

統計解析・確率論に関連 するパラドックス

イーピーエス株式会社
森岡 裕

一般に容認される前提から、反駁しがたい推論によって、一般に容認し難い結論を導く論説を逆理（パラドックスまたは逆説）という

— 「逆理」 日本数学会編『岩波 数学辞典 第4版』岩波書店、2007年（ISBN 978-4-00-080309-0）

常識的見解に矛盾するように見える見解、あるいは真理に矛盾するよう見えて、実はそうではない説

— 「パラドックス」 青本和彦（編集）ほか『岩波 数学入門辞典』岩波書店、2005年（ISBN 978-4-00-080209-3）

古代ギリシャのエウブリデスのパラドックスの一つ、古典的なパラドックス
ハゲ頭のパラドックス

「髪の毛が一本もない人はハゲである」(前提1)

「ハゲの人に髪の毛を一本足してもハゲである」(前提2)

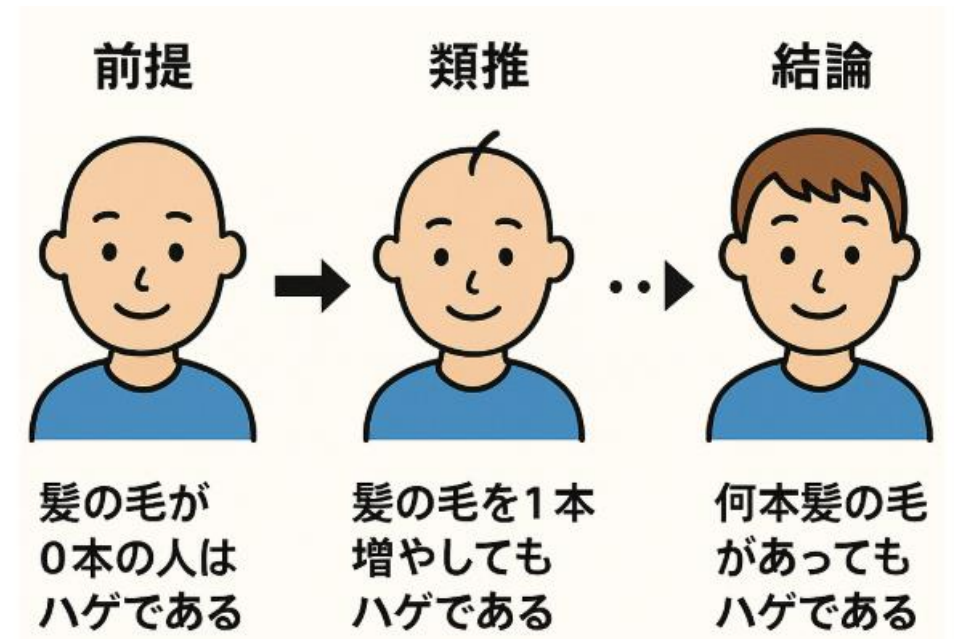
「髪の毛を1本増やしても、ハゲであることに変わりはない。」(類推)

「よって全ての人はハゲである」(結論)

※髪のある人から抜いてくパターンもあり

しかし、全ての人がハゲであるという
結論は一般には受け入れられない

なぜこんなことが起きるのか??



パラドックスの構造

正しそうに見える前提

- ・ 前提の中に正しくないものがある
- ・ 明示されていない前提が作用している

妥当に思える推論

- ・ 推論に誤りがある

受け入れがたい結論

- ・ 実は結論が正しい

問題点はどこか？

✗ 誤りは「類推」にある

漸進的な誤謬の観点で考えると・・・

各ステップ（1本増加）での変化が非常に小さいことから、
変化がないと「**みなしてしまう**」という錯覚が生じさせている
「微小な変化は無視してもよい」という前提を無限に繰り返すと、現実とは乖離した結論に陥る

○ 少しずつの変化は無視できる → ✗ 全体としても変化がない

形式論理で考えると…

形式的には、各ステップの論理（「1本増えても変わらない」）が正しいように思えるが、

「ハゲ」という概念自体が**連続量に対する曖昧な言語的分類**

境界が明確でないため、「変わらない」という主張が無限に続けられてしまう
ここが、**形式的推論と自然言語のギャップ**

シンプソンのパラドックス





VIEWTABLE: Work.Test

	Sex	Treatment	Response	Count
1	Male	No	Yes	4
2	Male	No	No	3
3	Male	Yes	Yes	8
4	Male	Yes	No	5
5	Female	No	Yes	2
6	Female	No	No	3
7	Female	Yes	Yes	12
8	Female	Yes	No	15

疫学調査のデータで，性別，ある治療の有無，奏効有無．

男女別，治療有無別に，奏効・非奏効のデータを集計する

```
* データの入力;
data test;
length Sex Treatment Response
$20.;
input Sex $ Treatment $ Response
$ Count;
datalines;
Male    No    Yes 4
Male    No    No  3
Male    Yes   Yes 8
Male    Yes   No  5
Female  No    Yes 2
Female  No    No  3
Female  Yes   Yes 12
Female  Yes   No  15
;
run;
```

```
title '性別別に見た治療とレスポンスの関係';
proc sql;
    select Sex, Treatment,
           sum(case when Response="Yes" then Count else 0 end) as Response,
           sum(case when Response="No" then Count else 0 end) as NoResponse,
           calculated Response / (calculated Response + calculated NoResponse) as ResponseRate
    format=percent8.2
    from test
    group by Sex, Treatment;
quit;

title '全体で見た治療と生存の関係';
proc sql;
    select Treatment,
           sum(case when Response="Yes" then Count else 0 end) as Response,
           sum(case when Response="No" then Count else 0 end) as NoResponse,
           calculated Response / (calculated Response + calculated NoResponse) as ResponseRate
    format=percent8.2
    from test
    group by Treatment;
quit;
```

性別別に見た治療とレスポンスの関係

Sex	Treatment	Response	NoResponse	ResponseRate
Female	No	2	3	40.00%
Female	Yes	12	15	44.44%
Male	No	4	3	57.14%
Male	Yes	8	5	61.54%

全体で見た治療とレスポンスの関係

Treatment	Response	NoResponse	ResponseRate
No	6	6	50.00%
Yes	20	20	50.00%

女性について

治療ありの方が奏効率が高い(40.00<40.44)

男性について

治療ありの方が奏効率が高い(57.14<61.55)

↓

✕ 故に、治療ありの方が奏効率が高い

全体でみると

治療有無による奏効率の差はない
(50.00 = 50.00)

性別別に見た治療とレスポンスの関係

Sex	Treatment	Response	NoResponse	ResponseRate
Female	No	2	3	40.00%
Female	Yes	12	15	44.44%
Male	No	4	3	57.14%
Male	Yes	8	5	61.54%

全体で見た治療とレスポンスの関係

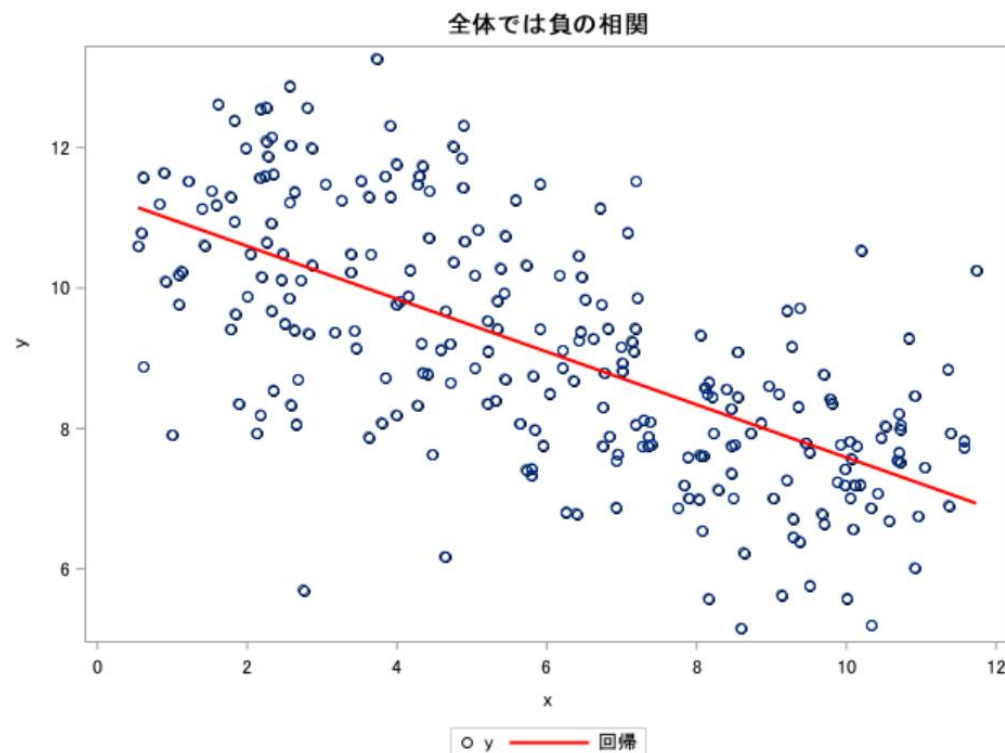
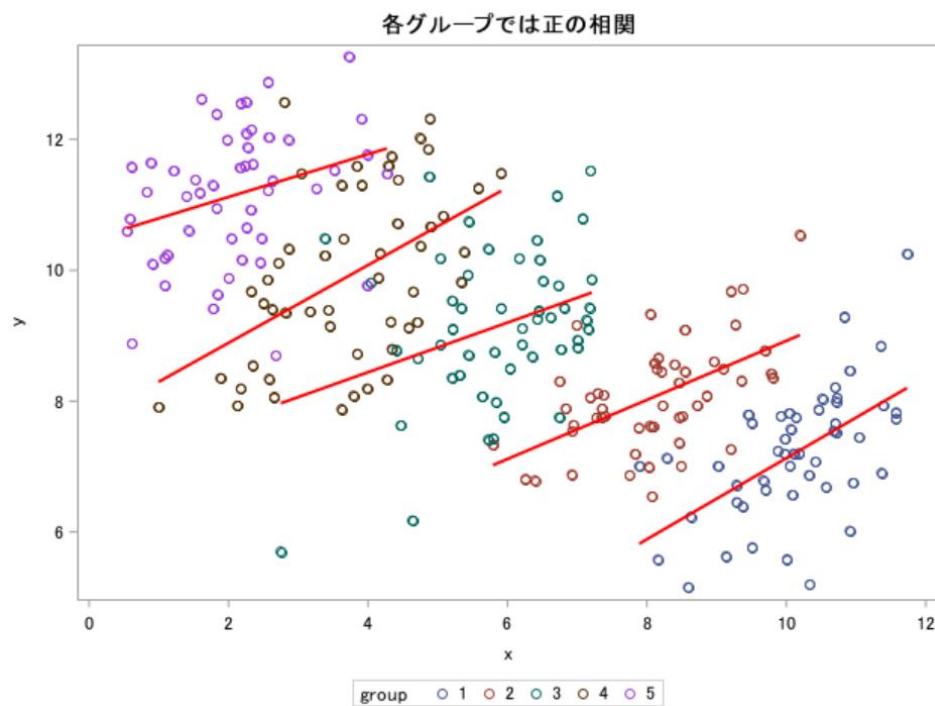
Treatment	Response	NoResponse	ResponseRate
No	6	6	50.00%
Yes	20	20	50.00%

