
4 Excel による統計処理

4.1 基本統計量

何かを調べて得られたデータは、そのすべてが同じ値になることは極めてまれで、普通はある範囲内で“分布”したものとなる。データがどのように分布しているかを表すための指標として、分布の中心を示す指標やばらつきを示す指標などがある。これらはデータの基本統計量と呼ばれる。“統計量”については、後に詳しく説明するので、ここでは、あるルールに従ってデータから計算される値と考えておけばよいであろう。

4.1.1 分布の中心を示す指標

(1) 算術平均値

“算術平均値”は、データの総和をデータの個数で割ったものであり、普通は“平均値”と呼ばれる。平均値は一般に $\bar{X}$ という記号で記述され、データ個数が n であり、データが $X_1, X_2, \dots, X_n$ であるとき $\bar{X} = (X_1 + X_2 + \dots + X_n)/n$ で計算されるものである。

Excel で平均値を求めるための関数は AVERAGE であり、数式は、次のようになる。

=AVERAGE(データの範囲)

(2) 中位数

データをその値の小さい順ないし大きい順に並べたとき、真ん中の順位にくるデータの値を“中位数（メディアン）”という。データの個数が奇数のときは真ん中の順位が存在するが、偶数のときは真ん中に相当する順位が存在しないので、中央の順位にくる2個の値の平均値をもって中位数とする。

Excel で中位数（メディアン）を求めるための関数は、MEDIAN であり数式は、次のようになる。

=MEDIAN(データの範囲)

(3) 最頻値

値が同じデータの個数を数えたとき、最も個数の多いデータの値を“最頻値（モード）”という。ただし、値が同じデータの個数が1や2と少ないときは意味のないものである。

Excel で最頻値（モード）を求めるための関数は、MODE であり数式は、次のようになる。

=MODE(データの範囲)

4.1.2 分布のバラツキの程度を示す指標

データがどのように分布するかを示す指標として、中心位置と共にデータのバラツキ方を数値化したものが用いられる。

(1) 範囲

データの中の最大値と最小値の差を“範囲”という。データの“最大値・最小値”と共に“範囲”は、分布のバラツキの程度を示す指標の1つとして用いられる。

Excel で最大値と最小値を求める関数はそれぞれ MAX と MIN であるので、範囲を求めるための数式は、次のようになる。

=MAX(データの範囲) - MIN(データの範囲)

(2) 分散

“分散”は、各データの“偏差（データの値 - 平均値）”を平方（2乗）したものの平均値である。分散を計算するためには、まず平均値を計算しなければならない。次に各データの偏差の平方を求め、その平均値として分散を計算する。データ個数が n 、データが $X_1, X_2, \dots, X_n$ であり、その平均値が $\bar{X}$

であるとき、分散 V は

$$V = \frac{(X_1 - \bar{X})^2 + (X_2 - \bar{X})^2 + \dots + (X_n - \bar{X})^2}{n}$$

で計算される。分散は、常に非負である。また、データにバラツキがない、すなわち、すべてのデータが同じ値のとき、平均値も同じ値になるので偏差がすべて 0 となり、分散も 0 となる。分散の値が大きいほどデータのバラツキが大きいことになる。

Excel では、このような面倒な計算を簡単に処理するための関数 VARP が用意されており、分散を求める数式は、次のようになる。

=VARP(データの範囲)

(3) 標準偏差

偏差はデータと同じ単位であるが、分散の単位は偏差すなわちデータの単位を平方（2 乗）したものとなる。そこで、データと同じ単位でのバラツキの尺度として、分散の平方根をとったものが考えられている。分散の平方根をとった指標は“標準偏差”と呼ばれる。

Excel で標準偏差を求める関数は、STDEVP であり数式は、次のようになる。

=STDEVP(データの範囲)

また、標準偏差は分散の平方根であり、平方根を求める関数が SQRT であるので、次の数式でもよいことになる。

=SQRT(VARP(データの範囲))

4.1.3 グラフによる分布の視覚化

(1) 散布図

データをドット（点）グラフで表したものが散布図であり、データの分布を視覚化するためにしばしば用いられる。

(2) ヒストグラム

“ヒストグラム”とは、データの度数分布表を棒グラフにしたものであり、データの分布を視覚化するために散布図と並んでよく用いられる。ヒストグラムの作成方法は先に説明したので省略する。

【例題 4.1】 100 家族の子供数を調査したところ、次の表のようであった。このとき、(1)平均値、(2)中位数、(3)最頻値、(4)分散、(5)標準偏差を求めよ。

子供数	0 人	1 人	2 人	3 人	4 人	5 人
家族数	5	11	40	22	17	5

《解答》

(1) 平均値は、データの総和をデータの個数で割ったものであるので

$$\begin{aligned} & \frac{0 + \dots + 0 + 1 + \dots + 1 + 2 + \dots + 2 + 3 + \dots + 3 + 4 + \dots + 4 + 5 + \dots + 5}{100} \\ &= \frac{0 \cdot 5 + 1 \cdot 11 + 2 \cdot 40 + 3 \cdot 22 + 4 \cdot 17 + 5 \cdot 5}{100} = \frac{250}{100} = 2.5 \end{aligned}$$

(2) 子供数（データ）を少ないものから順に並べると

5 11 40 22 17 5
0,0, ..., 0,1,1, ..., 1,2,2, ..., 2,3,3, ..., 3,4,4, ..., 4,5,5, ..., 5

となり、最初の 2 が 17 番目で最後の 2 が 56 番目となる。ここで、データ数は 100 であるので、中央は 50 番目と 51 番目となり、それらのデータの値は共に 2 人である。また、それらの平均値も 2 人となるので、中位数（メディアン）は 2 人の家族となる。

(3) この例題では最大個数 40 を示す 2 人の家族が最頻値となる。

(4) 平均値が 2.5 であるので、分散は次のようになる。

$$\frac{(0-2.5)^2 \cdot 5 + (1-2.5)^2 \cdot 11 + (2-2.5)^2 \cdot 40 + (3-2.5)^2 \cdot 22 + (4-2.5)^2 \cdot 17 + (5-2.5)^2 \cdot 5}{100} \\ = \frac{6.25 \cdot 5 + 2.25 \cdot 11 + 0.25 \cdot 40 + 0.25 \cdot 22 + 2.25 \cdot 17 + 6.25 \cdot 5}{100} = \frac{142}{100} = 1.42$$

(5) 標準偏差は、分散の平方根であるので

$$\sqrt{1.42} = 1.19 \quad \square$$

【例題 4.2】表 1 の都道府県別 1 住宅当たりの平均敷地面積データを用いて、(1)平均値、(2)中位数
(3)範囲、(4)分散、(5)標準偏差、(6)度数分布、(7)ヒストグラムを求めよ。

