

	結果が2択 「○・×」 「合格不合格」「あり・なし」	結果は程度（複数） 「数値」「点数」「金額」とか
どっちを使う？ 使う？使わない？ 2択		t 検定
量はくら どれい程 答え複数	結果がバラバラだから、割合は無理！ 結果の平均を比較してを検定	

A	9	13
B	12	13
C	10	9
D	10	7
E	12	16
F	11	7
G	13	11
H	11	15
I	11	8
J	11	11
平均	11	11

平均が同じでも

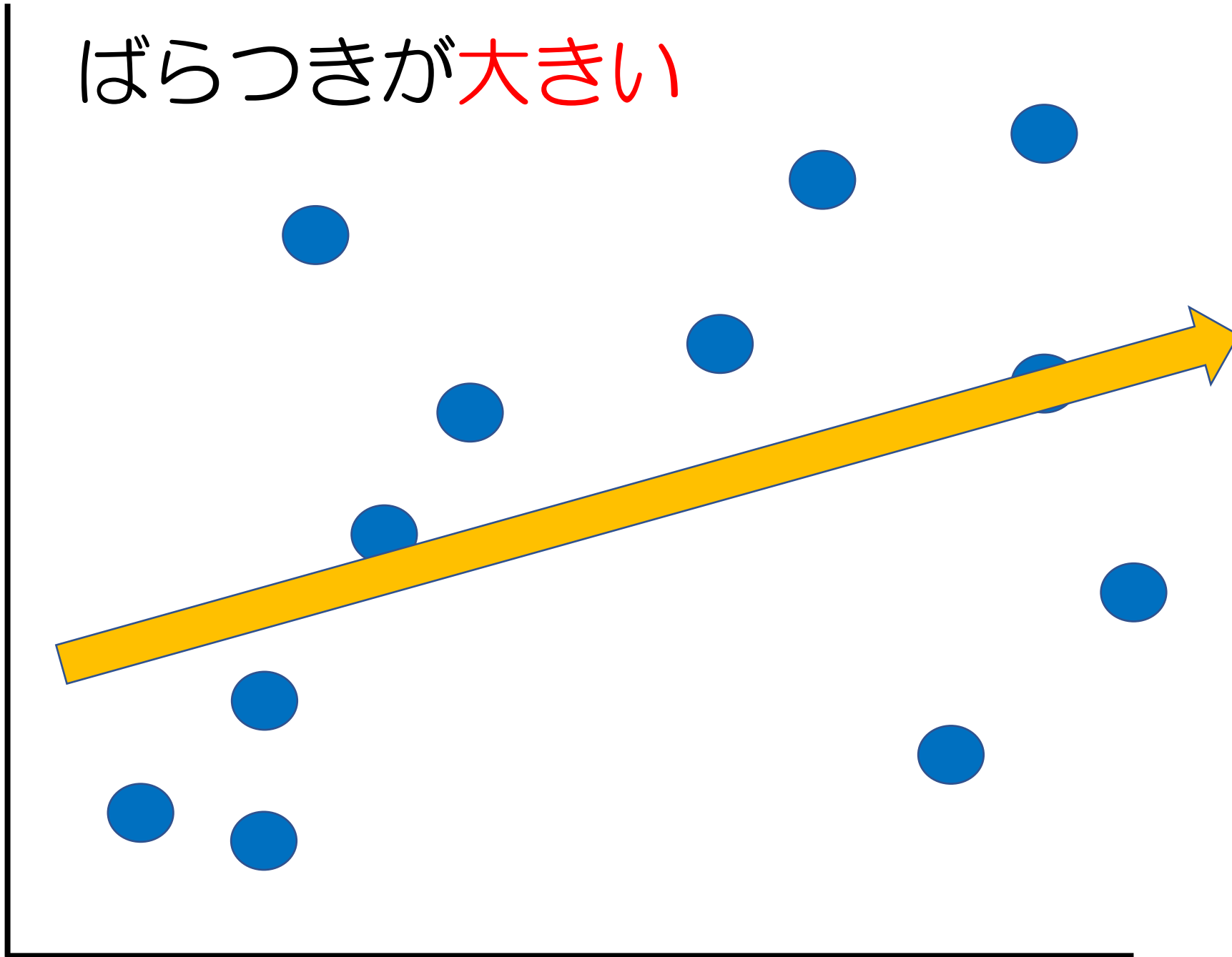
意味が違う！

ばらつきは

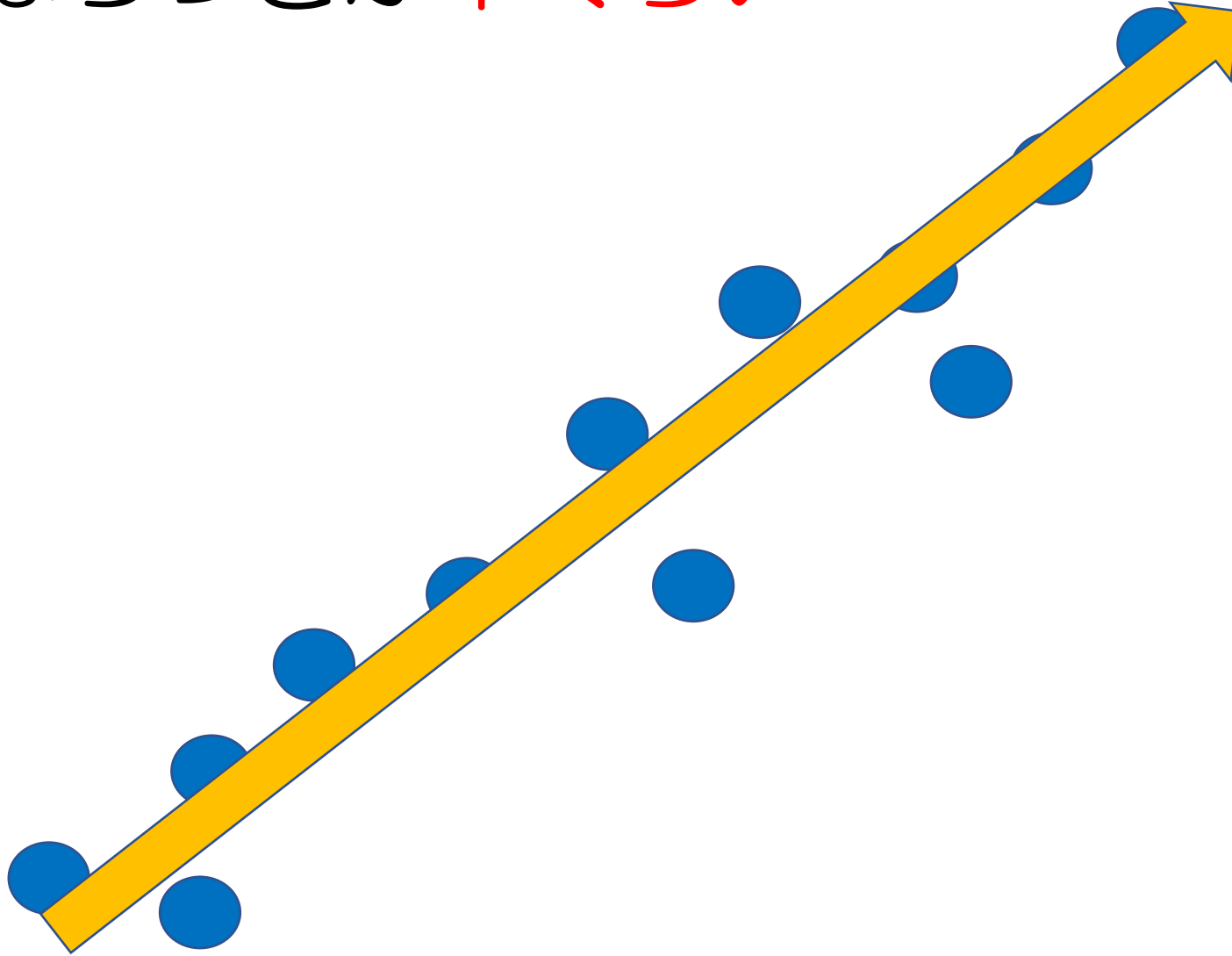
左：小さい

右：大きい

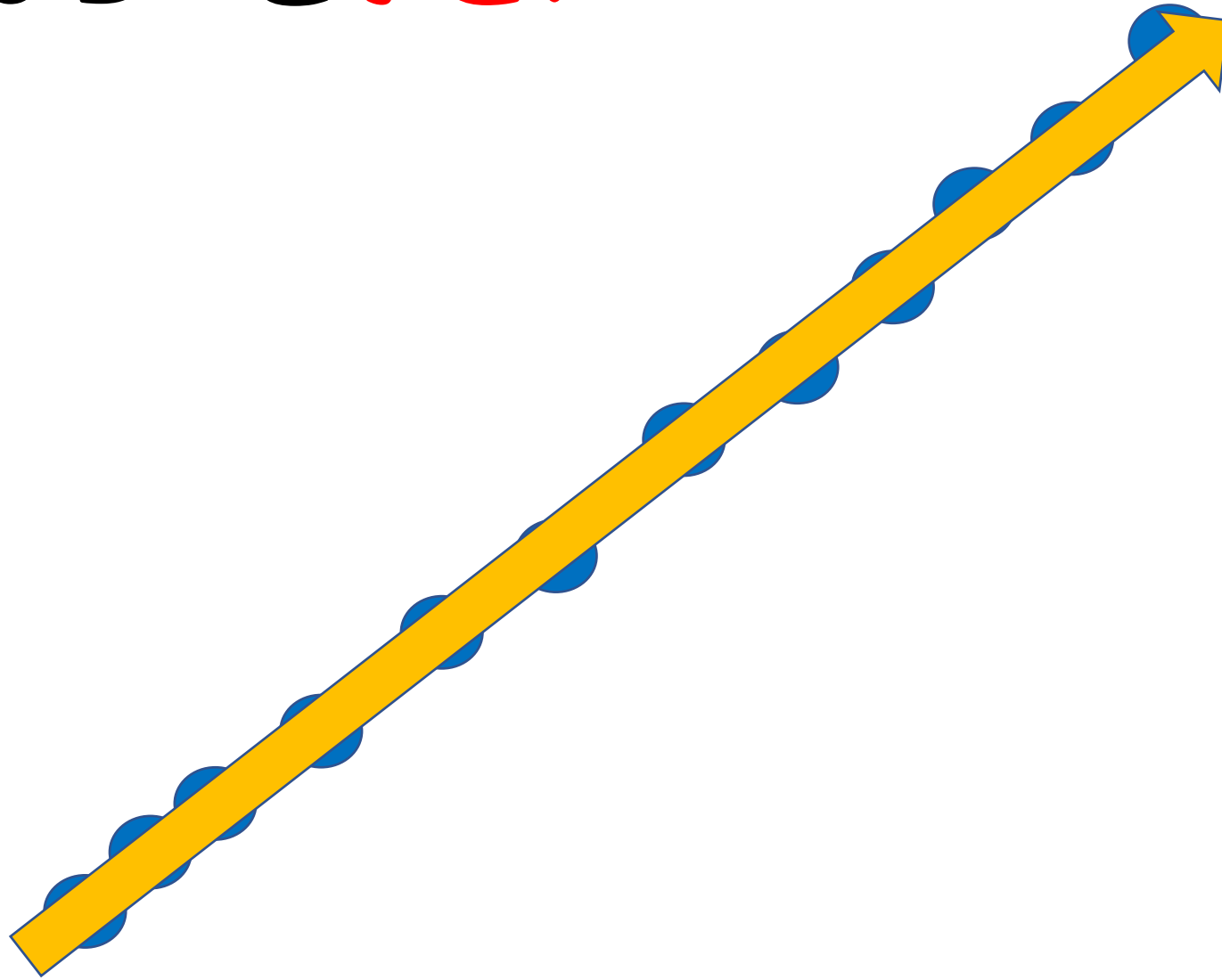
ばらつきが大きい



ばらつきが中くらい



ばらつき小さい



必要な標本数を考えるとき重要なのが
「誤差」（ばらつき）の割合

「誤差が少ない」

⇒ 少ない数でもだいたい正確

「誤差が多い」

⇒ 少ない数では不正確



その「効果」 （効果の大きさ：効果量 d ）

「効果量 d 」 介入によってどれだけ変化が起きたか？
を表す値（通常-3 から 3 の範囲）

例えば、新薬を飲んだグループと、
飲まないグループを比較して
薬の効果がどの程度あったか
を数値で示したのが効果量



その「効果」 （効果の大きさ：効果量 d ）

「効果量 d 」 介入によってどれだけ変化が起きたか？
を表す値（通常-3 から 3 の範囲）

例えば、

$d = 0.1$ 効果がほとんどない （2つの間に差がない）

$d = 1.0$ 効果がめっちゃあった（2つの間の差が大きい）

$d = 10$ ありえない結果 （データミスとか）

その「効果」 （効果の大きさ：効果量 d ）

「効果量 d 」 介入によってどれだけ変化が起きたか？
を表す値（通常-3 から 3 の範囲）

例えば、

$d = 0.1$ 効果がほとんどない （2つの間に差がない）

$d = 1.0$ 効果がめっちゃあった（2つの間の差が大きい）

$d = 10$ ありえない結果 （データミスとか）

その「効果」 （効果の大きさ：効果量 d ）

代表的な効果量の指標

「Cohen's d 」 「Hedges' g 」 「 r : 相関係数」

	Cohen's d	Hedges' g
効果量	記述的效果量	推定的效果量
意味	平均値の差を標準偏差で割った値	平均値の差を偏分散の平方根で割った値
特徴	サンプルの2つの差を記述した値	2つの母集団の差を推測する値