シンプソンのパラドックス (Simpson's paradox)

統計学者エドワード・H・シンプソン, 1951年

母集団での傾向と、母集団を分割した集団での傾向は、異なっている場合があるというパラドックス

つまり集団を分けた場合にある仮説が成立したとしても、集団全体では正反対の仮説が成立することがある。



治療の例では、性別 (C) は治療の有無 (A) と生死 (B) の共通の原因、すなわち交絡因子といえる

$$A \leftarrow C \rightarrow B$$

この場合でいうと、コ克蘭・マンツェル・ヘンツェルのような層別調整の手法が使えるかもしれない

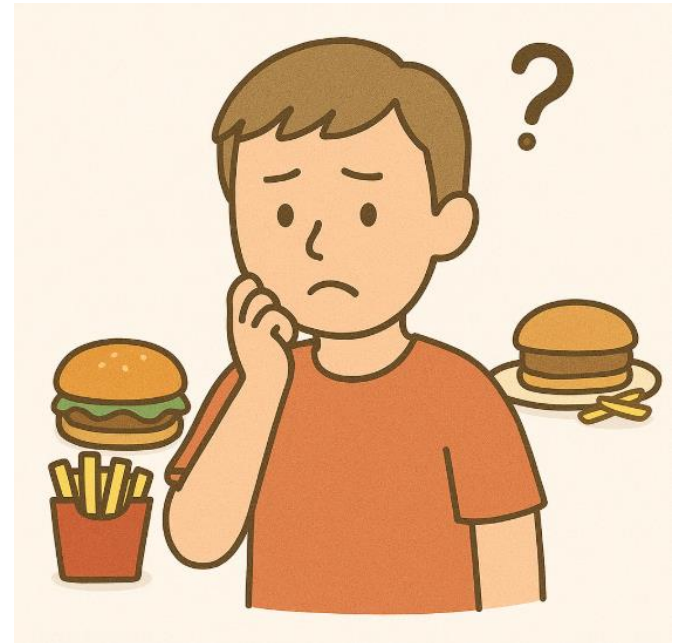
ただ

サブグループ解析がいいのか，層別解析がいいのか，全体で見たほうがいいのかは

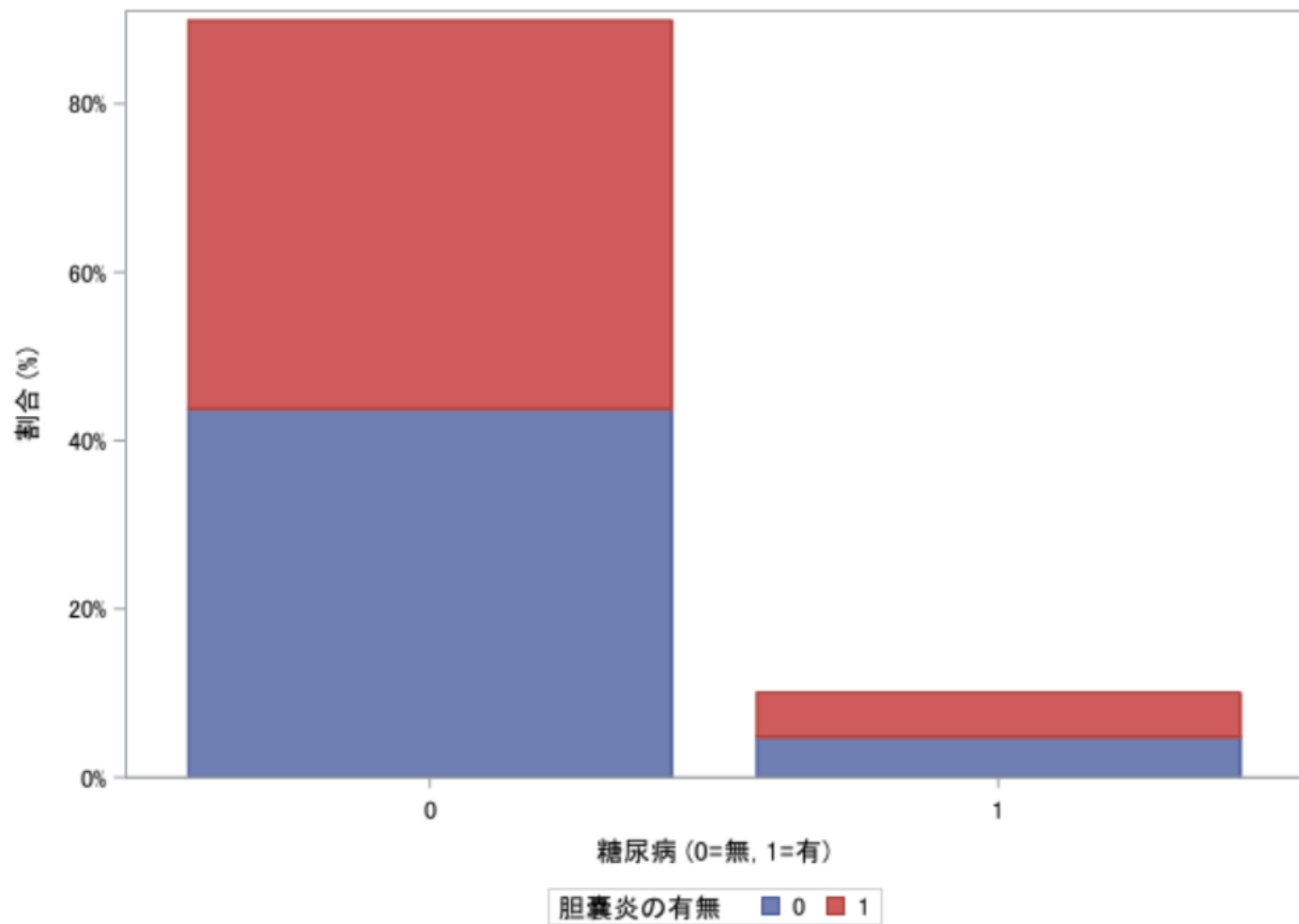
同じデータであっても因果構造によるため，常にこうして解消するというものではない

