

Problem porównywania średnich w k zależnych lub niezależnych grupach można potraktować jako szczególnie problem testowania hipotezy liniowej.

Model: $Y = \varphi + \varepsilon$

gdzie $Y \in R^n$ jest wektorem obserwacji
 $\varphi \in \Omega \subset R^n$ jest wektorem średnich, o którym wiadomo, że należy do pewnej
właściwej podprzestrzeni liniowej Ω przestrzeni R^n , tzn. $\Omega \subset R^n$ i
 $\dim(\Omega) = p < n$
 $\varepsilon \in R^n$ jest losowym wektorem błędów o rozkładzie $\varepsilon \sim N_n(\mathbf{0}, \sigma^2 \mathbf{I})$

Uwaga: $Y \sim N_n(\varphi, \sigma^2 \mathbf{I})$

Interesuje nas problem testowania hipotezy

$$H_0: \varphi \in \omega \subset \Omega \quad \text{przeciwko} \quad H_1: \varphi \in \Omega - \omega,$$

gdzie ω jest pewną **właściwą podprzestrzenią liniową** przestrzeni Ω ,

czyli $\omega \subset \Omega$ i $\dim(\omega) = p - r < \dim(\Omega)$.

Uwaga. Zbiory ω i Ω są **podprzestrzeniami liniowymi**, czyli nie mogą to być zbiory ograniczone np. kule w R^n .

Po wyborze parametryzacji (bazy w przestrzeni R^n) rozważany model liniowy można zapisać w postaci gaussowskiego modelu liniowego

$$Y = X\beta + \varepsilon$$

$$\varepsilon \sim N(\mathbf{0}, \sigma^2 \mathbf{I})$$

X jest tu nielosową macierzą wymiaru (n, p) zwaną **macierzą planu eksperymentu** o kolumnach x_i , $i=0, \dots, p-1$ przy czym $x_0 = \mathbf{1}$, liniowo niezależnych tzn. $\text{rz}(X) = p$.

Wektor $\varphi = E(Y) \in \Omega \subset R^n$ przy czym $\Omega = \text{span}\{x_i; i=0, \dots, p-1\} = \text{Im}(X)$

Testować będziemy hipotezę liniową

$$H_0: G\beta = \mathbf{0} \quad \text{przeciwko} \quad H_1: G\beta \neq \mathbf{0},$$

gdzie G jest macierzą (r, p) $r < p$ pełnego rzędu tzn. $\text{rz}(G) = r$, która specyfikuje podprzestrzeń

$\omega = \text{Im}(X|_{\text{Ker}G})$ ($\dim(\omega) = p - r$) będącą obrazem obcięcia przekształcenia X do jądra przekształcenia G .

Przykład. Porównywanie k średnich

Rozważmy k niezależnych prób prostych

$$\begin{array}{c} Y_{11}, \dots, Y_{1n_1} \\ \vdots \\ Y_{k1}, \dots, Y_{kn_k} \end{array} \quad \text{pochodzących odpowiednio z rozkładów}$$

$N(m_1, \sigma^2), \dots, N(m_k, \sigma^2)$. Można więc napisać

$$Y_{ij} = m_i + \varepsilon_{ij}, \quad i=1, \dots, k, \quad j=1, \dots, n_i,$$

przy czym ε_{ij} są niezależnymi zmiennymi losowymi o identycznych rozkładach normalnych $N(0, \sigma^2)$.

Testowana jest hipoteza $H_0: m_1 = \dots = m_k$ wobec alternatywy $H_1: \sim H_0$

W zapisie macierzowym

$$\begin{bmatrix} Y_{11} \\ \vdots \\ Y_{1n_1} \\ Y_{21} \\ \vdots \\ Y_{2k_2} \\ \vdots \\ Y_{k1} \\ \vdots \\ Y_{kn_k} \end{bmatrix} = m_1 \begin{bmatrix} 1 \\ \vdots \\ 1 \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix} + m_2 \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix} + \dots + m_k \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} + \begin{bmatrix} \varepsilon_{11} \\ \vdots \\ \varepsilon_{1n_1} \\ \varepsilon_{21} \\ \vdots \\ \varepsilon_{2k_2} \\ \vdots \\ \varepsilon_{k1} \\ \vdots \\ \varepsilon_{kn_k} \end{bmatrix}$$

$$\mathbf{Y} = m_1 \mathbf{X}_1 + m_2 \mathbf{X}_2 + \dots + m_k \mathbf{X}_k + \boldsymbol{\varepsilon}$$

$$\mathbf{Y} = \mathbf{X} \mathbf{m} + \boldsymbol{\varepsilon}$$

W tym problemie

$$\Omega = \text{span} \left\{ \begin{bmatrix} 1 \\ \vdots \\ 1 \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \dots, \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} \right\}, \dim \Omega = k, \omega = \text{span} \left\{ \begin{bmatrix} 1 \\ \vdots \\ 1 \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} \right\}, \dim \omega = 1$$

Inna parametryzacja w problemie porównywania k średnich

Niech $m = \frac{1}{k} \sum_{i=1}^k m_i$ będzie średnią arytmetyczną grupowych wartości oczekiwanych. Oznaczając

$\alpha_i = m_i - m$, $i=1, \dots, k$ odchylenie wartości oczekiwanej w i -tej grupie od m otrzymujemy ograniczenie

$$\sum_{i=1}^k \alpha_i = 0.$$

Przy tej parametryzacji model można zapisać w postaci

$$Y_{ij} = m + \alpha_i + \varepsilon_{ij}, \quad i=1, \dots, k, j=1, \dots, n_i, \quad \text{przy czym} \quad \sum_{i=1}^k \alpha_i = 0.$$

Jest to tzw. parametryzacja z sigma ograniczeniami

Parametr m reprezentuje średni poziom (poziom odniesienia) dla wszystkich obserwowanych zmiennych. Parametr α_i jest odchyleniem wartości oczekiwanej w i tej grupie od poziomu odniesienia. Problem testowania równości k średnich sprowadza się w tym przypadku do testowania hipotezy

$H_0: \alpha_1 = \dots = \alpha_k$ wobec alternatywy $H_1: \sim H_0$ przy czym $\sum_{i=1}^k \alpha_i = 0$ (parametr m nas nie interesuje)

Przy tej parametryzacji przestrzenie Ω i ω wyglądają następująco

$$\Omega = \text{span} \left\{ \begin{bmatrix} 1 \\ \vdots \\ 1 \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix}, \begin{bmatrix} 1 \\ \vdots \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \dots, \begin{bmatrix} 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} \right\} \quad \dim \Omega = k, \quad \omega = \text{span} \left\{ \begin{bmatrix} 1 \\ \vdots \\ 1 \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} \right\}, \quad \dim \omega = 1$$

Uwaga: pierwszy z $k+1$ wektorów jest sumą pozostałych.