北海道	283	埼 玉	239	岐 阜	283	鳥 取	310	佐 賀	311
青 森	339	千 葉	272	静 岡	261	島 根	285	長 崎	233
岩 手	350	東 京	150	愛 知	253	岡 山	258	熊 本	322
宮 城	348	神奈川	189	三 重	286	広 島	215	大 分	282
秋 田	386	新 潟	340	滋 賀	284	山 口	267	宮 崎	315
山 形	390	富 山	399	京 都	173	徳 島	282	鹿児島	290
福 島	360	石 川	288	大 阪	132	香 川	278	沖 縄	273
茨 城	423	福 井	321	兵 庫	199	愛 媛	225		
栃 木	393	山 梨	331	奈 良	235	高 知	183		
群 馬	350	長 野	335	和歌山	212	福 岡	267		

表 4.1 都道府県別 1 住宅当たりの平均敷地面積 (1993 年、単位: m²)

《解答》まず、図 1 のようにデータを入力する。

	A	B	C	D	E
1	都道府県名	平均敷地面積			
2	北海道	283		平均値	
3	青 森	339		中位数	
4	岩 手	350		最大値	
5	宮 城	348		最小値	
6	秋 田	386		範 囲	
7	山 形	390		分 散	
8	福 島	360		標準偏差	
9	茨 城	423			
10	栃 木	393	階 級	範 囲	
11	群 馬	350	1	150	
12	埼 玉	239	2	200	
13	千 葉	272	3	250	
14	東 京	150	4	300	
15	神奈川	189	5	350	
16	新 潟	340	6	400	
17	富 山	399	7	450	
18	石 川	288			
:	:	:			
48	沖 縄	273			

図 1

(1) 平均値を計算するためには、セル E2 に数式 =AVERAGE(B2:B48) を入力する。

(2) 中位数を計算するためには、セル E3 に数式 =MEDIAN(B2:B48) を入力する。

(3) 範囲を計算するために、最大値と最小値を別に計算することにする。最大値を計算するためにセル E4

に数式=MAX(B2:B48)、最小値を計算するためにセル E5 に数式=MIN(B2:B48)を入力し、範囲を計算するためにセル E6 に数式=E4-E5 を入力する。

(4) 分散を計算するためには、セル E7 に数式=VARP(B2:B48)を入力する。

(5) 標準偏差を計算するためには、セル E8 に数式=STDEVP(B2:B48)を入力する。

(6) 度数分布表を求めるためには、各度数の範囲を決めなければならない。ここでは、最大値と最小値を参考にして、図 1 に示すように 50m² 毎の 7 階級にする。まず、図 1 のようにセル E11 から E17 に各範囲の上限を入力する。つぎに、度数を計算するセル範囲 D11 から D17 を範囲指定し、数式=FREQUENCY(B2:B48,E11:E17)をキーインし、SHIFT キーと CTRL キーを同時に押しながらリターンキーを押す。

(7) ヒストグラムを求めるためには、D10 から E17 までを範囲指定し、棒グラフを描けばよい。□

4.2 確率分布

4.2.1 統計量

ある種の分布をする母集団と、その母集団から取られる大きさ n の標本（サンプル）を考え、母集団が従っている分布に関連した母数（母平均、母分散等）の集合を θ とし、標本データの集合を Ω で表すものとしよう。このとき、標本データおよび母数を用いることによって構成される数式 $f(\theta; \Omega)$ は、“統計量”と呼ばれる。

母数は、真なる唯一の値が存在するものである。しかし、標本データは、再度標本を取り直したとすれば、前回とは異なるデータが得られることとなり、このような標本抽出を繰り返したとすれば、多数の微妙に異なる標本データが得られることとなる。また、標本を抽出する毎に得られる標本データを代入して計算される数式 $f(\theta; \Omega)$ の値も、微妙に異なるものが多数得られることとなる。従って、標本抽出毎に得られる数式の値 $f(\theta; \Omega)$ は分布することになり、 $X=f(\theta; \Omega)$ としたとき、 X は確率変数となる。このことが、標本データを含む数式 $f(\theta; \Omega)$ が統計量と呼ばれる由縁である。

例として、ある母集団から大きさ n の標本を取ることを考えてみよう。このとき、 n 個のそれぞれの標本を $X_1, X_2, \dots, X_n$ とし、標本平均

$$\bar{X} = \frac{X_1 + X_2 + \dots + X_n}{n}$$

を考えると、 $\bar{X}$ は母数こそ含まないものの統計量となる。

4.2.2 不偏性

一般にデータを用いた統計分析を行うとき、母数は未知であることが多く、母数が未知のままであれば統計分析が行えないことが多い。そこで、データを用いて未知母数を推定し、統計分析を行うことになる。未知母数を推定するために適当な統計量を考えるが、このようにある値を推定するために用いる統計量は“推定量”と呼ばれる。推定量としては、どのようなものを用いてもよいが、推定量として望ましい性質がいくつか提案されており、できればこれらの性質を満たす推定量を用いるべきであろう。この望ましい性質の一つに“不偏性”と呼ばれるものがある。

推定量にデータの値を代入して得られる値は“推定値”と呼ばれるが、推定値はデータが取り直されれば、前回とは微妙に違った値になり、確率分布することになる。この確率分布の平均値が、推定しようとしている母数の値と一致することが保証されているとすれば、好ましいことになる。なぜならば、得られる推定値はその平均値の近くの値になることが多いと考えられるからである。このように推定値の平均値が母数と一致する性質を不偏性と呼ぶ。一般には“推定量が不偏性を満たす”と表現される。もしも、ある母数を推定するために採用された推定量が不偏性を満たさないとすれば、データを代入して得られる推定値が母数の近くの値になる保証がなく、適切な推定量とは言えないことになる。不偏性

に関して知られている性質のいくつかを以下に列記しておく。

ある母集団から無作為に抽出された大きさ n の標本が $X_1, X_2, \dots, X_n$ であるとき次の性質が成り立つ。

【性質 4.1】母平均の推定量として標本平均

$$\bar{X} = \frac{X_1 + X_2 + \dots + X_n}{n}$$

を考えたとき、標本平均は不偏性を満たす。

【性質 4.2】母平均 μ が既知であれば、母分散の推定量として

$$v^2 = \frac{(X_1 - \mu)^2 + (X_2 - \mu)^2 + \dots + (X_n - \mu)^2}{n}$$

を考えたとき、この推定量 v^2 は不偏性を満たす。

【性質 4.3】母平均 μ が未知のとき、母分散の推定量として標本分散

$$s^2 = \frac{(X_1 - \bar{X})^2 + (X_2 - \bar{X})^2 + \dots + (X_n - \bar{X})^2}{n}$$

を考えたとき、標本分散は不偏性を満たさない。

標本分散は母分散の不偏推定量（不偏性を満たす推定量）にならないことに注意すべきであろう。母分散の不偏推定量として、“不偏分散”なるものが考えられている。

【性質 4.4】母分散の推定量として不偏分散

$$u^2 = \frac{(X_1 - \bar{X})^2 + (X_2 - \bar{X})^2 + \dots + (X_n - \bar{X})^2}{n-1}$$