必要な標本数を考えるとき重要なのが
「誤差」（ばらつき）の割合

とも言えるやんね！

「誤差が少ない」 効果が大さい！

⇒ 少ない数でもだいたい正確

「誤差が多い」 効果が小さい

⇒ 少ない数では不正確



効果量 d とサンプルサイズの関係

効果量が多い：信号が強い状態

小さなサンプルサイズ（少ない数）でも信号を捉えることができる

効果量が少ない：信号が弱い

大きなサンプルサイズじゃないと信号を捉えることができない

サンプルサイズが小さい

ただ、どちらも可能性！

効果量が少ないと検出できない、大きいと検出できる

サンプルサイズが大きい

効果量が少なくても検出できる、大きいとより正確に検出できる

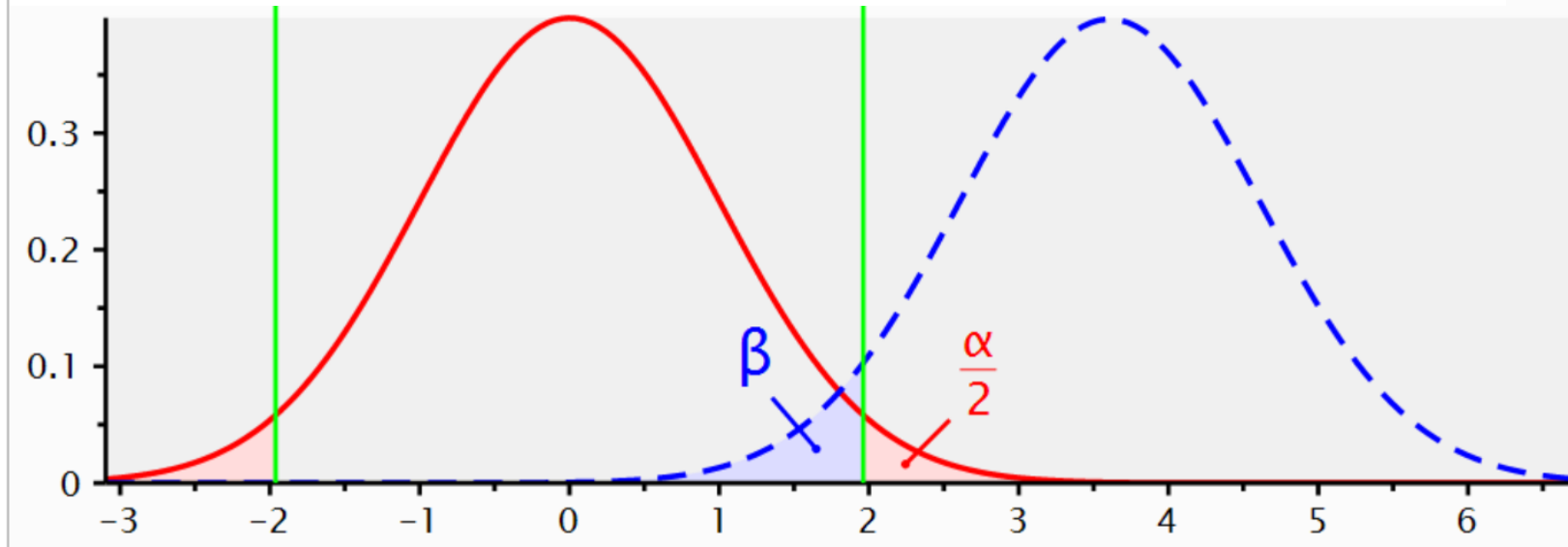
効果量 d とサンプルサイズの関係

計算はめんどくさいから、

統計ソフトに任せて



めちゃくちゃばらついてるとき



効果量 d が小さい

大きいサンプルサイズが必要！

A priori: Compute required sample size

Input Parameters

Tail(s) Two

Determine =>

Effect size d 0.1

α err prob 0.05

Power ($1 - \beta$ err prob) 0.95

Output Parameters

Noncentrality parameter δ 3.6083237

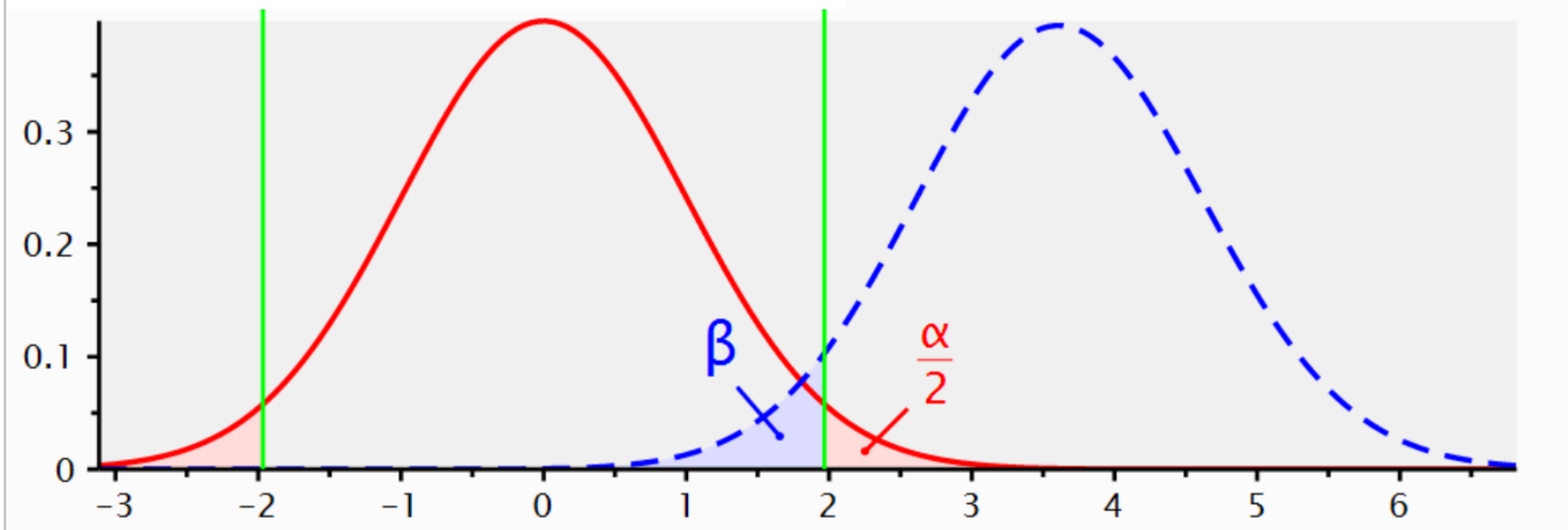
Critical t 1.9617891

Df 1301

Total sample size 1302

Actual power 0.9500867

ばらつきが大きい



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.2

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.6166283

Critical t

1.9672675

Df

326

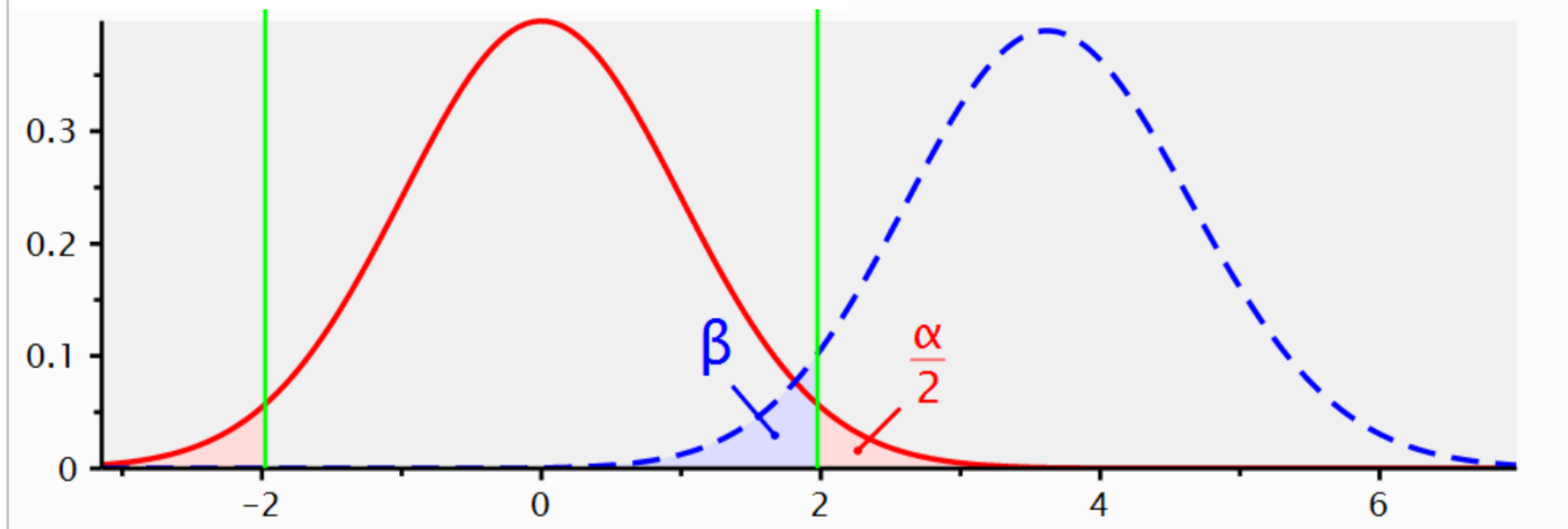
Total sample size

327

Actual power

0.9501171

ばらつきが大きい



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.3

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.6373067

Critical t

1.9763457

Df

146

Total sample size

147

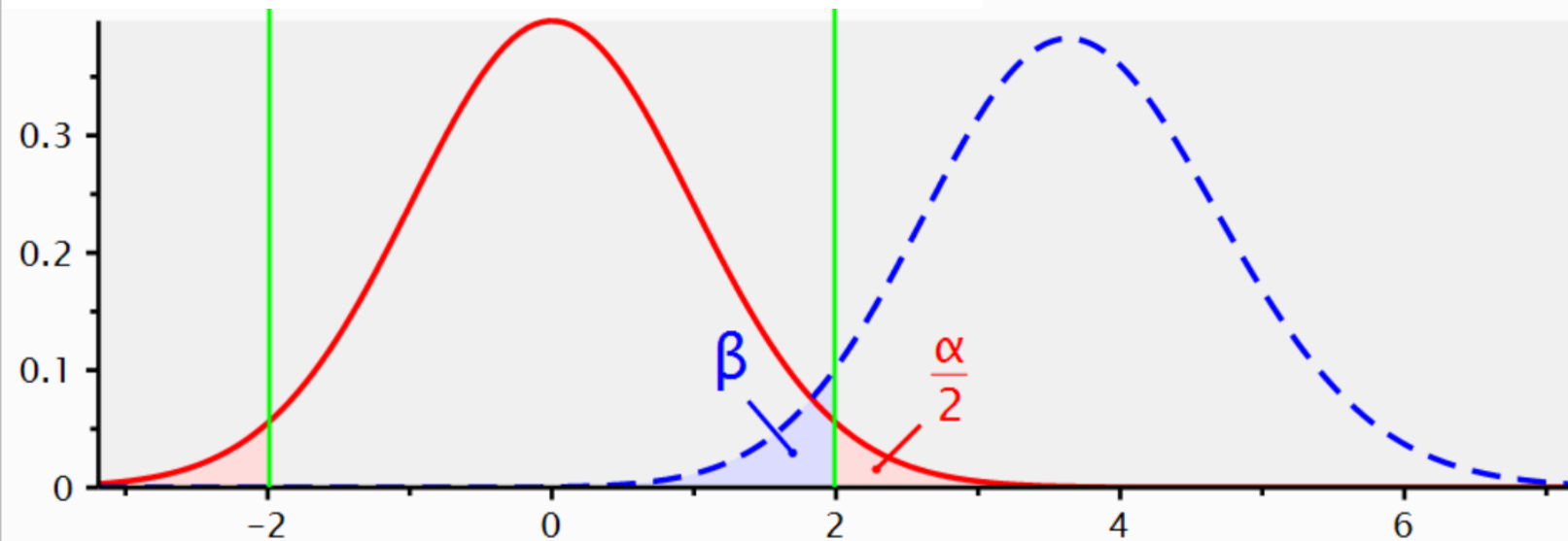
Actual power

0.9508665

ばらつきが中くらい

analyses

6



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.4

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.6660606

Critical t

1.9889598

Df

83

Total sample size

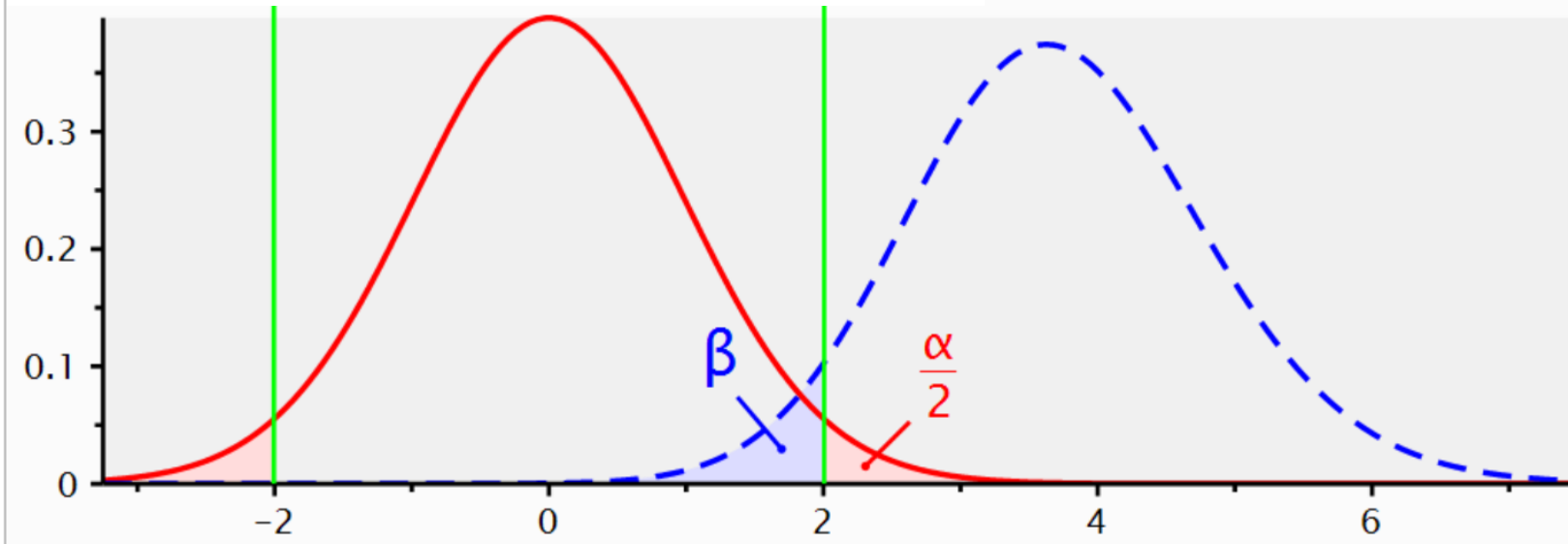
84

Actual power

0.9518794

ばらつきが中くらい

analyses



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.5

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.6742346

Critical t

2.0057460

Df

53