提唱されたのは70年以上前だが，いまだに大きな影響をもつパラドックス

バークソンのパラドックス



入院患者: 糖尿病と胆嚢炎の割合

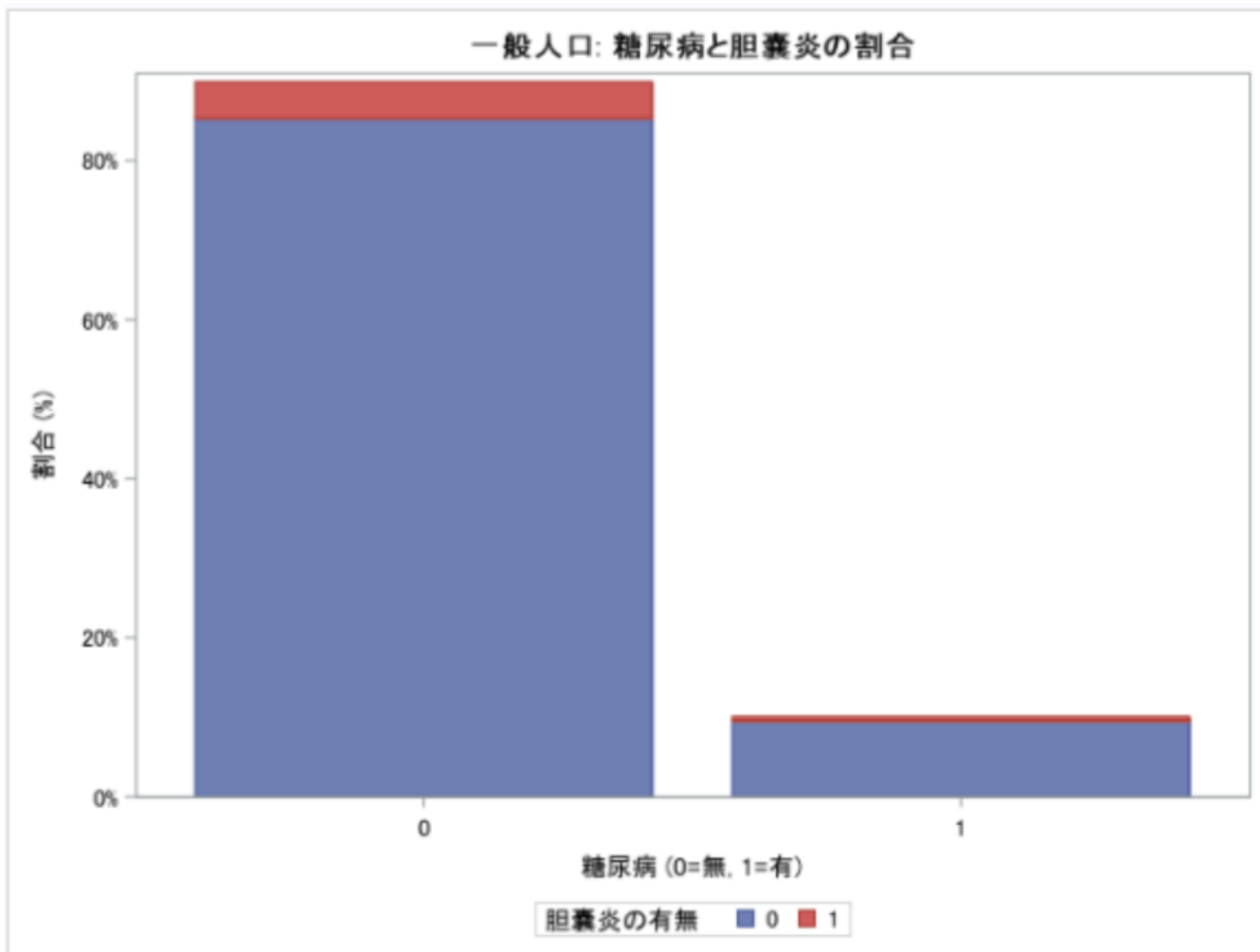


病院の入院患者から統計的に抽出された標本を用いて、ある疾患の危険因子を調査する後方視的研究を想定



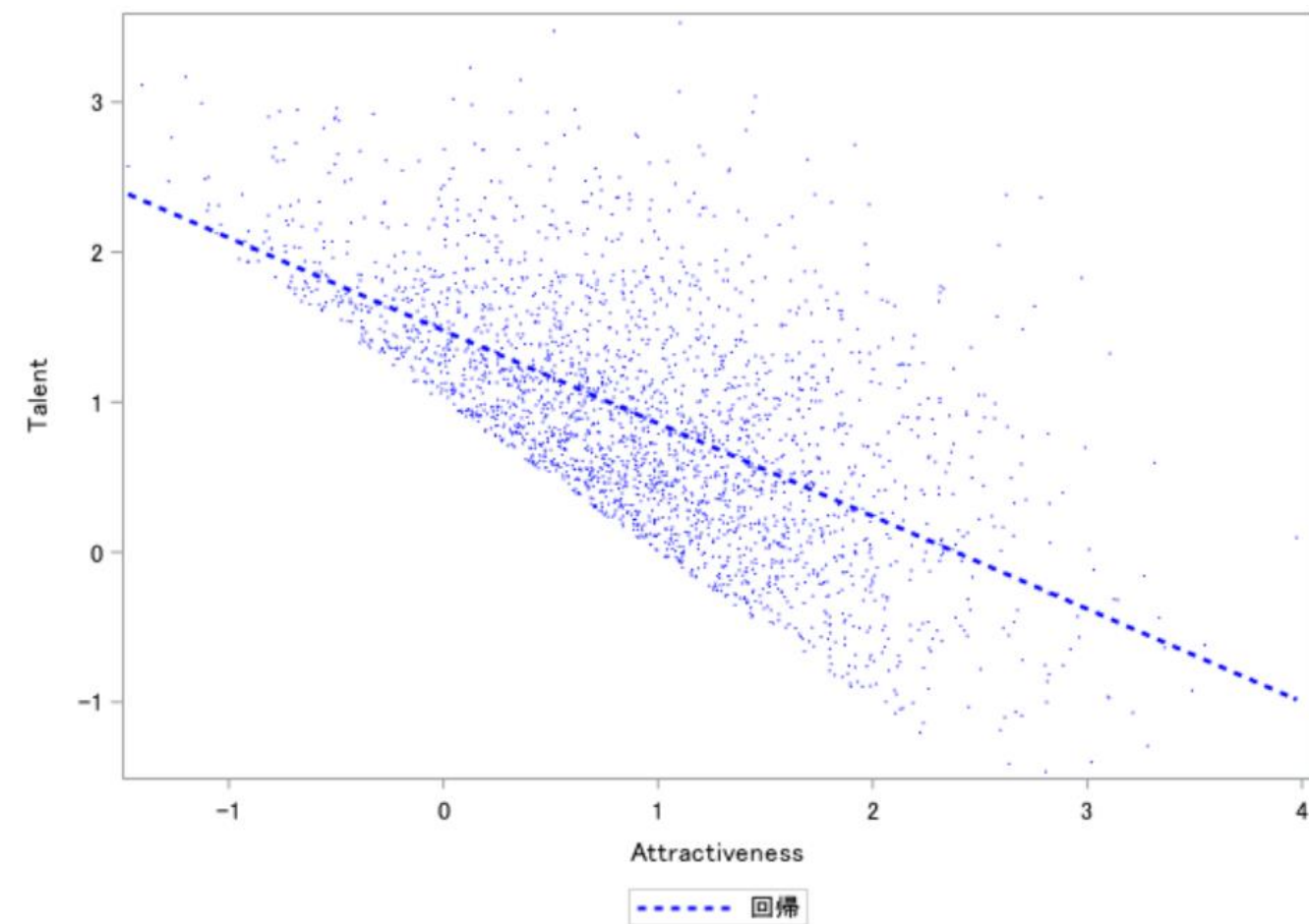
ここに注目して

糖尿病の半分程度が胆嚢炎を罹患しているので、この2つには関係があると考えるのは間違い



糖尿病を持たない入院患者は、糖尿病以外の入院理由（胆嚢炎の原因になるかもしれない）があるから、一般人に比べて胆嚢炎になる可能性が高いと考えられる

この結果は、一般人口において糖尿病と胆嚢炎の間に関連性があるかどうかに関わらず、得られるものである

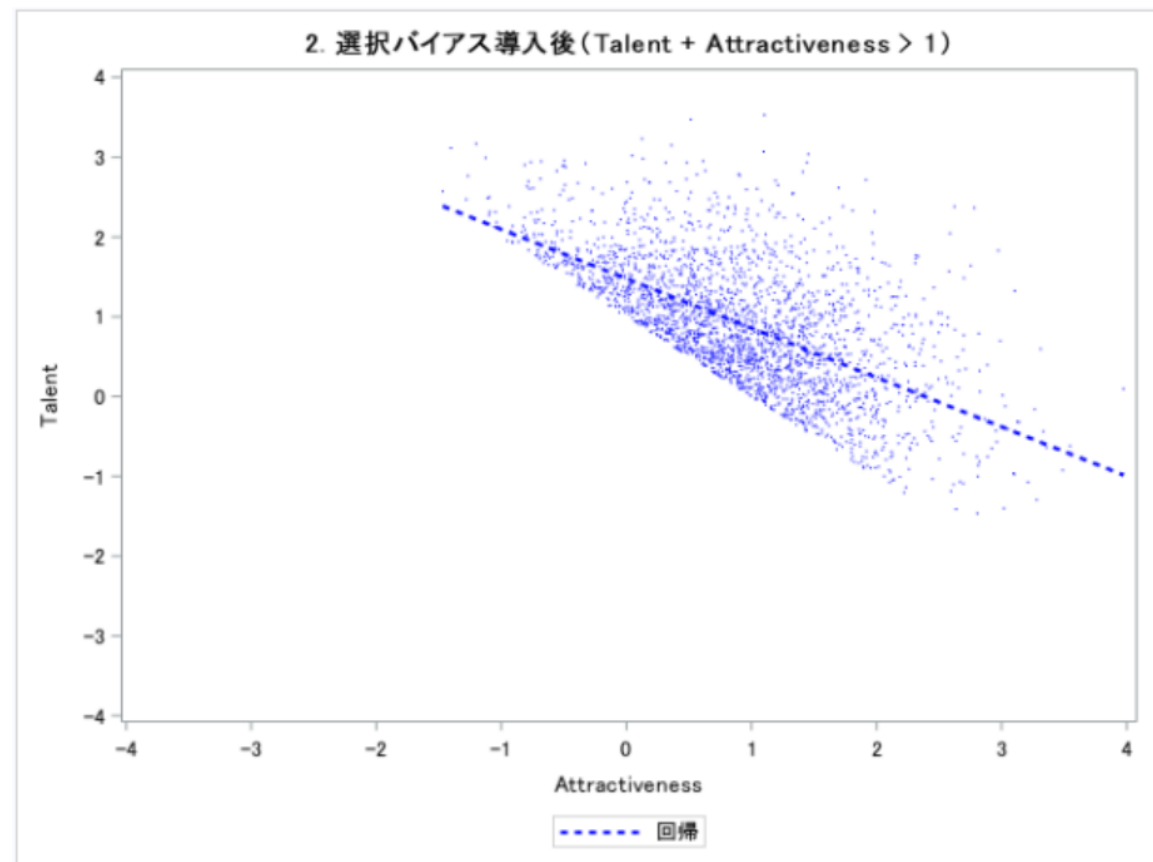
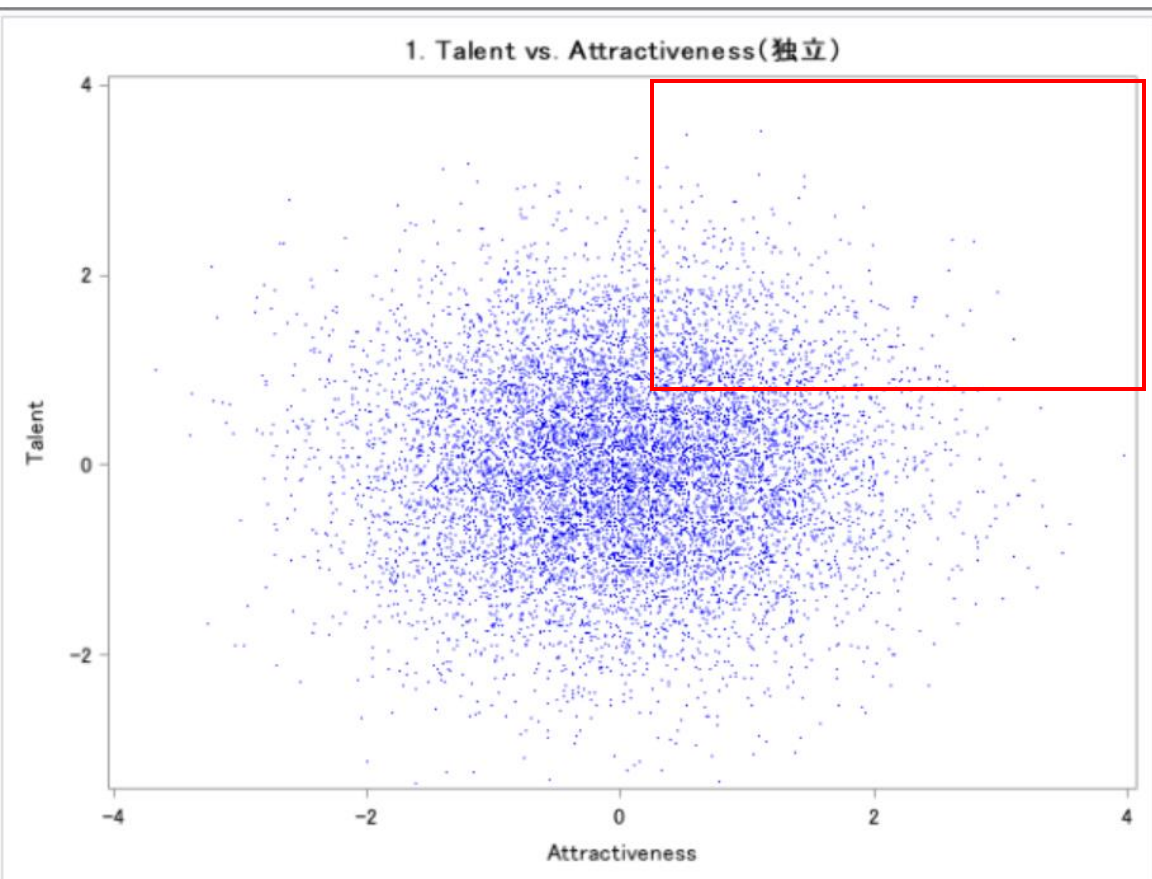