を考えたとき、不偏分散は不偏性を満たす。

標本分散と不偏分散の差は、分母が n か $n-1$ であるかの違いである。また、性質 4.2 の推定量 v^2 と不偏分散 u^2 は共に母分散の不偏推定量であるが、 v^2 が母平均 μ を含むのに対し、 u^2 は母数を含んでいない。一般には、母平均が未知であるので、母分散の推定量として不偏分散が使用される。

Excel において、標本分散 s^2 を求めるための数式が

$$=VARP(\text{データの範囲})$$

であり、標本標準偏差を求めるための数式が

$$=STDEVP(\text{データの範囲})$$

であることを示した。Excel は不偏分散 u^2 と不偏標準偏差 u を求めるための関数が用意されており、それぞれ VAR と STDEV で、数式は

$$=VAR(\text{データの範囲}) \quad \text{と} \quad =STDEV(\text{データの範囲})$$

である。

4.2.3 分布の基本的性質

母平均 μ 、母分散 σ^2 の母集団の分布を $A(\mu, \sigma^2)$ で表すものとする。このとき、この母集団から抽出された標本 X は、分布 $A(\mu, \sigma^2)$ に従うことになり、 c を定数としたとき、次の性質が成り立つ。

【性質 4.5】 $X+c$ は分布 $A(\mu+c, \sigma^2)$ に従う。

【性質 4.6】 cX は分布 $A(c\mu, c^2\sigma^2)$ に従う。

したがって、次の性質も成り立つ。

【性質 4.7】 $X - \mu$ は分布 $A(0, \sigma^2)$ に従う。

【性質 4.8】 $(X - \mu)/\sigma$ は分布 $A(0, 1)$ に従う。

n 個の母集団があり、それぞれの母平均と母分散が μ_i と σ_i^2 ($i=1, 2, \dots, n$) であるものとする。各母集団から 1 つずつ抽出された標本 X_1 の和 $X = X_1 + X_2 + \dots + X_n$ を考える。このとき、次の性質が成り立つ。

【性質 4.9】 $X = X_1 + X_2 + \dots + X_n$ の分布は、平均値が $\mu_1 + \mu_2 + \dots + \mu_n$ で分散が $\sigma_1^2 + \sigma_2^2 + \dots + \sigma_n^2$ の分布になる。

【性質 4.10】 $\mu_1 = \mu_2 = \dots = \mu_n = \mu$ かつ $\sigma_1^2 = \sigma_2^2 = \dots = \sigma_n^2 = \sigma^2$ であるならば、 X は平均 $n\mu$ 、分散 $n\sigma^2$ の分布に従うことになる。

【性質 4.11】 母平均 μ 、母分散 σ^2 の母集団より無作為に抽出された大きさ n の標本 $X_1, X_2, \dots, X_n$ の標本平均 $\bar{X}$ の分布を考える。このとき、標本平均 $\bar{X}$ は平均 μ 、分散 σ^2/n の分布に従うことになる。

4.2.4 正規分布

自然界に存在するものの多くが正規分布することと、次の定理（中心極限定理と呼ばれる）が成り立つことから、正規分布は最もよく用いられる。

【定理 4.1】 平均 μ と分散 σ^2 が存在する母集団において、母集団より無作為に抽出した大きさ n の標本の標本平均を $\bar{X}$ とする。このとき、母集団分布の形にかかわらず、 n が大きくなるとともに $\bar{X}$ は正規分布 $N(\mu, \sigma^2/n)$ に限りなく近づく。

性質 4.8 より、確率変数 X が、平均 μ 、分散 σ^2 の正規分布 $N(\mu, \sigma^2)$ に従うならば、統計量

$$Z = \frac{X - \mu}{\sigma}$$

は、平均 0、分散 1 の正規分布 $N(0, 1)$ に従う。平均 0、分散 1 の正規分布 $N(0, 1)$ は“標準正規分布”と呼ばれる。

性質 4.11 より、平均 μ 、分散 σ^2 の正規分布 $N(\mu, \sigma^2)$ に従う母集団から抽出された、大きさ n の標本の標本平均 $\bar{X}$ （これも統計量である）は、 $N(\mu, \sigma^2/n)$ に従う。また、統計量

$$Z = \frac{\bar{X} - \mu}{\sigma / \sqrt{n}}$$

は、標準正規分布 $N(0, 1)$ に従う。

Excel においては、標準正規分布 $N(0, 1)$ に関連した関数 NORMSDIST(z)、NORMSINV(p) と、一般の正規分布 $N(\mu, \sigma^2)$ に関連した関数 NORMDIST($z, \mu, \sigma, \text{TRUE}$ ないし FALSE) および NORMINV(p, μ, σ) が用意されている。

NORMSDIST(z) は、標準正規分布の累積分布関数の値（標準正規分布の確率密度関数を $-\infty$ から z まで積分した値）を計算する。計算結果は、標準正規分布に従う確率変数 X が z 以下となる確率（母集団全体の中で z 以下のものが占める割合）すなわち $\Pr\{X \leq z\}$ であり、 z 以上となる確率（母集団全体の中で z 以上のものが占める割合）すなわち $\Pr\{z < X\}$ は $1 - \text{NORMSDIST}(z)$ で計算できる。

NORMSINV(p)は、NORMSDIST 関数の逆関数であり、標準正規分布で z 以下となる確率（母集団全体の中で z 以下のものが占める割合）を p 、すなわち $\Pr\{X \leq z\}=p$ としたとき、 p を与えて z を計算するためのものである。

NORMDIST($z, \mu, \sigma, \text{TRUE}$ ないし FALSE)は、平均値 μ 、標準偏差 σ の正規分布 $N(\mu, \sigma^2)$ において、4 番目の引き数が TRUE のとき累積分布関数の値（母集団全体の中で z 以下のものが占める割合）すなわち $\Pr\{X \leq z\}$ 、 FALSE のとき確率密度関数の値を計算する。したがって、NORMDIST($z, 0, 1, \text{TRUE}$)は NORMSDIST(z)とまったく同じである。

NORMINV(p, μ, σ) は、NORMDIST 関数の逆関数であり、標準正規分布で z 以下となる確率を p 、すなわち $\Pr\{X \leq z\}=p$ としたとき、 p を与えて z を計算するためのものである。NORMINV($z, 0, 1$)は NORMSINV(z)とまったく同じである。

【例題 4.3】ある全国試験の結果は、平均点が 55 点、標準偏差が 10 点であった。この試験で、A 君の得点は 63 点であった。このとき、A 君より高得点の者は何%いたと考えられるか。また、B 君の得点が 49 点であったとき、B 君より高得点の者は何%いたか。

