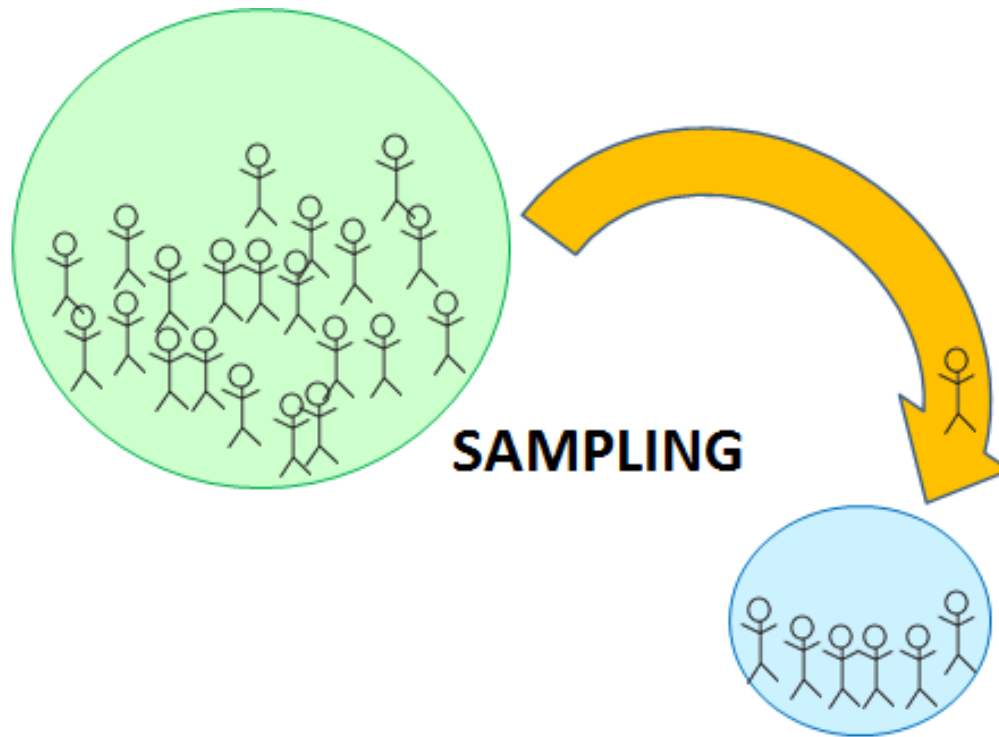


بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



Sampling and Sample Size in Medical Research

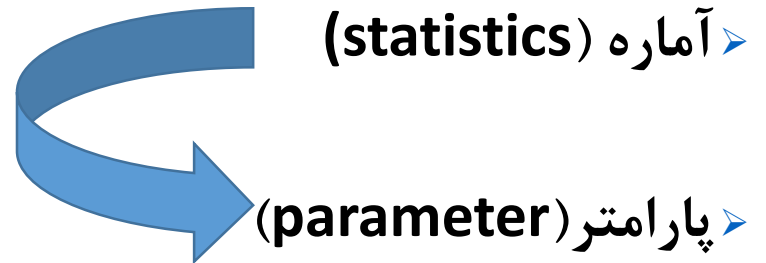


سؤال: چرا نمونه گیری می کنیم؟

- کاهش هزینه ها
- افزایش سرعت
- کاهش زمان
- غیر عملی بودن و محدودیت سرشماری در بعضی شرایط
- نتایج نمونه نسبت به سرشماری از تمام جمعیت اغلب از صحت و اعتبار بیشتری برخوردار است.
- نتایج نمونه نسبت به سرشماری از تمام جمعیت اغلب از ثبات و پایایی بیشتری برخوردار است.



اصطلاحات و مفاهیم نمونه گیری



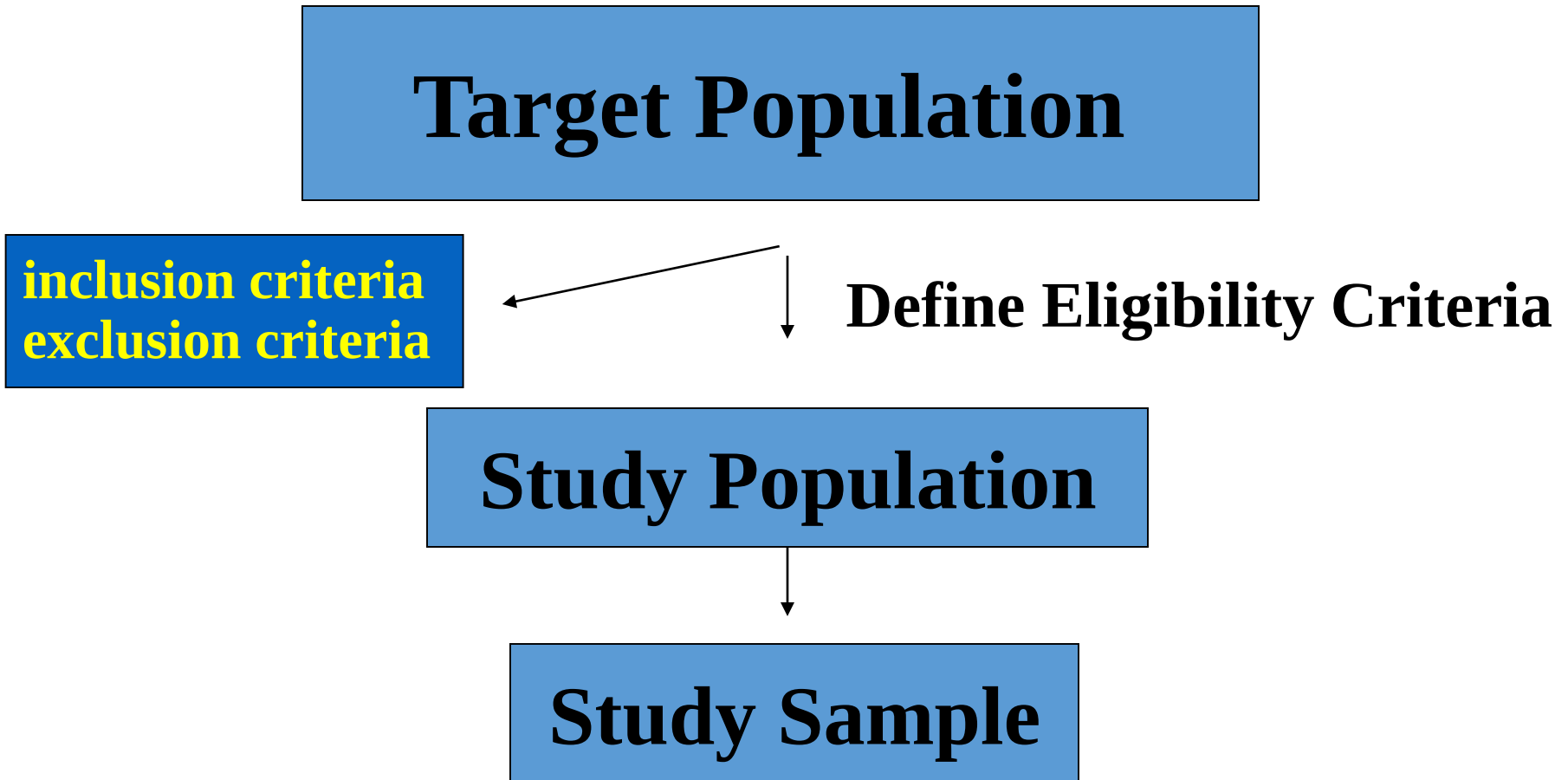
➤ چارچوب نمونه گیری (sampling frame): فهرستی از کل اعضای جامعه که

شامل کل واحدهای نمونه گیری (sampling unit) هستند.

