

Universidade Estadual de Santa Cruz
Departamento de Ciências Exatas e Tecnológicas
CET076 - Metodologia e Estatística Experimental
Curso de Agronomia

Notas de aulas expandidas.

Prof. José Cláudio Faria

Ilhéus – Bahia

Índice

NOTAS DO AUTOR	9
LITERATURA RECOMENDADA	10
RECURSOS DISPONÍVEIS NA WWW	10
LABORATÓRIOS VIRTUAIS DISPONÍVEIS NA INTERNET	10
SITE PARA ANÁLISES ON-LINE	10
EXEMPLOS DE RECURSOS DISPONÍVEIS NA WWW	11
SIMBOLOGIA ADOTADA NO CURSO	14
1. CALCULADORAS E APROXIMAÇÕES EM ESTATÍSTICA	15
1.1. CALCULADORA ADEQUADA	15
1.2. COMENTÁRIOS SOBRE OS RECURSOS BÁSICOS	15
1.3. APROXIMAÇÕES	15
1.4. UM TESTE	16
1.5. O QUE NÃO DEVE SER FEITO	17
2. REVISÃO DOS CURSOS PRELIMINARES	18
2.1. MÉDIA ARITMÉTICA	18
2.1.1. O QUE É	18
2.1.2. O QUE QUANTIFICA	18
2.1.3. SIMBOLOGIA E CÁLCULO	19
2.1.3.1. Cálculo	19
2.1.4. UNIDADE DE EXPRESSÃO	19
2.2. VARIÂNCIA	19
2.2.1. O QUE É	19
2.2.2. O QUE QUANTIFICA	19
2.2.3. SIMBOLOGIA E CÁLCULO	20
2.2.3.1. Cálculo	20
2.2.4. UNIDADE DE EXPRESSÃO	20
2.2.5. CONCEITO	20
2.2.6. FORMAS DE CÁLCULO	21
2.3. DESVIO PADRÃO	22
2.3.1. O QUE É	22
2.3.2. O QUE QUANTIFICA	22
2.3.3. SIMBOLOGIA E CÁLCULO	22
2.3.3.1. Cálculo	22
2.3.4. UNIDADE DE EXPRESSÃO	22
2.4. DESVIO PADRÃO RELATIVO E COEFICIENTE DE VARIAÇÃO	22
2.4.1. O QUE SÃO	22
2.4.2. O QUE QUANTIFICAM	23
2.4.3. SIMBOLOGIA E CÁLCULOS	23
2.4.3.1. Cálculos	23
2.4.4. JUSTIFICATIVAS PARA O USO E UNIDADES DE EXPRESSÃO	23
2.5. DEMONSTRAÇÕES	25
2.6. COVARIÂNCIA	27
2.6.1. O QUE É	27
2.6.2. O QUE QUANTIFICA	28

2.6.3.	SIMBOLOGIA E CÁLCULO	28
2.6.3.1.	Cálculo	28
2.6.4.	UNIDADE DE EXPRESSÃO	29
2.6.4.1.	Conceito	29
2.6.5.	EXEMPLOS DE CÁLCULO E VISUALIZAÇÃO DAS ASSOCIAÇÕES	30
2.6.5.1.	Variáveis com associação positiva e elevada	30
2.6.5.2.	Variáveis com associação negativa e elevada	30
2.6.5.3.	Variáveis não associadas	31
2.7.	TEOREMA CENTRAL DO LIMITE	31
2.7.1.	O QUE É	31
2.7.2.	O QUE SIGNIFICA	31
2.7.3.	COMO É USADO	32
2.8.	TESTE DE HIPÓTESES	33
2.8.1.	HIPÓTESE: O QUE É	33
2.8.2.	TESTE DE HIPÓTESES: O QUE É	33
2.8.3.	TIPOS DE HIPÓTESES	33
2.8.4.	TIPOS DE ERROS	33
2.9.	DISTRIBUIÇÃO F	34
2.9.1.	O QUE É	34
2.9.2.	O QUE SIGNIFICA	34
2.9.3.	COMO É USADA	37
2.9.4.	EXATIDÃO E PRECISÃO	38
2.9.5.	EXEMPLO BÁSICO DE APLICAÇÃO DA DISTRIBUIÇÃO F - COMPARAÇÃO DE PRECISÃO	39
2.9.5.1.	Mecanismo de decisão	40

3. ANÁLISE DE VARIÂNCIA **44**

3.1.	INTRODUÇÃO	44
3.2.	CONCEITOS E USO	44
3.2.1.	O QUE É?	44
3.2.2.	PARA QUE É USADA?	44
3.2.3.	QUAL DECISÃO É POSSÍVEL TOMAR?	44
3.2.4.	EXEMPLO	46
3.2.4.1.	Teste de hipóteses	46
3.2.4.2.	Procedimentos para a análise	46
3.2.5.	PRESSUPOSTOS DA ANÁLISE DE VARIÂNCIA	51
3.2.6.	DEMONSTRAÇÃO DA APLICAÇÃO DO TEOREMA CENTRAL DO LIMITE (TCL) NA ANOVA	52

4. NOÇÕES BÁSICAS DE EXPERIMENTAÇÃO **54**

4.1.	INTRODUÇÃO	54
4.2.	PÚBLICO	54
4.3.	PRINCIPAIS CONCEITOS	54
4.4.	A ORIGEM AGRÍCOLA	55
4.5.	PRINCÍPIOS BÁSICOS DA EXPERIMENTAÇÃO	56
4.5.1.	REPETIÇÃO	56
4.5.2.	CASUALIZAÇÃO	57
4.5.3.	CONTROLE LOCAL	58
4.6.	CONTROLE DE QUALIDADE DE EXPERIMENTOS	59
4.7.	TIPOS DE ERROS EM EXPERIMENTOS	60
4.7.1.	PRINCIPAIS FONTES DE ERRO E RESPECTIVOS CUIDADOS	61
4.7.1.1.	Heterogeneidade das condições ambientais	61
4.7.1.2.	Heterogeneidade do material experimental	61
4.7.1.3.	Condução diferenciada das unidades experimentais	61
4.7.1.4.	Competição intraparcelar	61

4.7.1.5.	Competição interparcelar	61
4.7.1.6.	Pragas, doenças e acidentes	61
4.8.	PLANEJAMENTO DE EXPERIMENTOS	61

5. DELINEAMENTO INTEIRAMENTE CASUALIZADO - DIC **62**

5.1.	INTRODUÇÃO	62
5.2.	PRINCÍPIOS UTILIZADOS	62
5.2.1.	REPETIÇÃO	62
5.2.2.	CASUALIZAÇÃO	62
5.2.3.	VANTAGENS E DESVANTAGENS	62
5.2.3.1.	Vantagens	62
5.2.3.2.	Desvantagens	63
5.3.	MODELO ESTATÍSTICO	63
5.4.	ESQUEMA DE CASUALIZAÇÃO DOS TRATAMENTOS	63
5.5.	COLETA DE DADOS	64
5.6.	ANÁLISE DE VARIÂNCIA	64
5.6.1.	ESQUEMA DA ANÁLISE DE VARIÂNCIA	64
5.6.2.	TESTE DE HIPÓTESES	64
5.7.	EXEMPLO COM UM MESMO NÚMERO DE REPETIÇÕES	65
5.7.1.	RESÍDUO	66
5.7.2.	O COEFICIENTE DE VARIAÇÃO E SUA INTERPRETAÇÃO	66
5.7.3.	TESTES DE COMPARAÇÃO DE MÉDIAS MÚLTIPLAS	67
5.7.4.	HIPÓTESES PARA OS CONTRASTES	67
5.7.5.	DESDOBRAMENTO DOS GL ASSOCIADOS A TRATAMENTOS EM CONTRASTES ORTOGONAIS	67
5.8.	EXEMPLO COM NÚMERO DIFERENTE DE REPETIÇÕES	68
5.8.1.	DESDOBRAMENTO DOS GL ASSOCIADOS A TRATAMENTOS EM CONTRASTES ORTOGONAIS	69
5.8.2.	ESTIMAÇÃO E TESTE DE HIPÓTESES PARA OS CONTRASTES	70
5.9.	CONSIDERAÇÕES FINAIS	71
5.10.	DEMONSTRAÇÕES E ILUSTRAÇÕES	72