たとえば ある有名人（人気ユーチューバーとか）で

縦軸に才能スコア，横軸に魅カスコア
（どうやって測るかはともかく）

をとると負の相関のように見える

つまり魅力がある人ほど，才能がない



才能か人間的魅力が大きい人が有名になりがちと考えられる

↓

その有名人の集団を切り取ってみると、仮に左図のように 才能と魅力が無相関でも右図のように相関関係が生じているかのように見える

バークソンのパラドックス (Berkson's paradox)

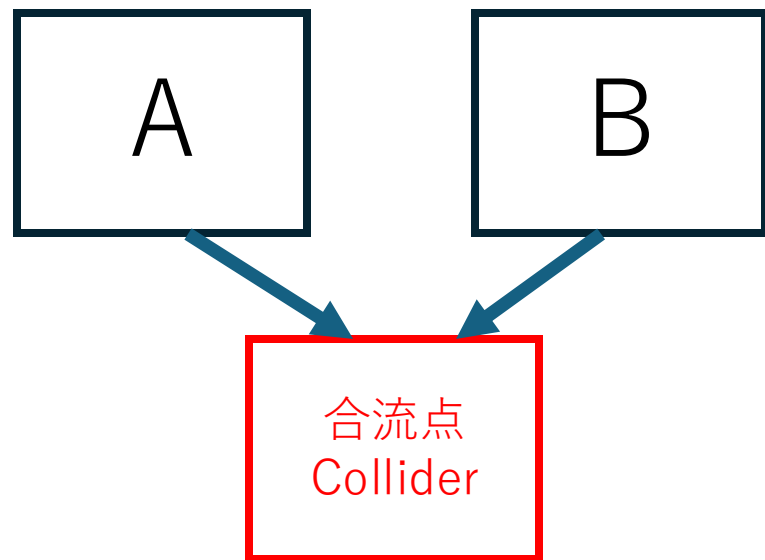
物理学者、医師、統計学者 ジョセフ・バークソン, 1946年

観察研究における選択バイアス, 確認バイアスの問題に起因するパラドックス

事象AまたはBの下で事象Aが起こる条件付き確率は
どちらも起こらない場合を除外しているため、事象A自体の確率よりも高くなる

AまたはBにおける、Bの下でのAの条件付き確率が母集団での条件付き確率と等しいときにパラドックスが生じる

AまたはBにおける、Aの無条件確率は、母集団全体での無条件確率と比べて高いので、AまたはBにおいて、Bが存在することでAの条件付き確率が減少する $\Rightarrow$ 全体の無条件確率に戻る



統計学や因果グラフにおいてある変数が2つ以上の変数から因果的な影響を受けている場合、その変数は合流点といわれる

合流点に影響を与える因果変数同士は、相互に関連しているとは限らない

誕生日のパラドックス



「何人集まれば、その中に誕生日が同一の2人（以上）がいる確率が、50%を超えるか？」



その問いの答えは理論的には **23**人なのですが、それが一般人の直観に反するというパラドックス

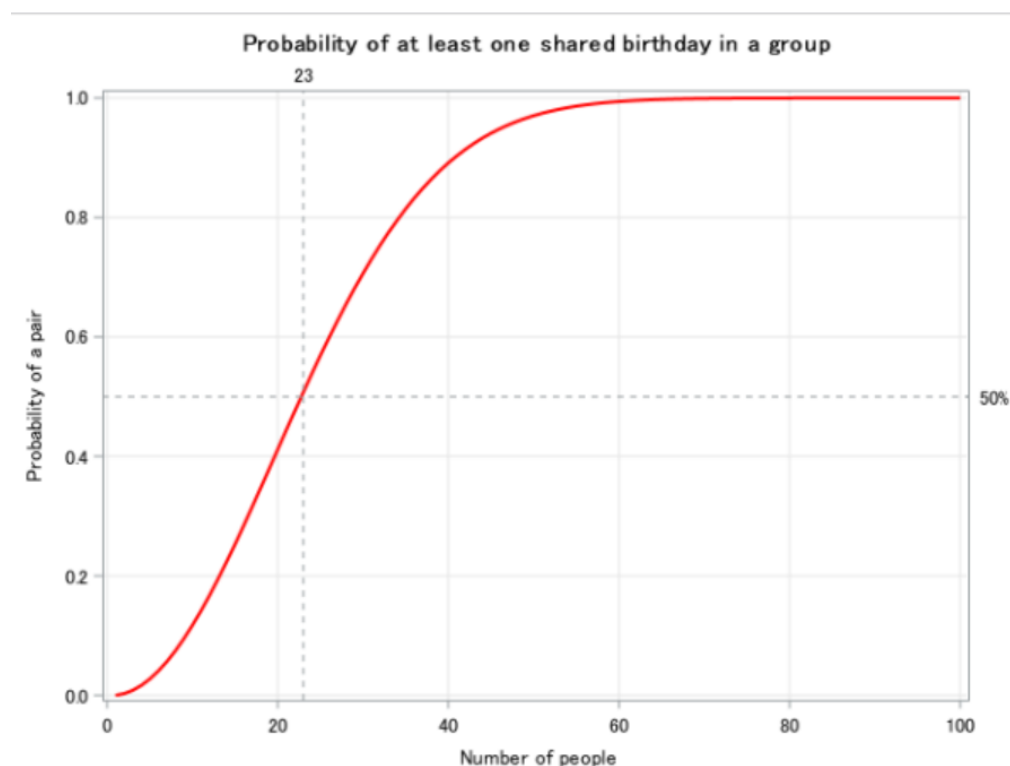
2人目が1人目と異なっている誕生日である確率は、 $364/365$.

次に、3人目が1人目2人目と異なる誕生日である確率は $363/365$ である

n人目は $(365-n+1)/365$ となる。つまり、n人の誕生日が全て異なる確率を考えて、それを1から引いて出せる

```
/* 1~100人の理論値を計算 */
```

```
data birthday_plot;  
  do n = 1 to 100;  
    p_no_match = 1;  
    do k = 1 to n-1;  
      p_no_match = p_no_match * (1 - k / 365);  
    end;  
    p_match = 1 - p_no_match;  
    output;  
  end;  
run;  
/* グラフ描画 */  
proc sgplot data=birthday_plot;  
  series x=n y=p_match / lineattrs=(color=red thickness=2);  
  refline 0.5 / axis=y lineattrs=(pattern=shortdash) label="50%";  
  refline 23 / axis=x lineattrs=(pattern=shortdash) label="23";  
  xaxis label="Number of people" grid;  
  yaxis label="Probability of a pair" grid;  
  title "Probability of at least one shared birthday in a group";  
run;
```



```
%let group_size = 23;      /* グループの人数 */
%let n_sim = 10000;        /* シミュレーションの回数 */
```

```
data birthday_sim;
  call streaminit(123); /* 乱数の初期化（再現性のため） */
  do sim = 1 to &n_sim;
    array bday[&group_size];
    do i = 1 to &group_size;
      bday[i] = ceil(365 * rand("uniform")); /* 1～365の誕生日 */
    end;
  end;
```

```
/* 重複チェック */
```

```
same_birthday = 0;
do i = 1 to &group_size-1 until (same_birthday = 1);
  do j = i+1 to &group_size;
    if bday[i] = bday[j] then same_birthday = 1;
  end;
end;
```

```
output;
```

```
end;
```

```
run;
```

```
proc means data=birthday_sim mean;
```

```
var same_birthday;
```

```
title "誕生日が一致する確率 (&group_size 人中、&n_sim 回のシミュレーション)";
```

```
run;
```