《解答》一般に、成績は正規分布をすると考えられている。正規分布は、平均値と分散（標準偏差の平方）が与えられれば、一意にその確率分布が決定されるものである。

ここでの問題では、平均値が 55、分散が 10^2 であるので、全受験者の成績は、 $N(55, 10^2)$ なる分布に従うと考えられる。したがって、無作為抽出される受験生の成績 X は、 $N(55, 10^2)$ なる分布に従う母集団からの標本と考えられる。なお、平均が 55、分散が 10^2 の正規分布（英語で normal distribution という）のことを記号で、 $N(55, 10^2)$ と表す。

正規分布 $N(55, 10^2)$ に従う母集団において、63 以下のものが占める割合を求めるための Excel の数式は

$$=\text{NORMDIST}(63, 55, 10, \text{TRUE})$$

であるので、63 以上のものが占める割合を求めるための Excel の数式は

$$=1 - \text{NORMDIST}(63, 55, 10, \text{TRUE})$$

となり、その計算結果は 0.211855 である。結局、問題の全国試験で 63 以上のものが全体に占める割合は 0.212 すなわち 21.2 % となる。

同様に、問題の全国試験で B 君の得点 49 以上のものが全体に占める割合は

$$=1 - \text{NORMDIST}(49, 55, 10, \text{TRUE})$$

で計算され、結果は 0.725747 すなわち 72.6 % となる。□

4.2.5 カイ 2 乗分布

【定義 4.1】 n 個の確率変数 $Z_1, Z_2, \dots, Z_n$ それぞれが標準正規分布 $N(0, 1)$ に従い、互いに独立であるとき、統計量

$$\chi^2 = Z_1^2 + Z_2^2 + \dots + Z_n^2$$

が従う確率分布を“自由度 n のカイ 2 乗分布”という。

カイ 2 乗分布する確率変数は、負の値をとらない。

$N(\mu, \sigma^2)$ に従う母集団から独立に抽出された、大きさ n の標本 $X_1, X_2, \dots, X_n$ を考える。このとき、標本平均 $\bar{X} = (X_1 + X_2 + \dots + X_n)/n$ は正規分布 $N(\mu, \sigma^2/n)$ に従うので、統計量

$$Z = \frac{\bar{X} - \mu}{\sigma / \sqrt{n}}$$

は標準正規分布 $N(0,1)$ に従い、統計量

$$\chi^2 = Z^2 = \frac{(\bar{X} - \mu)^2}{\sigma^2/n}$$

は自由度 1 のカイ 2 乗分布に従う。

また、統計量

$$Z_k = \frac{X_k - \mu}{\sigma}, \quad k=1,2,\dots,n$$

はそれぞれ標準正規分布 $N(0,1)$ に従うことになる。したがって、統計量

$$\chi^2 = Z_1^2 + Z_2^2 + \dots + Z_n^2 = \frac{(X_1 - \mu)^2 + (X_2 - \mu)^2 + \dots + (X_n - \mu)^2}{\sigma^2}$$

は、自由度 n のカイ 2 乗分布に従う。

【性質 4.1 2】統計量

$$\chi^2 = \frac{(X_1 - \bar{X})^2 + (X_2 - \bar{X})^2 + \dots + (X_n - \bar{X})^2}{\sigma^2}$$

は、自由度 $n-1$ のカイ 2 乗分布に従う。

標本分散を s^2 、不偏分散を u^2 としたとき、

$$(X_1 - \bar{X})^2 + (X_2 - \bar{X})^2 + \dots + (X_n - \bar{X})^2 = ns^2 = (n-1)u^2$$

であるので、統計量

$$\chi^2 = \frac{ns^2}{\sigma^2} = \frac{(n-1)u^2}{\sigma^2}$$

は、自由度 $n-1$ のカイ 2 乗分布に従うことになる。

Excel においては、カイ 2 乗分布に関連した関数 CHIDIST(y,m) と CHIINV(p,m) が用意されている。

CHIDIST(y,m) は、自由度 m のカイ 2 乗分布の確率密度関数を y から ∞ まで積分した値を計算する。計算結果は、自由度 m のカイ 2 乗分布に従う確率変数 X が y 以上となる確率（母集団全体の中で y 以上のものが占める割合）すなわち $\Pr\{y \leq X\}$ であり、 y 以下となる確率（母集団全体の中で y 以下のものが占める割合）すなわち $\Pr\{X \leq y\}$ は $1 - \text{CHIDIST}(y,m)$ で計算できる。

CHIINV(p,m) は、CHIDIST 関数の逆関数であり、自由度 m のカイ 2 乗分布で y 以上となる確率（母集団全体の中で y 以上のものが占める割合）を p 、すなわち $\Pr\{y \leq X\} = p$ としたとき、 p を与えて y を計算するためのものである。

4.2.6 t 分布

【定義 4.2】確率変数 Z が標準正規分布 $N(0,1)$ に従い、確率変数 Y が自由度 m のカイ 2 乗分布に従い、 Z と Y が独立のとき、統計量

$$T = \frac{Z}{\sqrt{Y/m}}$$

が従う確率分布を“自由度 m の t 分布”という。

t 分布は、平均が 0 であり、左右対称の釣り鐘形である。見た目は、標準正規分布と似ているが、自由度が小さいときは標準正規分布を上から押しつぶしたような形である。自由度が大きくなるに従って、標準正規分布に近づいていき、自由度が ∞ になれば標準正規分布に一致する。

平均 μ 、分散 σ^2 の正規分布 $N(\mu, \sigma^2)$ に従う母集団から抽出された大きさ n の標本の標本平均を $\bar{X}$ 、標本分散を s^2 (s は標本標準偏差) とすると、 $(\bar{X} - \mu) / (\sigma / \sqrt{n})$ が標準正規分布 $N(0,1)$ に従い、 ns^2 / σ^2 が

自由度 $n-1$ のカイ 2 乗分布に従うので、統計量

$$T = \frac{(\bar{X} - \mu) / (\sigma / \sqrt{n})}{\sqrt{ns^2 / \{\sigma^2 (n-1)\}}} = \frac{\bar{X} - \mu}{s / \sqrt{n-1}}$$

は、自由度 $n-1$ の t 分布に従う。なお、不偏分散を u^2 (u は不偏標準偏差) とすると、標本分散と不偏分散の間には $ns^2 = (n-1)u^2$ なる関係があるので、統計量

$$T = \frac{\bar{X} - \mu}{u / \sqrt{n}}$$

も、自由度 $n-1$ の t 分布に従う。

【性質 4.1 3】二つの母集団それぞれが、分散の等しい正規分布 $N(\mu_1, \sigma^2)$ と $N(\mu_2, \sigma^2)$ に従うものとする。このとき、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、標本平均が $\bar{X}_1$ と $\bar{X}_2$ 、標本標準偏差が s_1 と s_2 であるものとする。このとき

$$v = \sqrt{(n_1 s_1^2 + n_2 s_2^2) / (n_1 + n_2 - 2)}$$

とすると、統計量

$$T = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{v \sqrt{1/n_1 + 1/n_2}}$$