6. TESTES DE COMPARAÇÃO DE MÉDIAS MÚLTIPLAS **75**

6.1.	INTRODUÇÃO	75
6.2.	O FUNDAMENTO DOS TESTES	75
6.3.	OS TESTES	76
6.3.1.	TESTE DE DUNCAN	77
6.3.1.1.	Obtenção da dms	77
6.3.1.2.	Aplicação do teste	77
6.3.1.2.1.	Para contrastes que abrangem 4 médias	77
6.3.1.2.2.	Para contrastes que abrangem 3 médias	77
6.3.1.2.3.	Para testar contrastes que abrangem 2 médias	78
6.3.1.3.	Apresentação dos resultados e conclusão	78
6.3.2.	TESTE DE DUNNETT	79
6.3.2.1.	Obtenção da dms	79
6.3.2.2.	Aplicação do teste	79
6.3.2.3.	Apresentação dos resultados e conclusão	80
6.3.3.	TESTE DE TUKEY	80
6.3.3.1.	Obtenção da dms	80
6.3.3.2.	Aplicação do teste	81
6.3.3.3.	Apresentação dos resultados e conclusão	82
6.3.4.	TESTE DE STUDENT – NEWMAN – KEULS (SNK)	82
6.3.4.1.	Obtenção da dms	82
6.3.4.2.	Aplicação do teste	82
6.3.4.2.1.	Para contrastes que abrangem 4 médias	82

6.3.4.2.2.	Para contrastes que abrangem 3 médias	83
6.3.4.2.3.	Para contrastes que abrangem 2 médias	83
6.3.4.3.	Apresentação dos resultados e conclusão	84
6.3.5.	TESTE DE SCHEFFÉ	84
6.3.5.1.	Obtenção da dms	84
6.3.5.2.	Teste de Scheffé - médias de tratamentos	84
6.3.5.3.	Teste de Scheffé - grupos de médias de tratamentos	85
6.4.	EXEMPLO DE APLICAÇÃO EM EXPERIMENTOS DESBALANCEADOS	85
6.4.1.	TESTE DE DUNCAN	86
6.4.1.1.	Para contrastes que abrangem 4 médias: 4 vs. 4 repetições	86
6.4.1.2.	Para contrastes que abrangem 3 médias: 4 vs. 4 repetições	87
6.4.1.3.	Para contrastes que abrangem 3 médias: 4 vs. 5 repetições	87
6.4.1.4.	Para testar contrastes que abrangem 2 médias: 4 vs. 5 repetições	87
6.4.1.5.	Para testar contrastes que abrangem 2 médias: 4 vs. 4 repetições	88
6.4.2.	TESTE DE TUKEY	88
6.4.2.1.	Para testar contrastes que abrangem 2 médias: 5 vs. 4 repetições	89
6.4.2.2.	Para testar contrastes que abrangem 2 médias: 4 vs. 4 repetições	89

7. ESTUDO E APLICAÇÃO DE CONTRASTES **90**

7.1.	INTRODUÇÃO	90
7.2.	DEFINIÇÃO	90
7.3.	CONTRASTES ENTRE TOTAIS DE TRATAMENTOS COM UM MESMO NÚMERO DE REPETIÇÕES	91
7.3.1.	CÁLCULO DA SOMA DE QUADRADOS DOS DESVIOS	91
7.3.2.	ORTOGONALIDADE	91
7.4.	CONTRASTES ENTRE TOTAIS DE TRATAMENTOS COM NÚMERO DIFERENTES DE REPETIÇÕES	92
7.4.1.	CÁLCULO DA SOMA DE QUADRADOS DOS DESVIOS	92
7.4.2.	ORTOGONALIDADE	92
7.5.	REGRAS PARA OBTENÇÃO DE CONTRASTES ORTOGONAIS	93
7.5.1.	CONTRASTES COM UM MESMO NÚMERO DE REPETIÇÕES	93
7.5.2.	CONTRASTES COM NÚMERO DIFERENTE DE REPETIÇÕES	94
7.6.	VARIÂNCIA DE CONTRASTES	95
7.7.	COMPREENSÃO DO CÁLCULO AS SOMA DE QUADRADOS DOS DESVIOS DE CONTRASTES	96
7.7.1.	COM MÉDIAS DE TRATAMENTOS	96
7.7.2.	COM OS TOTAIS DE TRATAMENTOS	97

8. REFLEXÕES SOBRE A ANÁLISE DE VARIÂNCIA **98**

8.1.	INTRODUÇÃO	98
8.2.	REFLEXÕES	98
8.3.	BLOCO DE PERGUNTAS 1	105
8.4.	BLOCO DE PERGUNTAS 2	108
8.5.	ANÁLISE COMPUTACIONAL DE UM EXPERIMENTO	109
8.5.1.	PROGRAMA PARA A ANÁLISE	109
8.5.2.	RESULTADOS FORNECIDOS	110
8.5.2.1.	Análise de variância	110
8.5.2.2.	Testes de comparação de médias	110
8.5.2.2.1.	Teste de Tukey	110
8.5.2.2.2.	Teste de Duncan	111
8.5.2.2.3.	Teste de Dunnett	111
8.5.2.2.4.	Teste de Student – Newman – Keuls	111
8.6.	BLOCO DE PERGUNTAS 3	112