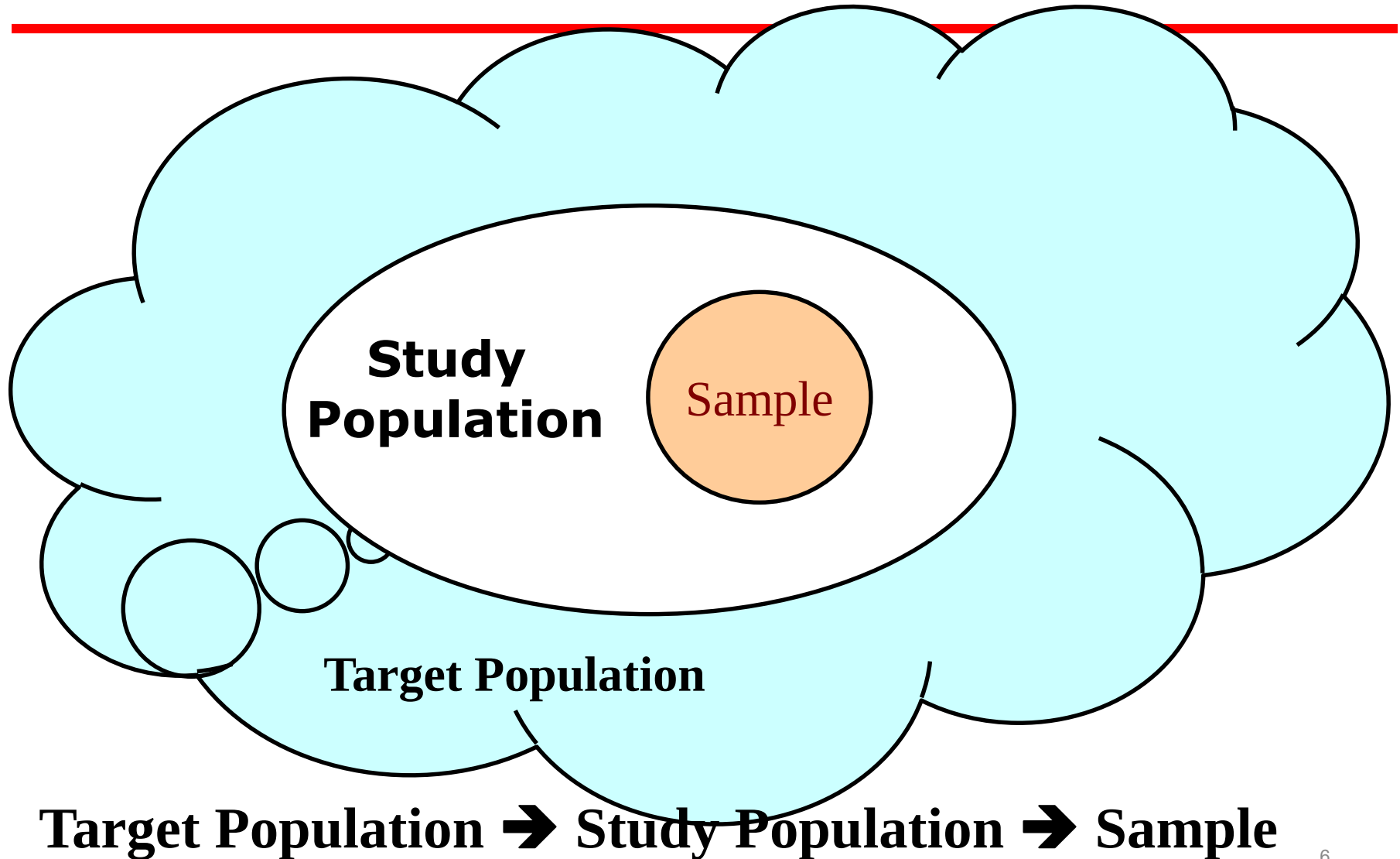
➤ طرح نمونه گیری (sampling design)

➤ معرف بودن نمونه (representativeness)

اصطلاحات و مفاهیم نمونه گیری



جمعیت هدف، جمعیت مطالعه، نمونه



روش های نمونه گیری

احتمالی (Probability sampling)

انتخاب افراد و واحدهای مطالعه به صورت تصادفی است. هر یک از عناصر جامعه، **شانسی** **مشخص و غیر صفر** برای انتخاب شدن در نمونه دارند و احتمال انتخاب همه عناصر جامعه **الزاماً برابر نمی باشد**.

- ✓ در این روش نمونه امکان **سوگیری انتخاب را کم می کند**
- ✓ اجازه می دهد که از **تئوری های آماری** در تحلیل ها استفاده شود
- ✓ **امکان تعمیم نتایج** و اطلاعات بدست آمده از مطالعه بر روی نمونه به جامعه اصلی وجود دارد.

غیر احتمالی (Non-probability sampling)

- انتخاب نمونه براساس **قوانین احتمالات صورت نمی گیرد**، نمونه ها به صورت تصادفی انتخاب نمی شوند و احتمال انتخاب شدن بستگی به نظر و **قضاوت پژوهشگر** دارد بنابراین
- ✓ نمونه ها **معرف واقعی جامعه نیستند**
 - ✓ **احتمال تورش انتخاب** وجود دارد
 - ✓ **نتایج قابل تعمیم به جمعیت هدف نمی باشد**

سوال

نمونه گیری تصادفی در چه نوع مطالعه ای اهمیت بیشتری دارد؟

استفاده از نمونه گیری های غیر احتمالی در **مطالعه های مقطعی**، که انتظار می رود مخرج کسر نماینده خوبی از جامعه باشد، اشتباه بزرگی است.

نمونه گیری غیر تصادفی در چه نوع مطالعاتی اهمیت پیدا می کند؟

در مطالعات **مقدماتی و پایلوت** مربوط به پیمایش های بزرگ کارآمد هستند. در چنین مواردی که هدف تعیین واریانس و تغییرات صفت مورد اندازه گیری (برای تعیین حجم نمونه و روش نمونه گیری) و... می باشد، احتیاجی به روش های نمونه گیری احتمالی نیست

گاهی اوقات نمونه گیری غیراحتمالی بهترین روش نمونه گیری می باشد. مانند زمانی که امکان **تهیه چارچوب**

نمونه گیری وجود نداشته باشد.

۱- نمونه گیری تصادفی ساده (Simple Random Sampling)

۲- نمونه گیری طبقه ای (Stratified sampling)

۳- نمونه گیری خوشه ای (Cluster sampling)

۴- نمونه گیری سیستماتیک (Systematic sampling)

۵- نمونه گیری چند مرحله ای (Multistage sampling)

احتمالی

روشهای
نمونه گیری

۱- نمونه گیری آسان یا در دسترس (Convenience sampling)

۲- نمونه گیری سهمیه ای (Quota sampling)

۳- روش هدفمند (Purposive sampling)

۴- نمونه گیری گلوله برفی (Snowball sampling)

غیراحتمالی

نمونه گیری تصادفی ساده (Simple Random Sampling)

- روش نمونه گیری تصادفی ساده، ساده ترین روش نمونه گیری احتمالاتی است.
- در این نوع نمونه گیری به صورت تصادفی ساده انجام خواهد شد و احتمال انتخاب افراد برابر با فراوانی نسبی آنها می باشد. به عبارت دیگر اگر حجم افراد جامعه N و حجم نمونه را n فرض کنیم، احتمال انتخاب هر فرد جامعه در نمونه مساوی n/N است.

■ روش اجرا:

- تهیه فهرست کل اعضای جامعه (چارچوب نمونه گیری)
- انتخاب تصادفی هر نمونه

انتخاب نمونه تصادفی ساده را به سه شیوه می توان انجام داد:

- ❖ شیوه اول به صورت قرعه کشی
- ❖ شیوه دوم با استفاده از جداول اعداد تصادفی
- ❖ و شیوه سوم با نرم افزارهای رایانه ای مثل Excel



مثالی از نمونه گیری تصادفی ساده

- می خواهیم شیوع پوسیدگی دندان (برمبنای DMF) را در بین ۱۲۰۰ نفر دانش آموزان یک مدرسه تعیین کنیم.
- فهرست دانش آموزان مدرسه را تهیه کنیم
- دانش آموزان را از ۱ تا ۱۲۰۰ شماره گذاری کنیم
- اگر حجم نمونه = ۶۰ نفر باشد، بایستی ۶۰ عدد تصادفی را بین ۱ تا ۱۲۰۰ انتخاب کنیم.

چگونه؟؟

جدول اعداد تصادفی : جدولی از اعداد ۵ رقمی که با حرکت از سطر و ستون های آن می توان به اعداد نمونه دست یافت .

57172	42088	70098	11333	26902	29959	43909	49607
33883	87680	28923	15659	09839	45817	89405	70743
77950	67344	10609	87119	15859	74577	42791	75889
11607	11596	01796	24498	17009	67119	00614	49529
56149	55678	38169	47228	49931	94303	67448	31286
80719	65101	77729	83949	83358	75230	56624	27549
93809	19505	82000	79068	45552	86776	48980	56684
40950	86216	48161	17646	24164	35513	94057	51834
12182	59744	65695	83710	41125	14291	74773	66391
13382	48076	73151	48724	35670	38453	63154	58116
38629	94576	48859	75654	17152	66516	78796	73099
60728	32063	12431	23898	23683	10853	04038	75246
01881	99056	46747	08846	01331	88163	74462	14551
23094	29831	95387	23917	07421	97869	88092	72201
15243	<u>21100</u>	48125	05243	16181	39641	36970	99522
53501	58431	68149	25405	23463	49168	02048	31522
07698	24181	01161	01527	17046	31460	91507	16050
22921	25930	79579	43488	13211	71120	91715	49881
68127	<u>00501</u>	37484	99278	28751	80855	02035	10910
55309	<u>10713</u>	36439	65660	72554	77021	46279	22705
92034	<u>90892</u>	69853	06175	61221	76825	18239	47687
50612	84077	41387	54107	09190	74305	68196	75634
81415	98504	32168	17822	49946	37545	47201	85224
38461	44528	30953	08633	08049	68698	08759	45611
07556	24587	88753	71626	64864	54986	38964	83534
60557	50031	75829	05622	30237	77795	41870	26300

Microsoft Excel

RANDBETWEEN

Returns a random number between the numbers you specify. A new random number is returned every time the worksheet is calculated.

Syntax

RANDBETWEEN (bottom,top)

Bottom is the smallest integer RANDBETWEEN will return.

Top is the largest integer RANDBETWEEN will return.