Total sample size

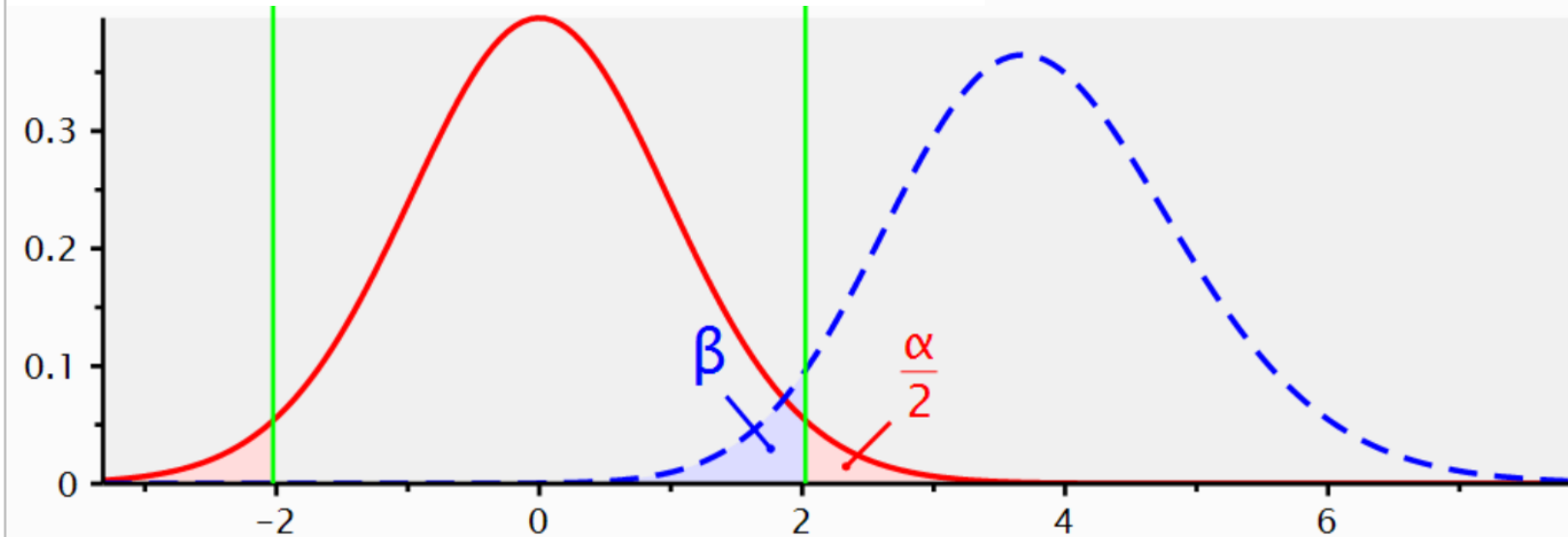
54

Actual power

0.9502120

ばらつきが中くらい

analyses



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.6

α err prob

0.05

Power (1- β err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.7469988

Critical t

2.0243942

Df

38

Total sample size

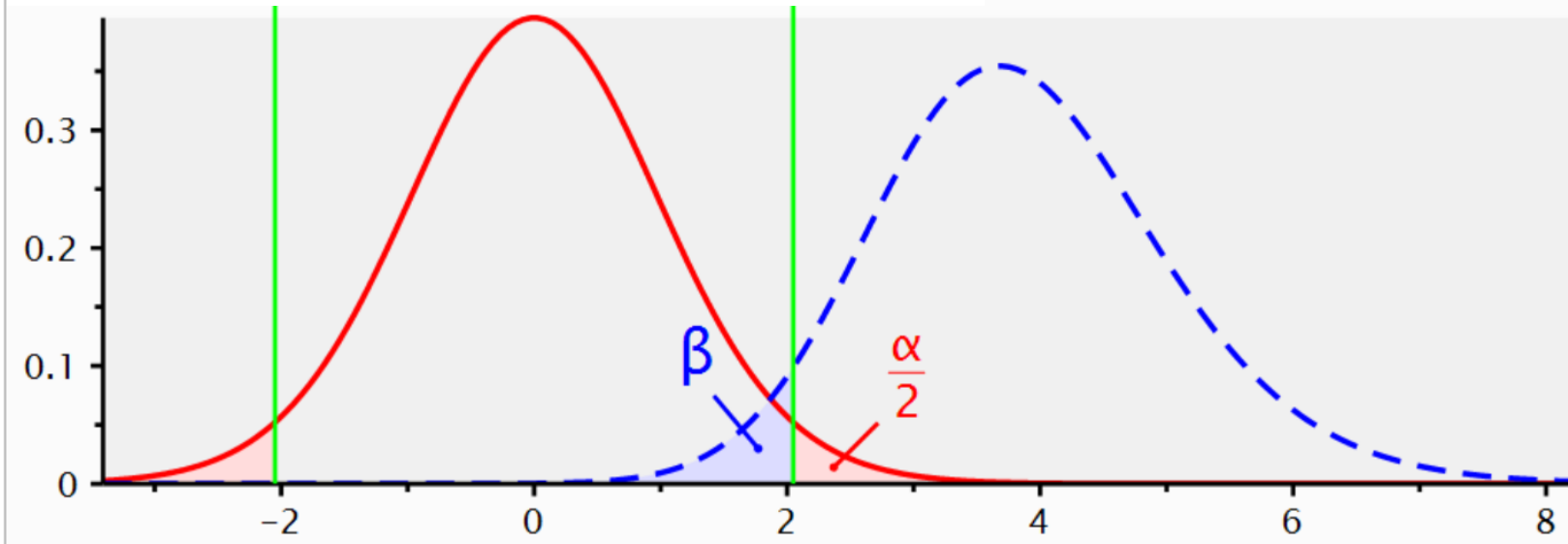
39

Actual power

0.9545588

ばらつきが中くらい

analyses



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.7

α err prob

0.05

Power (1- β err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.7696154

Critical t

2.0484071

Df

28

Total sample size

29

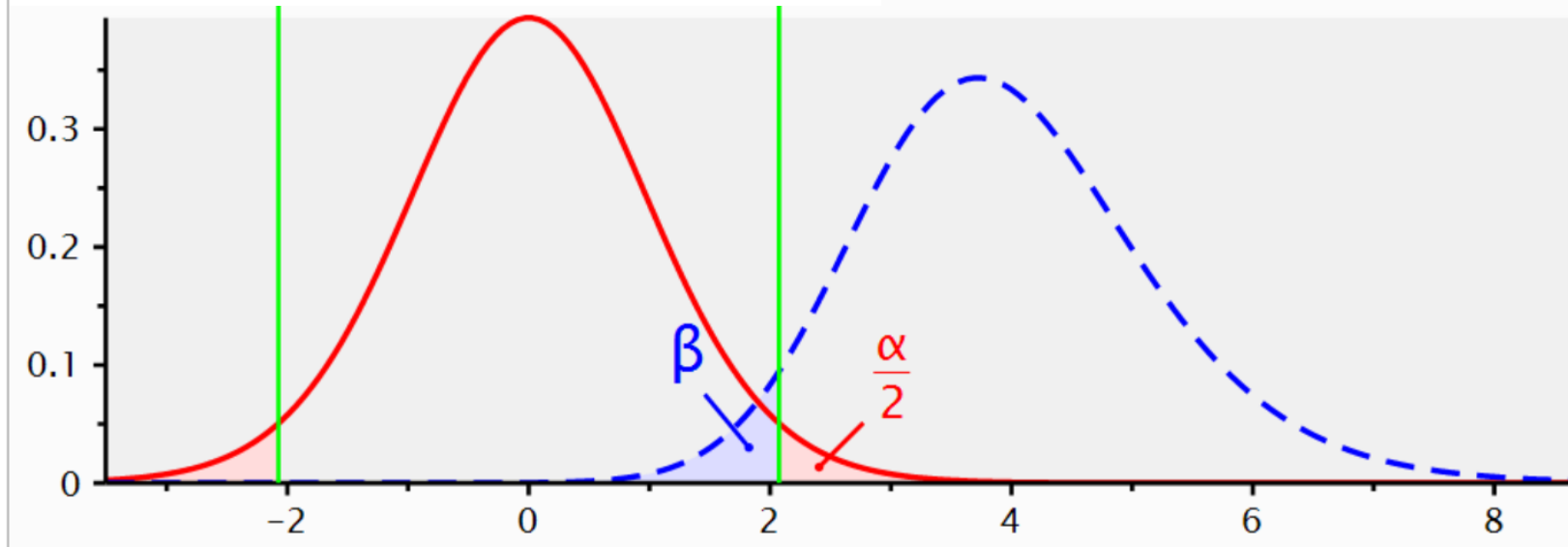
Actual power

0.9532752

ばらつきが小さい

power analyses

387



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.8

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.8366652

Critical t

2.0738731

Df

22

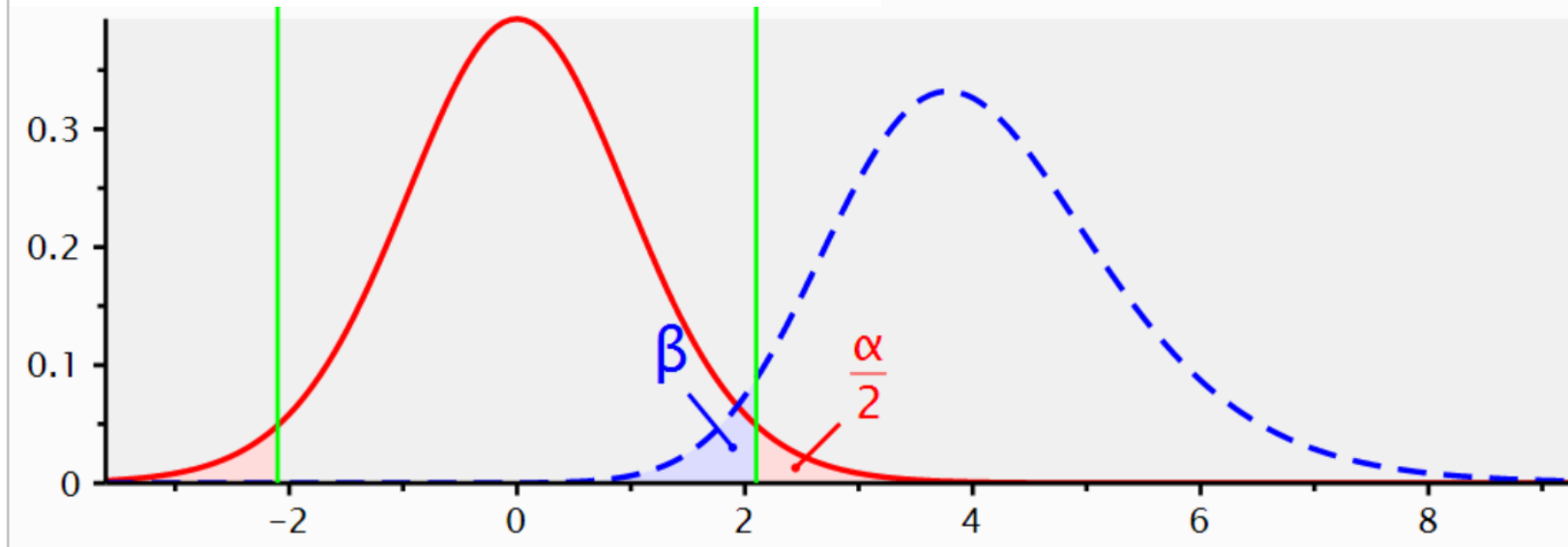
Total sample size

23

Actual power

0.9558497

ばらつきが小さい



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

0.9

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

3.9230090

Critical t

2.1009220

Df

18

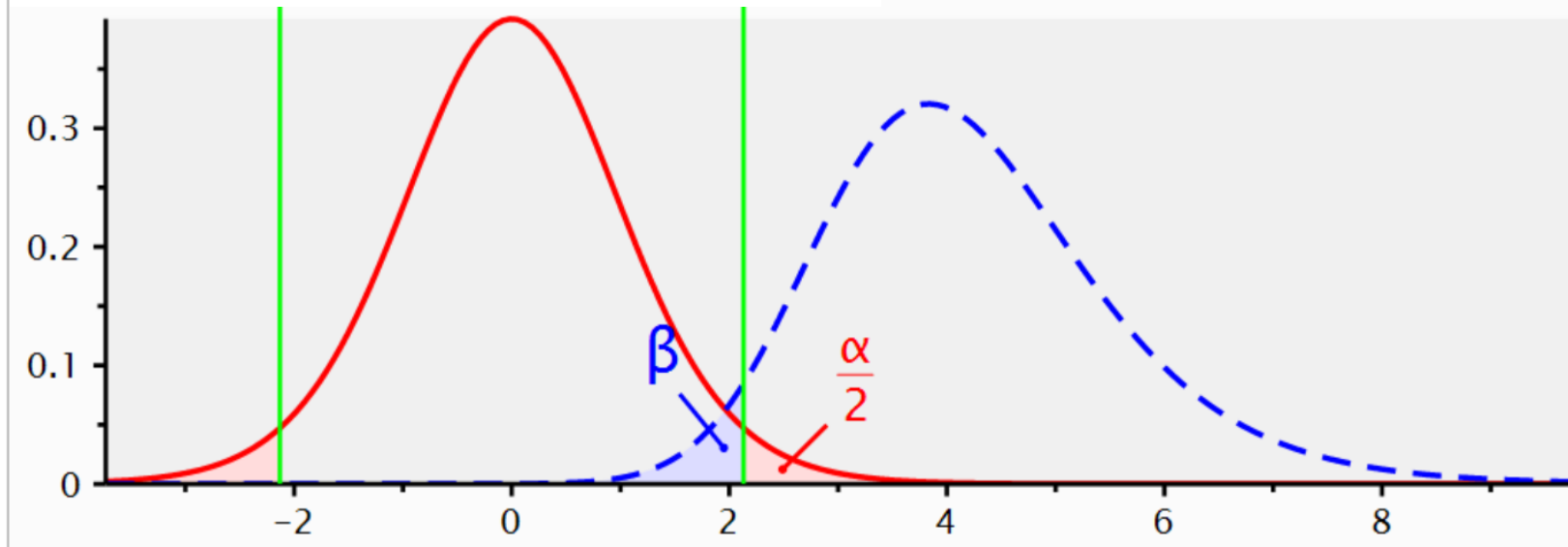
Total sample size

19

Actual power

0.9596224

ばらつきが小さい



Test family

t tests

Statistical test

Means: Difference from constant (one sample case)

Type of power analysis

A priori: Compute required sample size – given α , power, and effect size

Input Parameters

Tail(s) Two

Determine =>

Effect size d

1

α err prob

0.05

Power ($1-\beta$ err prob)

0.95

Output Parameters

Noncentrality parameter δ

4.0000000

Critical t

2.1314495

Df

15

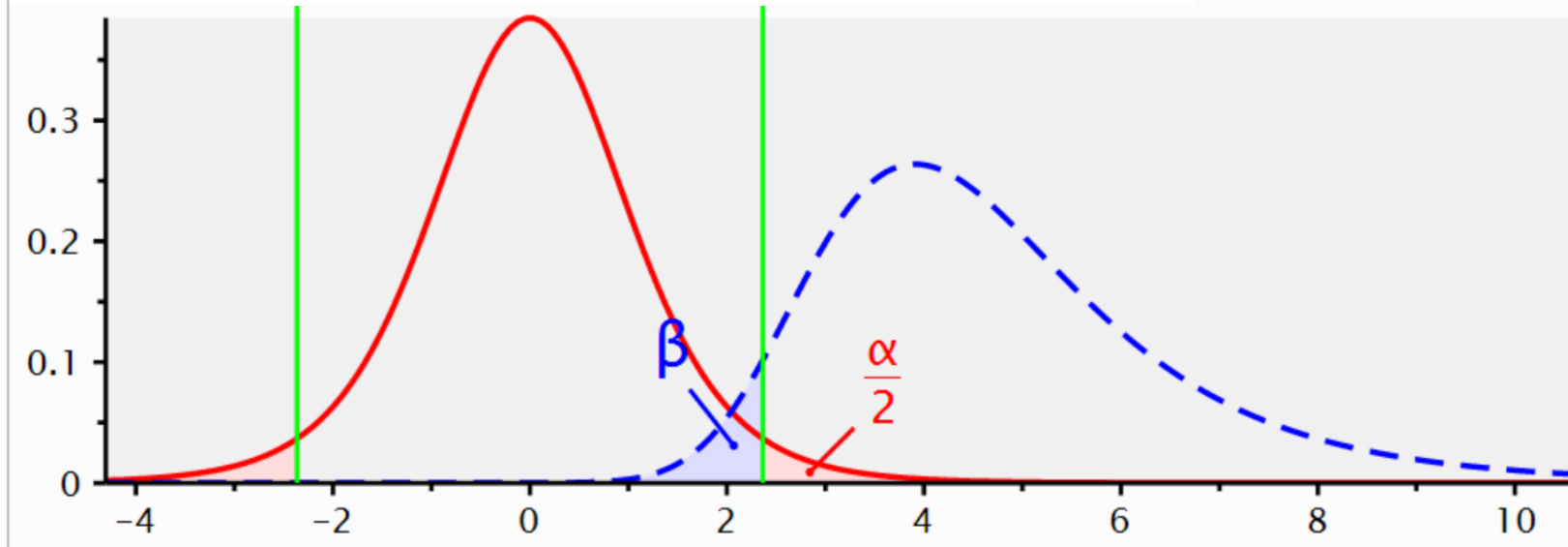
Total sample size

16

Actual power

0.9618851

ほとんどばらつきがない



効果量 d が大きい

小さいサンプルサイズでOK!