I jeszcze inna parametryzacja. Przypuśćmy, że jedna z grup (dla ustalenia uwagi pierwsza) jest wyróżniona np. jest to **grupa kontrolna**. Przyjmując parametry

$$\beta_1 = m_1$$

$$\beta_i = m_i - m_1, \quad i=2, \dots, k$$

$$\begin{bmatrix} X_{11} \\ \vdots \\ X_{1n_1} \\ X_{21} \\ \vdots \\ X_{2k_2} \\ \vdots \\ X_{k1} \\ \vdots \\ X_{kn_k} \end{bmatrix} = \beta_1 \begin{bmatrix} 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} + \beta_2 \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ \vdots \\ 1 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{bmatrix} + \dots + \beta_k \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 1 \\ \vdots \\ 1 \end{bmatrix} + \begin{bmatrix} \varepsilon_{11} \\ \vdots \\ \varepsilon_{1n_1} \\ \varepsilon_{21} \\ \vdots \\ \varepsilon_{2k_2} \\ \vdots \\ \varepsilon_{k1} \\ \vdots \\ \varepsilon_{kn_k} \end{bmatrix}$$

Testowanie hipotezy o równości średnich sprowadza się do testowania hipotezy $H_0: \beta_i = 0, \quad i=2, \dots, k$.

Parametr $\beta_1 = m_1$ wyznaczający „poziom odniesienia” nas nie interesuje.

Parametryzacja z grupą kontrolną jest przykładem **parametryzacji oszczędnej**. Taka parametryzacja jest stosowana w R

Zadanie1. Linie lotnicze zainteresowane są porównaniem przydatności 3 wyświetlaczy dla kontrolerów ruchu lotniczego. Każdy wyświetlacz był sprawdzany w 5 symulowanych warunkach zagrożenia. Trzydziestu wysoko wykwalifikowanych kontrolerów z podobnym doświadczeniem zawodowym zostało wybranych do badań. Kontrolerzy zostali losowo przypisani do klatek (wyświetlacz, warunki zagrożenia) po dwóch do każdej klatki. Mierzono czas w sekundach potrzebny do opanowania sytuacji awaryjnej. Wyniki przedstawia tabela

Wyświetlacz	Warunki zagrożenia				
	1	2	3	4	5
1	18	31	22	39	15
	16	35	27	36	12
2	13	33	24	35	10
	15	30	21	38	16
3	24	42	40	52	28
	28	46	37	57	24

Chcemy odpowiedzieć na pytania

- czy wyświetlacze różnią się między sobą?
- czy warunki zagrożenia wpływają na przydatność danego wyświetlacza?

#plik ANOVA.R

```
library(openxlsx)
dane <- read.xlsx("C:/Users/User/Documents/Dydaktyka/Modele
liniowe/Ćwiczenia/dane_anova.xlsx",colNames = TRUE)
```

```
dane1<-dane #przygotowanie danych do ANOVA
#colnames(data1)<-c("czas","wyświetlacz","warunki")
dane1$wyświetlacz <- factor (dane1$wyświetlacz )
levels(dane1$wyświetlacz) <- c("A","B","C")
dane1$warunki <- factor (dane1$warunki )
levels(dane1$warunki) <- c("I","II","III","IV","V")
attach(dane1)
```

#złe użycie funkcji aov

```
aov(czas~wyświetlacz+warunki,data=dane)
summary(aov(czas~wyświetlacz+warunki,data=dane))
summary(lm(czas~wyświetlacz+warunki,data=dane))
#funkcja nie rozpoznała zmiennych jakościowych (factor) i potraktowała je jako ilościowe
```

#poprawne użycie funkcji aov

```
wyn0<-(aov(czas~wyświetlacz+warunki,data=dane1)) #2 czynnikowa ANOVA efektów głównych (addytywny)
plot(wyn0) # wykresy diagnostyczne dla modelu efektów głównych
```

```
plot(czas~wyświetlacz+warunki,data=dane1)#wizualizacja średnich w grupach
summary(wyn0)
```

```
wyn<-aov(czas~wyświetlacz*warunki,data=dane1) #2 czynnikowa ANOVA z interakcją
plot(aov(czas~wyświetlacz*warunki,data=dane1))# wykresy diagnostyczne dla pełnego modelu czynnikowego
plot(czas~wyświetlacz*warunki,data=dane1)#wizualizacja średnich w grupach
```

```
summary(wyn)
```

#testowanie jednorodności wariancji

#Bartlett test of homogeneity of variances

```
bartlett.test(czas~ wyświetlacz,data=dane1)
bartlett.test(czas~ warunki,data=dane1)
bartlett.test(czas~interaction(wyświetlacz,warunki),data=dane1)
```

require(car)

#Levene's Test for Homogeneity of Variance (center = median)

```
leveneTest(czas~ wyświetlacz,data=dane1)
```

```
leveneTest(czas~ wyświetlacz,center=mean,data=dane1)
```

```
leveneTest(czas~ warunki,data=dane1)
```

```
leveneTest(czas~ wyświetlacz*warunki,data=dane1) # za mało obserwacji w klatkach
```

#Fligner-Killeen test of homogeneity of variances

```
fligner.test(czas~ wyświetlacz,data=dane1)
```

```
fligner.test(czas~ warunki,data=dane1)
```

```
fligner.test(czas~ interaction(wyświetlacz,warunki),data=dane1)
```

```
#interaction.plot(dane1$wyświetlacz,dane1$warunki,dane1$czas)
```

```
interaction.plot(wyświetlacz,warunki,czas)
```

```
interaction.plot(warunki,wyświetlacz,czas,xlab="warunki zagrożenia",ylab="czas reakcji")
```

#Analiza z użyciem funkcji lm()