Example

=RANDBETWEEN(1,1200)

Copy and Paste

نمونه گیری تصادفی ساده (Simple Random Sampling)

- مزایا نمونه گیری تصادفی ساده:

- سادگی روش
- خطای نمونه گیری به سادگی قابل اندازه گیری است.
- برای نمونه گیری از جوامع کوچک روشی مناسب و عملی است.
- سوگیری انتخاب در این روش کمتر وجود دارد.

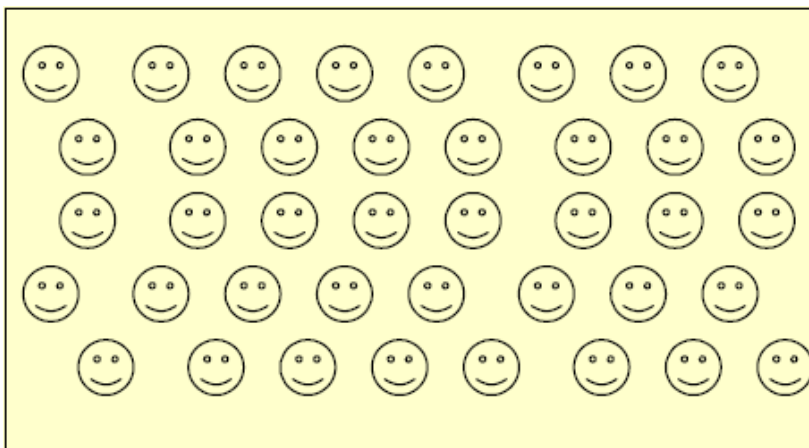
- معایب نمونه گیری تصادفی ساده :

- نیاز به لیست کامل و چارچوب نمونه گیری دارد.
- نمونه ممکن است بسیار پراکنده باشند و دسترسی به همه نمونه ها امکان پذیر نباشد.
- ممکن است نمونه بهترین معرف جامعه نباشد

نمونه گیری تصادفی منظم یا سیستماتیک

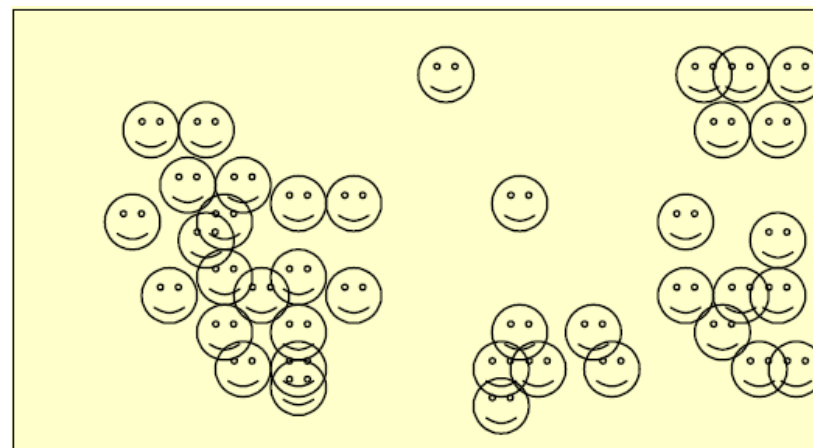
انتظار ما از نمونه گیری

تصادفی ساده



آنچه که در نمونه گیری تصادفی

ممکن است اتفاق بیافتد



نمونه گیری منظم (سیستماتیک) Systematic Sampling

منطبق بر یک قاعده و قانون مشخص

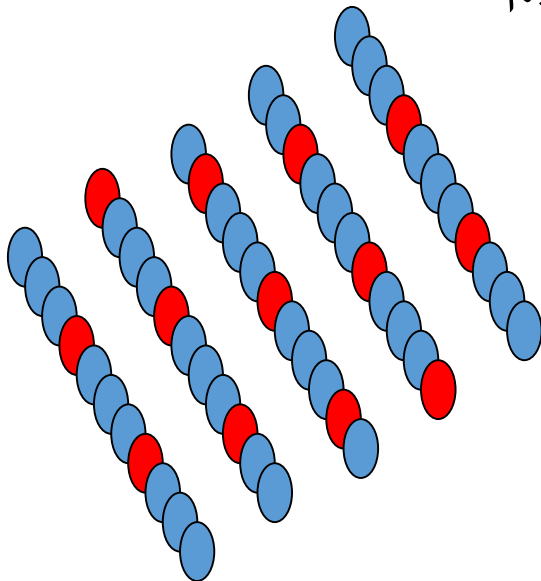
این روش همانند نمونه گیری تصادفی ساده است . با این تفاوت که در آن عدد کسر

$$k = \frac{N}{n}$$

نمونه گیری محاسبه و نمونه به فاصله معین انتخاب می شود .

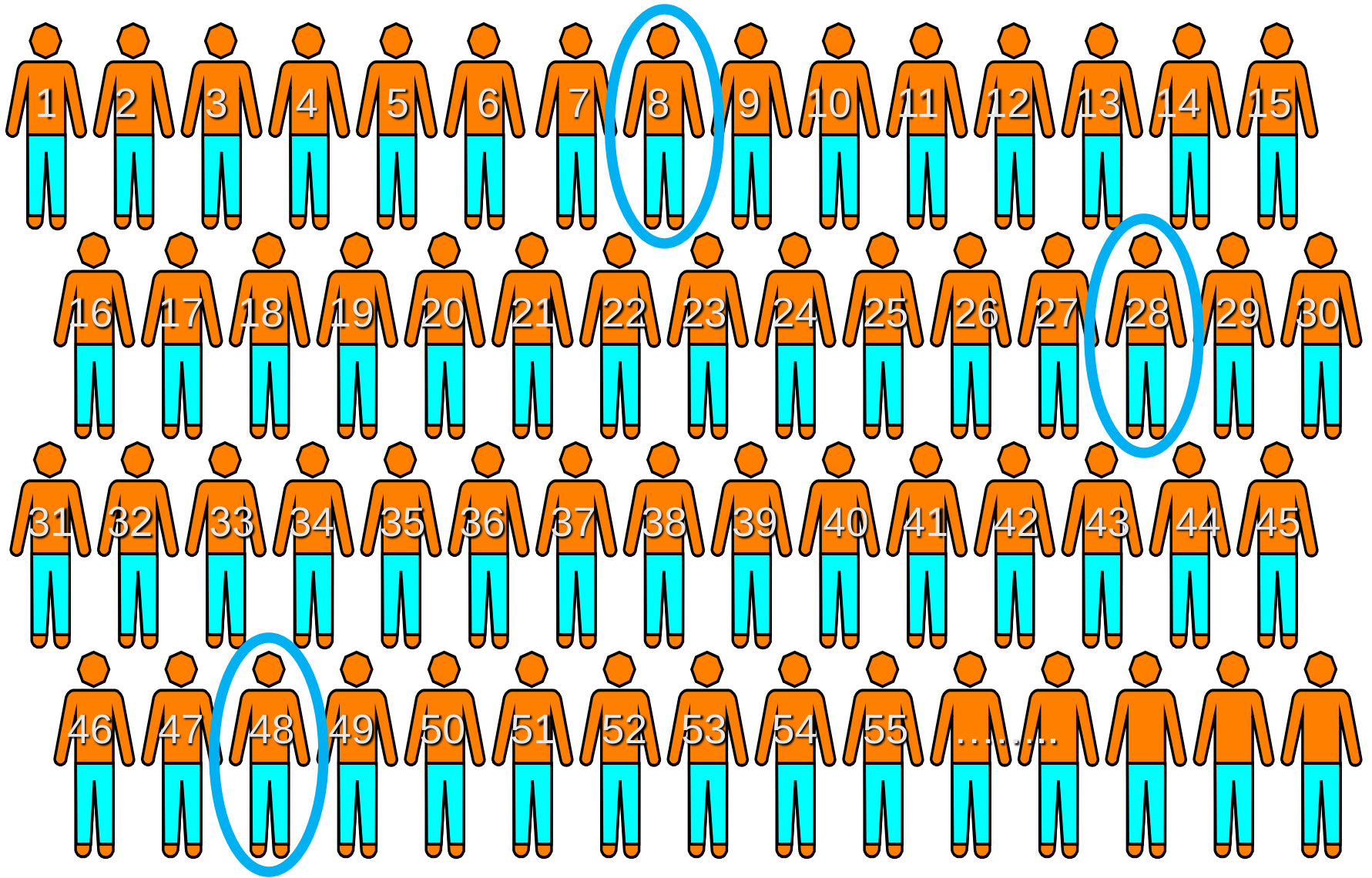
* تعداد اعضای جامعه (N) و نمونه (n) را مشخص می کنیم.

* عددی به صورت رانوم از ۱ تا k را انتخاب می کنیم .



نمونه گیری تصادفی منظم یا سیستماتیک

- مثال: می خواهیم شیوع پوسیدگی دندان (برمبنای DMF) را در بین ۱۲۰۰ نفر دانش آموزان یک مدرسه تعیین کنیم.
 - فهرست دانش آموزان مدرسه را تهیه کنیم
 - دانش آموزان را از ۱ تا ۱۲۰۰ شماره گذاری کنیم
 - اگر حجم نمونه ۶۰ نفر باشد، بایستی با توجه به $۱۲۰۰ / ۶۰ = ۲۰$ ، از هر ۲۰ نفر یکی انتخاب کنیم.
 - یک عدد تصادفی بین ۱ تا ۲۰ انتخاب می کنیم (مثلاً ۸)
 - سپس ۲۰ تا به عدد فوق (۸) اضافه می کنیم به این ترتیب نمونه اول فرد شماره ۸، بعدی ۲۸، بعدی ۴۸ و خواهند بود.