は、自由度 $n_1 + n_2 - 2$ の t 分布に従う。

【性質 4.1 4】二つの母集団それぞれが、分散が等しくない正規分布 $N(\mu_1, \sigma_1^2)$ と $N(\mu_2, \sigma_2^2)$ に従うものとする。このとき、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、その標本平均が $\bar{X}_1$ と $\bar{X}_2$ 、不偏標準偏差が u_1 と u_2 であるものとする。このとき、 $w_1 = u_1^2 / n_1$ 、 $w_2 = u_2^2 / n_2$ とすると、統計量

$$T = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{w_1 + w_2}$$

は、近似的に自由度 u の t 分布に従う。ただし、 u は

$$\frac{(w_1 + w_2)^2}{w_1^2 / (n_1 - 1) + w_2^2 / (n_2 - 1)}$$

に最も近い整数値である。

Excel においては、 t 分布に関連した関数 TDIST(t , m , 1 または 2) と TINV(p , m) が用意されている。

TDIST(t , m , 1) は、自由度 m の t 分布に従う確率変数 T が t 以上となる確率 (母集団全体の中で t 以上のものが占める割合) すなわち $\Pr\{T \geq t\}$ を計算するものである。ただし、 t は正でなければならない。 t 分布が左右対称であることから、自由度 m の t 分布に従う確率変数が $-t$ 以下となる確率も、TDIST(t , m , 1) で計算できる。

TDIST(t , m , 2) は、自由度 m の t 分布に従う確率変数が t 以上または $-t$ 以下となる確率 (母集団全体の中で t 以上のものと $-t$ 以下のものが占める割合) すなわち $\Pr\{T \leq -t \text{ または } T \geq t\}$ を計算するものである。ただし、 t は正でなければならない。TDIST(t , m , 2) は、 $2 * \text{TDIST}(t, m, 1)$ と同じであり、 $-t$ 以上 t 以下となる確率すなわち $\Pr\{-t \leq T \leq t\}$ は $1 - \text{TDIST}(t, m, 2)$ で計算できる。

TINV(p , m) は、TDIST 関数の逆関数であり、自由度 m の t 分布で t 以下または t 以上となる確率 (母集団全体の中で $-t$ 以下と t 以上のものが占める割合) を p 、すなわち $\Pr\{T \leq -t \text{ または } T \geq t\} = p$ としたとき、 p を与えて t を計算するためのものである。

4.2.7 F 分布

【性質 4.1 5】2つの確率変数 X_1 と X_2 が、それぞれ自由度 m と n のカイ 2 乗分布に従い、互いに独立であるとき、統計量

$$F = \frac{X_1/m}{X_2/n}$$

は自由度 m, n の F 分布に従う。逆に、 $F = (X_2/n)/(X_1/m)$ は自由度 n, m の F 分布に従う。

カイ 2 乗分布に従う確率変数が負の値をとらないので、F 分布に従う確率変数も負の値をとらない。

統計量 F が自由度 m, n の F 分布に従うとき、 $\Pr\{F \leq b\} = \alpha$ となるような b を求めたいとすれば、統計量 $G = 1/F$ が自由度 n, m の F 分布に従うことから、 $\Pr\{a \leq G\} = \alpha$ となるような a を求めれば $\Pr\{a \leq 1/F\} = \alpha$ であるので、 $\Pr\{F \leq 1/a\} = \alpha$ となり求める b は $1/a$ で与えられる。

【性質 4.1 5】二つの母集団それぞれが、分散の等しい正規分布 $N(\mu_1, \sigma^2)$ と $N(\mu_2, \sigma^2)$ に従うものとする。このとき、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、その不偏標準偏差が u_1 と u_2 であるものとする。このとき、統計量

$$F = \frac{u_1^2}{u_2^2}$$

は、自由度 $n_1 - 1, n_2 - 1$ の F 分布に従う。

Excel においては、F 分布に関連した関数 $\text{FDIST}(f, m, n)$ と $\text{FINV}(p, m, n)$ が用意されている。

$\text{FDIST}(f, m, n)$ は、自由度 m, n の F 分布の確率密度関数を f から ∞ まで積分した値を計算する。計算結果は、自由度 m, n の F 分布に従う確率変数 F が f 以上となる確率（母集団全体の中で f 以上のものが占める割合）すなわち $\Pr\{f \leq F\}$ であり、 f 以下となる確率（母集団全体の中で F 以下のものが占める割合）すなわち $\Pr\{F \leq f\}$ は $1 - \text{FDIST}(f, m, n)$ で計算できる。

$\text{FINV}(p, m, n)$ は、 FDIST 関数の逆関数であり、自由度 m, n の F 分布で f 以上となる確率（母集団全体の中で f 以上のものが占める割合）を p 、すなわち $\Pr\{f \leq F\} = p$ としたとき、 p を与えて f を計算するためのものである。

4.3 区間推定

確率変数 $X = f(\theta; \Omega)$ がある既知の分布に従い、この分布において $\Pr\{a \leq X \leq b\} = 1 - \alpha$ すなわち $\Pr\{a \leq f(\theta; \Omega) \leq b\} = 1 - \alpha$ であるものとしよう。このとき、 $a = f(\theta; \Omega)$ と $b = f(\theta; \Omega)$ を未知母数について解いた結果をそれぞれ $\theta = g_1(a; \Omega)$ と $\theta = g_2(b; \Omega)$ とすると $\Pr\{g_1(a; \Omega) \leq \theta \leq g_2(b; \Omega)\} = 1 - \alpha$ となる。したがって、 θ は

『信頼水準 $1 - \alpha$ で θ の信頼区間は、 $[g_1(b; \Omega), g_2(a; \Omega)]$ である。』

と区間推定されることになる。

母集団が $N(\mu, \sigma^2)$ すなわち平均 μ 、分散 σ^2 の正規分布に従うことが知られており、母数 μ と σ^2 は共に未知とする。このとき

$$T = \frac{\bar{X} - \mu}{u/\sqrt{n}}$$

は自由度 $n - 1$ の t 分布に従う。ただし、 n は標本の大きさであり、 $\bar{X}$ は標本平均、 u は不偏標準偏差すなわち

$$u^2 = \frac{(X_1 - \bar{X})^2 + (X_2 - \bar{X})^2 + \cdots + (X_n - \bar{X})^2}{n - 1}$$

である。自由度 $n - 1$ の t 分布において、 $\Pr\{-t_\alpha \leq T \leq t_\alpha\} = 1 - \alpha$ とすると

$$\Pr\left\{-t_{\alpha} \leq \frac{\bar{X} - \mu}{u/\sqrt{n}} \leq t_{\alpha}\right\} = 1 - \alpha$$

となり、解き直すことによって

$$\Pr\left\{\bar{X} - t_{\alpha} u/\sqrt{n} \leq \mu \leq \bar{X} + t_{\alpha} u/\sqrt{n}\right\} = 1 - \alpha$$