X<-model.matrix(wyn) #macierz planu eksperymentu czynnikowego

dane_do_lm<-data.frame(matrix(c(dane1\$czas,X),nrow=30)) #budujemy ramkę danych do modelu lm

dokładając na początku odpowiedź

```
colnames(dane_do_lm)<-
```

```
c("y","x1","x2","x3","x4","x5","x6","x7","x8","x9","x10","x11","x12","x13","x14","x15")
```

```
dane_do_lm
```

	y	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12	x13	x14	x15
1	18	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0
2	16	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0
3	13	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0
4	15	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0
5	24	1	0	1	0	0	0	0	0	0	0	0	0	0	0	0
6	28	1	0	1	0	0	0	0	0	0	0	0	0	0	0	0
7	31	1	0	0	1	0	0	0	0	0	0	0	0	0	0	0
8	35	1	0	0	1	0	0	0	0	0	0	0	0	0	0	0
9	33	1	1	0	1	0	0	0	1	0	0	0	0	0	0	0
10	30	1	1	0	1	0	0	0	1	0	0	0	0	0	0	0
11	42	1	0	1	1	0	0	0	0	1	0	0	0	0	0	0
12	46	1	0	1	1	0	0	0	0	1	0	0	0	0	0	0
13	22	1	0	0	0	1	0	0	0	0	0	0	0	0	0	0
14	27	1	0	0	0	1	0	0	0	0	0	0	0	0	0	0
15	24	1	1	0	0	1	0	0	0	0	1	0	0	0	0	0
16	21	1	1	0	0	1	0	0	0	0	1	0	0	0	0	0
17	40	1	0	1	0	1	0	0	0	0	0	1	0	0	0	0
18	37	1	0	1	0	1	0	0	0	0	0	1	0	0	0	0
19	39	1	0	0	0	0	1	0	0	0	0	0	0	0	0	0
20	36	1	0	0	0	0	1	0	0	0	0	0	0	0	0	0
21	35	1	1	0	0	0	1	0	0	0	0	0	1	0	0	0
22	38	1	1	0	0	0	1	0	0	0	0	0	1	0	0	0
23	52	1	0	1	0	0	1	0	0	0	0	0	0	1	0	0
24	57	1	0	1	0	0	1	0	0	0	0	0	0	1	0	0
25	15	1	0	0	0	0	0	1	0	0	0	0	0	0	0	0
26	12	1	0	0	0	0	0	1	0	0	0	0	0	0	0	0
27	10	1	1	0	0	0	0	1	0	0	0	0	0	0	1	0
28	16	1	1	0	0	0	0	1	0	0	0	0	0	0	1	0
29	28	1	0	1	0	0	0	1	0	0	0	0	0	0	0	1
30	24	1	0	1	0	0	0	1	0	0	0	0	0	0	0	1

```
g1<-lm(y~x1+x2+x3+x4+x5+x6+x7+x8+x9+x10+x11+x12+x13+x14+x15-1,dane_do_lm) #pełny model
czynnikowy
```

```
summary(g1)
```

```
g2<-lm(y~x1+x2+x3+x4+x5+x6+x7-1,dane_do_lm) #model efektów głównych
```

```
summary(g2)
```

```
anova(g2,g1) #testowanie istotności interakcji
```

```
g3<-lm(y~x1+x4+x5+x6+x7-1,dane_do_lm) #model efektów głównych bez wyświetlacza
summary(g3)
anova(g3,g2) # test hipotezy H0: wyświetlacze nie różnią się między sobą w modelu efektów głównych
```

Krótki przegląd konserwatywnych testów post hoc.

Wszystkie rozważane testy konserwatywne prowadzą do następującej reguły

Średnie m_i i m_j są istotnie różne jeżeli $|\bar{Y}_i - \bar{Y}_j| > \text{wartość progowa}$ (właściwa dla danego testu)

Rozważane testy kontrolują prawdopodobieństwa różnych błędów

Test Fishera NIR (najmniejsza istotna różnica) (inaczej *LSD* least significant difference).

Algorytm

1. Przeprowadzić ANOVA w celu przetestowania $H_0: m_1 = \dots = m_k$ przeciwko alternatywie H_1 : co najmniej dwie średnie różnią się między sobą
2. Jeżeli nie ma podstaw do odrzucenia H_0 kończymy analizę.
3. Jeżeli H_0 została odrzucona, to definiujemy najmniejszą istotną różnicę *NIR* (*LSD*) pomiędzy średnimi próbkowymi, którą należy zaobserwować, aby uznać odpowiadające im średnie w odpowiednich populacjach za istotnie różne.
4. Dla wyspecyfikowanej wartości α w celu porównania m_i z m_j obliczamy *NIR* wg wzoru

$$NIR = t_{1-\frac{\alpha}{2}, \sum_{i=1}^k n_i} \sqrt{S_w^2 \left(\frac{1}{n_i} + \frac{1}{n_j} \right)},$$
 gdzie $t_{1-\frac{\alpha}{2}, \sum_{i=1}^k n_i}$ jest kwantylem rzędu $1 - \frac{\alpha}{2}$ rozkładu *t*-Studenta o $\sum_{i=1}^k n_i$ stopniach swobody a S_w^2 jest sumą kwadratów wewnątrz grup (z ANOVA)
5. Porównujemy wszystkie pary średnich próbkowych. Jeżeli $|\bar{Y}_i - \bar{Y}_j| > NIR$, to uznajemy, że m_i i m_j są istotnie różne.
6. Dla wszystkich porównań par prawdopodobieństwo błędu I rodzaju jest ustalone na poziomie α

Komputerowe symulacje wykazały, że jeśli stosujemy test Fishera w połączeniu z ANOVA (tak jak opisano powyżej), to poziom istotności złożonego testu porównań wielokrotnych jest w przybliżeniu równy poziomowi istotności testu *F* (procedura Fisher's protected LSD).

Jeżeli stosujemy test *NIR* samodzielnie (bez ANOVA), to kontrolujemy jedynie błąd przy pojedynczych porównaniach (**per comparison**). Odpowiada to wielokrotnemu stosowaniu testu istotności różnicy dwóch średnich opartego na statystyce *t*-Studenta, przy czym estymator wariancji opieramy na całej próbie (z powodu jednorodności wariancji w grupach) a nie tylko na obserwacjach z aktualnie porównywanych grup.

Test W Tukey'a oparty jest na studentyzowanym rozstępie $\sqrt{n} \frac{\bar{Y}_{\max} - \bar{Y}_{\min}}{S}$ pomiędzy średnimi próbkowymi.

Algorytm (dla jednakowo licznych grup)

1. Uporządkować średnie próbkowe $\bar{Y}_i$, $i=1, \dots, k$
2. m_i i m_j są istotnie różne jeżeli $|\bar{Y}_i - \bar{Y}_j| \geq W$, gdzie $W = q_\alpha(k, df) \sqrt{\frac{S_w^2}{n}}$, df jest liczbą stopni swobody w S_w^2 , n ilość obserwacji w każdej grupie, $q_\alpha(k, df)$ prawostronna wartość krytyczna studentyzowanego rozstępu (tablice).
3. Prawdopodobieństwo zaobserwowania fałszywie istotnej różnicy dla porównań parami jest równe α .