معایب نمونه گیری سیستماتیک

- علاوه بر مشکلات نمونه گیری ساده ، دارای مشکل تناوب در چهارچوب نمونه گیری است .
- اگر فاصله نمونه گیری بر تغییرات منظم در درون جامعه منطبق شود، تورش انتخاب رخ می دهد. (selection bias)
- مثال : اگر در شکل زیر عدد تصادفی ۲ و نسبت نمونه گیری ۴ باشد . نمونه ها عبارت اند از :

۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰	۱۱	۱۲	۱۳	۱۴
پسر	دختر	پسر	دختر	پسر	دختر	پسر	دختر	پسر	دختر	پسر	دختر	پسر	دختر

- برای رفع آن بهتر است از روش نمونه گیری دیگر استفاده کنیم

مزایا نمونه گیری تصادفی منظم

□ آسان و سادگی اجرا دارد.

□ معمولاً پراکندگی تصادفی نمونه ها از تصادفی ساده بهتر است.

□ دقت نمونه گیری تصادفی سیستماتیک وقتی واحدهای جامعه به

طور تصادفی مرتب شده باشند (همانند قد افراد)، مساوی با نمونه

گیری تصادفی ساده و زمانی که واحدهای جامعه بر اساس صفتی

مرتبط با صفت مورد برآورد مرتب شده باشند، بهتر از نمونه گیری

تصادفی ساده و حتی بهتر از نمونه گیری طبقه ای است.

نمونه گیری طبقه بندی شده (Stratified Sampling)



نمونه گیری طبقه ای هنگامی مناسب است که بتوانیم جامعه آماری را نسبت به صفت مربوطه طوری تقسیم کنیم که واحدها در داخل طبقات از نظر صفت مربوطه شبیه به هم باشند.

مراحل:

ابتدا جامعه را به متغیر یا صفات مورد نظر طبقه بندی میکنیم.

چارچوب نمونه گیری و جدول توزیع فراوانی افراد جامعه را بر حسب طبقات مختلف بدست می آوریم.

نسبت درصد و سهم هر یک از طبقات را در کل جمعیت جامعه محاسبه مینماییم.

متناسب با سهم هر طبقه در جامعه ، سهم آن طبقه را در انتخاب نمونه ها نیز اعمال میکنیم.

$$P_k = N_k / N$$

با استفاده از نمونه گیری تصادفی ساده یا سیستماتیک افراد را از هر طبقه انتخاب مینماییم.

به طور مثال از یک جامعه آماری ۱۰۰۰۰ نفری که ۱۵ درصد آن دانشجو ، ۲۰ درصد کارمند اداری ، ۳۰ درصد کارگر و ۳۵ درصد کشاورز هستند می خواهیم ۴۰۰ نفر نمونه انتخاب کنیم . در مرحله اول تعداد مورد نیاز را در هر یک از این طبقات بر حسب درصدهای فوق معین می کنیم که به شرح زیر است :

تعداد افراد نمونه از بین دانشجویان $400 \times 0.15 = 60$

تعداد افراد نمونه از بین کارمندان $400 \times 0.20 = 80$

تعداد افراد نمونه از بین کارگران $400 \times 0.30 = 120$

تعداد افراد نمونه از بین کشاورزان $400 \times 0.35 = 140$

مجموع افراد انتخاب شده از طبقات ۴۰۰ نفر خواهد بود
در مرحله دوم از هر یک از طبقات جمعیت و با استفاده از روش تصادفی ساده یا سیستماتیک افراد نمونه را انتخاب می کنیم .

نکات نمونه گیری طبقه بندی شده (stratified)

- ❖ این نوع نمونه گیری ، شکل اصلاح شده ای از نمونه گیری تصادفی ساده و سیستماتیک است که هدف از آن ، رسیدن به نمونه های معرف تر و دقیق تر است.
- ❖ بین طبقه ها نباید هم پوشانی وجود داشته باشد.
- ❖ هیچ فردی از جامعه نبایستی بیرون از طبقه بندی قرار بگیرد.

هر چقدر داخل طبقات افراد بیشتر به هم شبیه (هموژن) باشند و بین طبقه ها تفاوت زیاد وجود داشته باشد، نمونه گیری طبقه ای معتبر تر است.

مزایای نمونه گیری تصادفی طبقه بندی شده

❖ افزایش دقت برآورد پارامترهای جامعه در صورت همگنی طبقات به علت کم بودن واریانس درون طبقات.

❖ نمونه انتخاب شده بر اساس متغیر طبقه بندی شده، نماینده واقعی جامعه مورد نظر است.

❖ کاهش خطای نمونه گیری

❖ هزینه اجرای این نوع نمونه برداری (به سبب راحت بودن آن) به مراتب کمتر از نمونه برداری تصادفی ساده است.

❖ شرایط یکسان بودن شانس انتخاب برای کل افراد درون طبقات تحقق پیدا می کند.

:(Sub Group Analysis)

❖ برای هر طبقه میتوان برآوردهای جداگانه از پارامترهای جامعه به دست آورد، بنابراین نتایج پژوهش را نه تنها برای کل جامعه بلکه برای بخشهای مختلف آن نیز میتوان بیان کرد.

معایب نمونه گیری تصادفی طبقه بندی شده

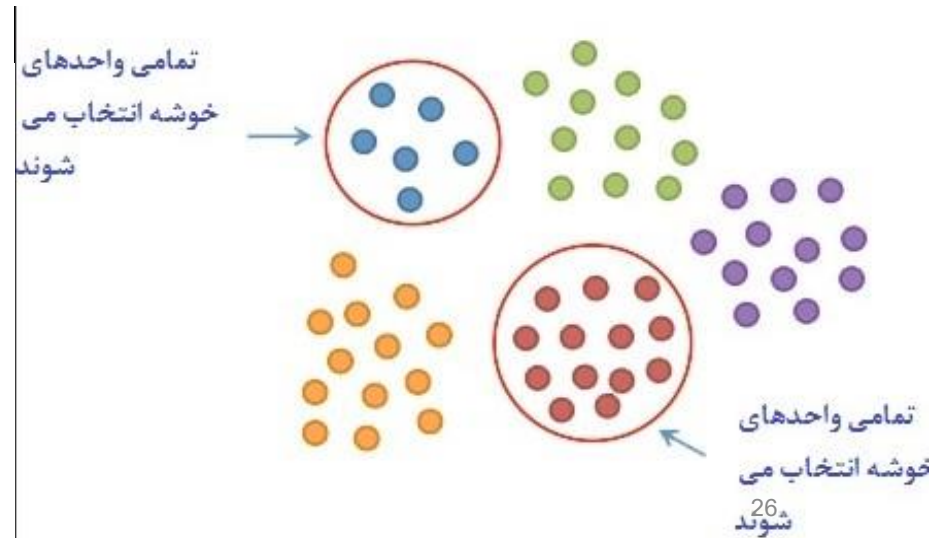
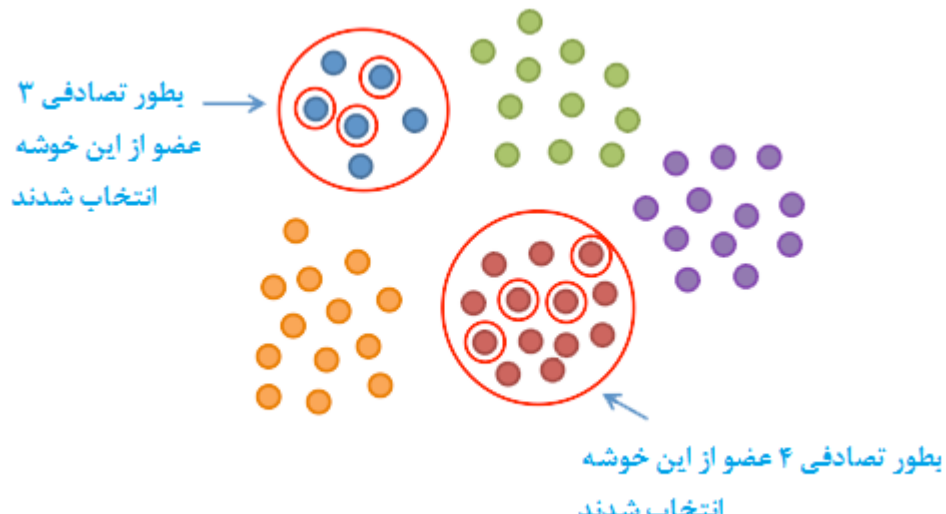
- ❖ عمده ترین مشکل این نمونه گیری ، در دست نبودن اطلاعات در مورد متغیر طبقه بندی است.
- ❖ همیشه طبقه بندی ممکن نیست.
- ❖ ممکن است تعداد نمونه در طبقات کم باشد.
- ❖ اگر صفت مورد نظر در داخل طبقه ها خیلی متفاوت باشد، این روش مناسب نیست.
- ❖ محاسبه خطای نمونه گیری (به دست آوردن SE) مشکل است.