となるので『信頼水準 $1 - \alpha$ で μ の信頼区間は $[\bar{X} - t_{\alpha} u/\sqrt{n}, \bar{X} + t_{\alpha} u/\sqrt{n}]$ と区間推定されることになる。Excel において、 t_{α} は TINV($\alpha, n-1$) で計算できるので、信頼区間は

$$= \text{AVERAGE}(\text{データの範囲}) - \text{TINV}(\alpha, n-1) * \text{STDEV}(\text{データの範囲}) / \text{SQRT}(n)$$

と

$$= \text{AVERAGE}(\text{データの範囲}) + \text{TINV}(\alpha, n-1) * \text{STDEV}(\text{データの範囲}) / \text{SQRT}(n)$$

で計算できる。

【例題 4.4】男性全体の平均身長を知る目的で、4 人の男性(標本)を調べたところ 165、170、170、175 であった。このデータから、男性全体の平均身長 μ を信頼水準 95 % で区間推定せよ。

《解答》Excel において、A1 に 165、A2 に 170、A3 に 170、A4 に 175 とデータを入力した場合を考える。この問題では、 $n=4$ 、 $\alpha=1-0.95=0.05$ であるので、信頼区間は

$$= \text{AVERAGE}(A1:A4) - \text{TINV}(0.05, 3) * \text{STDEV}(A1:A4) / \text{SQRT}(4)$$

と

$$= \text{AVERAGE}(A1:A4) + \text{TINV}(0.05, 3) * \text{STDEV}(A1:A4) / \text{SQRT}(4)$$

で計算でき、結果は 163.5039 と 176.4961 になる。このことから、信頼水準 95 % での μ の信頼区間は、163.5cm 以上 176.5cm 以下となる。□

4.4 統計的仮説検定

普通、仮説といえば真であることが切望されているものであるが、統計的仮説に限っては、偽であり否定されることが望まれるものである。そこで、この成り立たないことが願望される仮説を『帰無仮説』と呼んで、普通の仮説と区別している。

一方、統計的仮説検定とは、ある事象が真であるか偽であるかを統計的に判定することである。一般に、偽であることが望ましい事象を帰無仮説として設定し、標本データから、その帰無仮説を受け入れる(受容する)か、受け入れない(棄却する)かを判定する方法がとられる。

まず、帰無仮説が正しいと仮定して、手元の標本データが得られる確率を計算する。そして、この確率が非常に小さいならば、『このように小さな確率でしか起こらないはずのことが、たった一回の標本抽出で起こるのはおかしい。すなわち、どこかに矛盾があるはずだ。』との考え方から、帰無仮説を正しいとした仮定が誤りであったと判断し、帰無仮説を棄却する。

逆に、標本データが得られる確率が小さくないときは、正しいと仮定した状況(帰無仮説)は矛盾を含むものではなく、十分に有り得ることと判断し、帰無仮説を受容することになる。

ここで、確率が大きい小さいかを判定するための基準は、検定作業の最初に決定しなければならないが、一般に、5 % ないし 1 % が用いられる。すなわち、5 % ないし 1 % の確率は小さく、そう簡単には起こらないものとする。そして、これらの基準は『有意水準』と呼ばれる。また、帰無仮説を棄却するような判定が下されたときに、その判定が誤りである確率がまさにこの基準に等しいことから、『危険率』と呼ばれることもある。

帰無仮説が正しいにもかかわらず、帰無仮説を棄却するような判定を下す誤りを『第 1 種の過誤』と呼び、このような誤りを犯す確率が危険率である。一方、帰無仮説が正しくないにもかかわらず、帰無仮説を受容するような判定を下す誤りを『第 2 種の過誤』と呼ぶ。この第 2 種の過誤は、問題に応じて

異なる値となるが、第1種の過誤（危険率）を大きくすると小さくなり、逆に、危険率を小さくすると第2種の過誤が大きくなる。従って、両過誤のバランスをとることが必要となり、経験的に先に述べた5%ないし1%が良いとされている。

統計的仮説検定が行えるためには、帰無仮説が対象としている母数と標本データから構成される統計量が必要であり、また、この統計量が従う分布が既知でなければならない。このような条件が成り立つとき、帰無仮説が正しいと仮定することで、この統計量が従う分布が確定し、統計量がある値をとる確率が求まる。一般に、確率分布は平均値近辺の値がとられる確率が高く、平均値から大きく離れた値をとる確率は低い。このことから、めったに起こらない事象を、『平均値から大きく離れた値をとること』とするのは妥当な考え方であろう。

有意水準を α としたとき、ここでの統計量が従う分布から、 a 以下の値をとる確率が $\alpha/2$ 、 b 以上の値をとる確率が $\alpha/2$ となるような a と b を求める。すると、統計量が a 以下ないし b 以上の値をとる確率は α となり、このような事象はめったに起こらないということになる。

結局、統計的仮説検定は、手元にある標本データを用いて計算された統計量の値が、 a 以下ないし b 以上の値をとったとき、帰無仮説を棄却し、それ以外の値をとったとき帰無仮説を受容するものである。このことから、 a 以下ないし b 以上の範囲を『棄却域』と呼び、それ以外の範囲を『受容域』と呼ぶ。

【例題 4.5】内容量が 200cc と表示されている、あるメーカーの缶ジュースを愛飲しているが、どうも量が少ないように思われた。そこで、10 本の缶ジュースの量を調べたところ、平均が 196cc、不偏標準偏差が 5cc であった。この結果を基に、表示の真偽を検定しなさい。

《解答》缶ジュースの内容量は正規分布すると考えてよいであろう。そこで、その平均を μ 、分散を σ^2 とすると、10 本の標本の平均値は、 $N(\mu, \sigma^2/10)$ すなわち平均 μ 、分散 $\sigma^2/10$ の正規分布に従うことになる。従って、帰無仮説『 $\mu=200$ 』を検定することとなる。

統計量

$$T = \frac{\bar{X} - \mu}{u/\sqrt{n}}$$

が自由度 $n-1$ の t 分布に従うことが知られている。ただし、 n は標本の大きさであり、 $\bar{X}$ は標本平均、 u は不偏標準偏差である。いま、有意水準を $\alpha=0.05=5\%$ とすると、 $n=10$ であるので、棄却域は Excel の数式

$$=TINV(0.05,9)$$

で計算でき、計算結果 2.262159 より、棄却域は -2.262 以下ないし 2.262 以上と求まる。

一方、 $\bar{X}=196$ 、 $s=5$ であるので、統計量 T に標本データを代入して求まる値は Excel の数式

$$=(196-200)/(5/SQRT(10))$$

で計算でき、計算結果 -2.530 より、棄却域に入る。従って、帰無仮説は棄却される。すなわち、内容量の表示は誤りということになる。□

【例題 4.6】鉄筋を生産している会社がある。この会社の生産設備は正常に動いていれば、その直径が、平均 0.5 インチ、標準偏差 0.01 インチの正規分布に従った製品が生産できる。
あるとき、10 本の製品の平均直径を測定したところ、0.51 インチであった。設備に異常が発生したかどうかを有意水準 5% で検定せよ。