Jeżeli grupy nie są równoliczne, to zamiast W należy użyć $W_{ij} = q_\alpha(k, df) \sqrt{\frac{S_w^2}{2} \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}$

Test Tukey'a kontroluje błąd I rodzaju dla wszystkich porównań parami, tzn. prawdopodobieństwo (przy $H_0: m_1 = \dots = m_k$) zaobserwowania takiego układu średnich próbkowych $\bar{Y}_i$, $i=1, \dots, k$, dla którego przynajmniej jedna różnica pomiędzy średnimi m_i jest fałszywie uznana za istotną jest równe α . Jeżeli jest k grup, to test Tukey'a kontroluje łączny błąd I rodzaju dla $\binom{k}{2}$ porównań jednocześnie (**per experiment**)

Test Newmana-Keulsa jest modyfikacją testu Tukey'a wykorzystującą informację o ilości miejsc pomiędzy badanymi średnimi w uporządkowanym ciągu średnich.

Algorytm (dla jednakowo licznych grup)

1. Uporządkować średnie próbkowe $\bar{Y}_i$, $i=1, \dots, k$
2. Dla dwóch średnich $\bar{Y}_i$ i $\bar{Y}_j$ odległych o r miejsc odpowiadające im średnie m_i i m_j są istotnie różne jeżeli $|\bar{Y}_i - \bar{Y}_j| \geq W_r$, gdzie $W_r = q_\alpha(r, df) \sqrt{\frac{S_w^2}{n}}$, df jest liczbą stopni swobody w S_w^2 , n ilość obserwacji w każdej grupie, $q_\alpha(r, df)$ prawostronna wartość krytyczna studentyzowanego rozstępu (tablice). **Uwaga.** Przyjmujemy, że sąsiednie średnie odległe są o 2 a skrajne o k czyli $r_{ij} = |\text{ranga}(\bar{Y}_i) - \text{ranga}(\bar{Y}_j)| + 1$

Jeżeli grupy nie są równoliczne, to zamiast W_r należy użyć $W_{rij} = q_\alpha(r, df) \sqrt{\frac{S_w^2}{2} \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}$.

Test wielokrotnych rozstępów Duncana jest podobny do dwóch poprzednich, gdyż jest oparty na studentyzowanym rozstępie, lecz różni się od poprzednich tym, że zmienny jest poziom istotności przy poszczególnych porównaniach. Gdy średnie próbkowe zostaną uporządkowane i są odległe o r miejsc, to istotność różnicy odpowiadających im średnich w populacjach jest testowana na poziomie $1 - (1 - \alpha)^{r-1}$

Algorytm (dla jednakowo licznych grup)

1. Uporządkować średnie próbkowe $\bar{Y}_i$, $i=1, \dots, k$
2. Średnie m_i i m_j są istotnie różne jeżeli odpowiadające im średnie próbkowe odległe o r miejsc spełniają warunek $|\bar{Y}_i - \bar{Y}_j| \geq W'_r$, gdzie $W'_r = q'_\alpha(r, df) \sqrt{\frac{S_w^2}{n}}$, df jest liczbą stopni swobody

w S_w^2 , n ilość obserwacji w każdej grupie, $q'_\alpha(r, df)$ prawostronna wartość krytyczna testu Duncana (tablice).

Jeżeli grupy są w przybliżeniu równoliczne, to zamiast n należy użyć $\tilde{n} = \frac{k}{\frac{1}{n_1} + \frac{1}{n_2} + \dots + \frac{1}{n_k}}$.

Test Scheffe'go

1. Rozważmy dowolny kontrast $I_a = \sum_{i=1}^k a_i m_i$, $\sum_{i=1}^k a_i = 0$. Chcemy zweryfikować hipotezę

$$H_0: \forall \mathbf{a} \quad I_a = \sum_{i=1}^k a_i m_i = 0 \text{ wobec alternatywy } H_1: \exists \mathbf{a} \quad I_a \neq 0.$$

2. Rozważmy dowolny kontrast próbkowy $\hat{I}_a = \sum_{i=1}^k a_i \bar{X}_i$ dla którego nieobciążonym

$$\text{estymatorem wariancji jest } \hat{V}(\hat{I}) = \frac{1}{n-k} S_w^2 \sum_{i=1}^k \frac{1}{n_i} a_i^2$$

3. Kontrast I_a uznamy za istotny (istotnie różny od 0), gdy

$$\left| \frac{(n-k) \hat{I}_a^2}{(k-1) S_w^2 \sum_{i=1}^k \frac{a_i^2}{n_i}} \right| > F_{k-1, n-k, 1-\alpha},$$

gdzie $F_{k-1, n-k, 1-\alpha}$ jest kwantylem rzędu $1-\alpha$ rozkładu centralnego Snedecora Fishera F .

4. Prawdopodobieństwo zaobserwowania fałszywie istotnego kontrastu jest równe α .

Test Scheffe'go kontroluje łączny błąd większej liczby porównań niż test Tukey'a, więc jest bardziej konserwatywny (trudniej odrzucić H_0). W teście Scheffego na poziomie α prawdopodobieństwo zaobserwowania fałszywie istotnego kontrastu nie przekracza α .

Porównania testów. W poniższej tabeli zestawiono wartości progowe po przekroczeniu których różnice uznajemy za istotne. Rozważono porównanie 6 grup, przy czym próbka dla każdej grupy liczy $n=5$ elementów a $S_w^2=2451$.

Test	Liczba miejsc pomiędzy średnimi r				
	2	3	4	5	6
Fisher <i>NIR</i>	64.63	64.63	64.63	64.63	64.63
Tukey	96.75	96.75	96.75	96.75	96.75
Newman-Keuls	64.65	78.16	86.35	92.33	96.75
Duncan	64.65	67.97	69.74	71.29	72.62
Scheffe	196.31	196.31	196.31	196.31	196.31

Widać, że najbardziej konserwatywny jest test Scheffe'go a najmniej konserwatywny (czyli najbardziej czuły) test *NIR*. Z uwagi na kontrolę błędu godny polecenia jest test Tukey'a. Symulacje preferują raczej test Newmana-Keulusa. Z uwagi na czułość duże uznanie wśród praktyków wzbudził test Duncana.