Cluster Sampling

در صورتی که فهرست کامل افراد جامعه مورد مطالعه در دسترس نباشد می توان از نمونه گیری خوشه بندی استفاده نمود. در این روش ابتدا جامعه مورد بررسی را به بلوک ها (خوشه های) مختلف تقسیم میکنیم سپس به صورت تصادفی چند بلوک را انتخاب می کنیم. در این مرحله دو انتخاب اساسی داریم:

❖ در انتخاب اول کلیه اعضای خوشه های منتخب را مورد بررسی و ارزیابی قرار می دهیم. به این شیوه نمونه گیری نمونه گیری خوشه ای یک مرحله ای (Single stage cluster sampling) اطلاق می گردد.

❖ در انتخاب دوم از بین کلیه اعضای خوشه یا خوشه های انتخاب شده، تعدادی از اعضا به شیوه تصادفی ساده یا منظم انتخاب می گردند. به عبارتی، در این حالت در داخل خوشه ها نمونه گیری انجام می دهیم که به این شیوه، نمونه گیری خوشه ای دو مرحله ای (two stage cluster sampling) می گویند.



نمونه گیری خوشه ای

Cluster Sampling

❑ واحدهای نمونه گیری خوشه هایی شامل خانواده ها ، مدارس ، بیمارستان ها ، بلوکهای شهری ، دهکده ها و غیره هستند.

❑ هر چقدر داخل خوشه ها شباهت افراد کمتر و تفاوت بیشتر باشد (هتروژنیتی بیشتر) باشند نمونه گیری خوشه ای موفقتر است.

❑ چگونه نگرانی از مشابه نبودن خوشه ها را رفع کنیم؟

برای رفع این مشکل با استفاده از **اثر طرح** معمولاً تعداد نمونه لازم را در نمونه گیری خوشه ای $1/5$ تا 2 برابر می کنیم. به منظور عملی کردن میزان افزایش در اندازه نمونه در شرایط نمونه گیری درون خوشه ای، از ضریبی به نام شاخص اثر طرح یا **Design Effect** استفاده می نمایند.

A design effect (DEFF) is an adjustment made to find a survey sample size, due to a sampling method (**Cluster sampling and Multistage sampling**)

Design effect (DEEF) was calculated with the formula

$$DEEF = 1 + \delta(n-1)$$

“ δ ” or the interclass correlation coefficient (ICC)

“n” average size of clusters

- معمولاً این شاخص دارای اندازه‌ای بزرگ‌تر از یک می‌باشد.
- بزرگی این شاخص رابطه مستقیم با پراکندگی یا واریانس داخل خوشه‌ای دارد هر چه این پراکندگی بیشتر باشد، مقدار عددی اثر طرح بزرگ‌تر است.
- آنجا که در زمان طراحی مطالعه و تعیین روش نمونه‌گیری و حجم نمونه، مقدار دقیق اثر طرح نامشخص می‌باشد، معمولاً برآورد یا پیش‌بینی تقریبی برای شاخص اثر طرح را در محاسبه اندازه نمونه نهایی لحاظ می‌نمایند. در اغلب شرایط مقدار عددی **اثر طرح را حدود ۱/۵-۲** در نظر می‌گیرند و به عبارت دیگر عدد اندازه نمونه به دست آمده از فرمول‌های محاسبه اندازه نمونه را در این مقادیر تقریبی ضرب تا اندازه نمونه نهایی در شرایطی که روش نمونه‌گیری تصادفی خوشه‌ای است، به دست آید.

نمونه گیری خوشه ای

■ مزایا

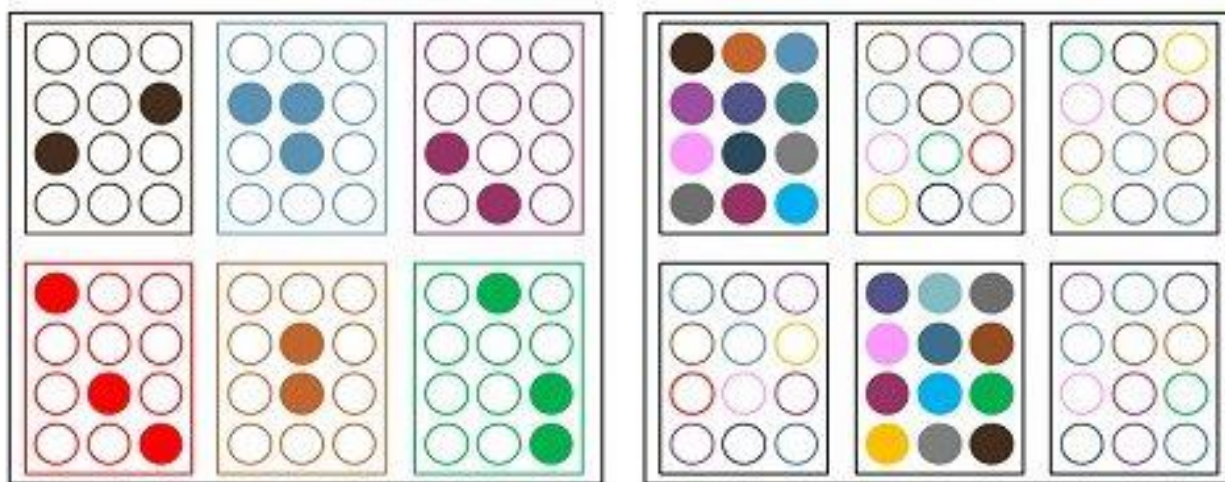
- در این روش نیازی به داشتن فهرست کامل واحدهای نمونه گیری (چارچوب نمونه گیری) نیست.
- هزینه و منابع مورد نیاز برای جمع آوری داده ها کمتر است.
- صرفه جویی در وقت

■ معایب

- اگر هموژنیتی صفت مورد نظر در داخل خوشه ها خیلی زیاد باشد، این روش مناسب نیست.
- نسبت به نمونه گیری تصادفی ساده، برای دستیابی به دقت (precision) مشابه به حجم نمونه بالاتری نیاز دارد.
- محاسبه خطای نمونه گیری (به دست آوردن SE) مشکل است

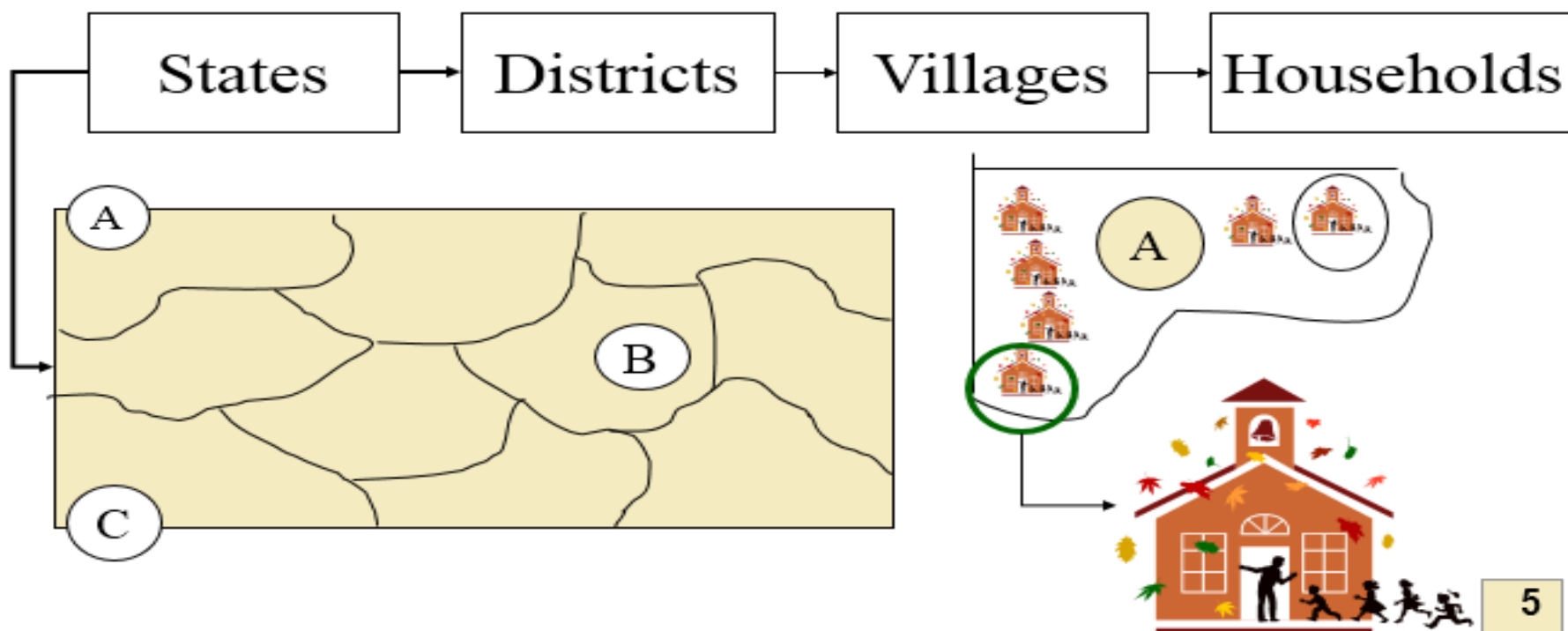
تفاوت بین نمونه گیری خوشه ای تک مرحله ای با نمونه گیری طبقه بندی شده

- الف-در نمونه گیری تصادفی طبقه ای از هر طبقه تعدادی از افراد را بر اساس فراوانی نسبی همان طبقه و به صورت تصادفی انتخاب می کنیم در صورتی که در نمونه گیری خوشه ای، تعدادی از خوشه ها را به تصادف انتخاب و سپس کل افراد درون خوشه را بعنوان نمونه مورد بررسی قرار می دهیم.
- ب- دقت نمونه گیری تصادفی طبقه ای در ارتباط مستقیم با همگنی درون طبقات و ناهمگنی بین طبقات است. اما دقت نمونه گیری تصادفی خوشه ای در ارتباط مستقیم با ناهمگنی درون خوشه ها و همگنی بین خوشه هاست.