モンテカルロ・シミュレーションで、独立した誕生日（双子とかじゃい）をもつ23人を集めて おんなじ誕生日がいるかを確認するって儀式を1万回繰り返して、いた回数を1万でわって、確率をだしてみる。理論的には約50.7%になるはずだが、実際どうだ？

50.8%となり、ほぼ理論値どおりの結果

誕生日が一致する確率(23 人、10000 回のシミュレーション)

MEANS プロシジャ

分析変数 : same_birthday	
平均	
0.5082000	

なぜ、誕生日のパラドックスが人間にとって、直観に反して パラドックス足りえるのか？？
人間以外にきいてみよう

ChatGPT 4o ▾

↑ 共有する



誕生日のパラドックスが 一般的に多くの人の直観に反するのは
何故ですか？？

とても良い質問です！

誕生日のパラドックス（Birthday Paradox）が**多くの**人にとって**直感に反する理由**には、いくつかの心理的・数学的な要因があります。以下に詳しく説明します。

1. 「比較の数」が直感で過小評価される

多く人は、23人の中に自分と同じ誕生日の人がいる確率を考えがちですが、誕生日パラドックスで考慮されているのは：

「任意の2人が同じ誕生日を持つペア」が1組でもいる確率

です。これは大きな違いです。

- 23人がいれば、全部で

$$\binom{23}{2} = \frac{23 \times 22}{2} = 253$$

組のペアがあります。

- つまり「253回」比較されるチャンスがあるので、思ったより一致が起きやすいのです。

2. 人は確率を「線形的」に考える傾向がある

人の直感では：

- 2人 → めったに一致しない
- 10人 → まだ珍しいだろう
- 20人 → 半分もいないし、まさか...

と、人数が増えるにつれて確率もゆっくり増えると思いがちですが、実際には **指数関数的に**増加します。

少人数でも爆発的に一致の確率が上がる
→ この非線形性が直感に反する

☒ 3. 人は「特定の誕生日」と「どれでもいい誕生日」を混同する

- 「私の誕生日と同じ人がいる確率」は、365分の1。
- でも「誰か2人が一致する確率」は全然別の話。

後者はどの誕生日でもいいので、一致すれば成立します。この“条件の緩さ”も直感では見落とされがちです。

🎲 4. 「希少イベント」に対する認識が弱い

たとえば、宝くじが当たる確率が小さいのに当たる人がいるのと同じで、

- 個別の一致確率は低くても
- 試行回数（ペア数）が多ければ起こる

という現象は、日常的に直感とズレて感じられます。

結論：人間の直感は「組合せの多さ」と「非線形な確率増加」に弱い

この2つの盲点により、「たった23人でそんなに高確率で一致するなんて...」と驚かされるのです。それがこのパラドックスの魅力でもあります。

モンティホール パラドックス



あるテレビ番組で、あなたはクイズに勝ち、3つのドアのうち1つを選ぶ権利を得ました。
3つのドアのうち、1つには「**車（当たり）**」、他の2つには「**ヤギ（ハズレ）**」が隠されています。





あなたが一つのドアを選択します

あなたが選ばなかった2つのドアのうち、1つを司会者があけて、そこにヤギ（はずれ）がいることを見せます。そして



ドアを選びなおしてもいいですよ。

この時、選択を変えない場合と変えた場合、どちらが得か？

正解は『ドアを変更する』

なぜなら、ドアを変更した場合には景品を当てる確率が2倍になるから



詳しくはwikipediaでモンティホール問題でみてもらえればですが、これはアメリカで大論争を巻き起こしました。

確率論から導かれる結果を説明されても、なお納得しない者が少なくないことから、**モンティ・ホール・ジレンマ**、**モンティ・ホール・パラドックス**とも称される。

「直感で正しいと思える解答と、論理的に正しい解答が異なる問題」

つまり、実は結論は正しいパターンのパラドックス

1. 最初の選択で車を当てている確率は 1/3

なぜなら、プレイヤーはランダムにドアを選んでいるため：

- ドア①に車がある確率 → 1/3
- ドア②に車がある確率 → 1/3
- ドア③に車がある確率 → 1/3

2. 最初の選択でヤギを選んでいる確率は 2/3

- ドア①がヤギ（失敗）である確率 → 2/3

このとき、モンティは必ず「もう1つのヤギのドアを開ける」ので、
残ったドア（＝選ばなかった＆開けられなかったドア）には車がある確率が2/3となります

隠されている車の位置	プレイヤーの最初の選択	モンティが開けるドア	切り替えた場合の結果
ドア①（車）	ドア①（正解）	ドア② or ③（ヤギ）	ハズレ（ヤギ）
ドア②（車）	ドア①（不正解）	ドア③（ヤギ）	当たり（車）
ドア③（車）	ドア①（不正解）	ドア②（ヤギ）	当たり（車）

- このように、「最初に正解していた場合（1/3）」だけが、選び直すと損。
- それ以外の「最初に間違っていた場合（2/3）」は、**選び直すと当たる**。

👉 **結論：選び直すことで当たる確率は 2/3。**

最初の選択が正解である確率： $1/3 \rightarrow$ 変更すると損
最初の選択が不正解である確率： $2/3 \rightarrow$ 変更すると得
よって、選択を変更するほうが2倍の確率で当たる



モンティが必ずヤギのドアを開けるというルールがあるため、
プレイヤーの行動に「バイアスされた情報」が加わる
これがただの運任せ（モンティが無作為に開ける）とは全く違う点であり、
「条件付き確率」が成立する決定的な理由となる

大数学者ポール・エルデシュはこの説明をうけても納得しなかった

弟子のアンドリュー・ヴァージョニが、本問題を自前のパーソナルコンピュータでモンテカルロ法を用いて数百回のシミュレーションを行って、変えたほうが得になることを証明し、エルディッシュが納得したというエピソードがあるので

アンドリュー・ヴァージョニに倣ってシミュレーションで考えてみる

当時はコンピューター性能的に数百回だったそうですが、私たちは贅沢に100万回まわしてみましようか

```

/* シミュレーション回数 */
%let num_simulations = 1000000;
data monty_hall_simulation;
  do simulation_id = 1 to &num_simulations;
    /* 1. 車のドアをランダムに決定 (1, 2, 3) */
    car_door = ceil(3 * ranuni(0));

    /* 2. 参加者が最初に選ぶドアをランダムに決定 */
    initial_choice = ceil(3 * ranuni(0));

    /* 3. モンティがヤギのドアを開ける */
    /* 参加者の選択と新車のドア以外のドアを選ぶ */
    do monty_opens = 1 to 3 until (monty_opens ne initial_choice and monty_opens ne car_door);
    end;

    /* 4. ドアを変更する場合の選択 */
    /* 最初の選択とモンティが開けたドア以外のドアを選ぶ */
    do switch_choice = 1 to 3 until (switch_choice ne initial_choice and switch_choice ne monty_opens);
    end;
  output;
end;
keep initial_choice switch_choice car_door;
run;

```

```

proc sql ;
  title "ドアを変更した場合の的中率";
  select mean(switch_choice = car_door)
  into :win_rate_switch
  from monty_hall_simulation;

  title "ドアを変更しない場合の的中率";
  select mean(initial_choice = car_door)
  into :win_rate_stay
  from monty_hall_simulation;
quit;

```

ドアを変更した場合の的中率

0.66735

ドアを変更しない場合の的中率

0.33265

ドアを変更するという行為は、最初の「2/3の確率でハズレを選んでいた」という事象に賭け直すことになるので、変えたほうが得ということが実際、プログラムにしてみても、私は初めて理解できた。プログラムでのパラドックスの検証は私はとても有用だと思います

ロードのパラドックス



たとえば，ダイエット療法における体重変化を男女で比較したい場合を考える
ダイエット開始前の体重と後の体重と性別があるとして

```
data weight_data;  
length sex $20.;  
call streaminit(67890);  
  
/* データ生成 */  
/* 男子グループ */  
do sex = "Male";  
do i = 1 to 100;  
/* PreとPostは相関0.8の多変量正規分布に従う */  
z1 = rand('NORMAL', 0, 1);  
z2 = rand('NORMAL', 0, 1);  
pre_weight = 70 + 5 * z1; /* Pre: 平均70kg、SD=5 */  
post_weight = 72 + 5 * (0.8 * z1 + sqrt(1 - 0.8**2) * z2); /* Post:  
平均72kg、相関0.8 */  
diff_weight = post_weight - pre_weight;  
output;  
end;  
end;  
  
/* 女子グループ */  
do sex = "Female";  
do i = 1 to 100;  
z1 = rand('NORMAL', 0, 1);  
z2 = rand('NORMAL', 0, 1);  
pre_weight = 60 + 5 * z1; /* Pre: 平均60kg、SD=5 */  
post_weight = 62 + 5 * (0.8 * z1 + sqrt(1 - 0.8**2) * z2); /* Post:  
平均62kg、相関0.8 */  
diff_weight = post_weight - pre_weight;  
output;  
end;  
end;  
  
drop z1 z2 i;  
run;
```

VIEWTABLE: Work.Weight_data				
	sex	pre_weight	post_weight	diff_weight
1	Male	66.282817254	68.895898334	2.61308108
2	Male	71.539042141	69.486584605	-2.052457536
3	Male	75.263995845	72.27089341	-2.993102435
4	Male	78.296556591	76.701380054	-1.595176537
5	Male	72.488166599	73.982077541	1.4939109413
6	Male	70.958265535	72.053936986	1.0956714511
7	Male	66.700757509	66.855978706	0.1552211974
8	Male	66.596821253	70.761912028	4.1650907728
9	Male	72.847174226	78.798112195	5.9509379696
10	Male	70.088464431	68.499433794	-1.589030638
11	Male	77.698769422	77.342872205	-0.355897217
12	Male	77.361840524	78.0744381	0.7125975759
13	Male	67.366496507	68.777560981	1.4110644736
14	Male	69.923458532	73.672037819	3.7485792875
15	Male	73.939132775	71.742952566	-2.196180209
16	Male	72.141802723	78.669125499	6.527322776
17	Male	69.98821047	72.131772261	2.1435617913
18	Male	87.858243022	91.635961114	3.7777180919
19	Male	66.844128763	68.288986282	1.444857519
20	Male	71.022512414	75.359594576	4.3370821617