Na uwagę zasługuje jeszcze test **Dunetta** wielokrotnych porównań z **wyróżnioną grupą kontrolną**

#testy post-hoc


```

TukeyHSD(wyn,"wyświetlacz")#test HSD dla czynnika wyświetlacz
plot(TukeyHSD(wyn,"wyświetlacz")) #test HSD dla czynnika wyświetlacz
TukeyHSD(wyn)#test HSD dla wszystkich czynników i ich interakcji
plot(TukeyHSD(wyn))

# testy z pakietu agricolae
library(agricolae)
??agricolae
# ustawienie parametru group=TRUE (default) powoduje otrzymanie jednorodnych grup
#parametr console = TRUE powoduje wyświetlenie wyników testu
out<-HSD.test(wyn0,"wyświetlacz",alpha=0.05,group=TRUE,console = TRUE) #z pakietu agricolae
#jednorodne grupy można przedstawić graficznie
plot(out) #z zakresem
plot(out,variation = "IQR") # z IQR
plot(out,variation = "SD") # z SD
HSD.test(wyn0,"wyświetlacz",alpha=0.05,group=FALSE,console = TRUE) #zwraca istotne różnice

out.SNK<-SNK.test(wyn,"wyświetlacz",alpha=0.05,group=TRUE,console = TRUE)#zwraca na konsolę
jednorodne grupy
plot(out.SNK)# rysunek jednorodnych grup
(SNK.test(wyn,"wyświetlacz",alpha=0.05,group=FALSE,console = TRUE))# zwraca istotne różnice

#test LSD Fishera z pakietu agricolae
out.LSD<-LSD.test(wyn,"wyświetlacz",p.adj="holm")
out.LSD #zwraca jednorodne grupy
plot(out.LSD)#st Studenta-Newmana-Keulsa z pakietu agricolae
out.LSD1<-LSD.test(wyn,"wyświetlacz",p.adj="holm",group=FALSE)
out.LSD1 #zwraca p-wartości dla porównań

#test Scheffego z pakietu agricolae
out.Scheffe<-scheffe.test(wyn,"wyświetlacz")
out.Scheffe
plot(out.Scheffe)
out.Scheffe1<-scheffe.test(wyn,"wyświetlacz",group=FALSE,console = TRUE)

out.duncan<-duncan.test(wyn, "wyświetlacz",group=TRUE,console=TRUE)
plot(out.duncan)
plot(out.duncan,variation = "IQR")
plot(out.duncan,variation = "SD")
duncan.test(wyn, "wyświetlacz",console=TRUE,group=FALSE)
print(duncan.test(wyn,"warunki",console=TRUE,group=FALSE))

```

Testowanie sekwencyjne i brzegowe w modelach liniowych.

Rozważmy model 2 czynnikowej ANOVA z czynnikami **A** i **B**. Używając składni języka R model efektów głównych z odpowiedzią **y** można zapisać w postaci **y~A+B** natomiast pełny model czynnikowy w postaci **y~A*B** co jest równoważne zapisowi **y~A+B+A:B**. Po przetestowaniu hipotezy H_0 o nieistotnym wpływie na odpowiedź wszystkich czynników polegającym na porównaniu modeli **y~A*B** lub **y~A+B** z modelem pustym **y~1** i odrzuceniu hipotezy zerowej pojawia się problem testowania istotności poszczególnych czynników i ich interakcji. **Funkcjonują 2 zasadnicze strategie testowania:**

- testowanie typu I – testowanie sekwencyjne (*sequential*),**
- testowanie typu III –testowanie brzegowe (*marginal*).**

Testowanie typu I (sekwencyjne) polega na sprawdzeniu istotności kolejnych zmiennych po dołączeniu do coraz większego modelu. W pierwszym kroku testowana jest istotność czynnika A przez porównanie (testem F ANOVA) modelu $y \sim A$ z modelem pustym $y \sim 1$. W drugim kroku sprawdzana jest istotność drugiego czynnika poprzez porównanie modelu $y \sim A+B$ z modelem $y \sim A$. Wreszcie w trzecim kroku sprawdzana jest istotność interakcji $A:B$ przez porównanie modelu $y \sim A+B+A:B$ z modelem $y \sim A+B$. Wyniki testowania mogą zależeć (i zwykle zależą) od kolejności zmiennych. Testowanie sekwencyjne dokonywane jest z użyciem funkcji `anova()`

Testowanie typu III (brzegowe) polega na sprawdzeniu istotności zmiennych przez porównanie modelu pełnego z modelem po usunięciu efektu danej zmiennej. Istotność czynnika A jest testowana przez porównanie (testem F) modelu $y \sim A+B+A:B$ z modelem $y \sim B+A:B$. podobnie istotność czynnika B jest sprawdzana przez porównanie modelu $y \sim A+B+A:B$ z modelem $y \sim A+A:B$. Wreszcie istotność interakcji jest testowana przez porównanie modelu $y \sim A+B+A:B$ z $y \sim A+B$. Wyniki testowania nie zależą od kolejności testowania zmiennych.

Testowanie brzegowe dokonywane jest przy użyciu funkcji `summary()`

Eksperyment numeryczny- wygenerujemy 100 obserwacji z 3 wymiarowego rozkładu normalnego

$N(\mathbf{m}, \Sigma)$ z $\mathbf{m} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}$ i $\Sigma = \begin{bmatrix} 1 & 0.25 & 0.25 \\ 0.25 & 1 & 0.75 \\ 0.25 & 0.75 & 1 \end{bmatrix}$. Macierz kowariancyjna jest w tym przypadku równa

macierzy korelacyjnej. Wybraliśmy macierz korelacyjną tak aby 2 i 3 składowa były silnie skorelowane.

Dokonując dekompozycji Choleskiego macierzy $\Sigma = \mathbf{U}^T \mathbf{U} = \mathbf{L} \mathbf{L}^T$, gdzie $\mathbf{U}$ jest macierzą trójkątną górną a $\mathbf{L}$ trójkątną dolną za pomocą funkcji `chol()` przekształcamy 100 elementową próbkę prostą z rozkładu $N(\mathbf{0}, \mathbf{I})$ w 100 elementową próbkę prostą z rozkładu $N(\mathbf{m}, \Sigma)$ za pomocą transformacji $\mathbf{Y} = \mathbf{U}^T \mathbf{X}$.

```
R<-matrix(c(1,0.25,0.25,0.25,1,0.75,0.25,0.75,1),nrow=3)
U<-chol(R) # dekompozycja Choleskiego macierzy korelacyjnej R
t(U)%*%U
#wyznaczanie pierwiastka symetrycznego z macierzy R=P.D.t(P) zamiast dekompozycji Choleskiego
ds<-eigen(R)
P<-ds$vectors
D<-diag(sqrt(ds$values),3,3)
#D<-diag(ds$values,3,3)
UU<-P%*%D%*%t(P)
UU
UU%*%UU

XYZ<-rnorm(300) # generujemy 300 realizacji zmiennej N(0,1)
XYZ<-matrix(XYZ,ncol=3) # przekształcamy wektor o 300 elementach w macierz typu (100,3)
XYZ<-t(U)%*%t(XYZ) # przekształcamy macierz przez transformację Y=t(U)X
XYZ<-data.frame(t(XYZ)) # przekształcamy macierz w ramkę danych
colnames(XYZ)<-c("Y", "X", "Z")
```

```
attach(XYZ)
cor(XYZ) # próbkowa macierz kowariancji z zaokrągleniem do 3 miejsc
summary(lm(Y~X+Z),data=XYZ) # testowanie typu III – brzegowe
anova(lm(Y~X+Z)) #testowanie typu I – sekwencyjne
anova(lm(Y~Z+X))# testowanie typu I - sekwencyjne w innej kolejności
```

```
#testowanie sekwencyjne istotności wpływu zmiennej Z na odpowiedź Y
g0<-lm(Y~1,data=XYZ) # tworzymy model pusty
gz<-lm(Y~Z,data=XYZ)#tworzymy model z udziałem zmiennej Z
summary(gz)
anova(g0,gz)
```

```
gzx<-lm(Y~Z+X,data=XYZ)
summary(gzx)
anova(gz,gzx)
anova(g0,gzx)
anova(gzx)
```