9. DELINEAMENTO EM BLOCOS CASUALIZADOS - DBC	114
9.1. INTRODUÇÃO	114
9.2. PRINCÍPIOS UTILIZADOS	114
9.2.1. REPETIÇÃO	114
9.2.2. CASUALIZAÇÃO	114
9.2.3. CONTROLE LOCAL	114
9.2.4. EXEMPLOS DE CONTROLE LOCAL	114
9.3. VANTAGENS E DESVANTAGENS	115
9.3.1. VANTAGENS	115
9.3.2. DESVANTAGENS	115
9.4. MODELO ESTATÍSTICO	115
9.5. ESQUEMA DE CASUALIZAÇÃO DOS TRATAMENTOS	115
9.6. COLETA DE DADOS	116
9.7. ANÁLISE DE VARIÂNCIA	116
9.7.1. ESQUEMA DA ANÁLISE DE VARIÂNCIA	116
9.7.2. TESTE DE HIPÓTESES	117
9.8. EXEMPLO COM UM MESMO NÚMERO DE REPETIÇÕES	117
9.8.1. TESTES DE COMPARAÇÃO DE MÉDIAS MÚLTIPLAS	118
9.8.2. DESDOBRAMENTO DOS GL ASSOCIADOS A TRATAMENTOS EM CONTRASTES ORTOGONAIS	118
9.9. CONSIDERAÇÕES FINAIS	119
10. DELINEAMENTO EM QUADRADO LATINO - DQL	120
10.1. INTRODUÇÃO	120
10.2. PRINCÍPIOS UTILIZADOS	120
10.2.1. REPETIÇÃO	120
10.2.2. CASUALIZAÇÃO	120
10.2.3. CONTROLE LOCAL	120
10.2.4. EXEMPLOS DE CAUSAS DE VARIAÇÃO CONTROLADAS POR ESTE DELINEAMENTO	120
10.3. VANTAGENS E DESVANTAGENS	121
10.3.1. VANTAGENS	121
10.3.2. DESVANTAGENS	121
10.4. MODELO ESTATÍSTICO	121
10.5. ESQUEMA DE CASUALIZAÇÃO DOS TRATAMENTOS	122
10.6. COLETA DE DADOS	122
10.7. ANÁLISE DE VARIÂNCIA	123
10.7.1. ESQUEMA DA ANÁLISE DE VARIÂNCIA	123
10.7.2. TESTE DE HIPÓTESES RELATIVAS AOS TRATAMENTOS	123
10.8. EXEMPLO COM UM MESMO NÚMERO DE REPETIÇÕES	123
10.8.1. TESTES DE COMPARAÇÃO DE MÉDIAS MÚLTIPLAS	125
10.8.2. DESDOBRAMENTO DOS GL DE TRATAMENTOS EM CONTRASTES ORTOGONAIS	125
10.9. CONSIDERAÇÕES FINAIS	126
11. EXPERIMENTOS FATORIAIS	127
11.1. INTRODUÇÃO	127
11.2. CLASSIFICAÇÃO DOS EFEITOS	128
11.2.1. EFEITO PRINCIPAL	128
11.2.2. EFEITO DA INTERAÇÃO	128
11.3. VANTAGENS E DESVANTAGENS	130
11.3.1. VANTAGENS	130
11.3.2. DESVANTAGENS	130
11.4. MODELO ESTATÍSTICO	130
11.5. COLETA DE DADOS	130

11.6. ANÁLISE DE VARIÂNCIA	131
11.6.1. ESQUEMA DA ANÁLISE DE VARIÂNCIA	131
11.6.2. TESTES DE HIPÓTESES	131
11.7. EXEMPLOS	131
11.7.1. EXPERIMENTO MONTADO NO DIC COM INTERAÇÃO NÃO SIGNIFICATIVA	131
11.7.2. EXPERIMENTO MONTADO NO DIC COM INTERAÇÃO SIGNIFICATIVA	134
11.7.3. EXPERIMENTO MONTADO NO DBC COM INTERAÇÃO SIGNIFICATIVA	138
11.7.4. EXPERIMENTO MONTADO NO DIC COM INTERAÇÃO SIGNIFICATIVA	145
<u>12. EXPERIMENTOS EM PARCELAS SUBDIVIDIDAS</u>	<u>151</u>
12.1. INTRODUÇÃO	151
12.2. FATORIAL VS. PARCELA SUBDIVIDIDA	151
12.3. CLASSIFICAÇÃO DOS EFEITOS	152
12.3.1. EFEITO PRINCIPAL	152
12.3.2. EFEITO DA INTERAÇÃO	152
12.4. VANTAGENS E DESVANTAGENS	153
12.4.1. VANTAGENS	153
12.4.2. DESVANTAGENS	153
12.5. MODELO ESTATÍSTICO	153
12.6. COLETA DE DADOS	154
12.7. ANÁLISE DE VARIÂNCIA	154
12.7.1. TESTE DE HIPÓTESES	154
12.8. EXEMPLO: PARCELA SUBDIVIDIDA NO ESPAÇO	155
12.8.1. TESTE DE TUKEY APLICADO AOS EFEITOS PRINCIPAIS	157
12.9. EXEMPLO: PARCELA SUBDIVIDIDA NO TEMPO	158
12.9.1. DESDOBRAMENTO DA INTERAÇÃO	161
<u>13. CORRELAÇÃO LINEAR SIMPLES</u>	<u>168</u>
13.1. INTRODUÇÃO	168
13.2. DEFINIÇÃO	168
13.3. CONCEITOS E COMPREENSÃO A PARTIR DE UM EXEMPLO	169
13.4. PRESSUPOSIÇÕES DA CORRELAÇÃO	173
<u>14. INTRODUÇÃO AO ESTUDO DE REGRESSÃO LINEAR SIMPLES</u>	<u>176</u>
14.1. INTRODUÇÃO	176
14.1.1. CRITÉRIOS PARA SE AJUSTAR UMA RETA	183
14.1.2. AJUSTANDO UMA RETA	184
14.2. ANÁLISE DE VARIÂNCIA DA REGRESSÃO	192
14.2.1. CÁLCULOS ALTERNATIVOS DA SOMA DE QUADRADOS DOS DESVIOS	195
14.2.2. COEFICIENTE DE DETERMINAÇÃO DA REGRESSÃO	196
14.2.3. RELAÇÃO ENTRE O COEFICIENTE DE DETERMINAÇÃO E O COEFICIENTE DE CORRELAÇÃO	196
14.2.4. OBSERVAÇÕES A RESPEITO DA REGRESSÃO	197
14.2.5. ANÁLISE DE REGRESSÃO DE DADOS PROVENIENTES DE DELINEAMENTOS EXPERIMENTAIS	197
14.3. CRITÉRIOS PARA DECISÃO DE UM MODELO AJUSTADO E CONSIDERAÇÕES FINAIS	199
14.4. EXEMPLO DE ANÁLISE COMPLETA DE UM EXPERIMENTO	200
<u>15. TRANSFORMAÇÃO DE DADOS</u>	<u>208</u>
15.1. INTRODUÇÃO	208
15.2. TRANSFORMAÇÃO ANGULAR	208

15.2.1.	PRESSUPOSIÇÕES	208
15.2.2.	USO	208
15.2.3.	RECOMENDAÇÕES	209
15.1.	TRANSFORMAÇÃO RAIZ QUADRADA	209
15.1.1.	PRESSUPOSIÇÕES	209
15.1.2.	USO	209
15.1.3.	RECOMENDAÇÕES	209
18.1.1.	DICAS ÚTEIS	209
15.2.	TRANSFORMAÇÃO LOGARÍTMICA	210
15.2.1.	PRESSUPOSIÇÕES	210
15.2.2.	USO	210
15.2.3.	RECOMENDAÇÕES	210
15.2.1.	DICAS ÚTEIS	210

16. TABELAS ESTATÍSTICAS **I**

12ª edição

Estas anotações contêm, entre outras informações, as transparências utilizadas em sala de aula no curso de CET076 – Metodologia e Estatística Experimental do curso de Agronomia da Universidade Estadual de Santa Cruz, Ilhéus, Bahia.

Sua reunião, no formato de uma apostila, tem como objetivo fornecer aos estudantes as informações essenciais discutidas em sala de aula, evitando as anotações excessivas, assim como, servir como material de referência para as necessárias consultas à literatura.

Em hipótese alguma este material deve ser considerado como suficiente para os estudos durante o transcorrer do curso, além do que, deve ser complementado de forma pessoal por anotações decorrentes das discussões em sala de aula.