2つの分析法を行う

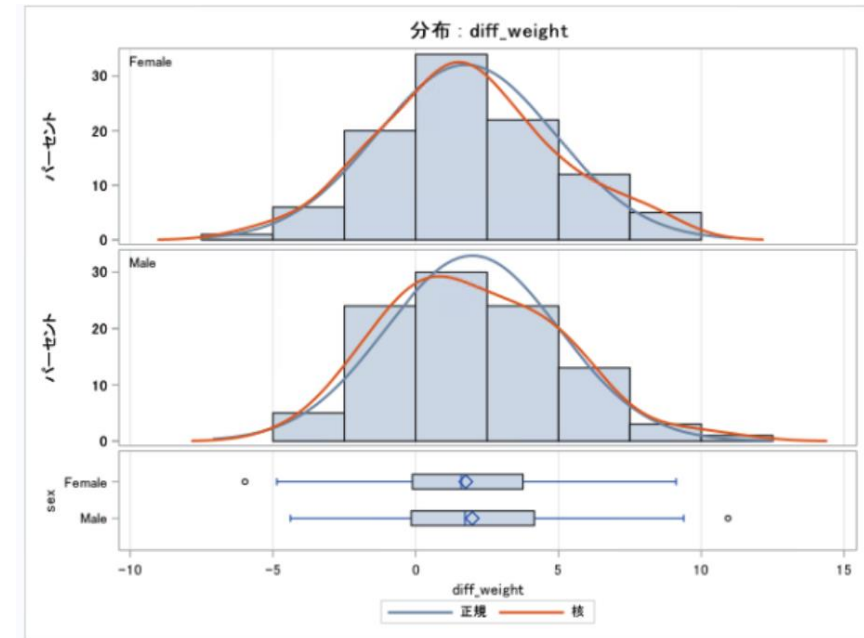
```
/* 分析1：差の分析（体重変化の比較） */  
proc ttest data=weight_data;  
    class sex;  
    var diff_weight;  
    title '差の分析：体重変化（Post - Pre）のグループ間比較';  
run;
```

```
/* 分析2：共分散分析（ANCOVA） */  
proc glm data=weight_data;  
    class sex;  
    model post_weight = sex pre_weight / solution;  
    lsmeans sex / adjust=tukey;  
    title '共分散分析：Preを共変量としたPostのグループ間比較';  
run;  
quit;
```

t検定

sex	手法	平均	平均の 95% 信頼限界		標準偏差	標準偏差の 95% 信頼限界	
Female		1.7568	1.1402	2.3734	3.1075	2.7284	3.6099
Male		1.9816	1.3806	2.5826	3.0290	2.6594	3.5187
Diff (1-2)	Pooled	-0.2248	-1.0805	0.6310	3.0685	2.7937	3.4036
Diff (1-2)	Satterthwaite	-0.2248	-1.0805	0.6310			

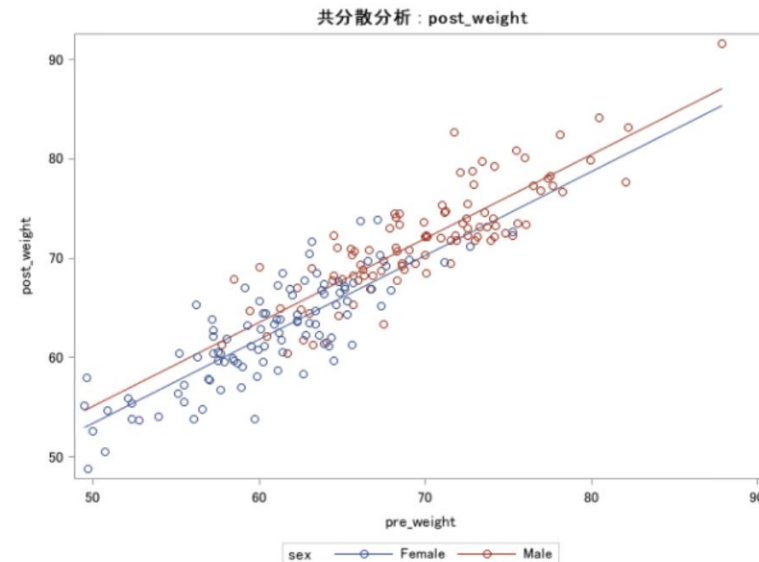
手法	分散	自由度	t 値	Pr > t
Pooled	Equal	198	-0.52	0.6051
Satterthwaite	Unequal	197.87	-0.52	0.6051



共分散分析

要因	自由度	Type III 平方和	平均平方	F 値	Pr > F
sex	1	78.754914	78.754914	8.98	0.0031
pre_weight	1	4044.783812	4044.783812	461.04	<.0001

パラメータ	推定値	標準誤差	t 値	Pr > t
Intercept	12.82130742	B 2.76949365	4.63	<.0001
sex Female	-1.66768321	B 0.55661126	-3.00	0.0031
sex Male	0.00000000	B .	.	.
pre_weight	0.84506914	0.03935699	21.47	<.0001



ロードのパラドックス (Lord's paradox)

統計学者フレデリック・マザー・ロード, 1967年

統計学者が既存の差異を調整するかどうかによって異なる結論に達する可能性がある例をパラドックスとして示した

先の例としては 単純な変化量分析であれば有意にならず, ベースラインを調整した共分散分析では有意になった

変化量分析 は個々の参加者の変化そのものに注目していて

共変量分析 は初期値が同じであったとした場合に、最終値に差があるかどうか注目している

これも、昨今の因果推論の枠組みの中でいまだ ホットな議論になっている。
一番最初に説明したシンプソンのパラドックスの連続量版に近いという見方もある

平均への回帰

あんまりパラドックスとはいわない，ただの現象だけど
ロードのパラドックスに絡んでるので一応

1000人の参加者について、ある評価をして、スコアをとります（その評価の母平均は100で、標準偏差は15だとします）。人間なので、測定ごとに、その人の本来のスコアにノイズが入るとします

得られたスコアのうち、1回目のスコアで、上位5%の高得点集団と下位5%の低得点集団を定義してみます

```
%let n = 1000;
/* 真のスコアとノイズを作成し、1回目と2回目の観測スコアを
生成 */
data regression_to_mean;
  call streaminit(123);
  do id = 1 to &n;

    /* 真のスコア（平均100, SD=15） */
    true_score = rand("normal", 100, 15);

    /* ノイズを加えた1回目と2回目の観測スコア */
    noise1 = rand("normal", 0, 10);
    observed1 = true_score + noise1;

    noise2 = rand("normal", 0, 10);
    observed2 = true_score + noise2;

    output;
  end;
run;

/* 観測値のパーセンタイル（5%と95%）を計算 */
proc univariate data=regression_to_mean noprint;
  var observed1;
  output out=percentiles pctlpre=P_ pctlpts=5
95;
run;

data regression_to_mean;
  if _n_ = 1 then set percentiles;
  set regression_to_mean;

  if observed1 >= P_95 then high = 1;
  else high = 0;

  if observed1 <= P_5 then low = 1;
  else low = 0;
run;
```