Podsumowując, w rozważanym przykładzie jest zależność pomiędzy Y a X i Y i Z. Zależność ta jest widoczna, gdy porównujemy model pusty z modelem z tylko jedną zmienną objaśnianą. Słabiej ta zależność jest widoczna, gdy model pusty porównujemy z modelem z dwiema zmiennymi objaśniającymi (model pełny) ponieważ każda ze zmiennych objaśniających wyraża tą samą część zmienności zmiennej objaśnianej Y. Widać to w testach brzegowych, które pokazują że zmienna X ma nieistotny wpływ na Y jeżeli jest uwzględniona w modelu zmienna Z (i symetrycznie). W sytuacji gdy zmienne objaśniane są od siebie zależne należy starannie wybrać model

ANCOVA

W pewnym eksperymencie rolniczym porównywano plony Y trzech różnych odmian kukurydzy. Mierzono również opady deszczu X. Zakładając, że odmiany zostały losowo przypisane do pól, zbadano, czy odmiany istotnie różnią się między sobą wysokością plonu.

Odmiana I		Odmiana II		Odmiana III	
X	Y	X	Y	X	Y
14	68,5	30	81,5	15	70,0
36	83,0	33	85,2	11	84,0
31	83,0	39	87,1	42	75,2
27	66,5	26	69,3	14	67,5
15	58,1	43	73,5	15	66,0
17	82,4	31	65,5	42	79,1
		18	73,4	26	75,5
		14	56,1		

Poniżej jest osadzony plik z danymi

	1	2	3	4	5
	X	Y	Odmiana	Korekta	Ykor
1	14	68,5	1		76,25367
2	36	83	1		76,13247
3	31	83	1		79,45547
4	27	66,5	1		65,61387
5	15	58,1	1		65,18907
6	17	82,4	1		88,15987
7	30	81,5	2		78,60508
8	22	85,2	2		80,20008

Wykonanie

Zmienna zależna: Y (Plon)

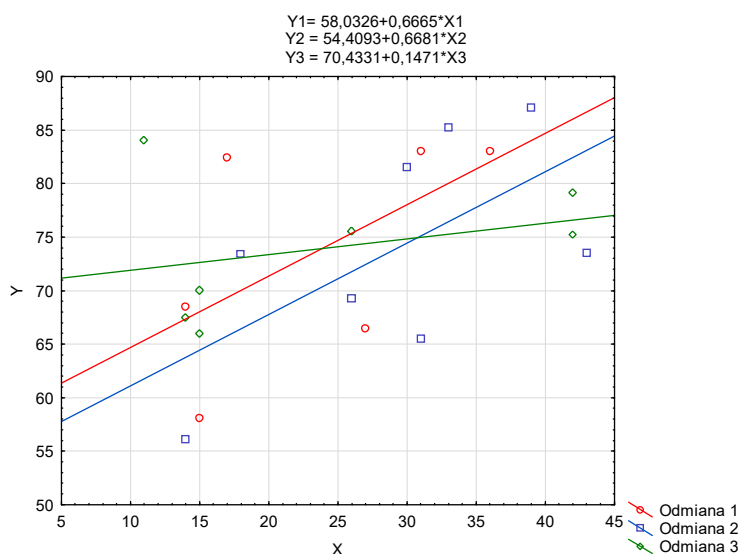
Pred. jakościowe: Odmiana

Pred. ilościowe: X (nawodnienie)

Klikamy przycisk **Efekty międzygrupowe–użyj efektów dostosowanych dla układu międzygrupowego** i wybieramy oba czynniki Odmiana, X (nawodnienie), klikamy przycisk **pełny układ czynnikowy**

Najpierw zaczniemy od modelu jednakowych nachyleń

Efekt	Jednowymiarowe testy istotności, wielkości efektów i moce dla Y (Arkusz in C:\User Parametryzacja z sigma-ograniczeniami Dekompozycja efektywnych hipotez							
	SS	Stopnie swobody	MS	F	p	Eta-kwadrat częściowe	Niecentralność	Moc obserwowana (alfa=0,05)
Wyraz wolny	11706	1	11706	162,42	0,0000	0,9052	162,42	1,0000
X	366	1	366	5,08	0,0377	0,2300	5,08	0,5657
Odmiana	22	2	11	0,15	0,8601	0,0176	0,30	0,0698
Błąd	1225	17	72					



Etykieta	Etykiety kolumn (Arkusz in C:\Documents and Settings\Adam Ćmiel Etykiety kolumn macierzy eksperymentu X						
	Kolumna	Zm.	Poziom Zm.	wzgl. Poziom	Zm.	Poziom Zm.	wzgl. Poziom
Wyraz wolny	1						
Odmiana	2	Odmiana	1	3			
Odmiana	3	Odmiana	2	3			
X	4	X					
Odmiana*X	5	Odmiana	1	3	X		
Odmiana*X	6	Odmiana	2	3	X		

Efekt	Jednowymiarowe testy istotności, wielkości efektów i moce dla Y (Arkusz in Dokl Parametryzacja z sigma-ograniczeniami Dekompozycja efektywnych hipotez							
	SS	Stopnie swobody	MS	F	p	Eta-kwadrat cząstkowe	Niecentr alność	Moc obserwowana (alfa=0.05)
Wyraz wolny	9582	1	9582	133,4	0,0000	0,899	133,4	1,000
Odmiana	153	2	77	1,1	0,3692	0,124	2,1	0,202
X	463	1	463	6,4	0,0227	0,300	6,4	0,660
Odmiana*X	148	2	74	1,0	0,3815	0,121	2,1	0,196
Błąd	1078	15	72					