Esta edição passou por uma ampla revisão, tendo-se empregado esforços no sentido de padronizar a notação usada, adequar o índice, as fórmulas e as ilustrações, assim como, na correções de erros.

O autor agradece quaisquer sugestões que possam contribuir para o aprimoramento do conteúdo.

José Cláudio Faria, 15/04/2006.

emails:

joseclaudio.faria@terra.com.br

jc_faria@uesc.br

jc_faria@uol.com.br

Literatura recomendada

- BANZATTO, D.A & KRONKA, S.N. **Experimentação agrícola**. Jaboticabal: FUNEP, 1989. 247p.
- COCHRAN, W.G & COX, G.M. **Experimental design**. 2. Ed. New York: John Wiley, 1957. 462p.
- KACHIGAN, S.K. **Statistical analysis: an interdisciplinary introduction to univariate & multivariate methods**. New York: Radius Press. 1986. 589p.
- STORK, L.; GARCIA, D.C; LOPES, S.J. ESTEFANEL,V . **Experimentação vegetal**. Santa Maria: Ed. UFSM, 2000. 198p.
- ZAR, J.H. **Biostatistical analysis**. 4 ed. New Jersey: Prentice Hall. 1999. 663p. app 1-205.

Observações:

A literatura recomendada está listada por ordem alfabética dos autores.

Em caso da opção para aquisição textos de referência na língua portuguesa, para compor a biblioteca pessoal, recomenda-se BANZATTO, D.A & KRONKA, S.N, e ou, STORK et al.

ZAR, J.H. possui a seguinte referência na biblioteca da UESC:

- 574.015195
- Z 36 bio

Recursos disponíveis na WWW

Em função dos recursos didáticos avançados, recomenda-se que os laboratórios virtuais de estatística disponíveis na WWW sejam regularmente usados, pois são de inestimável valia para o aprendizado da estatística.

Os laboratórios indicados, além das experiências virtuais disponíveis, disponibilizam programas e links que permitem análises de dados em tempo real, podendo ser usados para o aprendizado, resoluções de exercícios e avaliações.

Laboratórios virtuais disponíveis na Internet

<http://www.ruf.rice.edu/~lane/rvls.html>

<http://www.kuleuven.ac.be/ucs/java/>

<http://www.stat.vt.edu/~sundar/java/applets/>

<http://www.isds.duke.edu/sites/java.html>

Site para análises on-line

<http://www.stat.sc.edu/webstat/>

Exemplos de recursos disponíveis na WWW

Distribuições amostrais

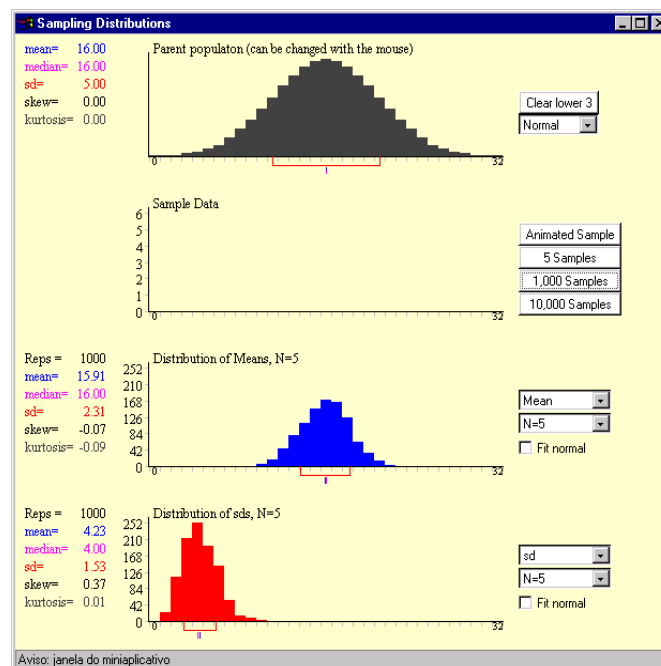


Figura 0.1 - Excelente para entender o teorema central do limite.

http://www.ruf.rice.edu/~lane/stat_sim/sampling_dist/index.html

Distribuição normal

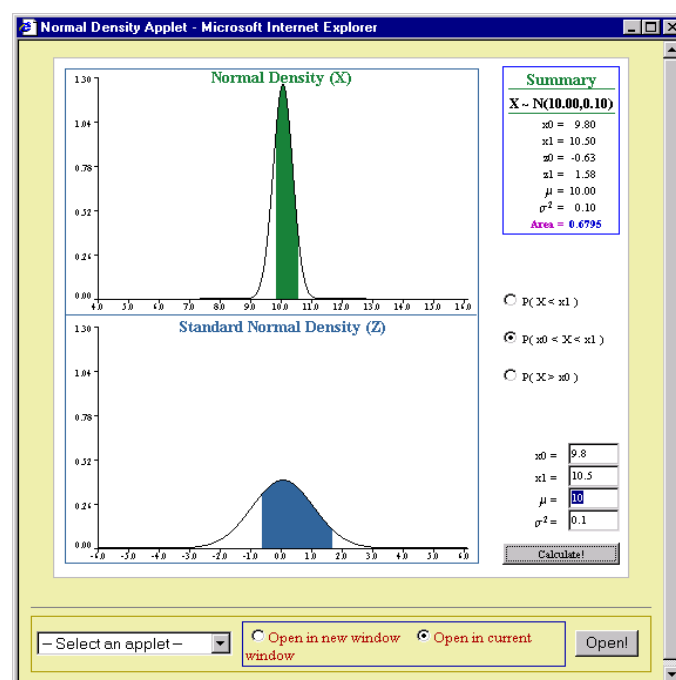


Figura 0.2 - Permite entender e realizar cálculos da distribuição normal.

<http://www.stat.vt.edu/~sundar/java/applets/>

Intervalo de confiança para a média populacional

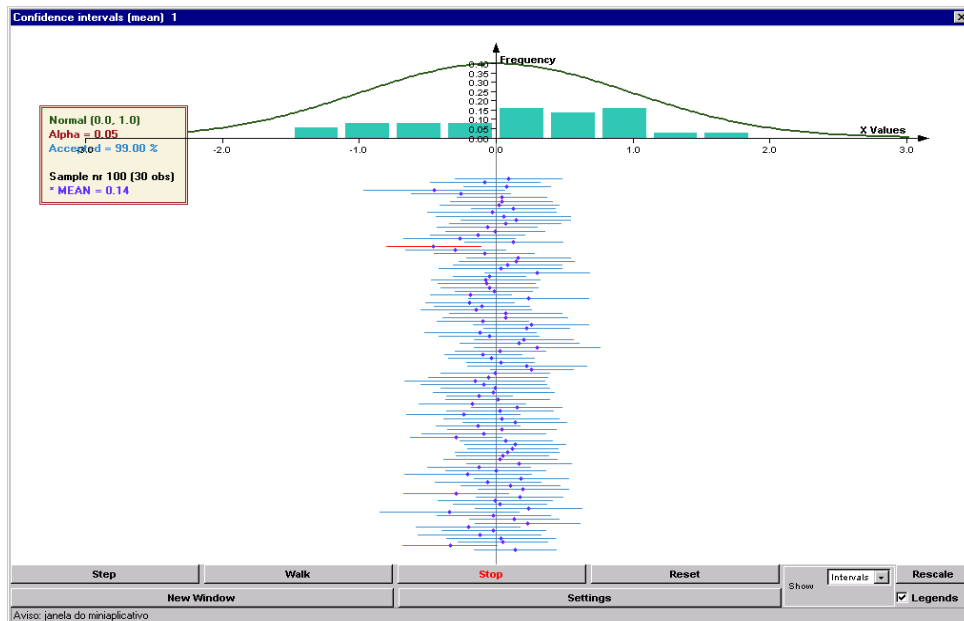


Figura 0.3 – Permite gerar populações, definir o tamanho das amostras e as variáveis que influenciam no intervalo de confiança para a média populacional.