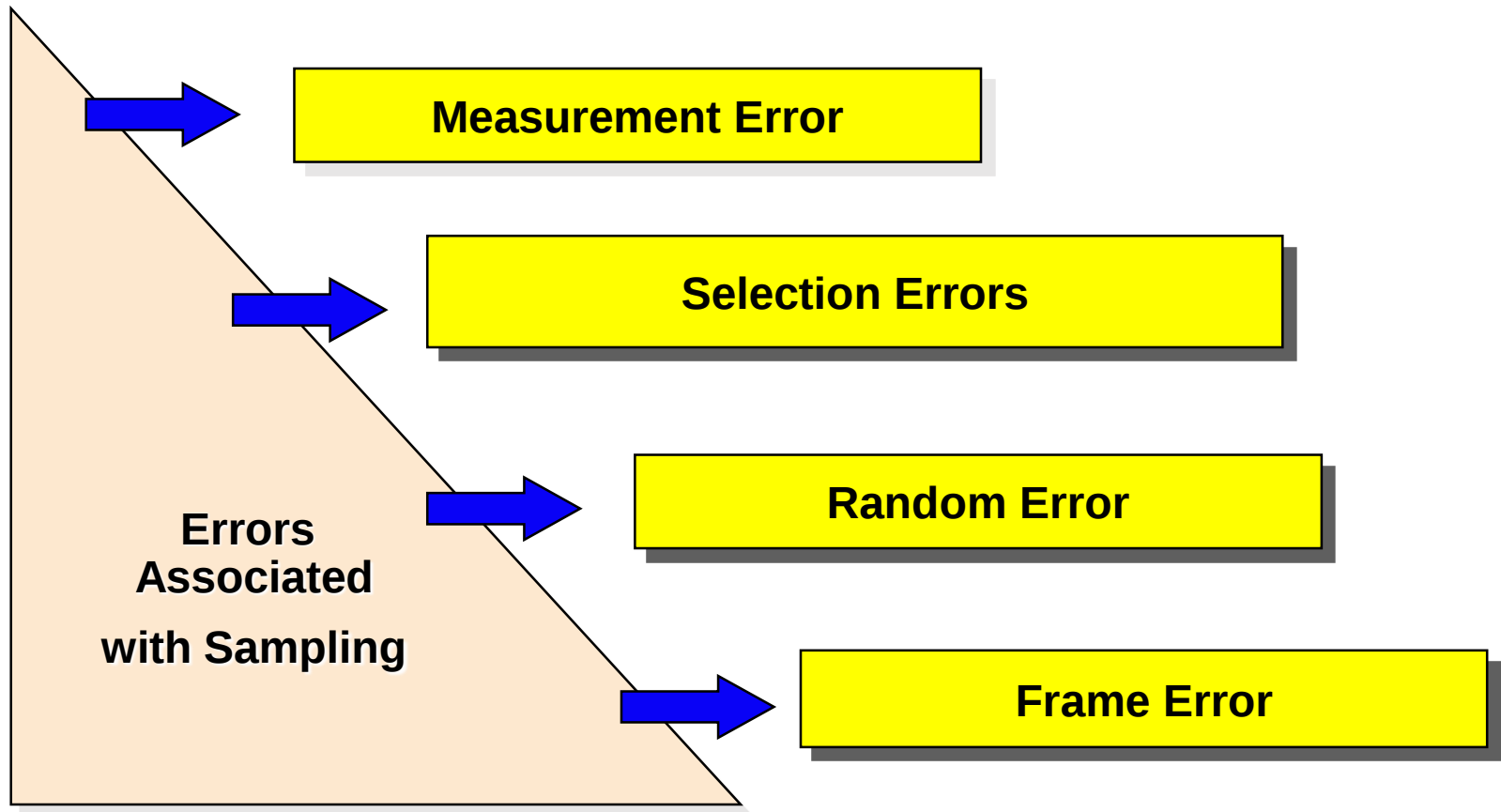


۵- نمونه گیری چندمرحله ای (Multistage sampling)

- در این روش نمونه گیری با استفاده از روشهای نمونه گیری احتمالی واحد یا سطح نمونه گیری مرحله به مرحله کوچکتر و محدودتر می شود و در آخرین مرحله نمونه گیری به واحد اصلی دسترسی پیدا می کنیم .



Types of Sampling Errors



Sample size

اخذ تصمیم درباره حجم نمونه، از لحاظ تامین میزان دقت نتایج نمونه گیری و صرفه جویی در مقدار وقت و هزینه، از اهمیتی خاص برخوردار است. بدیهی است که بزرگ بودن حجم نمونه موجب صرف هزینه و وقت زیاد، و کوچک بودن حجم نمونه موجب عدم دقت کافی برآوردها می شود.

برای تعیین اندازه نمونه، چند راه پیش رو داریم:

✓ تعیین اندازه نمونه بر اساس جدول های آماده مانند جدول موران

در این جدول ها با توجه به اندازه جامعه (N)، اندازه نمونه (n) تعیین شده است.

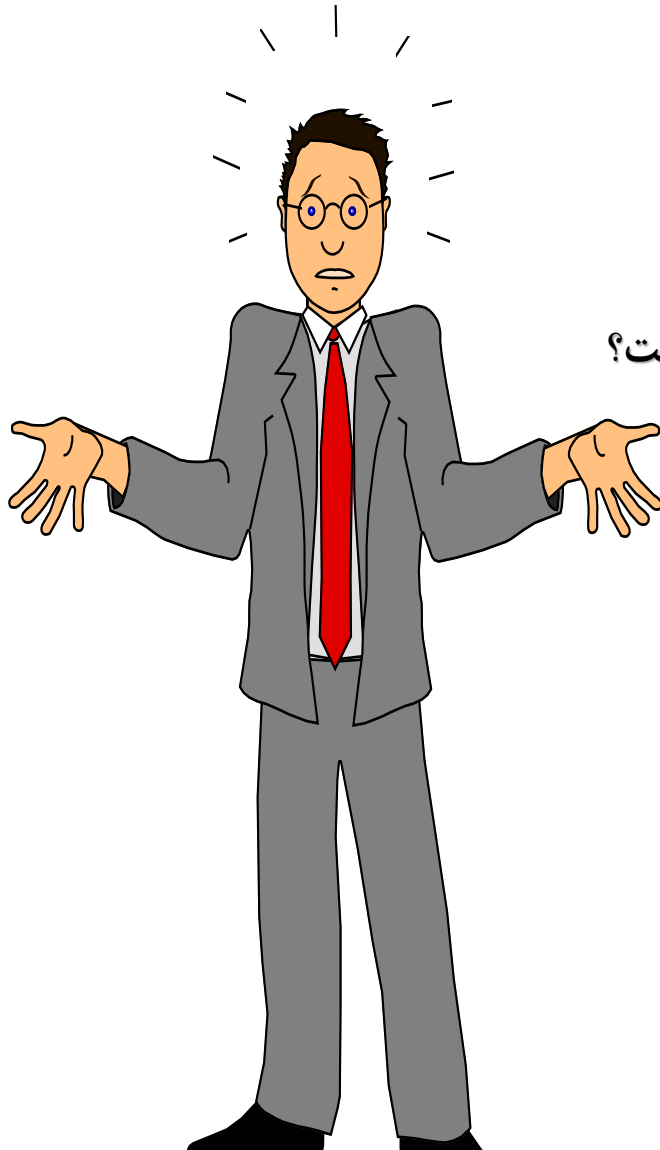
✓ تعیین حجم نمونه بر اساس نظر پژوهش گر

در این روش پژوهش گر با در نظر گرفتن عامل هایی مثل بودجه، زمان، نیروی انسانی ماهر، امکانات، شیوع بیماری و در دسترس بودن موارد درصدی از جامعه را به عنوان نمونه انتخاب می کند. برخی از پژوهشگران حداقل اندازه نمونه را ۱۰ درصد اندازه جامعه ذکر کرده اند.

✓ تعیین اندازه نمونه بر اساس محاسبات آماری

متداول ترین روش محاسبه اندازه نمونه، استفاده از محاسبات آماری است.

برای یافتن روش تعیین حجم نمونه نیاز است که محقق در ابتدا به چند سؤال پاسخ دهد:



• هدف اصلی پژوهشگر و نوع مطالعه (توصیفی یا تحلیلی) چیست؟

• متغیر اصلی مورد پژوهش کمی است یا کیفی؟

• چه سطح اطمینان و در صورت لزوم چه توانی برای نیل به هدف کافی است؟

$(\alpha, \beta, 1-\alpha, 1-\beta)$

• مقدار تغییرات **متغیر وابسته** (انحراف معیار و شیوع) چقدر است؟

• محقق چه مقدار خطا را می‌پذیرد؟ (Precision)

• بودجه طرح چقدر است؟

Sample Size and Population Size

- Where is **N** (size of the population) in the sample size determination formula?