Efekt	Oceny parametrów (Arkusz in C:\Documents and Settings\Adam Ćmiel\Moje d Parametryzacja z sigma-ograniczeniami									
	Poziom Efekt	Kolumna	X Param.	X Bł. std.	X t	X p	-95,00 % Gr.ufn.	+95,00 % Gr.ufn.	X Beta (β)	X Bł.Std. β
Wyraz wolny		1	60,9583	5,278	11,55	0,000	49,7	72,21		
Odmiana	1	2	-2,9257	7,888	-0,37	0,716	-19,7	13,89	-0,264	0,712
Odmiana	2	3	-6,5490	7,833	-0,84	0,416	-23,2	10,15	-0,635	0,759
X		4	0,4939	0,195	2,54	0,023	0,1	0,91	0,600	0,236
Odmiana*X	1	5	0,1726	0,306	0,56	0,581	-0,5	0,82	0,402	0,713
Odmiana*X	2	6	0,1742	0,271	0,64	0,530	-0,4	0,75	0,483	0,752

Model $Y_{ij} = m + \alpha_i + \beta X_{ij} + \gamma_i X_{ij} + \varepsilon_{ij}$

$$\sum_{i=1}^3 \alpha_i = 0, \quad \sum_{i=1}^3 \gamma_i = 0$$

$$\varepsilon_{ij} \sim N(0, \sigma^2)$$

Parametryzacja programu Statistica

$$\begin{bmatrix} Y_1 \\ Y_2 \\ Y_3 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} m + \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} \alpha_1 + \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix} \alpha_2 + \begin{bmatrix} X_1 \\ X_2 \\ X_3 \end{bmatrix} \beta + \begin{bmatrix} X_1 \\ 0 \\ -X_3 \end{bmatrix} \gamma_1 + \begin{bmatrix} 0 \\ X_2 \\ -X_3 \end{bmatrix} \gamma_2 + \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \varepsilon_3 \end{bmatrix}$$

m	α_1	α_2	β	γ_1	γ_2
60,95835	-2,92574	-6,54900	0,49387	0,17259	0,17419

$$\bar{X} = 25,66667$$

$$Y_{kor} = (Y - (\beta + \gamma_1) * (X - \bar{X} \cdot \mathbf{1})) * (\text{Odm}=1) + (Y - (\beta + \gamma_2) * (X - \bar{X} \cdot \mathbf{1})) * (\text{Odm}=2) + \\ + (Y - (\beta + \gamma_1 + \gamma_2) * (X - \bar{X} \cdot \mathbf{1})) * (\text{Odm}=3)$$

Po wyeliminowaniu wpływu zmiennej towarzyszącej wykonujemy procedurę ANOVA dla skorygowanej zmiennej Y_{kor}

Jednowymiarowe testy istotności, wielkości efektów i moce dla Y_{kor} (Arkusz in C:\Docume Parametryzacja z sigma-ograniczeniami Dekompozycja efektywnych hipotez								
Efekt	SS	Stopnie swobody	MS	F	p	Eta-kwadrat cząstkowe	Niecentralność	Moc obserwowana (alfa=0,05)
Wyraz wolny	112302,3	1	112302,3	1875,940	0,000	0,990	1876	1,000
Odmiana	49,8	2	24,9	0,416	0,666	0,044	1	0,107
Błąd	1077,6	18	59,9					

i stwierdzamy, że odmiany nie różnią się istotnie między sobą.

```
library(openxlsx)
dane <- read.xlsx("C:/Users/User/Documents/Dydaktyka/Modele liniowe/Ćwiczenia/ANCOVA.xlsx", colNames = TRUE)
dane1 <- dane #przygotowanie danych do ANCOVA
colnames(dane1) <- c("X", "Y", "Odmiana")
dane1$Odmiana <- factor(dane1$Odmiana)
levels(dane1$Odmiana) <- c("I", "II", "III")
attach(dane1)
Dane zostały przygotowane do użycia
```

```
lapply(split(dane1, dane1$Odmiana), summary)
Wyznaczenie podstawowych charakterystyk zmiennych X i Y dla Odmian
```

```
$I
      X      Y  Odmiana
Min. :14.00 Min. :58.10 I :6
1st Qu.:15.50 1st Qu.:67.00 II :0
Median :22.00 Median :75.45 III:0
Mean :23.33 Mean :73.58
3rd Qu.:30.00 3rd Qu.:82.85
Max. :36.00 Max. :83.00
```

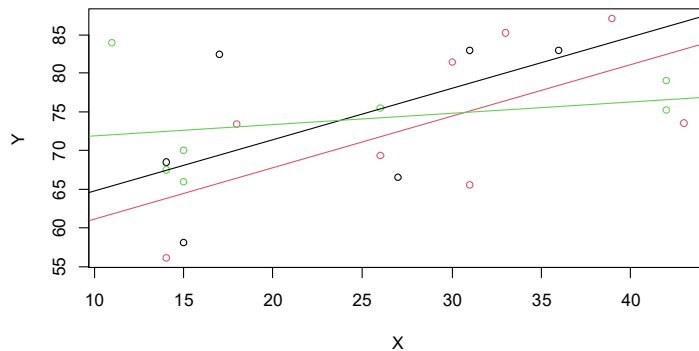
```
$II
      X      Y  Odmiana
Min. :14.00 Min. :56.10 I :0
1st Qu.:24.00 1st Qu.:68.35 II :8
Median :30.50 Median :73.45 III:0
Mean :29.25 Mean :73.95
3rd Qu.:34.50 3rd Qu.:82.42
Max. :43.00 Max. :87.10
```

```
$III
      X      Y  Odmiana
Min. :11.00 Min. :66.00 I :0
1st Qu.:14.50 1st Qu.:68.75 II :0
Median :15.00 Median :75.20 III:7
Mean :23.57 Mean :73.90
3rd Qu.:34.00 3rd Qu.:77.30
Max. :42.00 Max. :84.00
```

Wyznaczymy trzy linie regresji $y \sim x$ dla poszczególnych odmian z użyciem funkcji `subset(zbiór, kryterium, zmienne)` - składnia funkcji