<http://www.kuleuven.ac.be/ucs/java/>

Distribuição da variância

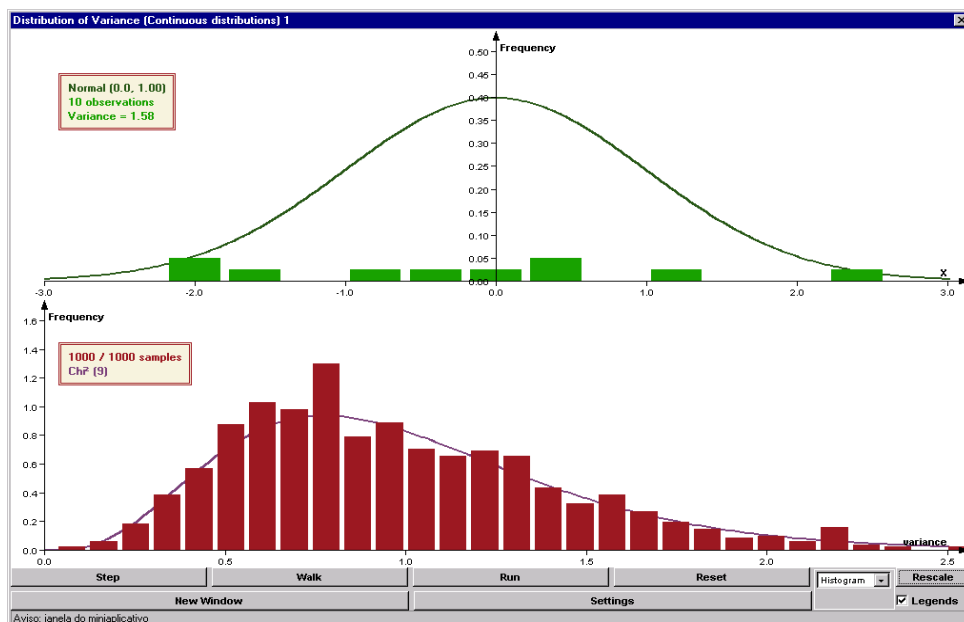


Figura 0.4 - Excelente para entender distribuição do Qui-quadrado.

<http://www.kuleuven.ac.be/ucs/java/>

Análise de variância – ANOVA

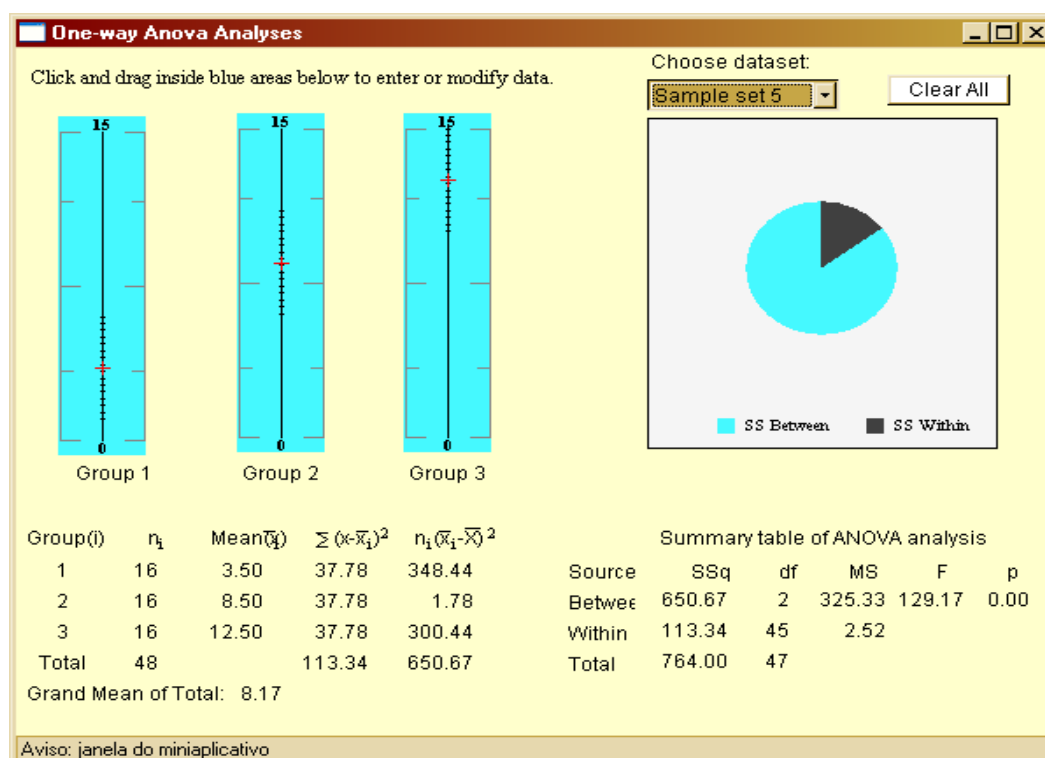


Figura 0.4 – Indispensável para entender os fundamentos da ANOVA permitindo a simulação de dados com o uso do mouse.

http://www.ruf.rice.edu/~lane/stat_sim/one_way/index.html

Medida	Populacional	Amostral (estimativa ou estatística)
Média	μ	m
Mediana	Md	md
Moda	Mo	mo
Variância	σ^2	s^2
Desvio padrão	σ	s
Desvio padrão relativo	DPR	dpr
Coeficiente de variação	CV	cv
Número de elementos	N	n
Correlação	ρ	r
Covariância	COV	cov
Parâmetro genérico	θ	$\hat{\theta}$

Variável	Valor observado	Valor estimado
Variável aleatória	Y	$\hat{Y}$

Sigla/Símbolo	Significado
GL , gl ou j	Graus de liberdade
SQD	Soma do quadrado dos desvios em relação à média
QMD	Quadrado médio dos desvios em relação à média

O termo **parâmetro** (θ) refere-se a toda e qualquer característica medida em populações, enquanto a **estimativa do parâmetro** ($\hat{\theta}$) é o correlato obtido em amostras representativas.

O termo **grau de liberdade** (GL, gl ou j) geralmente nos informa sobre o tamanho da amostra a partir da qual alguma estimativa ou estatística foi calculada. Na análise de contrastes a cada um é atribuído 1 GL e o mesmo é feito na análise de regressão onde cada parâmetro estimado no modelo recebe também 1 GL.

1. Calculadoras e aproximações em estatística

A experiência no ensino da estatística tem mostrado que uma parte considerável das dificuldades no aprendizado e no rendimento acadêmico relaciona-se ao uso de calculadoras inadequadas, a subutilização dos recursos de calculadoras adequadas e a problemas de aproximações de valores intermediários em cálculos sequencias comuns em estatística.

