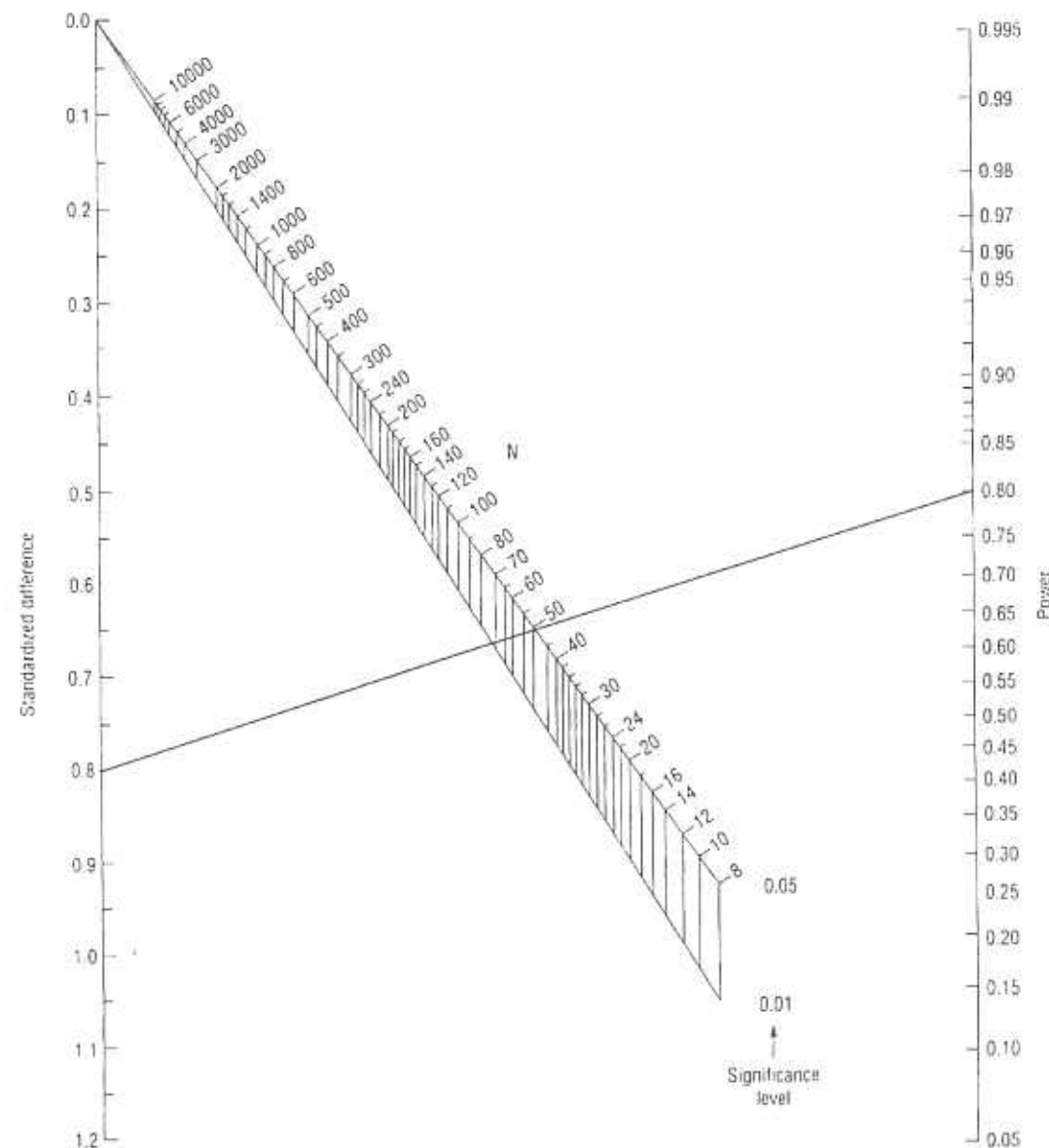
یوست ۲. جدول توزیع Z (ادامه)

[illegible]

منابع:

- 1) Hully S.B;et al, Designing Clinical Research, 2nd ed., USA, Lippincott, 2001.
- 2) Polgaris.;Thomas S., Introduction to Research in the Health Sciences, 4th ed., UK,Churchill Living Stone,2000.
- 3) Levy P.;Lemeshow S.,Sampling of populations methods and application, 3rd ed., USA,Wiley interscience,1999.
- 4) Wingo P.A.;et al, An Epidemiologic Approach to Reproductive Health, CDC & FHI & WHO, 1991.
- 5) Munro B.H., Statistical Methods for Health Care Research,4th ed., USA, Lippincott, 2001.
- 6) Riegelman R.K., Studing a Study & Testing a Test. How to Read Medical Evidence, 4th ed., USA, Lippincott, 2000.
- 7) Petrie A.;Sabin C., Medical Statistics at a Glance, USA, Blackwell Science, 2000.
- 8) Jolly J.; Mitchell M., Research Design explaind, 4th ed., USA, Harcourt, 2001.
- 9) Knapp R.; Miller M., Clinical Epidemiology and Biostatistics, 1st ed., USA ,Williams and Wilkins, 1992.
- 10) Schlesselman J., Case-Control Studies, New-York, Oxford University Press, 1982.
- 11) Andrew G., Manual of EPI Info,Ver. 6, USA, CDC, 1997.
- 12) Armitage P., Berry G., Statistical Methods in Medical Research, 2nd ed., Blackwell Scientific Publications, Oxford, 1987.
- 13) Lwana S.; Tye C.; Ayeni O., Teaching Health Statistics, 2nd ed., Switzerland, WHO, 1999.

پیوست ۵. نوموگرام آلتمن (Altman's Nomogram)



- 14) Mully S.; Cummings S.; Browner W., Designing Clinical Research, 2nd ed., USA, Williams & Wilkins, 2001.
- 15) Polit D.; Reck C.; Mungler B., Essentials of Nursing Research, 5th ed., USA, Lippincott, 2001.
- 16) Peck R.; Olsen C.; Devore J., Introduction to Statistics and Data Analysis, 1st ed., DUXBURY, USA, 2001.
- 17) Sackett D.; Bias in Analytic Research, J Chro Dis., 1979, vol 32, pp 51-63
- 18) Simel D.; Samsa G.; Matcher D., Likelihood Ratios With Confidence: Sample Size Estimation For Diagnostic Test Studies, J Clin Epid., vol 44, No.8, pp 763-770, 1991.
- 19) [http://www.members.aol.com/john_p71/postpower.html\(80/6/12\)](http://www.members.aol.com/john_p71/postpower.html(80/6/12))
- 20) <http://www.trochim.human.cornel.edu/kb>, visited at 17 may 2001.
- 21) <http://www.rvc.ac.uk/epivetnet/manual>, visited at 17 may 2001.

۲۲) سرایی، حسن، "مقدمه‌ای بر نمونه‌گیری در تحقیق"، چاپ اول، تهران، انتشارات سمت، ۱۳۷۲

۲۳) داوسون، ب، "آمار پزشکی پایه- بالینی"، ترجمه سرافراز، علی اکبر؛ غفارزادگان، کامران، چاپ دوم، مشهد، انتشارات دانشگاه علوم پزشکی مشهد، ۱۳۷۷

۲۴) لووازگا، اس کا؛ لمتو اس؛ "تعیین حجم نمونه در مطالعات بهداشتی"، ترجمه محمد کاظم؛ صائمی، سیدحسن، انتشارات معاونت پژوهشی وزارت بهداشت، درمان و آموزش پزشکی، ۱۳۷۱

۲۵) ساده، مهدی، "روشهای تحقیق با تاکید بر جنبه‌های کاربردی آن"، چاپ اول، مؤلف، ۱۳۷۵

۲۶) موزر، س؛ کالتون، ج؛ "روش تحقیق"، ترجمه ایزدی، کاظم، چاپ سوم، سازمان انتشارات کیهان، ۱۳۷۱

۲۷) بشردوست، نصرا،؛ اردلان، علی، "طراحی انواع مطالعات اپیدمیولوژیک"، چاپ اول، انتشارات طب گستر، ۱۳۷۸

۲۸) مازنی، ج؛ بان، آ، "اصول اپیدمیولوژی"، ترجمه ملک‌افضلی، حسین؛ ناصری، کیومرث، چاپ اول، مرکز نشر دانشگاهی، ۱۳۶۳