観測スコア1回目 vs 2回目: 平均への回帰の視覚化 1回目高得点集団

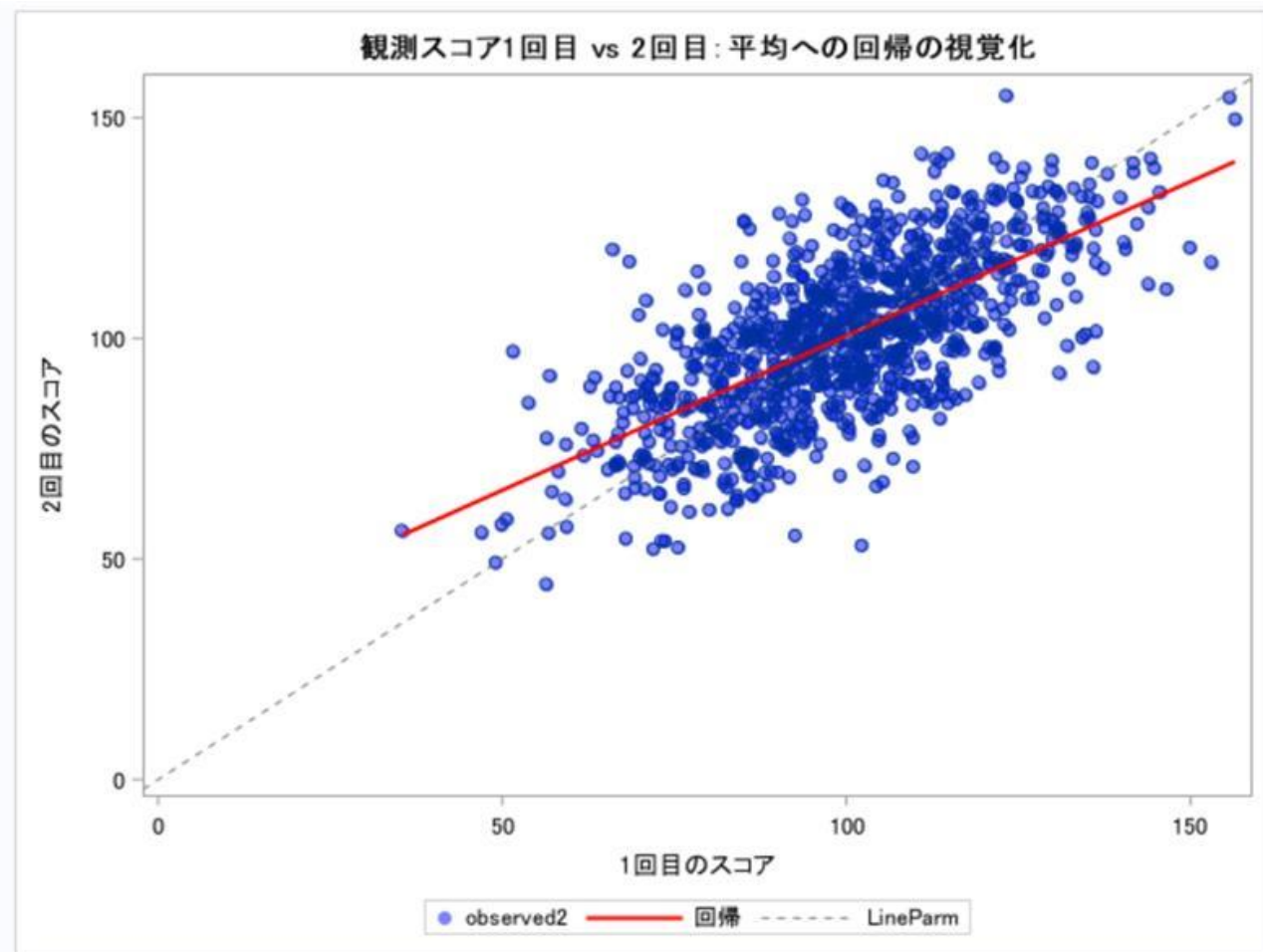
MEANS プロシジャ

high	Obs 数	変数	平均
1	50	observed1	137.53
		observed2	123.58

観測スコア1回目 vs 2回目: 平均への回帰の視覚化 1回目低得点集団

MEANS プロシジャ

low	Obs 数	変数	平均
1	50	observed1	62.71
		observed2	76.74

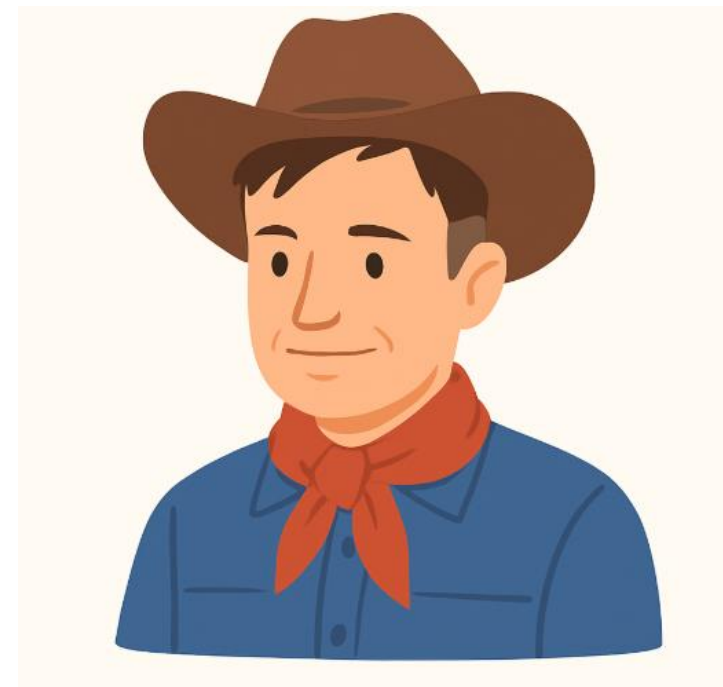


1回目のスコアの高得点集団の2回目のテストは前より低くなり

1回目のスコアの低得点集団の2回目のテストは前より高くなる

グラフの回帰直線は、完全な $y=x$ の直線より傾きが小さく、これも「**極端な値が平均に戻る**」傾向を示している

ウィル・ロジャース現象



もしオクラホマ州の出稼ぎ労働者がカリフォルニア州に移動したら、両方の州の知的レベルが上がるだろう

アメリカ合衆国のコメディアンであるウィル・ロジャースが、1930年代の世界恐慌の際に言ったコメントが由来。

現代でこれ言ったら一発NGレベルな問題発言っぽいが、

つまり
ある集合の中の数を別の集合に移した結果、両方の平均が高く（低く）なる現象のことである
この現象は、次の2つの条件がともに成り立つ時に起こる

移動する要素は移動前の集合の平均よりも小さい。
移動する要素は移動先の集合の平均よりも大きい。



VIEWTABLE: Work.Initial		
	Group	Score
1	A	40
2	A	45
3	A	50
4	A	55
5	A	60
6	B	65
7	B	70
8	B	75
9	B	80
10	B	85

Aにおけるスコア60は
Aの平均より高く
Bの平均より小さい

これがAからBに動くと
両群の平均が下がる

```
/* 平均の確認（初期状態） */
```

```
proc means data=initial mean;
```

```
class Group;
```

```
var Score;
```

```
title "Initial Group Means";
```

```
run; /* Step 2: スコアが 60 の人を Group A から B に移動 */
```

```
data moved;
```

```
set initial;
```

```
if Group="A" and Score=60 then Group="B";
```

```
run;
```

```
/* 平均の確認（移動後） */
```

```
proc means data=moved mean;
```

```
class Group;
```

```
var Score;
```

```
title "Group Means After Moving Score=60 from A to B";
```

```
run;
```

Initial Group Means

MEANS プロシジャ

分析変数 : Score		
Group	Obs 数	平均
A	5	50.0000000
B	5	75.0000000

Group Means After Moving Score=60 from A to B

MEANS プロシジャ

分析変数 : Score		
Group	Obs 数	平均
A	4	47.5000000
B	6	72.5000000

ウィル・ロジャース現象は 実際の医療データで起きている

グループ Healthy（診断前・健康）と Cancer（診断済み）に分類。
診断技術（例：FDG-PETスキャン）の向上により、「境界的な人（移行期）」が
Healthy → Cancer に再分類される。結果的に、両グループの平均寿命が向上して見える。

	Group	Lifespan
1	Healthy	80
2	Healthy	82
3	Healthy	78
4	Healthy	75
5	Healthy	65
6	Cancer	60
7	Cancer	58
8	Cancer	55
9	Cancer	52
10	Cancer	50

診断技術が発達して
この、今まで健康人扱いだった寿命65年の患者が
癌患者と判定されるようになった場合

カテゴリーの付け替えによって
健康人も、癌患者も 平均寿命が延びる

なにも変わってないのに、検査精度があがっただけで
寿命が延びている錯覚を起こす

Initial Mean Lifespan by Group

MEANS プロシジャ

分析変数 : Lifespan		
Group	Obs 数	平均
Cancer	5	55.00000000
Healthy	5	76.00000000

両グループの平均寿命が上がったように見える
→ 診断が良くなっただけで、実際に寿命が延びたわけではない。

ウィル・ロジャース現象は
疾患の評価法や、診断基準が変わった場合に起きうるので、

オンコロジー領域などでの評価・診断法変更時には要注意

Mean Lifespan After Reclassification of Transitional Case

MEANS プロシジャ

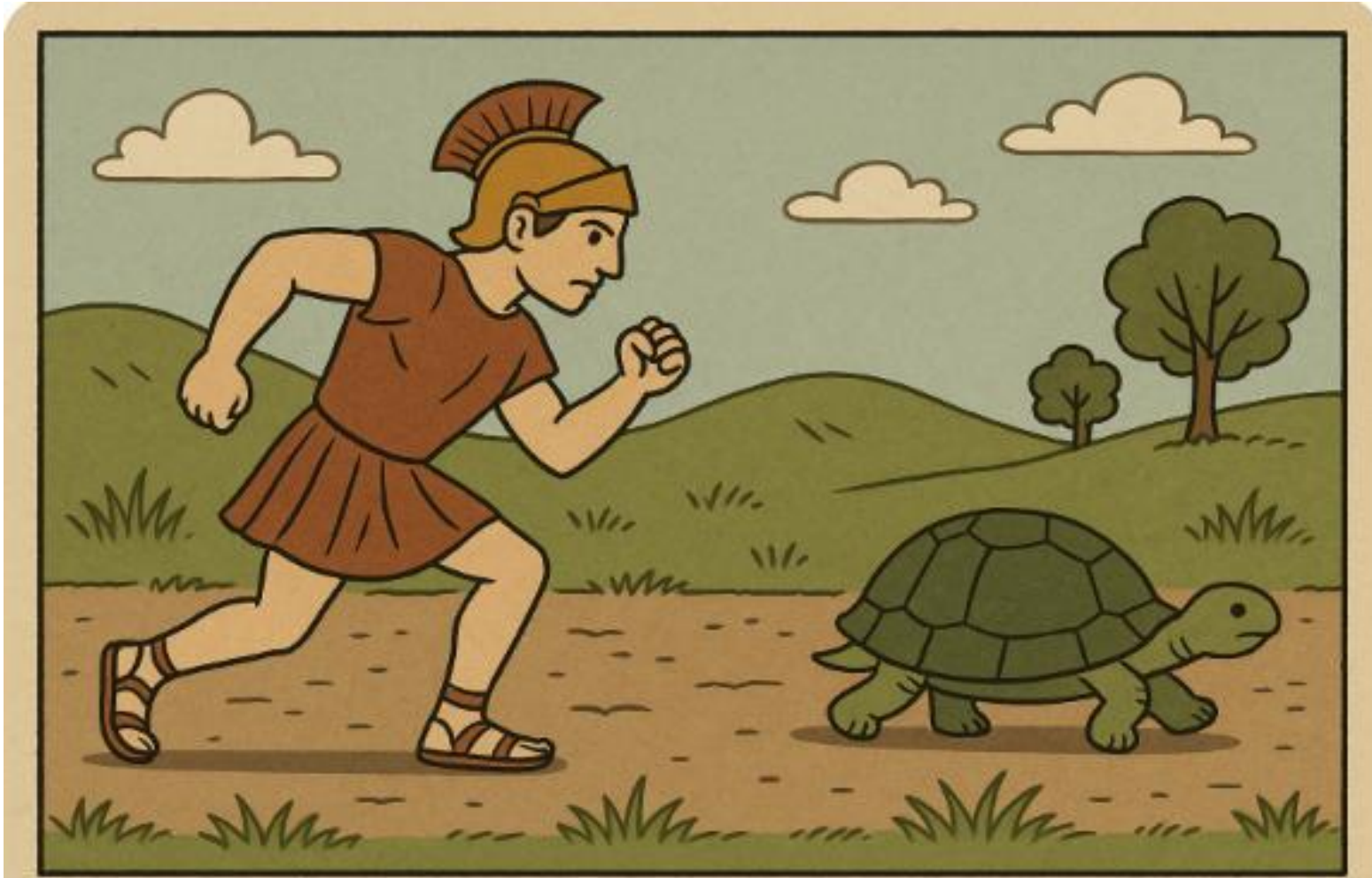
分析変数 : Lifespan		
Group	Obs 数	平均
Cancer	6	56.66666667
Healthy	4	78.75000000