O objetivo destas considerações iniciais é esclarecer previamente o tipo de calculadora científica necessária, o uso adequado dos recursos básicos e as aproximações normalmente usadas em estatística.

1.1. Calculadora adequada

Uma calculadora adequada, não somente para os cursos de estatística, mas para o decorrer das disciplinas dos cursos de graduação, deve conter, no mínimo os seguintes recursos:

Medidas estatísticas básicas: média, variância, e ou, desvio padrão.

Somatórios básicos: $\sum x$ $\sum y$ $\sum x^2$ $\sum y^2$ $\sum xy$

Permitir a edição da série de dados armazenada na memória estatística.

Endereços de memória para armazenar de 5 a 10 resultados parciais.

Trabalhar com listas de números.

1.2. Comentários sobre os recursos básicos

Medidas estatísticas: são muito usadas e suas determinações, com calculadoras comuns, embora possível, são trabalhosas.

Somatórios básicos: são necessários em várias determinações.

Edição de dados: calculadoras que não possuem este recurso dificultam o trabalho com séries extensas de dados, pois depois de inseridos na memória estatística, não é possível conferi-los nem corrigi-los, o que ocasiona incerteza dos resultados e fadiga desnecessária devido à necessidade de repetição da digitação.

Endereços de memória: são muito usados para o armazenamento e recuperação de resultados intermediários que são usados em cálculos sucessivos.

Trabalhar com listas: permite que uma mesma operação seja feita em uma lista de dados, ao invés de elemento por elemento.

Exemplo:

$$\{4 \ 3 \ 5 \ 6\} - 3 = \{1 \ 0 \ 2 \ 3\}^2 = \{1 \ 0 \ 4 \ 9\} \square \sum_{\text{lista}} \Rightarrow = 14$$

1.3. Aproximações

Os cálculos estatísticos, embora simples, são em geral seqüenciais, de modo que resultados parciais são usados em novas determinações e assim por diante. Desta forma, o resultados intermediários devem ser sempre armazenados em variáveis de memória com todos os decimais possíveis e usados dessa forma. **Apenas no fim dos cálculos é que o resultado final deve ser aproximado, para o número de casas decimais suficiente para o problema numérico.** Se estes cuidados não forem tomados, as

aproximações sucessivas levam a distorções consideráveis no resultado final, podendo levar a conclusões equivocadas.

Em geral 2 ou 3 casas decimais são suficientes para a maioria dos problemas acadêmicos. Imagine que você está analisando algo que foi medido em metro (m), por exemplo 1 m, com uma casa decimal você estaria dando importância a um decímetro (1,0 m), com duas casas decimais você já estaria fazendo o mesmo com a um centímetro (1,00 m), com 3 casas decimais ao milímetro (1,000 m) e assim por diante. Bem, na grande maioria dos casos, quando estamos medindo algo em metro, aproximações finais em nível de centímetro ou milímetro são satisfatórias. Mais que isto, por exemplo, 1,000000000 m, poderia ser considerado desnecessário pois você estaria dando importância ao nanômetro, visível apenas com o auxílio de microscópios potentes.

1.4. Um teste

Vamos supor duas séries de dados com 15 elementos cada uma:

$A = \{12,31 \ 14,52 \ 13,23 \ 14,71 \ 16,82 \ 19,33 \ 14,99 \ 17,98 \ 13,67 \ 14,16 \ 14,85 \ 14,63 \ 13,24 \ 17,65 \ 13,26\}$ $B = \{14,13 \ 16,94 \ 11,55 \ 13,36 \ 18,17 \ 13,28 \ 14,19 \ 16,28 \ 12,17 \ 18,46 \ 12,55 \ 11,34 \ 12,13 \ 14,22 \ 18,11\}$
--

Os seguintes procedimentos são necessários:

a. Calcular a média aritmética simples de cada série

$m_A = 15,02$ $m_B = 14,46$

b. Diminuir cada valor das séries de suas respectivas médias

$A = \{(12,31 - 15,02) \ (14,52 - 15,02) \ \dots \ (13,26 - 15,02)\}$ $B = \{(14,13 - 14,46) \ (16,94 - 14,46) \ \dots \ (18,11 - 14,46)\}$

c. Para cada série elevar ao quadrado as diferenças e efetuar o somatório

$A = \{(-2,71)^2 + (-0,50)^2 + \dots + (-1,77)^2\}$ $B = \{(-0,33)^2 + (2,48)^2 + \dots + (3,65)^2\}$

d. Dividir cada resultado da etapa anterior (c) por 14

$A = \frac{57,40}{14} = 4,10$ $B = \frac{87,91}{14} = 6,28$

e. Dividir o maior pelo menor valor dos encontrados na etapa anterior (d) e expressar o resultado final com duas casas decimais

$$\frac{6,28}{4,10} = 1,53$$

Este é o resultado trabalhando com todos os resultados intermediários em variáveis de memória. Deve-se realizar o teste acima considerando que afastamentos do valor indicado (1,63) implicaram na adoção de procedimentos inadequados que necessitam ser revistos e melhorados.

1.5. O que não deve ser feito

- Não armazenar os valores das médias em variáveis de memória.
- Subtrair os valores das médias aproximadas (15,02 e 14,46) e não dos valores reais (15,02333... e 14,458666...).
- Redigitar as diferenças aproximadas para elevar ao quadrado e depois redigitar novamente os valores para efetuar o somatório.
- Redigitar novamente os resultados anteriores para efetuar a divisão por 14.
- Redigitar os valores aproximados anteriores para efetuar a divisão final.

É fácil perceber que devido às aproximações de resultados intermediários pode-se chegar a resultados bem diferentes do real. Adicionalmente, as digitações ocasionam erros (adicionais aos das aproximações) além da fadiga desnecessária.

Alguns estudantes realizam cálculos armazenando os valores das médias em variáveis de memória, digitam cada valor da série, que é subtraído da média, elevado e armazenado na memória de soma (M+). Posteriormente a soma final é recuperada e dividida por 14. Embora seja um paliativo, este procedimento encontra-se muito aquém do uso eficiente dos recursos disponíveis. Nas resoluções de exercícios toma muito tempo e via de regra compromete as avaliações.

Existem varias formas alternativas de realizar os cálculos anteriores utilizando os recursos das calculadoras científicas. A mais simples e usual é informar o valor de cada série na memória estatística e solicitar a medida estatística de dispersão dos dados em torno da média (variância amostral), armazenar cada valor (4,10 e 6,28) em variáveis de memória e posteriormente realizar a divisão entre elas.

Outra forma interessante é trabalhar com as séries na forma de listas.

Exemplo:

$$\{12,31 \ 14,52 \dots 13,26\} - 15,02 = \{-2,71 \ -0,50 \dots -1,76\}^2 = \{7,36 \ 0,25 \dots 3,11\} \quad \Sigma \text{Lista} \quad \frac{57,40}{14} = 4,10$$

Deve-se ter em mente que, além da necessidade da calculadora dispor dos recursos necessários, é importante saber usá-los adequadamente. Assim, cada usuário deve estudar o manual de instruções de sua calculadora pessoal a fim de que possa ter clareza e domínio sobre os recursos disponíveis.

2. Revisão dos cursos preliminares

O objetivo deste capítulo é o nivelamento básico dos conceitos já vistos em disciplinas consideradas pré-requisitos para o curso de Metodologia e Estatística Experimental.