Population Size	e=±3% Sample Size	e=±4% Sample Size
10,000	1,067	600
100,000	1,067	600
1,000,000	1,067	600
100,000,000	1,067	600

In almost all cases, the sample size is independent of the size of the population.

TABLE FOR DETERMINING SAMPLE SIZE FROM A GIVEN POPULATION or [Cochran's Sample Size Formula](#).

	فرض صفر درست	فرض صفر غلط
رد فرض صفر	خطای نوع اول α	توان آزمون $1-\beta$
پذیرش فرض صفر	خطایی رخ نداده $1-\alpha$	خطای نوع دوم β

The quantities α , β , Δ (Effect Size) and N are all interrelated.

$$\text{POWER} = 1 - \beta = P\{\text{accept } H_A | H_A \text{ is true}\}$$

Holding all other values constant, what happens to the power of the study if

- α decreases? Power ↓
- N increases? Power ↑
- variability increases? Power ↓
- Δ (Effect Size) increases? Power ↑

Note: Typical error rates are $\alpha = .05$ and $\beta = .1$ or $.2$ (80 or 90% power). Why is α often smaller than β ?

$$\hat{p} \pm 1.96 * \sqrt{\frac{pq}{n}}$$

point estimate $\pm$ margin of error

$$\text{margin of error} = 1.96 * \sqrt{\frac{pq}{n}}$$

$$(\text{margin of error})^2 = (1.96)^2 * pq/n$$

$$n = (1.96)^2 * pq/d^2$$

محاسبه حجم نمونه با فرمول کوکران

- یکی از روشهای پرکاربرد در تعیین حجم نمونه فرمول کوکران است. فرمول کوکران به صورت زیر محاسبه می شود:

$$n = \frac{\frac{z^2 pq}{d^2}}{1 + \frac{1}{N} \left(\frac{z^2 pq}{d^2} - 1 \right)}$$

- $P=0.5$ $q=.5$ $z=1.96$ $d=0.05$

جدول مورگان برای محاسبه حجم نمونه

مورگان چیزی جز همان کاربرد فرمول کورکران نیست. و به عبارتی دیگر هر یک از اعداد این جدول را در فرمول کوکران بگذارید همین حجم نمونه را مشاهده خواهید کرد.

حجم نمونه	جامعه آماری	حجم نمونه	جامعه آماری	حجم نمونه	جامعه آماری
291	1200	140	220	10	10
297	1300	144	230	14	15
302	1400	148	240	19	20
306	1500	152	25	24	25
310	1600	155	260	28	30
313	1700	159	270	32	35
317	1800	162	280	36	40
320	1900	165	290	40	45
322	2000	169	300	44	50
327	2200	175	320	48	55
331	2400	181	340	52	60
335	2600	186	360	56	65
338	2800	191	380	59	70
341	3000	196	400	63	75
346	3500	201	420	66	80
351	4000	205	440	70	85
354	4500	210	460	73	90
357	5000	214	480	76	95
361	6000	217	500	80	100
364	7000	226	550	86	110
367	8000	234	600	92	120
368	9000	242	650	97	130
370	10000	248	700	103	140
375	15000	254	750	108	150
377	20000	260	800	113	160
379	30000	265	850	118	170
380	40000	269	900	123	180
381	50000	274	950	127	190
382	75000	278	1000	132	200
384	100000	285	1100	136	210

اندازه نمونه در مطالعات توصیفی

- در مطالعات توصیفی، معمولاً هدف برآورد میانگین یک صفت کمی یا نسبت یک صفت کیفی است.

متغیر کیفی (برآورد نسبت)

متغیر کمی (برآورد میانگین)

- حجم نمونه در مطالعات توصیفی

محاسبه حداقل حجم نمونه برای برآورد نسبت (صفت کیفی)

$$n \geq \frac{(z_{1-\alpha/2})^2 * p(1-p)}{d^2}$$

α : احتمال خطای نوع اول؛

اگر $\alpha = 0.05$ باشد $z_{1-\alpha/2}$ برابر 1.96 است

p : تخمین نسبت یا شیوع (prevalence) صفت مورد نظر

d : خطای قابل قبول در برآورد نسبت مورد نظر

محاسبه حداقل حجم نمونه برای برآورد میانگین (صفت کمی)

$$n = \left(\frac{Z_{1-\alpha/2} * \sigma}{d} \right)^2$$

α : احتمال خطای نوع اول؛

اگر $\alpha = 0.05$ باشد $Z_{1-\alpha/2}$ برابر 1.96 است

σ : انحراف معیار (SD) صفت مورد نظر

d : خطای قابل قبول در برآورد میانگین

اندازه نمونه در مطالعات تحلیلی

در مواردی که پژوهش تحلیلی است یعنی محقق بدنبال یافتن اختلاف بین گروه‌ها با آزمون ارتباط بین متغیرها است.

نکته : علاوه بر مقوله اطمینان به نتایج، توان آزمون مورد استفاده در کشف یا نشان دادن حقیقت نیز مطرح می‌شود.

متغیر کیفی (مقایسه نسبتها)

• حجم نمونه در مطالعات تحلیلی

متغیر کمی (مقایسه میانگینها)

محاسبه حداقل حجم نمونه برای مقایسه نسبت ها در دو جامعه مستقل

$$n = \frac{(z_{1-\alpha/2} + z_{1-\beta})^2 * [p_1(1-p_1) + p_2(1-p_2)]}{(p_1 - p_2)^2}$$

α : احتمال خطای نوع اول؛ اگر $\alpha = 0.05$ باشد $Z_{1-\alpha/2}$ برابر 1.96 است

β : احتمال خطای نوع دوم؛ اگر $\beta = 0.2$ باشد $Z_{1-\beta}$ برابر 0.84 است

P1: نسبت یا سهم پیامد مورد نظر در گروه اول (مورد - مواجهه یافته یا مداخله)

P2: نسبت یا سهم پیامد مورد نظر در گروه دوم (شاهد، غیر مواجهه یافته یا مقایسه)

تفاوت دو نسبت در دو گروه در واقع همان دقت برآورد در مطالعه است.

محاسبه حداقل حجم نمونه برای مقایسه میانگین ها در دو جامعه مستقل

$$n = \frac{(Z_{1-\alpha/2} + Z_{1-\beta})^2 * (\sigma_1^2 + \sigma_2^2)}{(\bar{X}_1 - \bar{X}_2)^2}$$

α : احتمال خطای نوع اول؛ اگر $\alpha = 0.05$ باشد $Z_{1-\alpha/2}$ برابر 1.96 است

β : احتمال خطای نوع دوم؛ اگر $\beta = 0.2$ باشد $Z_{1-\beta}$ برابر 0.84 است

σ_1 : انحراف معیار صفت در جامعه اول (مورد - مواجهه یافته یا مداخله)

σ_2 : انحراف معیار صفت در جامعه دوم (شاهد، غیر مواجهه یافته یا مقایسه)

X_1 : میانگین صفت در جامعه اول (مورد - مواجهه یافته یا مداخله)

X_2 : میانگین صفت در جامعه دوم (شاهد، غیر مواجهه یافته یا مقایسه)

DETERMINING of SAMPLE SIZE PARAMETERS

it is appropriate to have a precision of 5% if the prevalence of the disease is going to be between 10% and 90%.

However, when the prevalence is going to be below 10% or more than 90%, the precision of 5% seems to be inappropriate.

Table 1 95% CI of rare diseases ($P \leq 0.05$) and common diseases ($P \geq 0.95$) with a precision (d) set at 0.05

<i>P</i>	<i>n</i>	95% CI	
		Lower	Upper
0.01	16	-0.04	0.06
0.02	31	-0.03	0.07
0.03	45	-0.02	0.08
0.04	60	-0.01	0.09
0.05	73	0	0.10
0.95	73	0.90	1.00
0.96	60	0.91	1.01
0.97	45	0.92	1.02
0.98	31	0.93	1.03
0.99	16	0.94	1.04

For example, if the prevalence is 1% (in a rare disease) the precision of 5% is obviously crude and it may cause problems. The obvious problem is that 95% CIs of the estimated prevalence will end up with irrelevant negative lower-bound values or larger than 1 upper bound values as seen in the Table 1

DETERMINING of SAMPLE SIZE PARAMETERS

- Therefore, we recommend d as a half of P if P is below 0.1 (10%) and if P is above 0.9 (90%), d can be $\{0.5(1-P)\}$.
- For example, if P is 0.04, investigators may use $d=0.02$, and if P is 0.98, we recommend $d=0.01$.
- However, if there is a resource limitation, investigators may use a larger d . In case of a preliminary study, investigators may use a larger d (e.g. >10%). However, justification for the selection of d should be stated clearly (e.g. limitation of resources) in their research proposal.