Input Parameters

Determine =>

Tail(s) Two

Effect size d 1.5

α err prob 0.05

Power ($1-\beta$ err prob) 0.95

Output Parameters

Noncentrality parameter δ 4.2426407

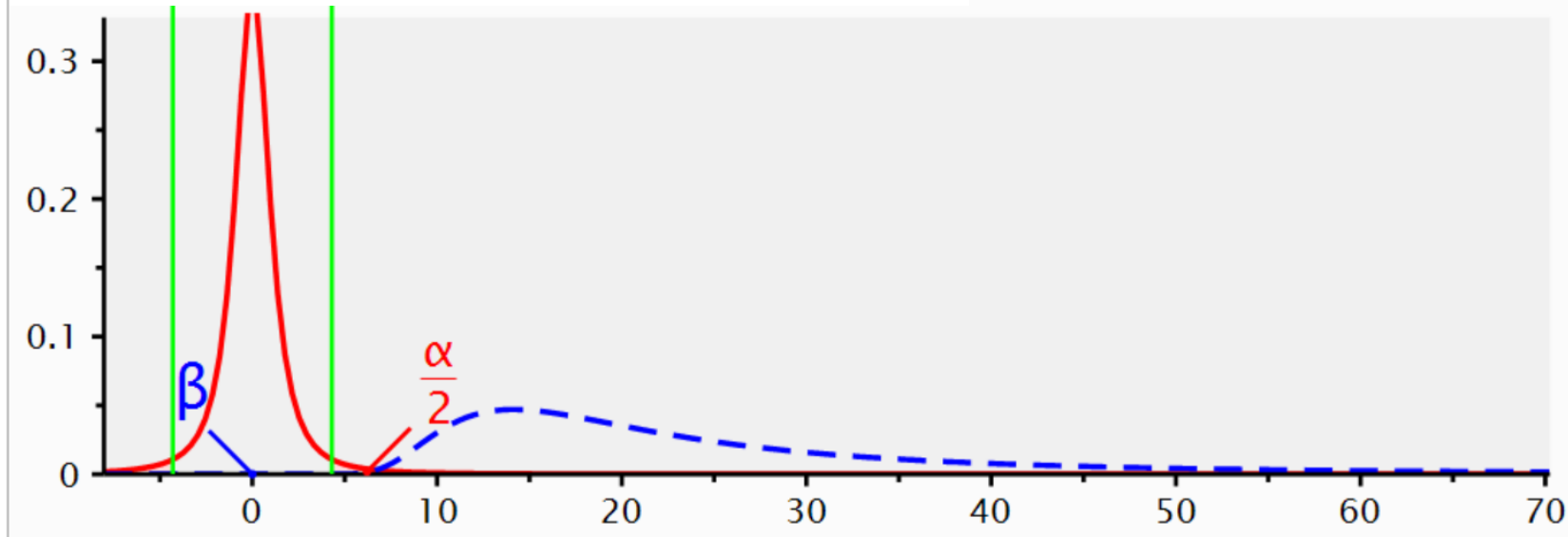
Critical t 2.3646243

Df 7

Total sample size 8

Actual power 0.9509518

ばらつきが全くない



効果量 d がありえない

Test family: Statistical test:
 in constant (one sample case)
 A priori: Compute required sample size – given α , power, and effect size

Input Parameters		Output Parameters	
Determine =>	Tail(s) Two	Noncentrality parameter δ	
Effect size d	10	Critical t	4.3020527
α err prob	0.05	Df	2
Power ($1-\beta$ err prob)	0.95	Total sample size	3
		Actual power	0.9999996

3個でOK!