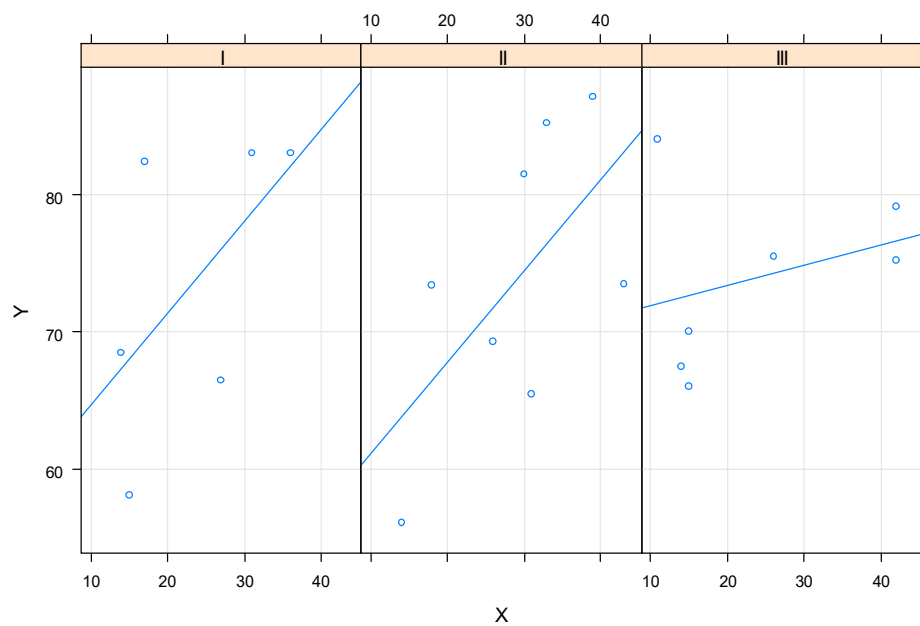
```
(m1<-lm(Y~X,data=subset(dane1,Odmiana=="I")))  
(m2<-lm(Y~X,data=subset(dane1,Odmiana=="II")))  
(m3<-lm(Y~X,data=subset(dane1,Odmiana=="III")))
```

```
plot(Y~ X,data=dane,col=Odmiana)  
abline(coef(m1),col=1)  
abline(coef(m2),col=2)  
abline(coef(m3),col=3)
```

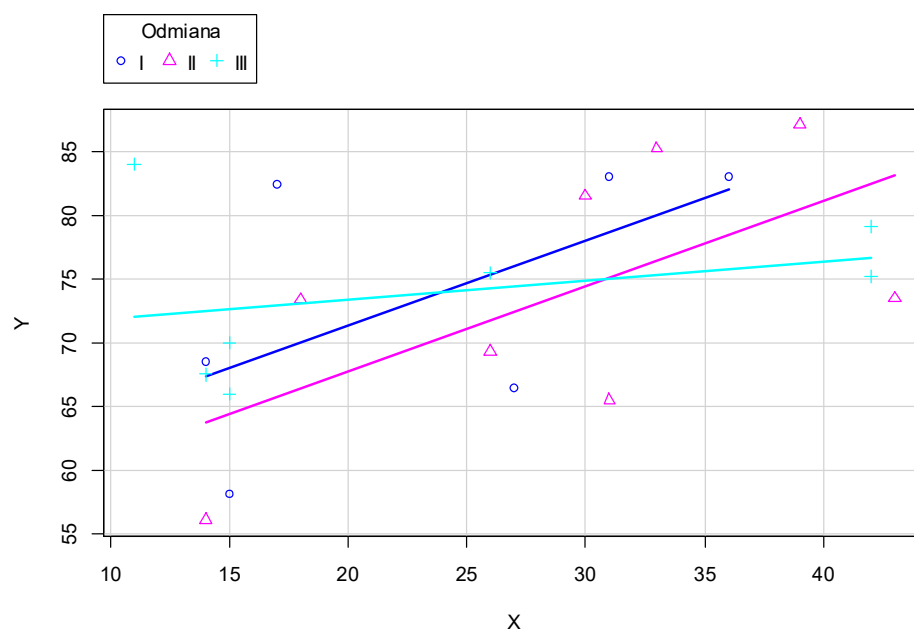


```
wyn0<-aov(Y~Odmiana*X,data=dane1)# ANCOVA różnych nachyleń  
plot(wyn0) # wykresy diagnostyczne dla modelu efektów głównych  
plot(Y~Odmiana*X,data=dane1)#wizualizacja średnich w grupach  
summary(wyn0)
```

```
library(lattice)  
xyplot(Y~X|Odmiana,data=dane1,type=c("p","g","r"))
```



```
library(car)
#scatterplot(Y~X|Odmiana,data=dane1,smooth=FALSE)
sp(Y~X|Odmiana,data=dane1,smooth=FALSE) # udoskonalona funkcja scatterplot
```




```
XX<-model.matrix(wyn0)#macierz planu eksperymentu czynnikowego
dane_do_lm<-data.frame(matrix(c(dane1$Y,XX),nrow=21))#budujemy ramkę danych do modelu lm dokładając
na początku odpowiedź
colnames(dane_do_lm)<-
c("Y","Intercept","Odmiana_II","Odmiana_III","X","Odmiana_II:X","Odmiana_III:X")
```

```
g0<-lm(Y~1,dane_do_lm) #model pusty
summary(g0)
```

```
g1<-lm(Y~Odmiana_II+Odmiana_III+X+Odmiana_II:X+Odmiana_III:X,dane_do_lm)#pełny model
czynnikowy
summary(g1)
```

```
g2<-lm(Y~Odmiana_II+Odmiana_III+X,dane_do_lm)
summary(g2)
anova(g2,g1)# brzegowy test istotności interakcji
```

```
g3<-lm(Y~Odmiana_II+Odmiana_III,dane_do_lm) #model jednakowych nachyleń
summary(g3)
anova(g3,g2)# brzegowy test istotności czynnika X w modelu jednakowych nachyleń
```

```
g4<-lm(Y~Odmiana_II+Odmiana_III+Odmiana_II:X+Odmiana_III:X,dane_do_lm)
summary(g4)
anova(g4,g1)#brzegowy test istotności dla czynnika X- porównać z testem t-Studenta
```

```
g5<-lm(Y~X+Odmiana_II:X+Odmiana_III:X,dane_do_lm)
summary(g5)
anova(g5,g1)#brzegowy test istotności dla czynnika Odmiana OK
```

```
g6<-lm(Y~X,dane_do_lm)
summary(g6)
anova(g0,g6)# test istotności dla zmiennej X
```

```
(model.matrix(g1))
```

```
#korekta eliminacja wpływu zmiennej towarzyszącej
Y_kor<-Y
n=length(Y)
for (i in 1:n) {
  if(Odmiana[i]=="I") Y_kor=Y-(coef(m1)[2]*(X-mean(X)))
  if(Odmiana[i]=="II") Y_kor=Y-(coef(m2)[2]*(X-mean(X)))
  if(Odmiana[i]=="III") Y_kor=Y-(coef(m2)[2]*(X-mean(X)))
}
dane2<-data.frame(Y_kor,Odmiana)
#ramka danych dla skorygowanej zmiennej Y
```