Os conceitos discutidos são essenciais para o entendimento das técnicas de análise que serão tratadas neste curso. Assim, caso necessário, recomenda-se o aprofundamento do entendimento através da literatura pertinente.

Medidas estatísticas são números utilizados para resumir ou sintetizar as propriedades de uma série de dados.

2.1. Média aritmética

2.1.1. O que é

A média (ou esperança matemática) é uma medida estatística de tendência central.

É definida como a razão entre soma de todos os valores $\sum y_i$, e o número de elementos da série, **N** para populações ou **n** para amostras.

2.1.2. O que quantifica

Em uma série, quantifica a posição central, o ponto de equilíbrio ou o centro de gravidade:

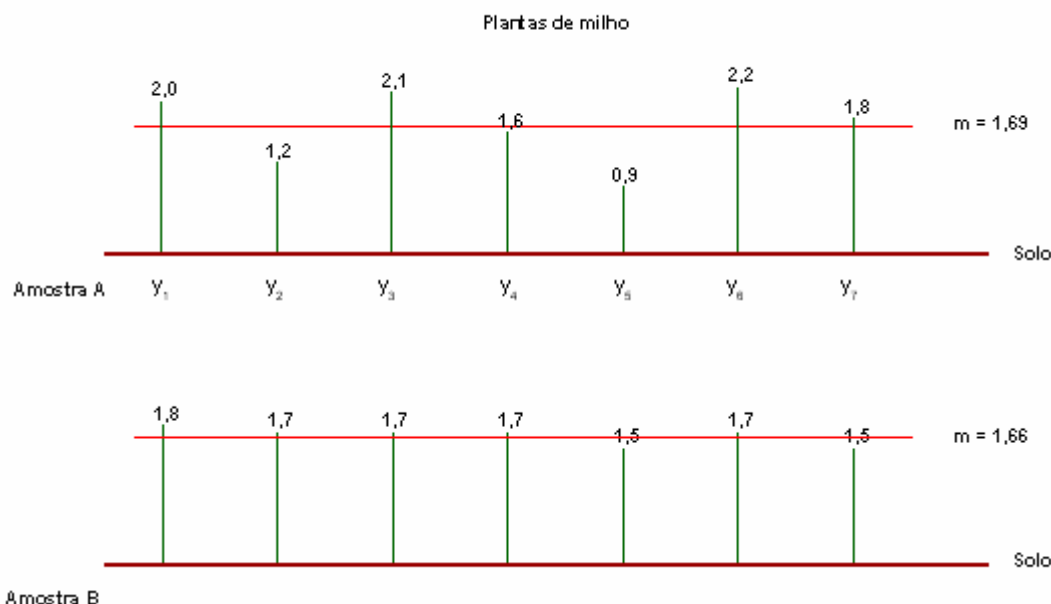


Figura 2.1 – Ilustração da média aritmética da altura de plantas.

2.1.3. Simbologia e cálculo

É simbolizada por α para populações e m para amostras.

2.1.3.1. Cálculo

$$\alpha = \frac{\sum y}{N} \qquad m = \frac{\sum y}{n}$$

Amostra A:

$$m(A) = \frac{\sum y}{n} = \frac{(2,0 + \dots + 1,8)}{7} = 1,69 m$$

Amostra B:

$$m(B) = \frac{\sum y}{n} = \frac{(1,8 + \dots + 1,5)}{7} = 1,66 m$$

2.1.4. Unidade de expressão

A unidade de expressão é a mesma da variável aleatória em questão. Para o exemplo dado na Figura 2.1, altura de plantas, a unidade é o metro, m:

$$\alpha \text{ ou } m = \frac{\sum y}{N \text{ ou } n} = \frac{m + \dots + m}{\text{número}} = m$$

2.2. Variância

2.2.1. O que é

É uma medida estatística da dispersão dos dados em relação à média aritmética.

É definida como a esperança matemática da soma de quadrados dos desvios em relação à média aritmética, ΣD^2 .

2.2.2. O que quantifica

Quantifica a dispersão dos dados em relação à média aritmética.

Permite distinguir séries de dados em relação à homogeneidade:

Séries homogêneas $\Rightarrow$ menor valor da variância

Séries heterogêneas $\Rightarrow$ maior valor da variância

2.2.3. Simbologia e cálculo

É simbolizada por σ^2 para populações e s^2 para amostras.

2.2.3.1. Cálculo

i. Populações:

$$\sigma^2 = \frac{\sum D^2}{N} \quad \text{onde } D = y - \alpha \text{ ou} \quad \sigma^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{N}}{N}$$

ii. Amostras:

a. α é conhecido (caso raro):

$$\sigma^2 = \frac{\sum D^2}{n} \quad \text{onde } D = y - \alpha \text{ ou} \quad s^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n}$$

b. α é desconhecido (caso comum):

$$s^2 = \frac{\sum d^2}{n-1} \quad \text{onde } d = y - m \quad \text{ou} \quad s^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1}$$

2.2.4. Unidade de expressão

A unidade de expressão é a mesma da variável aleatória em questão, porém, elevada ao quadrado. Para o exemplo dado na Figura 2.2, altura de plantas, a unidade é o metro elevado ao quadrado, m^2 :

$$\sigma^2 \text{ ou } s^2 = \frac{\sum D^2 \text{ ou } \sum d^2}{N \text{ ou } (n-1)} = \frac{m^2 + \dots + m^2}{\text{número}} = m^2$$

2.2.5. Conceito

É muito comum a dificuldade do estudante compreender o significado das medidas absolutas de dispersão (variância e do desvio padrão). Ou seja, compreender o conceito, o fundamento, antecedendo a qualquer cálculo:

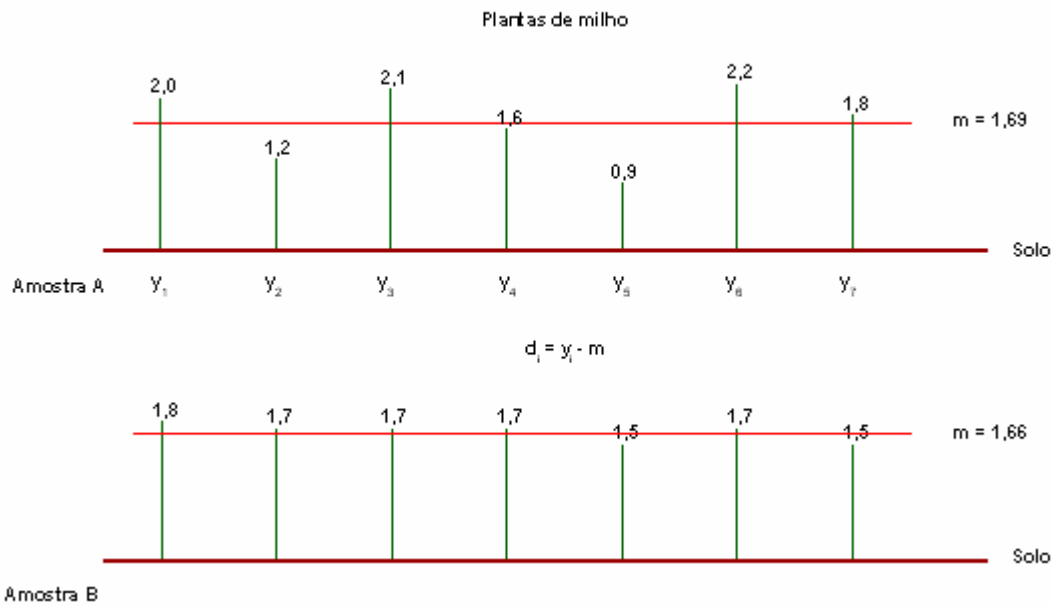


Figura 2.2 – Ilustração do significado da variância²s As barras verdes representam a altura das plantas de milho em relação ao solo e d representa o desvio da altura de uma planta em relação à média da série.