詳しい原理はわからんけど…

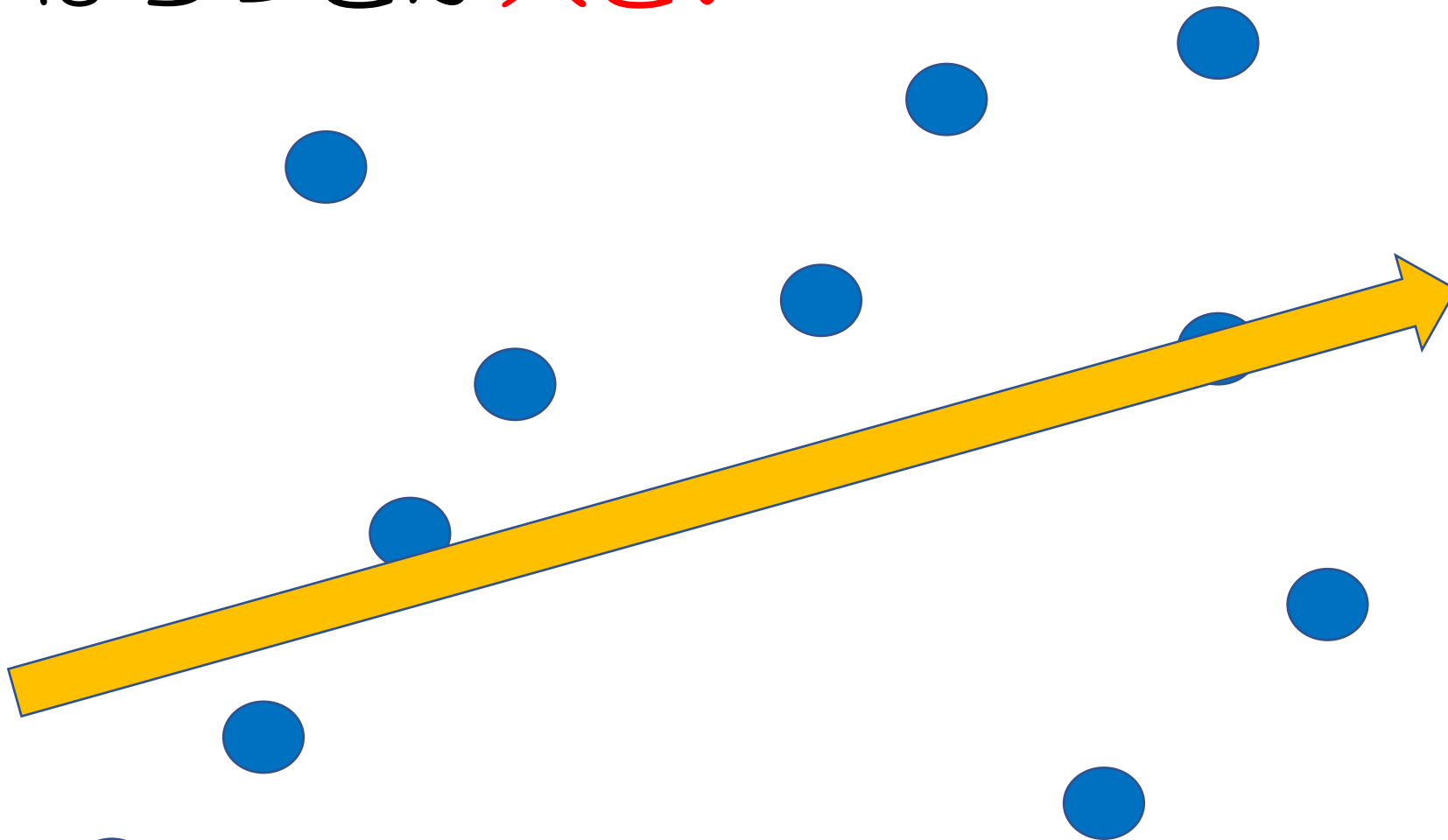
「ばらつき大」：330くらい

「ばらつき中」：50くらい

「ばらつき小」：20くらい

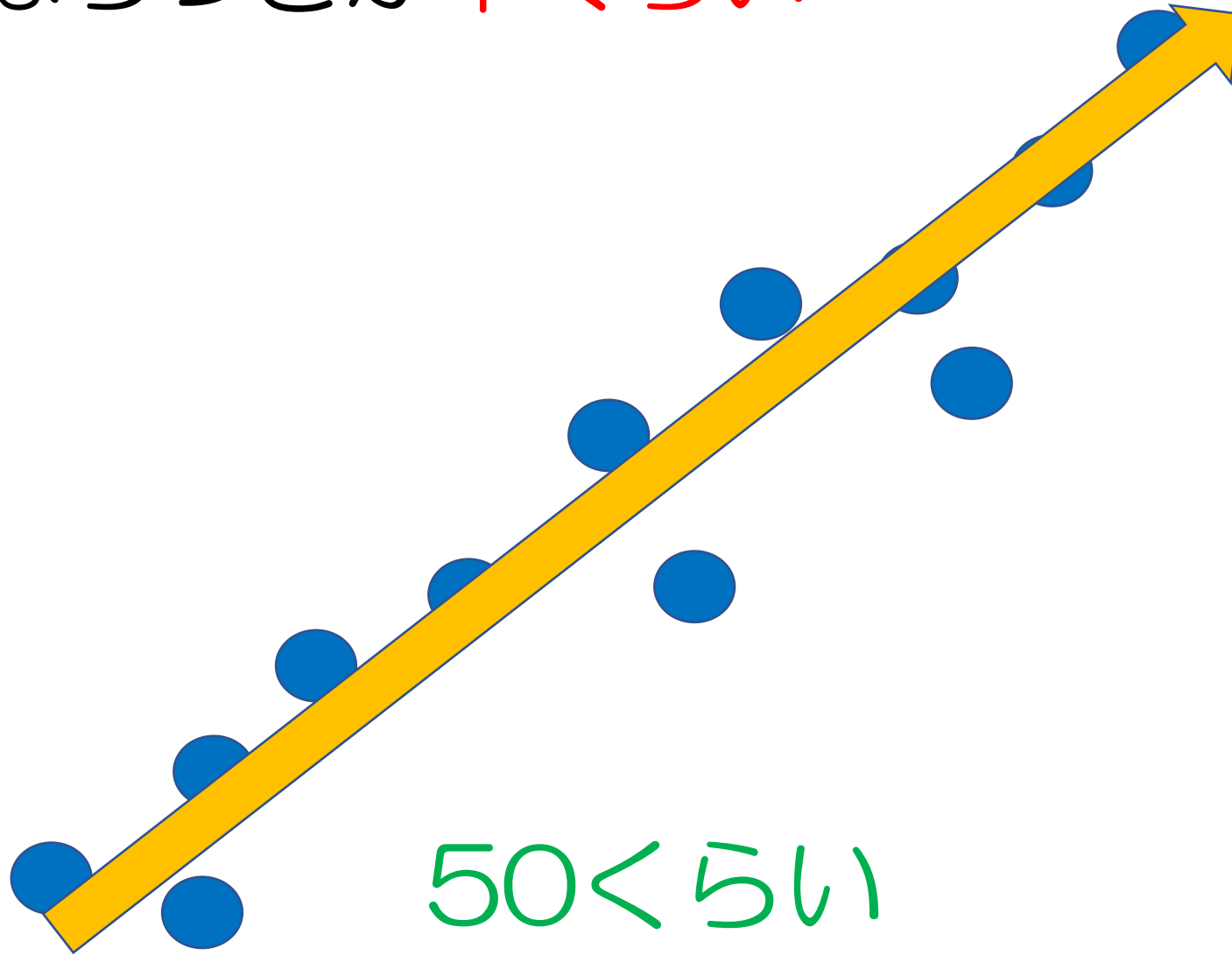
のところが大体の目安

ばらつきが大きい



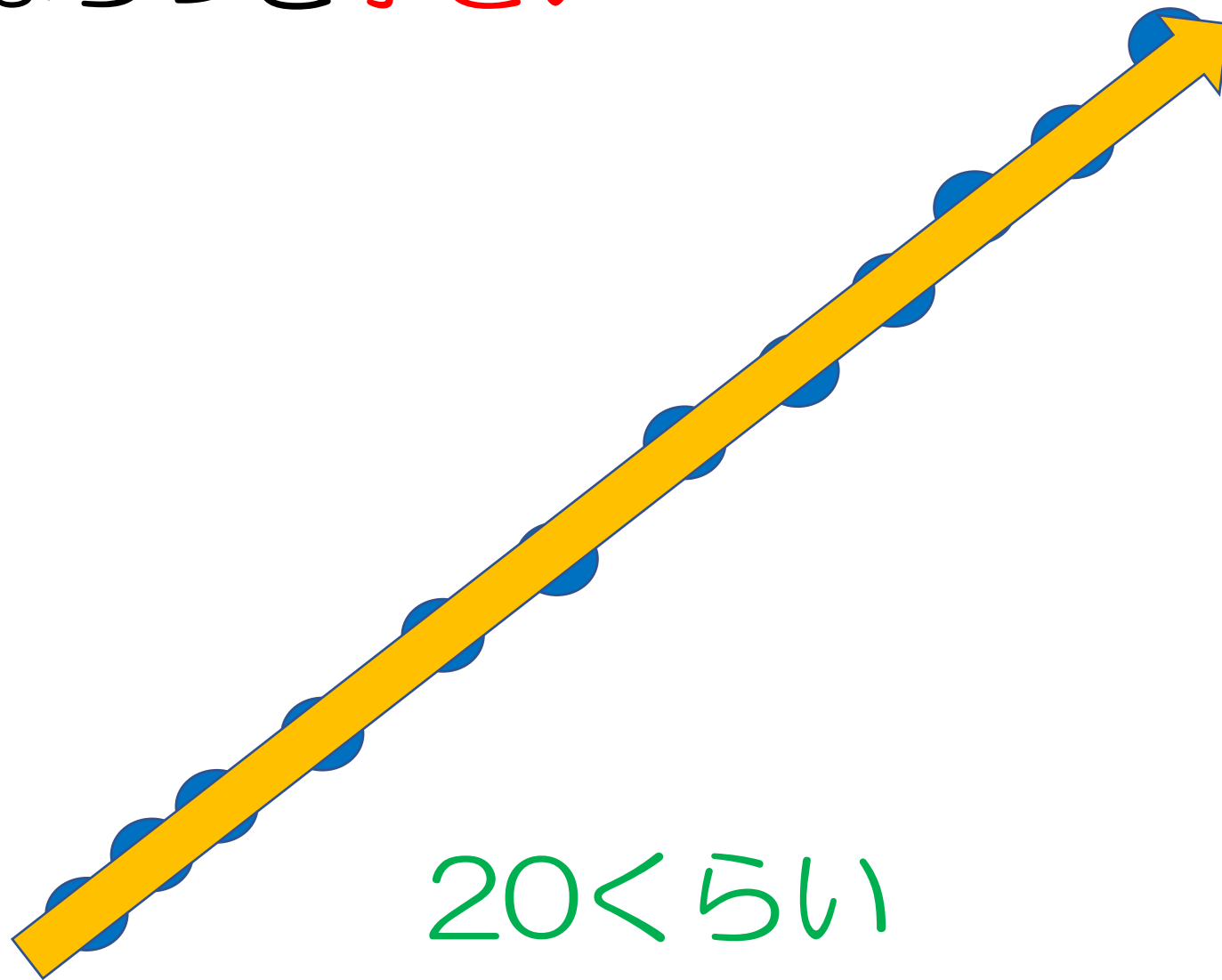
330くらい

ばらつきが中くらい



50くらい

ばらつき小さい

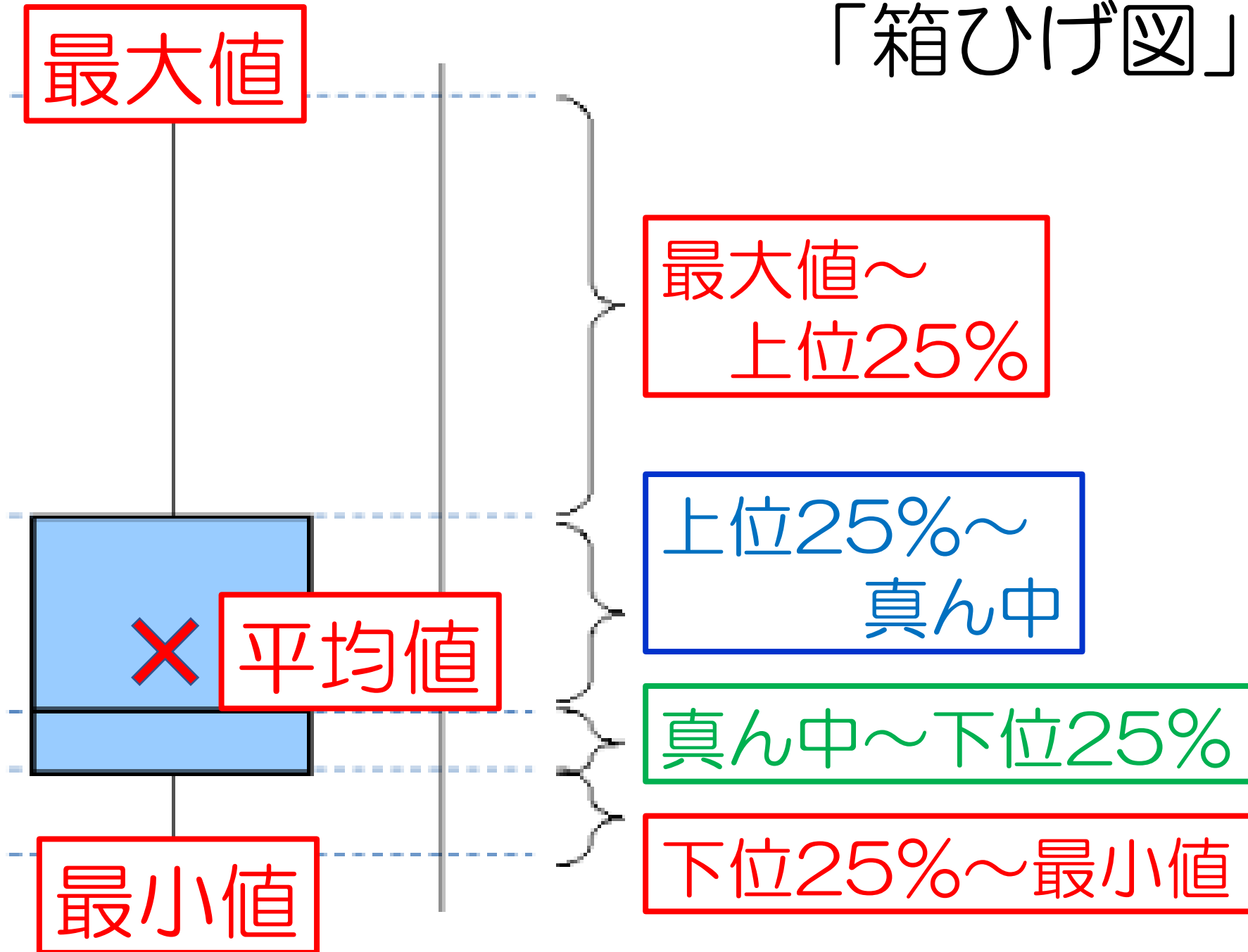


20くらい

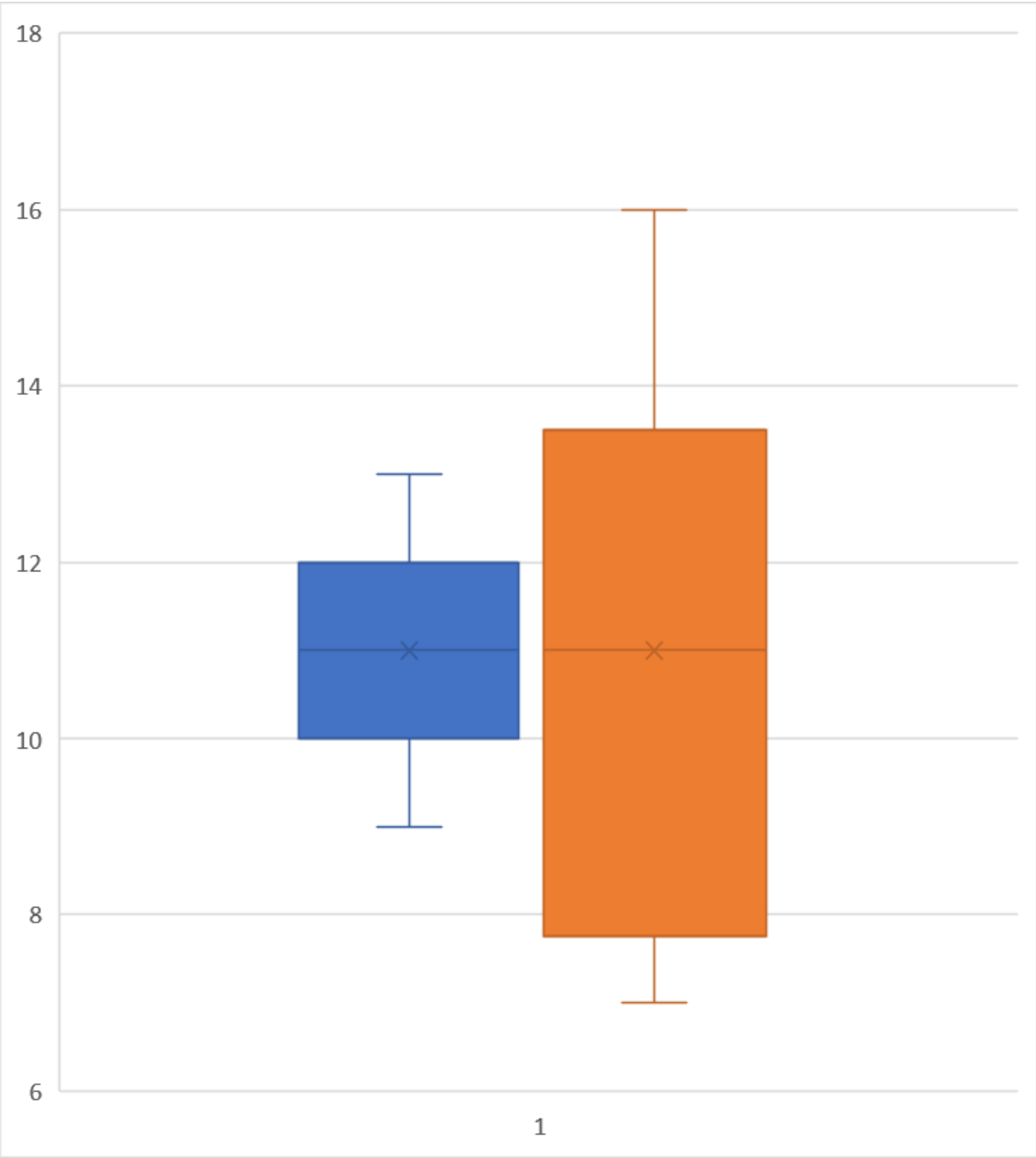
そのばらつきをどうやって表すのか？

「箱ひげ図」

「箱ひげ図」



A	9	13
B	12	13
C	10	9
D	10	7
E	12	16
F	11	7
G	13	11
H	11	15
I	11	8
J	11	11
平均	11	11



仮説検定（統計的検定）

1 カイ二乗検定

2 t 検定

3 回帰分析

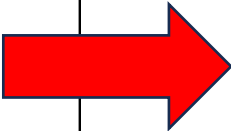


	結果が2択 「○・×」 「合格不合格」「あり・なし」	結果は程度（複数） 「数値」「点数」「金額」とか
どっちを使う？ 使う？使わない？ 2択	×二乗検定	t 検定
量はどれくらい？ どれくらい程度？ 答え複数	ロジスティック回帰	回帰分析

	結果が2択 「○・×」 「合格不合格」「あり・なし」	結果は程度（複数） 「数値」「点数」「金額」とか
どっちを使う？ 使う？使わない？ 2択	×二乗検定 	
量はどれくらい？ どれくらい程度？ 答え複数	赤と青のお皿で、食べるかどうかの 割合 を検定	

	結果が2択 「○・×」 「合格不合格」「あり・なし」	結果は程度（複数） 「数値」「点数」「金額」とか
どっちを使う？ 使う？使わない？ 2択		t 検定 
量はくら どれ い程 答え複数	結果がバラバラだから、割合は無理！ 結果の平均を比較してを検定	

	結果が2択 「○・×」 「合格不合格」「あり・なし」	結果は程度（複数） 「数値」「点数」「金額」とか
どっちを 使う？		
使う かな	これは前にやった「相関係数」と ちよつと似てる	
2		
量はこれ くらい？		
どれくら い程度？		
答え複数		



回帰分析

	結果が2択「○・×」 「合格不合格」「あり・なし」	結果は程度（複数） 「数値」「点数」「金額」とか
どっちを使う？ 使う？使わない？ 2択	↑ 結果は変わる？ ← これによって 原因	結果
量はどれくらい？ どれくらい程度？ 答え複数		