《解答》 $N(\mu, \sigma^2)$ に従って分布する母集団から取られた、大きさ n の標本の平均値 $\bar{X}$ は、 $N(\mu, \sigma^2/n)$ に従って分布することとなる。すなわち、大きさ n の標本が取られる度に計算される標本平均 $\bar{X}$ は、標本毎に微妙に異なった値となり、分布することとなる。

一方、有意水準 5 %での両側検定とは、帰無仮説（ここでは、「設備が正常である」とする仮説）が正しいとしたときに、ある統計量が従う分布の両側 5 %、すなわち小さい方 2.5 %ないし大きい方 2.5 %の部分（この部分を『棄却域』と呼ぶ）を考え、この統計量の値がこれらの領域に入ったとき、帰無仮説を棄却する（「設備に異常が発生した」と判断する）検定法である。

ここでの問題では、直径は $N(0.5, 0.01^2)$ に従って分布するので、 $n=10$ ずつ取り出して調べた標本の平均直径は、設備が正常であれば、 $N(0.5, 0.01^2/10)$ に従って分布することになる。正規分布 $N(0.5, 0.01^2/10)$ において両側 5 %となる部分（棄却域）は、Excel の数式

$$=NORMINV(0.025, 0.5, 0.01/SQRT(10))$$

と

$$=NORMINV(0.975, 0.5, 0.01/SQRT(10))$$

で計算され、0.4938 以下と 0.5062 以上の領域となる。すなわち、設備が正常に動いているときには、10 本の標本の平均がこの領域の値をとることはめったにないということになり、ある標本平均がこの領域の値をとったとすれば、帰無仮説が正しいとすると矛盾が生ずると判断し、帰無仮説を棄却すなわち設備は正常でないと判定される。

ここでの問題では、標本平均が 0.51 であり、棄却域に入るので、設備が正常であるとした帰無仮説が棄却され、設備に異常が発生したと判定される。□

【例題 4.7】ある蛍光灯の平均寿命は 1,600 時間といわれている。ところが、100 本のサンプルを調べたところ、その平均が 1,570 時間、標準偏差が 120 時間であった。このとき、「平均寿命が 1,600 時間である」という主張を有意水準 5 %と 1 %で両側検定せよ。

《解答》母集団が $N(\mu, \sigma^2)$ に従うとき、この母集団からとられた大きさ n の標本の標本平均 $\bar{X}$ は $N(\mu, \sigma^2/n)$ に従い、統計量 $Z=(\bar{X}-\mu)/(\sigma/\sqrt{n})$ は標準正規分布 $N(0,1)$ に従う。しかし、 σ が未知であるので、 $T=(\bar{X}-\mu)\{s/(n-1)^{1/2}\}$ を考えると、統計量 T は自由度 $n-1$ の t 分布に従うこととなる。

ここでの蛍光灯の寿命は、正規分布 $N(\mu, \sigma^2)$ に従って分布すると考えられる。そして、帰無仮説は当然「 $\mu=1,600$ 」である。

このとき、大きさ n の標本の標本平均を $\bar{X}$ 、標本標準偏差を s とすると

$$T = \frac{\bar{X}-1600}{s/\sqrt{n-1}}$$

は、自由度 $n-1$ の t 分布に従うことになる。

t 分布における棄却域は、自由度が $100-1=99$ であることから、有意水準 5 %のときは Excel 数式

$$=TINV(0.05, 99)$$

での計算結果より、-1.984 以下と 1.984 以上の領域になり、有意水準 1 %のときは Excel 数式

$$=TINV(0.01, 99)$$

での計算結果より、-2.626 以下と 2.626 以上の領域になる。

また、上式に標本平均 1570 と標本標準偏差 120 を代入すると Excel 数式

$$=(1570-1600)/(120/SQRT(99))$$

での計算結果は -2.487 となり、この値は有意水準 5 %のときは棄却域に入るが、有意水準 1 %のときは棄却域に入らない。したがって、帰無仮説（すなわち「全蛍光灯の平均寿命が 1,600 時間である」という仮説）は、有意水準 5 %のときは棄却されるが、有意水準 1 %のときは棄却されない。□

このように、有意水準によって棄却されたり棄却されなかったりする。このことから、分析後に有意水準を決める方法では恣意的になり易いので、一般に、有意水準は分析前に決めておかなければならない。

【例題 4.8】A 社の電球 100 個を選んで寿命時間を調べたところ、標本標準偏差 100 時間を得た。B 社の電球 75 個を選んで同様に調べたところ、標本標準偏差 105 時間を得た。両社の電球の寿命のばらつき、すなわち母分散に差はあるか。

《解答》二つの母集団それぞれが、分散の等しい正規分布 $N(\mu_1, \sigma^2)$ と $N(\mu_2, \sigma^2)$ に従うものとする。このとき、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、その不偏標準偏差が u_1 と u_2 であるものとする、統計量 $\chi^2_1 = (n_1 - 1)u_1^2 / \sigma^2$ と $\chi^2_2 = (n_2 - 1)u_2^2 / \sigma^2$ は、それぞれ自由度 $n_1 - 1$ と $n_2 - 1$ のカイ 2 乗分布に従うことになる。したがって

$$F = \frac{\{(n_1 - 1)u_1^2 / \sigma^2\} / (n_1 - 1)}{\{(n_2 - 1)u_2^2 / \sigma^2\} / (n_2 - 1)} = \frac{u_1^2}{u_2^2}$$

は、自由度 $n_2 - 1, n_1 - 1$ の F 分布に従う。

ここでは、両社の電球の寿命の分散に差がないという帰無仮説を設定する。また、電球の寿命は正規分布に従うと考えられる。したがって、上の F 統計量に関する性質が利用できる。

ここでの問題では、 $n_F = 100$ 、 $n_2 = 75$ 、 $s_1 = 100$ 、 $s_2 = 105$ であるので

$$u_1^2 = \frac{n_1 s_1^2}{n_1 - 1} = \frac{100 \cdot 100^2}{99} = 10101.0$$

$$u_2^2 = \frac{n_2 s_2^2}{n_2 - 1} = \frac{75 \cdot 105^2}{74} = 11174.0$$

となる。ここで、 $u_1^2 < u_2^2$ であるので、統計量 $F = u_2^2 / u_1^2$ を考えると、この統計量は自由度 74,99 の F 分布に従うことになる。また、有意水準（片側）5% で、自由度 74,99 の F 分布における棄却値は、Excel の数式

$$=FINV(0.05, 74, 99)$$

によって 1.42 と計算される。一方、標本データによる統計量 F の値は

$$F = \frac{u_2^2}{u_1^2} = \frac{11174.0}{10101.0} = 1.106$$

となる。したがって、ここでの帰無仮説は棄却されず、受容される。このことから、ここでの両母集団の母分散は等しいとみなせることになる。□

【例題 4.9】A 校の 16 人の学生の I.Q は、標本平均 107、標本標準偏差 10 であり、B 校の 14 人の学生の I.Q は、標本平均 112、標本標準偏差 8 であった。両校の学生の I.Q には差があるか。