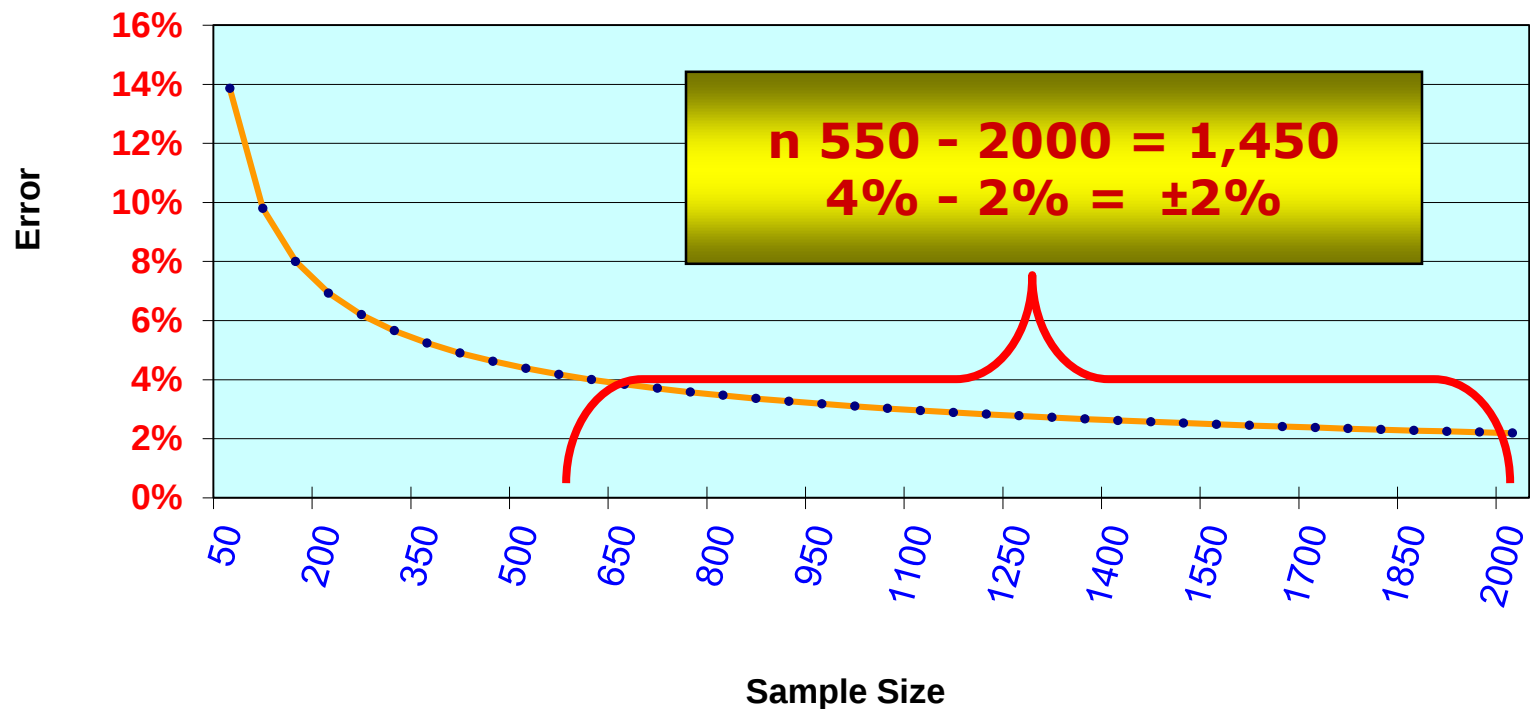
Estimating P

- Preferably, P from the studies with **pilot study, similar study design** and study population from the most recent studies would be most preferable.
- about the value of P , it is best to towards **50%** as it would lead to a **larger** sample size. Setting $P=0.5$ is provide the biggest sample size.
- If we have a range of P , for instance, 20% to 30%, we should use 30% as it will give a larger sample size .If the range is 60% to 80%, we should use 60% as it will give a larger sample size. If the range is 40% to 60%, 50% will give a larger sample size.

Sample Size

- **A larger sample size is needed to detect the smallest meaningful difference.**
- **A larger sample size is needed when there is much variability in the population**
- **A larger sample size is required to increase the power of a study.**

Sample Size and Accuracy



Probability sample error (accuracy) can be calculated with a simple formula, and expressed as a ± % number.

متغیرها

تعریف :

مشخصه یک فرد ، شیء ، یا پدیده که قابل اندازه گیری بوده و بتواند مقادیر مختلفی را بپذیرد.(از فردی به فرد دیگر تغییر می کند)

متغیرها

اسمی (Nominal)

ترتیبی یا رتبه ای (Ordinal)

فاصله ای (Interval)

نسبتی (Ratio)

واژه NOIR را بخاطر بسپارید که مخفف همه انواع متغیرهاست.

نوع متغیرها (مقیاس سنجش متغیرها)

کیفی

- اسمی
- رتبه ای

کمی

- فاصله ای
- نسبتی

انواع متغیر ها

- **متغیر کمی:** متغیری است که با عدد نمایش داده می شود این متغیر به دو دسته متغیر **گسسته و پیوسته** تقسیم خواهد شد که متغیر پیوسته مقادیر کسری را هم می پذیرد ولی گسسته این امکان را ندارد.
- **متغیر کیفی:** این متغیری است که کیفیت صفات با آن معرفی می شود.

🔗 نوع متغیر (سندش متغیر) :

۱- اسمی (NOMINAL) :

متغیری که تنها می توان نامی را بر آن نهاد ولی هیچگونه دلیلی بر برتر بودن یکی بر دیگری نداریم بنابراین تنها می توان آنها را دسته بندی و درصد گذاری کرد.



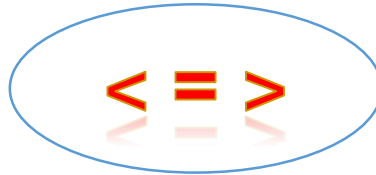
در این متغیرها هیچ یک از عملیات ریاضی را نمی توان انجام داد.

مانند: گروههای خونی، رنگ چشم ، نژاد ، جنسیت، ملیت و...

متغیر رتبه ای

در این مقیاس اندازه گیری علاوه برداشتن نامی برای متغیر مورد بررسی می توان رتبه آنها را نسبت به یکدیگر بیان نمود مانند درجه گواتر (۰، ۱، ۲ و ۳) و شدت درد (کم، متوسط، زیاد) ، سطح تحصیلات (ابتدایی، راهنمایی، دبیرستان، دانشگاه)

قابل محاسبه هستند.



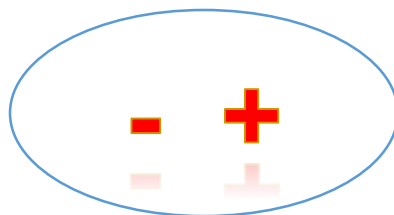
در این متغیرها عملیات ریاضی

نکته : دو نوع متغیر اسمی و رتبه ای را متغیر کیفی می گویند.

متغیر فاصله ای

متغیری است که به صورت مقدار و عدد بیان می شود اما فاقد صفر ذاتی است و تنها بصورت صفر قراردادی مشخص می شوند. در این متغیرها علاوه بر دسته بندی و رتبه بندی متغیرها فاصله بین متغیرها را نیز می توان محاسبه کرد.

قابل محاسبه هستند.



در این متغیرها عملیات ریاضی

اعشار پذیر نیستند و اعداد صحیح را می پذیرند.

در علوم پزشکی درجه حرارت یک مقیاس فاصله ای است. مثالهای دیگر: مبداء سالهای شمسی، قمری، میلادی

متغیر نسبتی

متغیری که دارای واحد مشخص، معین، تعریف شده و نیز دارای صفر ذاتی می باشد. صفر نشان دهنده فقدان واقعی خاصیت موردنظر است.
فشارخون، جرم، میزان هموگلوبین، صدا، میزان مرگ و میر و ...

قابل محاسبه هستند.



در این متغیرها عملیات ریاضی

نکته : دو متغیر فاصله ای و نسبتی را متغیر کمی می گویند .

انواع متغیر ها (از لحاظ نقش)

❑ **متغیر مستقل:** آن متغیری است که محقق قصد دارد تاثیر آن را بر سایر متغیرها مورد سنجش قرار می دهد.

❑ **متغیر وابسته:** متغیری است که متغیر مستقل بر روی آن اثر می کنند.

بررسی تاثیر ورزش بر روی فشارخون

❑ **متغیر زمینه ای** (جمعیت شناسی یا دموگرافیک): این متغیرها خصوصیات جامعه مورد مطالعه را به نحوه مطلوبی توصیف می کنند و به شناخت بهتر موضوع کمک می کنند.

سن، جنس، وضعیت تاهل، محل سکونت، محل اقامت

❑ **متغیر مداخله گر:** متغیری است که ارتباط بین متغیر مستقل و وابسته را مداخله و تحریف می کند.

شروط: (۱) با متغیر وابسته رابطه علیتی داشته باشد (۲) با متغیر مستقل رابطه علی یا غیر علی داشته باشد (۳) در مسیر بین متغیر مستقل و وابسته نباشد.

❑ **متغیر واسطه گر:** متغیر است که در مسیر بین متغیر مستقل و وابسته قرار می گیرد.



THE END

By Sajjad Rahimi
Ph.D in Epidemiology



Thank You!