「相関分析と回帰分析の違い」

相関分析

相関係数がメイン。1対1のみ
XとYの関係性を見るだけ
XとYの関係式は出てこない
どっちが原因とかは考えない

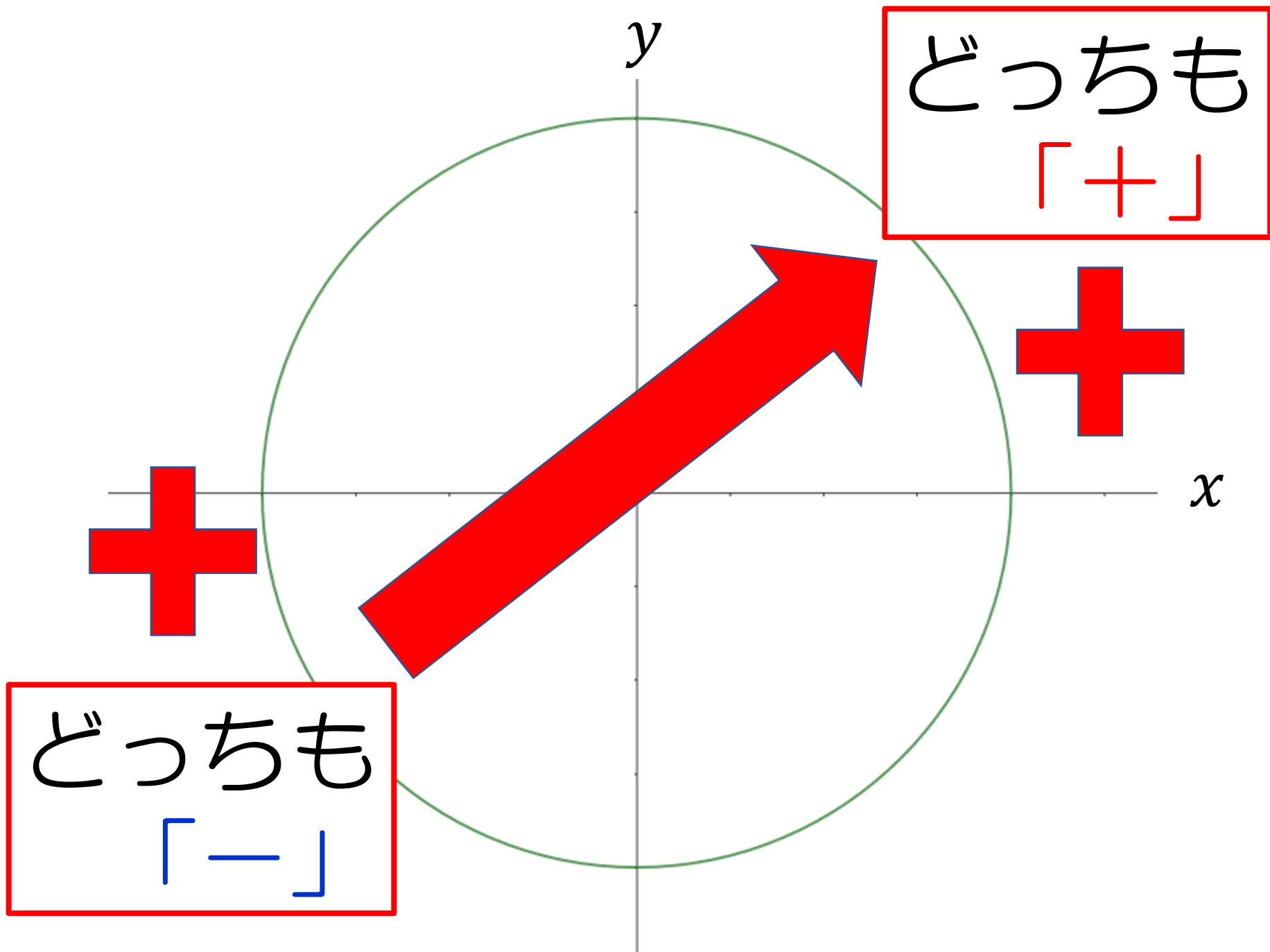
回帰分析

1対1でも、1対複数でも！
XからみたYの変化を見る
数式 $Y = aX + b$ を立てて予測
Xが原因でYが結果として進める

回帰分析は意外と簡単！

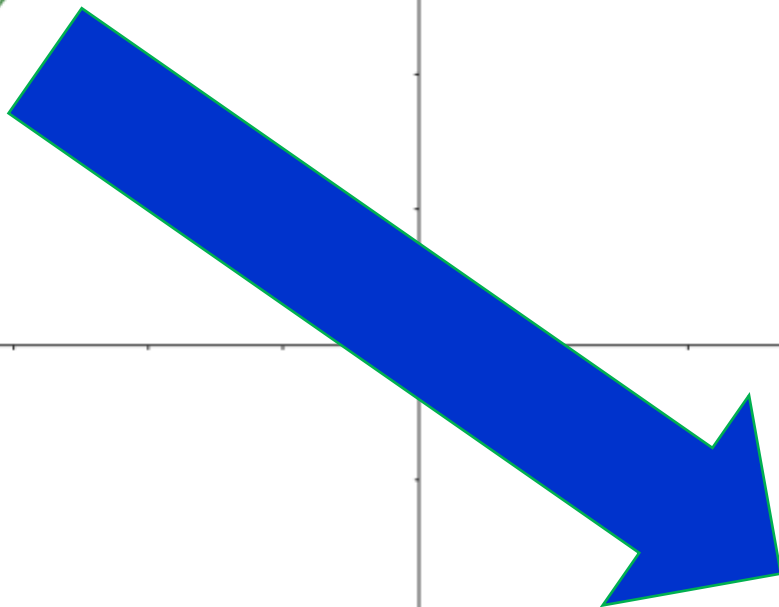
いろいろ考えずにとりあえず
やってみよう！
何となくわかってくると思う！





「 $-$ 」と「 $+$ 」

$-$



「 $+$ 」と「 $-$ 」

$-$

これと同じようなの
どっかで見た覚えはない？



相関係数

2種類のデータの関係性の強さを

「 -1 から $+1$ 」

の間の値で表した数

「 r 」で表されることが多い

相関の強さ

「Xが大きくなると
「Y」も大きくなる

$$0.7 \leq r \leq 1.0$$

強い正の相関

$$0.4 \leq r \leq 0.7$$

正の相関

$$0.2 \leq r \leq 0.4$$

弱い正の相関

$$-0.2 \leq r \leq 0.2$$

相関なし

$$-0.4 \leq r \leq -0.2$$

弱い負の相関

$$-0.7 \leq r \leq -0.4$$

負の相関

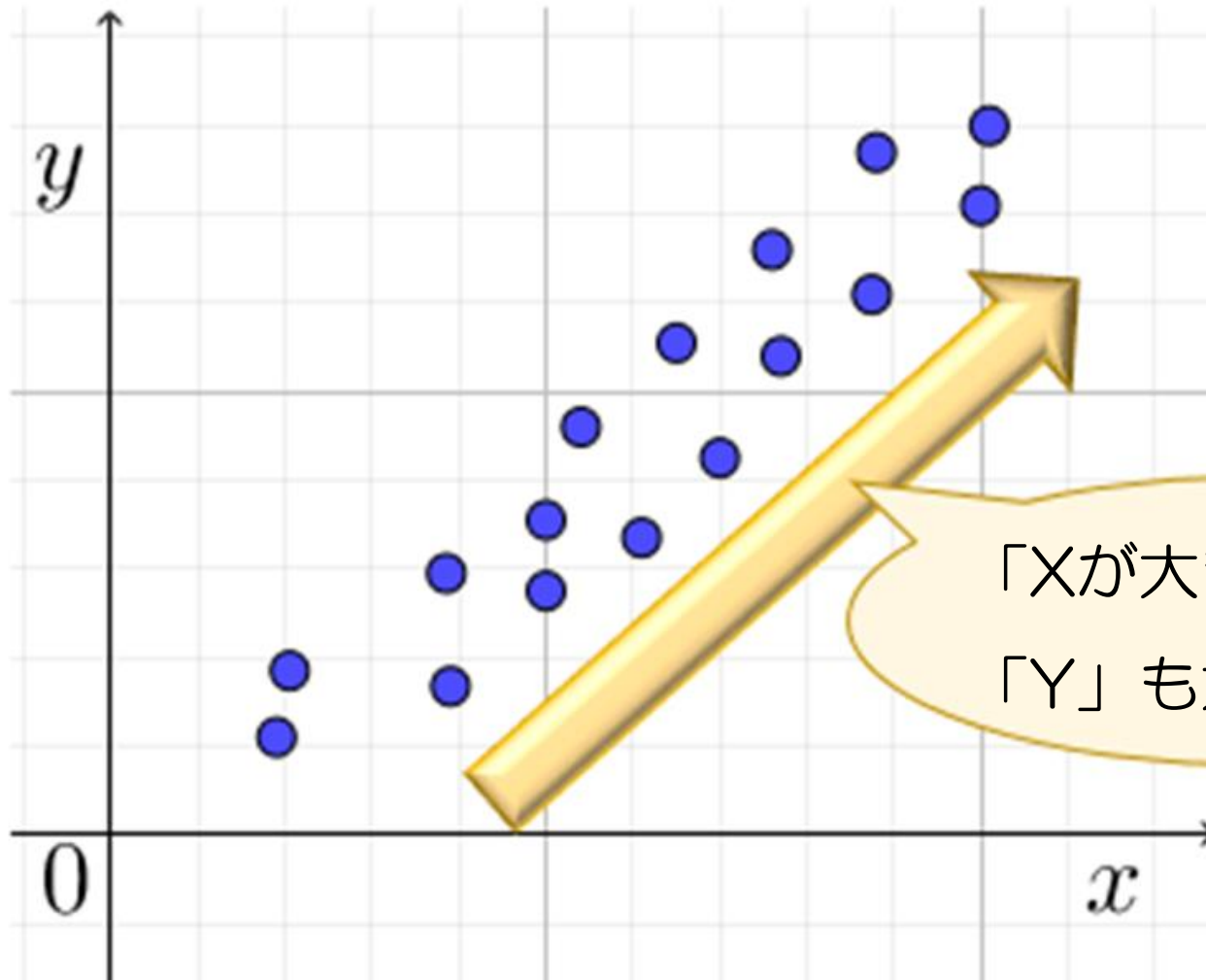
$$-1.0 \leq r \leq -0.7$$

強い負の相関

「Xが大きくなると
「Y」は小さくなる

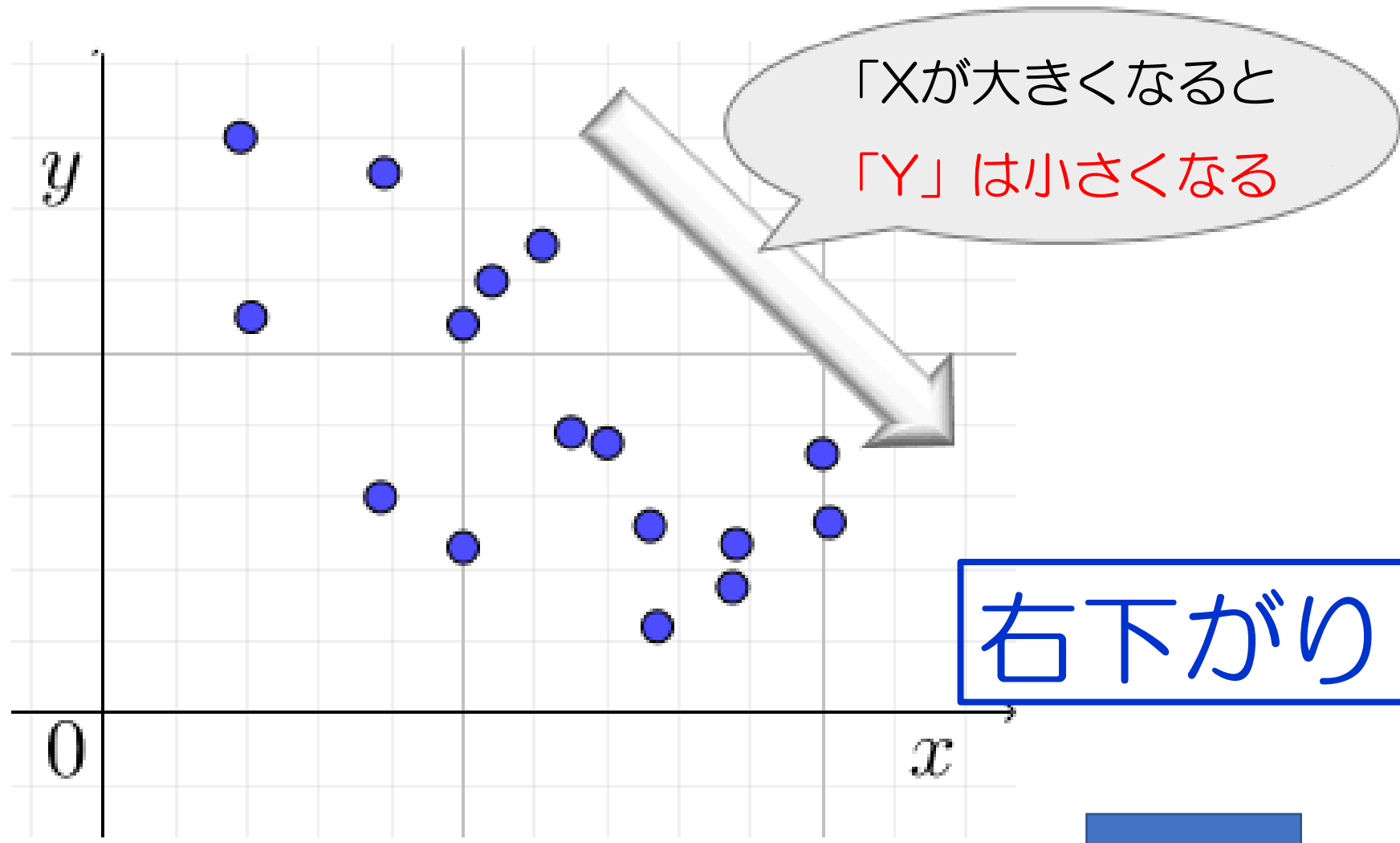
相関係数：0.94

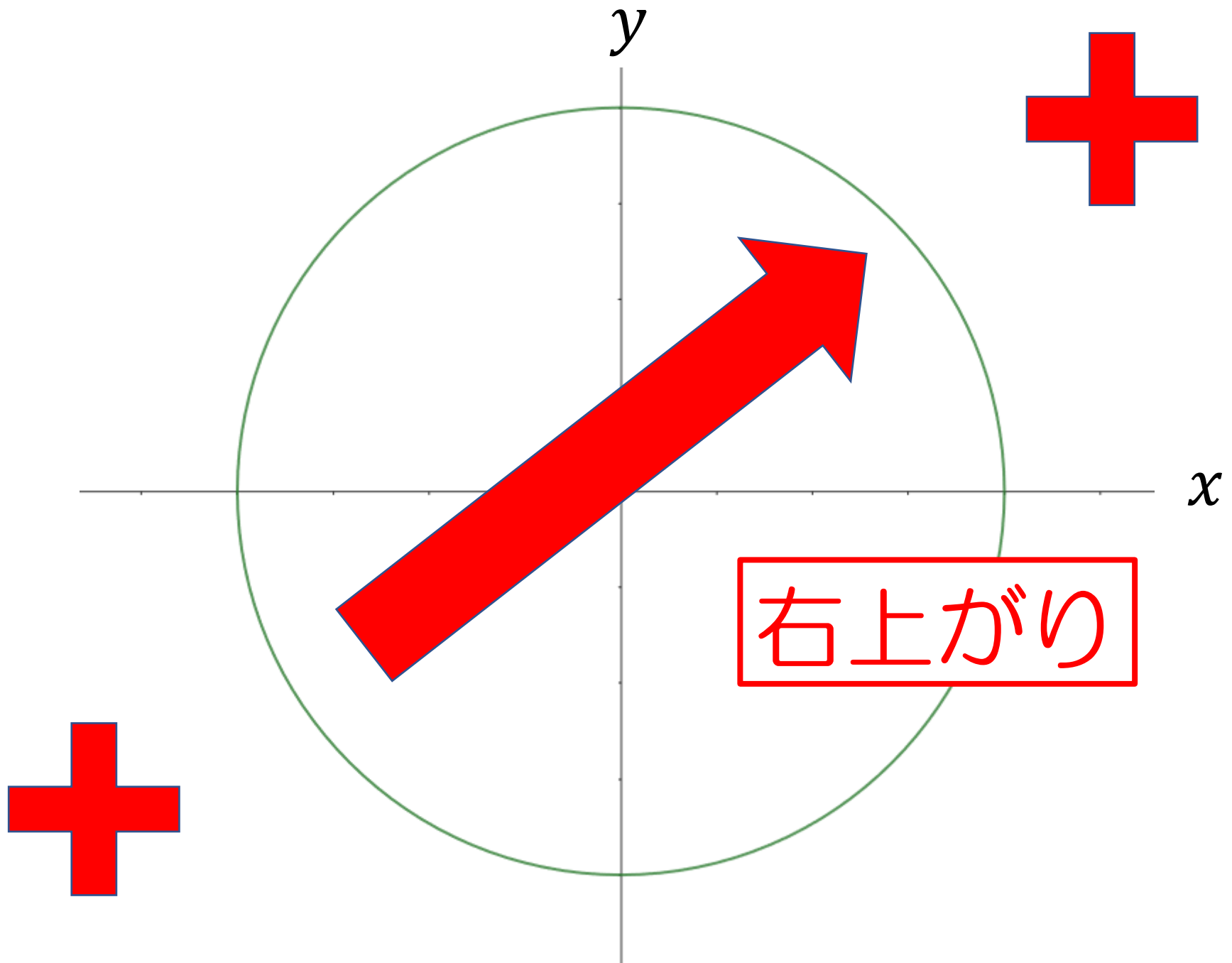
右上がり



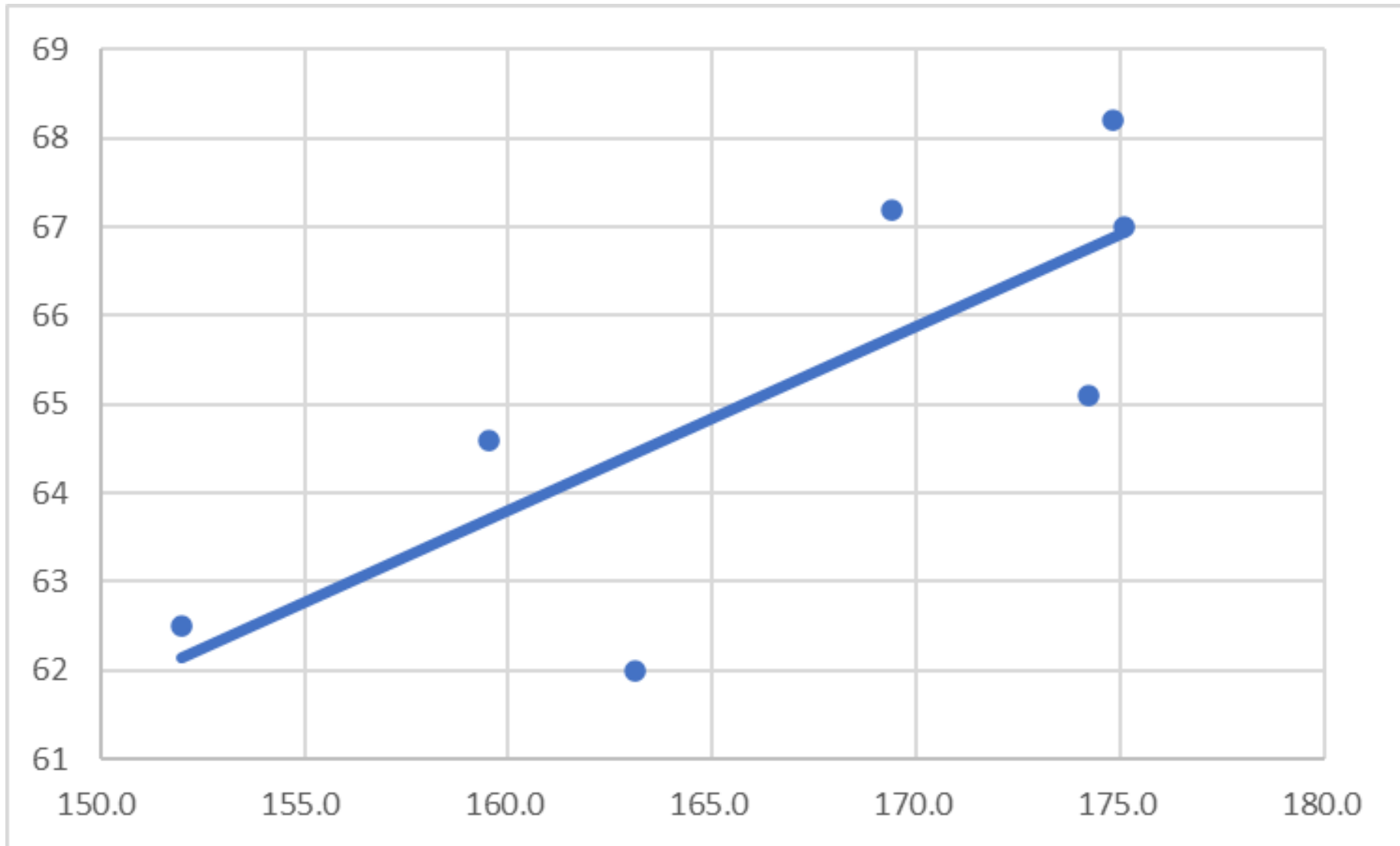
「Xが大きくなると
「Y」も大きくなる

相関係数：-0.74





散布図・近似式で見ても「右上がり」



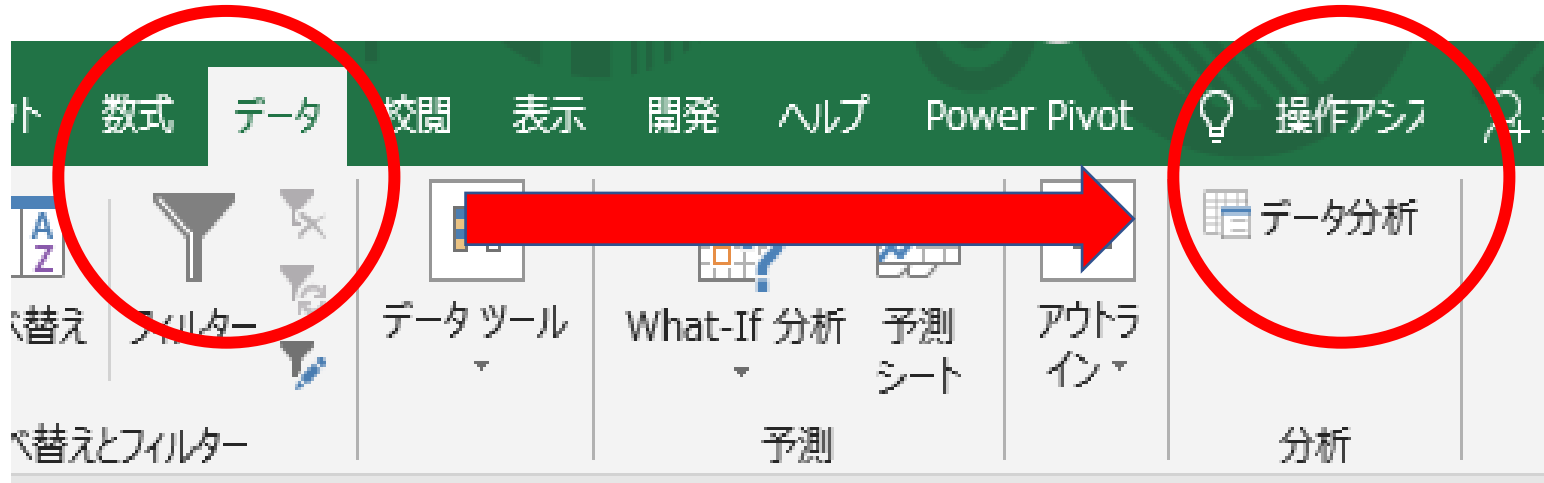
共分散：「平均の差」×「平均の差」の事

今わかったのは

「右上がりの関係にある」

ってことだけ

Excelだともっと簡単



Excelの「データ」タブの中にある

「データ分析」 → 「共分散」

こんなの出るはず

	身長	体重
身長	68.63	
体重	14.22	4.876

—瞬！

...
4.63
40.58
14.22
(平均の差×平均の差)の平均

同じようにやってみて！

シート「共分散2」

- 1 年齢と最高血圧
- 2 年齢と最低血圧
- 3 年齢と年収
- 4 最高血圧と最低血圧
- 5 最高血圧と年収
- 6 最低血圧と年収 の相関は？

こんなの出るはず

	年齢	最高血圧	最低血圧	年収
年齢	174.0			
最高血圧	121.1	113.1		
最低血圧	70.9	52.2	38.8	
年収	1491.4	796.8	604.7	17678.0

「共分散」は向きがわかるだけ…

どれくらい強い関係か

調べるために行うのが

「相関係数の有意差検定」

前に「相関係数」を調べたときは

$$0.7 \leq r \leq 1.0$$

強い正の相関

$$0.4 \leq r \leq 0.7$$

正の相関

$$0.2 \leq r \leq 0.4$$

弱い正の相関

$$-0.2 \leq r \leq 0.2$$

相関なし

$$-0.4 \leq r \leq -0.2$$

弱い負の相関

$$-0.7 \leq r \leq -0.4$$

負の相関

$$-1.0 \leq r \leq -0.7$$

強い負の相関

強さを

おおまかに

調べた

相関係数の基準値

$$r_0 = \sqrt{\frac{4}{n+2}}$$

n : データ数

シート「相関1」

AグループとBグループの相関係数

それぞれ求めてみて

A	身長	体重
身長	1	
体重	0.778	1