ステージマイグレーションのパラドクスともいうらしい

生物統計も確率論も関係ないけど…

おまけ

アキレスと亀のパラドックス



ゼノンのパラドックスの一つ

走ることの最も遅いものですら最も速いものによって決して追いつかれ
ないであろう。なぜなら、追うものは、追いつく以前に、逃げるものが
走りはじめた点に着かなければならず、したがって、より遅いものは常
にいくらかずつ先んじていなければならないからである

•初期設定:

- アキレスの速度 = 10 m/s、 亀の速度 = 1 m/s
- 亀はスタート時に 100m 先にいる

•ロジック:

- 各ステップで、アキレスは亀のその時点での位置に追いつくのに必要な時間を計算
- その時間だけ進んだときの両者の位置を更新
- この繰り返しは、理論的には無限に続くが、ここでは15ステップで打ち切る

•結果:

- アキレスがいくら進んでも、常に亀が先にいる

```
/* ゼノンのステップを計算: アキレスが亀のいた位置に到達する点 */
data zeno_steps;
  /* 初期条件 */
  v_a = 10;
  v_t = 1;
  d = 100;
  /* 最初の15ステップを計算 */
  do step = 0 to 15;
    /* ステップnでの時間:  $t_n = d / (v_a - v_t) * (1 - (v_t/v_a)^n)$  */
    time = (d / (v_a - v_t)) * (1 - (v_t / v_a)**step);
    /* アキレスの位置 */
    achilles_pos = v_a * time;
    /* 亀の位置 */
    tortoise_pos = d + v_t * time;

    /*アキレスと亀の距離*/
    diff=tortoise_pos-achilles_pos;
    output;
  end;
format diff 30.28;
run;
```

step	time	achilles_pos	tortoise_pos	diff
0	0	0	100	100.000000000000000000000000000000
1	10	100	110	10.000000000000000000000000000000
2	11	110	111	1.000000000000000000000000000000
3	11.1	111	111.1	0.099999999999999900000000000000
4	11.11	111.1	111.11	0.010000000000000510000000000000
5	11.111	111.11	111.111	0.000999999999999056000000000000
6	11.1111	111.111	111.1111	0.000100000000000331900000000000
7	11.11111	111.1111	111.11111	0.000010000000000317410000000000
8	11.111111	111.11111	111.111111	0.000000999999997475240000000000
9	11.1111111	111.111111	111.1111111	0.000000099999994063182000000000
10	11.11111111	111.1111111	111.11111111	0.000000010000007932831000000000
11	11.111111111	111.11111111	111.111111111	0.000000001000003635454050000000
12	11.1111111111	111.111111111	111.1111111111	0.000000000100001784630876000000
13	11.11111111111	111.1111111111	111.11111111111	0.000000000010018652574217400000
14	11.111111111111	111.11111111111	111.111111111111	0.0000000000009947598300641400
15	11.1111111111111	111.111111111111	111.1111111111111	0.0000000000000994759830064140

限りなく 0 に近づくが、けっして逆転できない

追い越せることは自明なので、それにどう説明をつけて解消するか

/* データセットの作成：時間ごとのアキレスと亀の位置 */

data zeno_paradox;

/* 初期条件 */

v_a = 10; /* アキレスの速度 (m/s) */

v_t = 1; /* 亀の速度 (m/s) */

d = 100; /* 亀の初期位置 (m) */

/* 時間軸：0秒から12秒まで0.1秒刻み */

do time = 0 to 12 by 0.1;

/* アキレスの位置: $x_a = v_a * time$ */

achilles_pos = v_a * time;

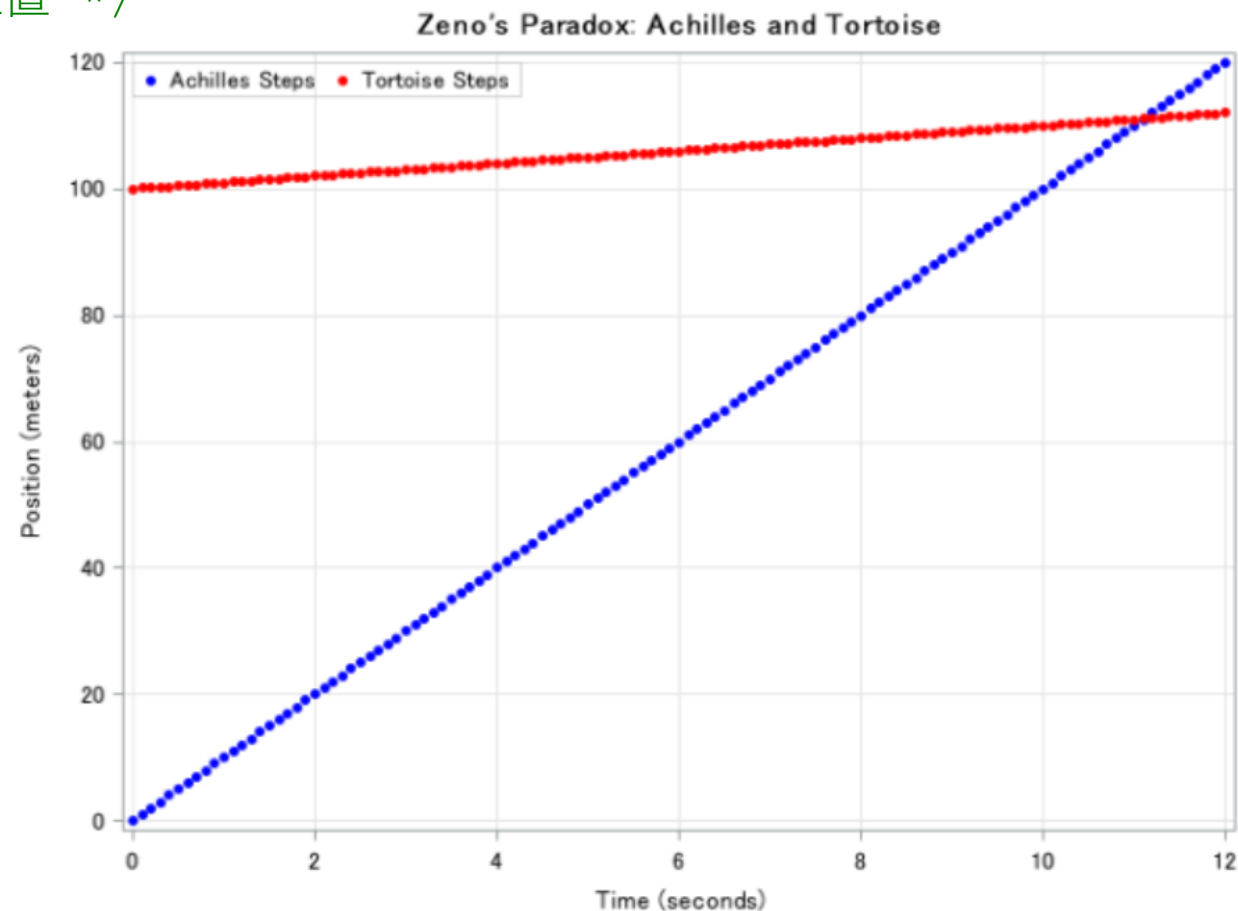
/* 亀の位置: $x_t = d + v_t * time$ */

tortoise_pos = d + v_t * time;

output;

end;

run;



総合的な解消のポイント

ゼノンのパラドックスは、無限の概念や連続性の直感的な理解が未熟だった時代に、運動や分割の不可能性を強調するために作られた。これらは現代数学・物理学で次のように解消される

1. **無限級数の収束:** 無限に分割された距離や時間も、数学的に有限の値に収束する。
2. **微積分学:** 無限小の概念を扱うことで、連続的な変化を正確に記述。
3. **空間と時間の連続性:** 物理学的に、空間と時間は無限分割の概念を超えて連続的に存在。

「無限の回数が無限の時間で行われる」

👉 この実際にはおかしい論理を（わざと）自然に仕込ませている点がこのパラドックスを成立させている

まとめ

統計解析におけるパラドックスの重要性への所感

統計的思考の深化

パラドックスを理解することで、「数字」だけでなく、その背景や条件を考慮する統計的思考が養われる

科学的判断の精度向上

パラドックスを通して、「何がデータから本当に言えるのか？」という視点が得られる

前提・推論・結果の解釈 のいずれかに誤解や見落としがないかを考えることは

解析の計画段階、 実施段階、 結果の考察 どの段階においても重要になる

教育や啓発への応用

パラドックスは、初学者から専門家まで幅広く使える教材となりえる

クイズ形式で導入することで興味を引き、データの「意味」を深く理解させるきっかけになる

ご清聴ありがとうございました！