《解答》正規分布 $N(\mu_1, \sigma_1^2)$ と $N(\mu_2, \sigma_2^2)$ に従う 2 つの母集団があり、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、その標本平均が $\bar{X}_1$ と $\bar{X}_2$ 、不偏標準偏差が u_1 と u_2 であるものとする。

このとき、まずすべき分析は、両母集団の分散に差があるかないかを検定することである。このために、帰無仮説： $\sigma_1^2 = \sigma_2^2$ を仮説検定しなければならない。

それぞれの不偏標準偏差が u_1 と u_2 であるものとする、統計量 $F = u_1^2 / u_2^2$ は、自由度 $n_1 - 1, n_2 - 1$ の F 分布に従う。

ここでの問題では、 $n_F = 16$ 、 $n_2 = 14$ 、 $s_1 = 10$ 、 $s_2 = 8$ であるので

$$F = \frac{u_1^2}{u_2^2} = \frac{n_1 s_1^2 / (n_1 - 1)}{n_2 s_2^2 / (n_2 - 1)} = \frac{16 \cdot 100 / 15}{14 \cdot 64 / 13} = 1.548$$

となる。また、有意水準 5% で、自由度 15,13 の F 分布の臨界値は Excel 数式

$$=FINV(0.05, 15, 13)$$

で計算され、結果は 2.53113 となるので、棄却域は 2.53 以上となる。したがって、ここでの帰無仮説は

棄却されず、受容される。このことから、ここでの両母集団の母分散は等しいとみなせることになる。

両母集団の母分散が等しいとみなせるとき、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、標本平均が $\bar{X}_1$ と $\bar{X}_2$ 、標本標準偏差が s_1 と s_2 であるものとする、統計量

$$T = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{\sqrt{(n_1 s_1^2 + n_2 s_2^2) / (n_1 + n_2 - 2)}} \sqrt{1/n_1 + 1/n_2}$$

は、自由度 $n_1 + n_2 - 2$ の t 分布に従う。

ここでの問題では、 $n_1=16$ 、 $n_2=14$ 、 $\bar{X}_1=107$ 、 $\bar{X}_2=112$ 、 $s_1=10$ 、 $s_2=8$ であり、帰無仮説として $\mu_1=\mu_2$ すなわち両校の学生の平均 I.Q.には差がないという仮説を設定する。このとき、統計量 T の値

$$T = \frac{107 - 112}{\sqrt{(16 \cdot 10^2 + 14 \cdot 8^2) / (16 + 14 - 2)}} \sqrt{1/16 + 1/14}$$

は Excel 数式

$$= (107 - 112) / (\text{SQRT}((16 \cdot 100 + 14 \cdot 64) / (16 + 14 - 2)) * \text{SQRT}(1/16 + 1/14))$$

で計算され、結果は -1.447 となる。一方、有意水準 5 % で、自由度 28 の t 分布の臨界値は Excel 数式

$$= \text{TINV}(0.05, 28)$$

で計算され、結果は 2.048409 となり、棄却域は -2.048 以下と 2.048 以上となるので、帰無仮説は棄却されない。すなわち、両校の学生の平均 I.Q.には差がないと判定される。□

【例題 4.10】A 社と B 社から販売されている刃物の硬度を調べたところ、次表の結果を得た。この結果より、両社の刃物の平均硬度は等しいといえるか。

ロックウエル硬度	40	41	42	43	44	45	46	47	48	計
A 社	0	0	2	5	11	7	1	0	0	26
B 社	0	2	4	4	5	4	4	1	1	25

《解答》まず、基本統計量を求めると

$$\text{A 社 : } n_1=26, \quad u_1=25, \quad \bar{X}_1=44.00, \quad u_1^2=0.96, \quad s_1^2=0.92$$

$$\text{B 社 : } n_2=25, \quad u_2=24, \quad \bar{X}_2=44.04, \quad u_2^2=3.46, \quad s_2^2=3.32$$

となる。母分散の同異を検定すると

$$F = u_2^2 / u_1^2 = 3.46 / 0.96 = 3.60$$

となり、有意水準 1 % で自由度 $u_2=24$ 、 $u_1=25$ の F 分布の臨界値は Excel 数式

$$= \text{FINV}(0.01, 24, 25)$$

で計算され、結果は 2.620254 となるので。棄却域は 2.62 以上となる。したがって、ここでの帰無仮説(母分散が等しいとする仮説)は棄却され、ここでの両母集団の母分散は等しくないと判断される。

両母集団の母分散が等しくないとき、両母集団からそれぞれ大きさ n_1 と n_2 の標本がとられ、標本平均が $\bar{X}_1$ と $\bar{X}_2$ 、不偏標準偏差が u_1 と u_2 であるとし、 $w_1=u_1^2/n_1$ 、 $w_2=u_2^2/n_2$ すると、統計量

$$T = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{\sqrt{w_1 + w_2}}$$

は、近似的に自由度

$$u = \frac{(w_1 + w_2)^2}{w_1^2 / (n_1 - 1) + w_2^2 / (n_2 - 1)}$$

の t 分布に従う。

帰無仮説として $\mu_1=\mu_2$ すなわち両社の刃物の平均硬度には差がないという仮説を設定する。このとき、 $w_1=u_1^2/n_1=0.96/26=0.0369$ 、 $w_2=u_2^2/n_2=3.46/25=0.1384$ となるので

$$\frac{(w_1+w_2)^2}{w_1^2/(n_1-1)+w_2^2/(n_2-1)} = \frac{(0.0369+0.1384)^2}{0.0369^2/25+0.1384^2/24}$$

は Excel 数式

$$=(0.0369+0.1384)^2/(0.0369^2/25+0.1384^2/24)$$

で計算され、結果 36.05 より、自由度は近似的に 36 となる。有意水準 5 %で、自由度 36 の t 分布の臨界値は Excel 数式

$$=TINV(0.05,36)$$

で計算され、結果は 2.028091 となるので、棄却域は、-2.03 以下ないし 2.03 以上であり

$$T = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{\sqrt{w_1 + w_2}} = \frac{44.00 - 44.04}{\sqrt{0.0369 + 0.1384}}$$

は Excel 数式

$$=(44.00-44.04)/SQRT(0.0369+0.1384)$$

で計算され、結果は-0.0955 となるので、帰無仮説は棄却されず、受容される。すなわち、両社の刃物の平均硬度には差がないと判定される。□