B	身長	体重
身長	1	
体重	0.922	1

基準値は

$$r_0 = \sqrt{\frac{4}{7+2}} = 0.667$$

A	身長	体重	B	身長	体重
身長	1		身長	1	
体重	0.778	1	体重	0.922	1

「0.667」より大きいとき

統計的に「相関関係にある」と言える

シート「相関2」

平均気温とビール消費量の相関係数

を求めてみて

基準値は

$$r_0 = \sqrt{\frac{4}{12 + 2}} = 0.535$$

	平均気温	ビール消費量
平均気温	1	
ビール消費量	0.510	1

「0.535」より小さいから

統計的に「相関関係にない」と言える

ここでちょっと考えてみる

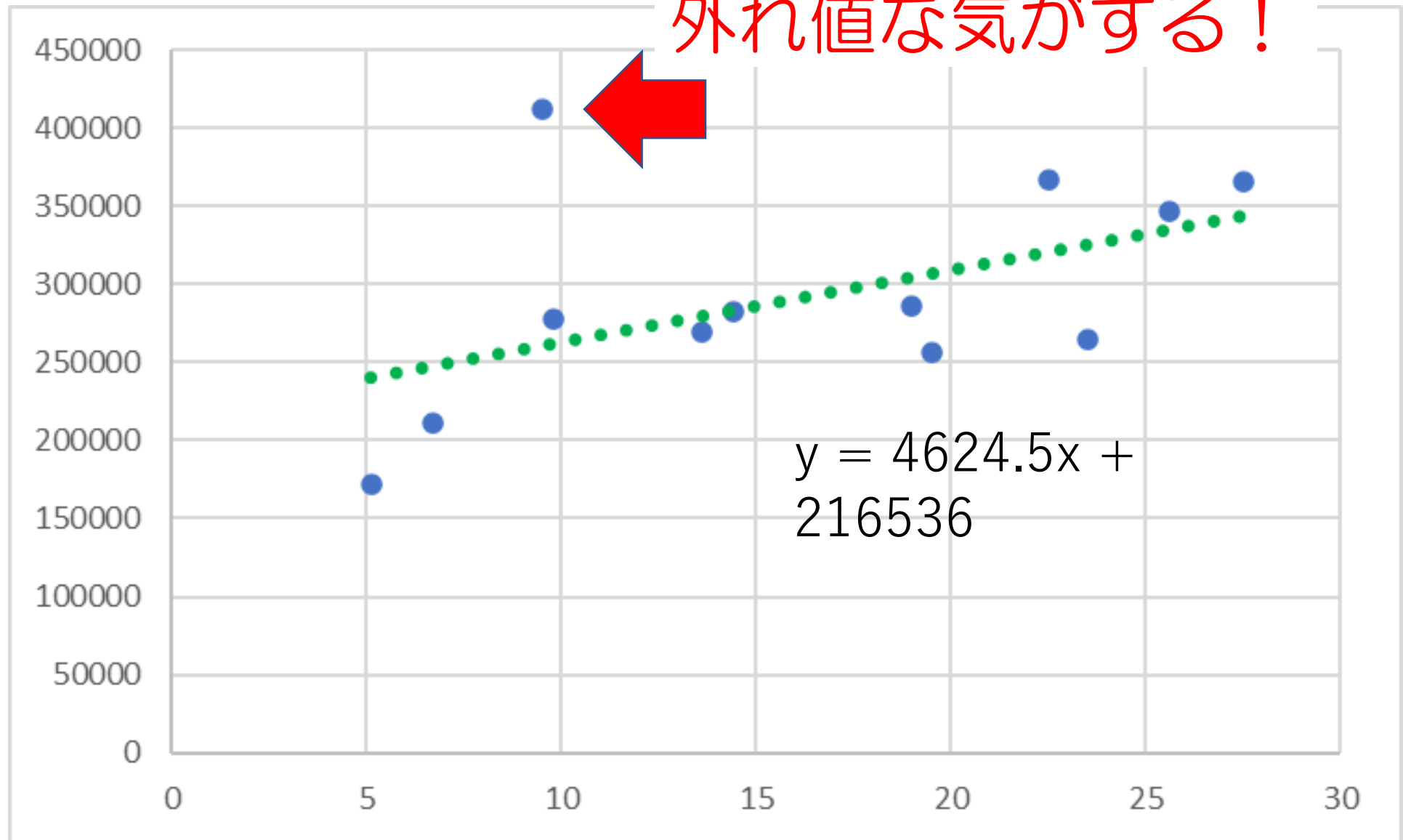
前やった時も「外れ値」があると
相関係数がむちゃくちゃになった。



「散布図」を使って確認してみる

こんなものになるはず

外れ値な気がする！



「外れ値」のデータを抜いて

もう一度相関係数を調べる

基準値は $r_0 = \sqrt{\frac{4}{11 + 2}} = 0.555$

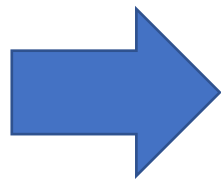
$$r_0 = \sqrt{\frac{4}{12 + 2}} = 0.555$$

	平均気温	ビール消費量
平均気温	1	
ビール消費量	0.829	1

「0.555」より大きいから

統計的に「相関関係にある」と言える

「外れ値」以外を調べると
相関関係にあった



「外れ値」になった理由が
推測できるかも…

「1 2月は忘年会などで気温に関係なく
ビールの消費量が増えたのではないか」
と推測できる！

相関関係の調べ方はわかったけど

いつになったら

「回帰分析」が始まるのか？

2択

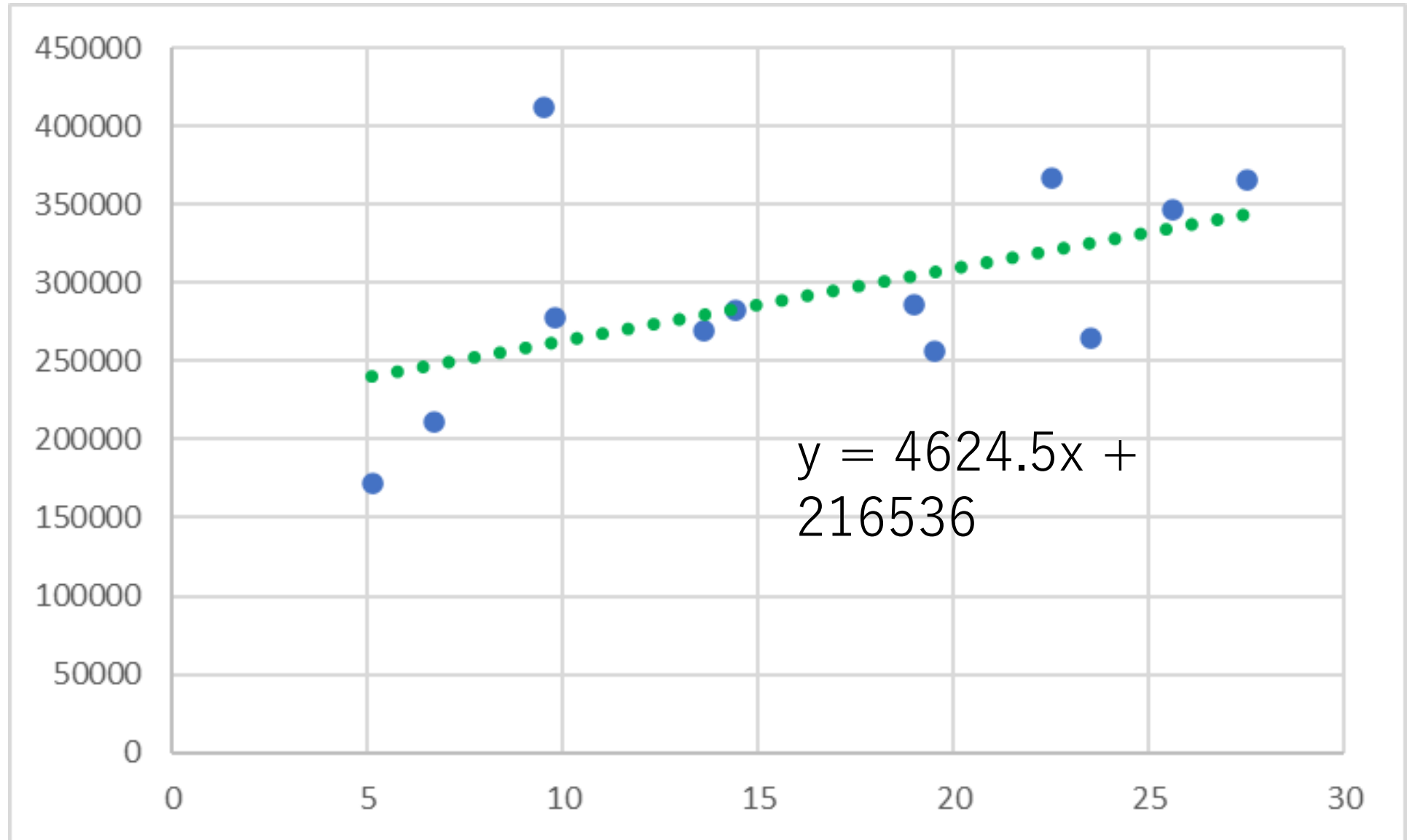
量はどれ
くらい？

どれくら
い程度？

答え複数

回帰分析

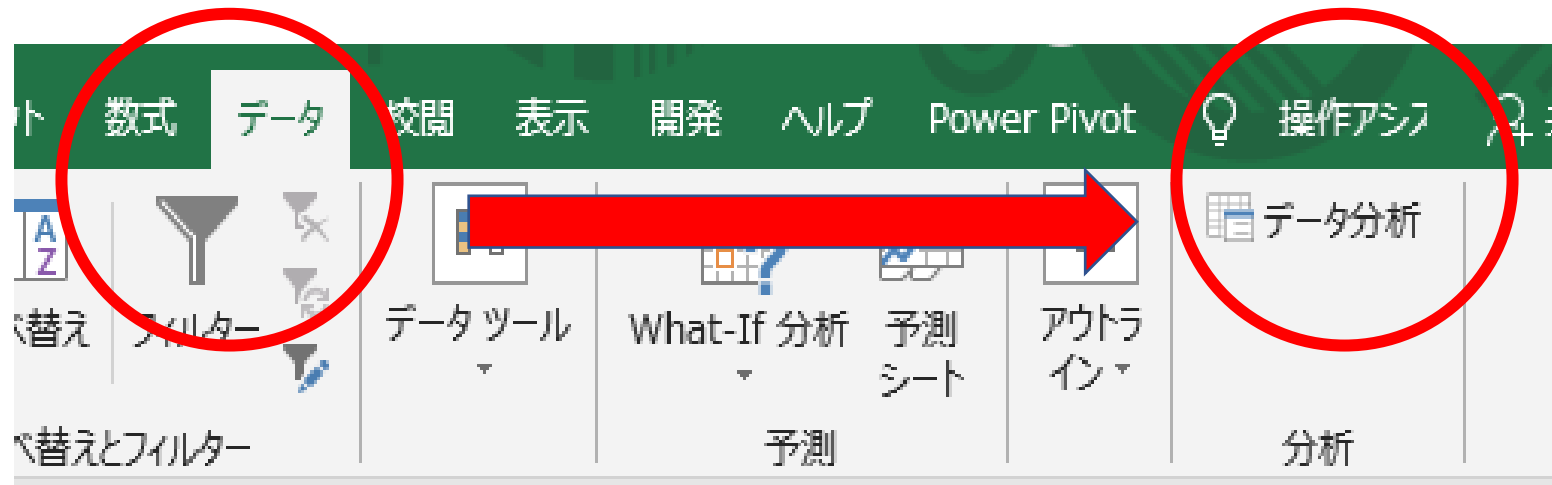
実はもう「**回帰分析**」をやってる！



も一回さっきの

「平均気温とビール消費量」に戻って

「回帰分析」をやってみよう



「回帰分析」

入力Y範囲：ビール消費量(11月まで)

入力X範囲：平均気温 (11月まで)

The image shows a screenshot of a regression analysis software dialog box. It contains several sections with checkboxes and radio buttons. A red oval highlights the top-left section containing the following options:

- ☒ ラベル(L)
- ☒ 有意水準(Q)

To the right of these options is another section with the following options:

- ☐ 定数に 0 を使用(Z)
- 95 %

Below these is the "出力オプション" (Output Options) section, which contains three radio buttons:

- ☒ 一覧の出力先(S):
- ☐ 新規ワークシート(P):
- ☐ 新規ブック(W)

At the bottom is the "残差" (Residuals) section, which contains four checkboxes:

- ☐ 残差(R)
- ☐ 標準化された残差(I)
- ☐ 残差グラフの作成(Q)
- ☒ 観測値グラフの作成(I)

A red rectangle with the handwritten text "ここを チェック" (Check here) is placed over the "観測値グラフの作成(I)" checkbox. A red oval also highlights this checkbox.

こんなものになるはず

回帰統計		
重相関 R	0.829	
重決定 R2	0.688	
補正 R2	0.653	
標準誤差	35696.891	
観測数	11.000	
分散分析表		
	自由度	変動
回帰	1.000	25271620194.279
残差	9.000	11468412424.267
合計	10.000	36740032618.546
	係数	標準誤差
切片	169580.739	27329.437
平均気温	6573.653	1476.117

さっきと同じ

「相関分析と回帰分析の違い」

相関分析

相関係数がメイン。1対1のみ
XとYの関係性を見るだけ
XとYの関係式は出てこない
どっちが原因とかは考えない

回帰分析

1対1でも、1対複数でも！
XからみたYの変化を見る
数式 $Y = aX + b$ を立てて予測
Xが原因でYが結果として進める

「相関分析と回帰分析の違い」

回帰分析

1対1でも、1対複数でも！

XからみたYの変化を見る

数式 $Y = aX + b$ を立てて予測

Xが原因でYが結果として進める

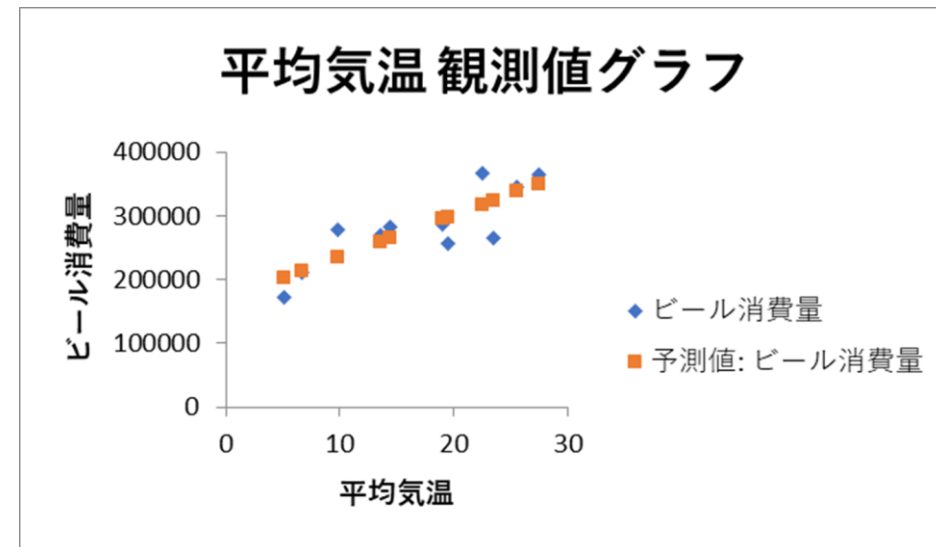
「回帰分析」

新たなデータXからデータY

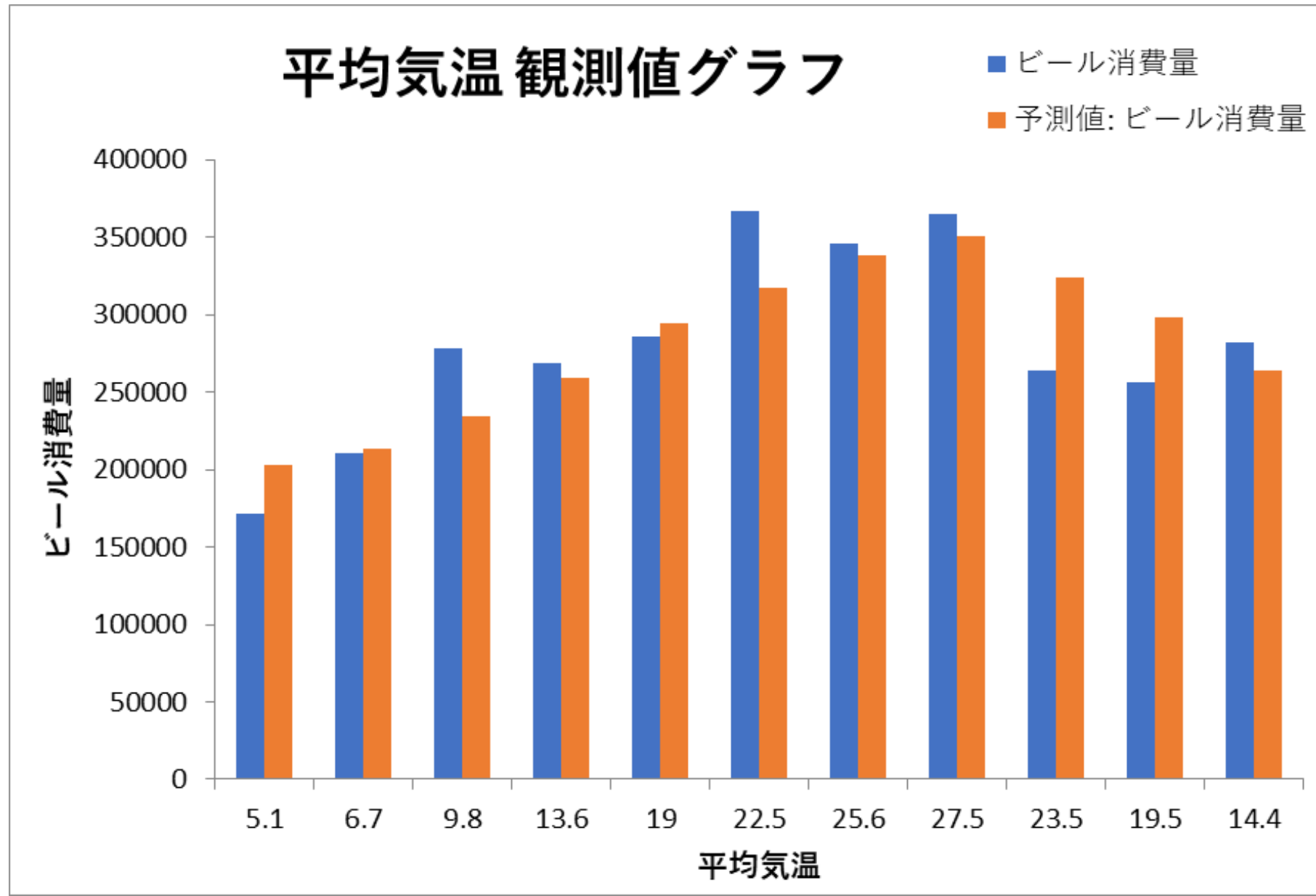
を予測する

$$Y = aX + b$$

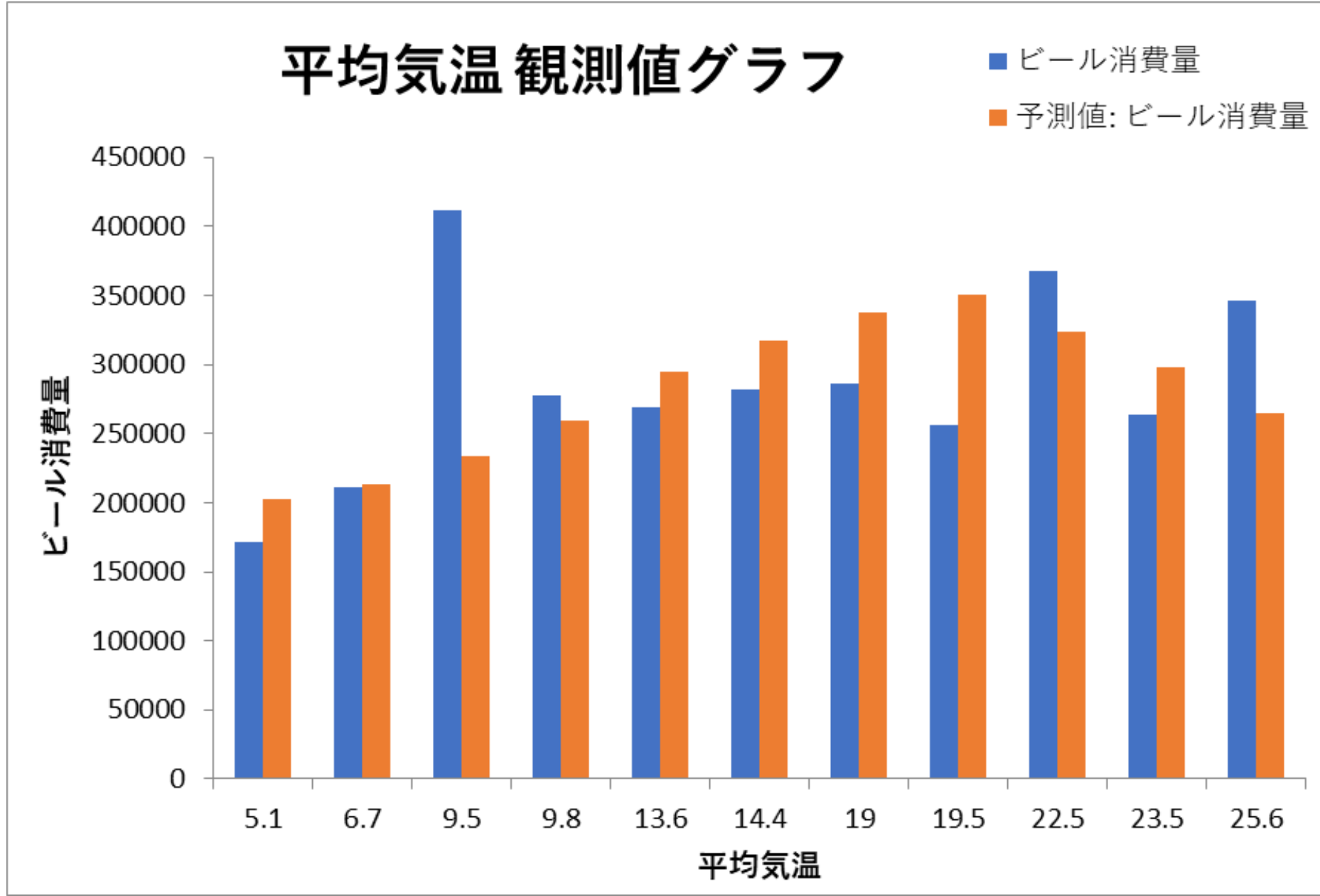
横のほうに出てる
グラフ
を見てみて



グラフを棒グラフに変えると



元データを昇順にすると完成！



「相関分析と回帰分析の違い」

回帰分析

1対1でも、1対複数でも！

XからみたYの変化を見る

数式 $Y = aX + b$ を立てて予測

Xが原因でYが結果として進める

「回帰分析」

新たなデータXからデータY
を予測する

$$Y = aX + b$$

次はこの式の意味を考える！

さっきの表に戻って

次に見るのはココ

回帰統計		
重相関 R	0.829	
重決定 R2	0.688	
補正 R2	0.653	
標準誤差	35696.891	
観測数	11.000	
分散分析表		
	自由度	変動
回帰	1.000	25271620194.279
残差	9.000	11468412424.267
合計	10.000	36740032618.546
	係数	標準誤差
切片	169580.739	27329.437
平均気温	6573.653	1476.117

Yを「ビール消費量」

Xを「平均気温」にしたから

これを代入すると...

$$Y = aX + b$$

$$Y = 6573 \times X + 169580$$

Y : 「ビール消費量」

X : 「平均気温」

予測できること

- 平均気温が1° 上がると

消費量が6573増える

- 平均気温が0° では消費量は169580

「相関分析と回帰分析の違い」

回帰分析

1対1でも、1対複数でも！

XからみたYの変化を見る

数式 $Y = aX + b$ を立てて予測

Xが原因でYが結果として進める

さっきやったのは

「（単）回帰分析」（1対1）

シート「重回帰」やってみて
やることは一緒！

問題設定は

「家賃」と「面積・築年数・時間」
の関係を調べたい

これをメインに調べたい

「データ分析」 ⇒ 「重回帰分析」

入力Y範囲：家賃

入力X範囲：面積・築年数・時間全部

The screenshot shows the 'Data Analysis' dialog box in Excel, specifically the 'Regression' sub-dialog. Two red circles highlight specific options: the top-left group of checkboxes and the bottom-right group of checkboxes.

Options highlighted by red circles:

- ☒ ラベル(L)
- ☒ 有意水準(O) 95 %
- ☐ 残差(R)
- ☒ 観測値グラフの作成(I)

Other visible options and settings:

- ☐ 定数に 0 を使用(Z)
- 出力オプション
 - ☒ 一覧の出力先(S): \$E\$2
 - ☐ 新規ワークシート(P):
 - ☐ 新規ブック(W)
- 残差
 - ☐ 残差(R)
 - ☐ 標準化された残差(I)
 - ☐ 残差グラフの作成(D)

	係数
切片	40900
面積	642.32
築年数	-213.3
駅までの時間	-507.3

これになった？

$$Y = 642X_1 - 213X_2 - 507X_3 + 40899$$

Y : 家賃 X_1 : 面積

X_2 : 築年数 X_3 : 時間

$$Y = 642X_1 - 213X_2 - 507X_3 + 40899$$

Y : 家賃 X_1 : 面積

X_2 : 築年数 X_3 : 時間

家賃は、基本が40899円で

面積が1 増えると642円上がり

築年数が1 増えると213円安くなり

時間が1 増えると507円安くなる

ここから先で大事な話

「パラメトリック」

「ノンパラメトリック」

特徴	パラメトリック	ノンパラメトリック
データの分布	特定の分布を仮定 (通常は正規分布)	特定の分布を仮定しない
分析対象	平均、標準偏差など	順位
代表的な手法	t検定	マン・ホイットニーのU検定 ウィルコクソンの符号順位検定
メリット	小さな差も検出できる	適用範囲が広い

「パラメトリック」

母集団の分布（正規分布とか）

が事前にわかってるとき

この分布に従ってパラメータを

統計的に推測する（ t 検定とか）



「ノンパラメトリック」

母集団の分布がわからないとき

分布がわからないとして推測

(どんな分布でもOK)

- 母集団に正規性がなく、

サンプルサイズが小さいとき

- 極端な外れ値があるけど無視できないとき

- マンホイットニーのU検定、
- ウィルコクソンの符号順位検定

「使い分け」

まず、パラメトリックが無理かを考える



ダメなときはノンパラメトリックな手法を用いる

(パラメトリックな時に、
ノンパラメトリックを使っても
いいけど検出力は不利になる)

これで

「医療統計」 終わり

テストはPCを使ってやります

授業でやった

記述統計・推測統計の中で

「提出」 してないやつを重点的に！