A variância, para uma variável aleatória em estudo, nada mais é que uma medida da totalidade dos desvios em relação à média.

Intuitivamente, portanto, a amostra A deve apresentar um maior valor da variância da altura das plantas de milho que a amostra B, pois os dados, em A, encontram-se mais dispersos em relação à média.

$$s_A^2 = \frac{\sum d^2}{n-1} = \frac{(2,0 - 1,69)^2 + (1,2 - 1,69)^2 + \dots + (1,8 - 1,69)^2}{7-1} = 0,23 \, m^2$$

$$s_B^2 = \frac{\sum d^2}{n-1} = \frac{(1,8 - 1,66)^2 + (1,7 - 1,66)^2 + \dots + (1,5 - 1,66)^2}{7-1} = 0,01 \, m^2$$

2.2.6. Formas de cálculo

Amostra A:

$$s_A^2 = \frac{\sum d^2}{n-1} = \frac{d_1^2 + \dots + d_7^2}{n-1} = \frac{(2,0 - 1,69)^2 + \dots + (1,8 - 1,69)^2}{7-1} = \frac{(0,31)^2 + \dots + (0,11)^2}{6} = 0,23 \, m^2$$

$$s_A^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1} = \frac{21,30 - \frac{(11,80)^2}{7}}{6} = 0,23 \, m^2$$

Amostra B:

$$s_B^2 = \frac{\sum d^2}{n-1} = \frac{d_1^2 + \dots + d_7^2}{n-1} = \frac{(1,8 - 1,66)^2 + \dots + (1,5 - 1,66)^2}{7-1} = \frac{(0,14)^2 + \dots + (-0,16)^2}{6} = 0,01 \text{ m}^2$$

$$s_A^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1} = \frac{19,30 - \frac{(11,60)^2}{7}}{6} = 0,01 \text{ m}^2$$

2.3. Desvio padrão

2.3.1. O que é

É uma medida estatística da dispersão dos dados em relação à média aritmética.
É definido como a raiz quadrada da variância.

2.3.2. O que quantifica

Quantifica a dispersão dos dados em relação à média aritmética.

2.3.3. Simbologia e cálculo

É simbolizada por σ para populações e s para amostras.

2.3.3.1. Cálculo

i. **Populações:**

$$\sigma = \sqrt{\sigma^2}$$

ii. **Amostras:**

$$s = \sqrt{s^2} \quad \therefore \quad s_A = \sqrt{s_A^2} = \sqrt{0,23 \text{ m}^2} = 0,48 \text{ m}$$

2.3.4. Unidade de expressão

A unidade de expressão é a mesma da variável aleatória em questão. Para o exemplo dado, a unidade é o metro, m:

$$\sigma \text{ ou } s = \sqrt{\text{m}^2} = \text{m}$$

2.4. Desvio padrão relativo e coeficiente de variação

2.4.1. O que são

São medidas estatísticas relativas da dispersão dos dados em relação à média.
São definidas como a razão entre o desvio padrão e a média aritmética.

2.4.2. O que quantificam

Quantificam a dispersão relativa dos dados em relação à média aritmética.

2.4.3. Simbologia e cálculos

O desvio padrão relativo é simbolizado por **DPR** para populações e **dpr** para amostras, o coeficiente de variação por **CV** para populações e **cv** para amostras.

2.4.3.1. Cálculos

i. Populações:

$$DPR = \frac{\sigma}{\bar{x}} \quad CV = \frac{\sigma}{\bar{x}} \cdot 100$$

ii. Amostras:

$$dpr = \frac{s}{\bar{x}} \quad cv = \frac{s}{\bar{x}} \cdot 100$$

2.4.4. Justificativas para o uso e unidades de expressão

Freqüentemente em trabalhos de pesquisa são necessárias comparações em situações nas quais as medidas estatísticas das variáveis em estudo foram feitas usando-se unidades distintas. Por exemplo: um pesquisador usou o metro, m, e outro o centímetro, cm.

Como as medidas absolutas de dispersão (variância e desvio padrão) são influenciadas pela unidade de medida das variáveis em estudo, a comparação entre os trabalhos fica dificultada.

Por serem adimensionais, é conveniente determinar uma das medidas relativas de dispersão, sendo a mais usada o coeficiente de variação.

Considerando que a unidade de medida das variáveis estudadas foi o metro, m:

i. População:

$$DPR = \frac{\sigma}{\bar{x}} = \frac{m}{m} = \text{adimensional} \quad CV = \frac{\sigma}{\bar{x}} \cdot 100 = \frac{m}{m} \cdot 100 = \% \text{ (adimensional)}$$

ii. Amostra:

$$dpr = \frac{s}{\bar{x}} = \frac{m}{m} = \text{adimensional} \quad cv = \frac{s}{\bar{x}} \cdot 100 = \frac{m}{m} \cdot 100 = \% \text{ (adimensional)}$$

Desta forma pode-se saber, independentemente da influência das unidades usadas, qual estudo apresentou maior ou menor dispersão.

Exemplo:

Considerando os dados da Figura 2.2:

i. Tomados em metro (m):

Amostra A:

$$cv = \frac{s}{m} \cdot 100 = \frac{0,48}{1,69} \cdot 100 = 28,74\%$$

Amostra B:

$$cv = \frac{s}{m} \cdot 100 = \frac{0,11}{1,66} \cdot 100 = 6,84\%$$

ii. Tomados em outras unidades de medida:

a. Amostra A em milímetro (mm):

$$cv = \frac{s}{m} \cdot 100 = \frac{484,52}{1.685,71} \cdot 100 = 28,74\%$$

b. Amostra B em centímetro (cm):

$$cv = \frac{s}{m} \cdot 100 = \frac{11,34}{165,71} \cdot 100 = 6,84\%$$

2.5. Demonstrações

i. Fórmula para cálculo da estimativa da variância:

$$\begin{aligned}
 s_Y^2 &= \frac{1}{n-1} \sum d^2 \\
 s_Y^2 &= \frac{1}{n-1} \sum (y - m)^2 \\
 s_Y^2 &= \frac{1}{n-1} \sum (y^2 - 2ym + m^2) \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - 2m \sum y + \sum m^2 \quad \therefore \quad \sum K \cdot y = K \sum y \\
 &\quad \text{se } m = \frac{\sum y}{n} \quad \text{então} \quad \sum y = n \cdot m \\
 &\quad \sum m^2 = n \cdot m^2 \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - (2m)(n \cdot m) + n \cdot m^2 \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - 2n \cdot m^2 + n \cdot m^2 \quad \therefore \quad 2a - a = a \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - n \cdot m^2 \quad \therefore \quad m = \frac{\sum y}{n} \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - n \cdot \left(\frac{\sum y}{n} \right)^2 \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - n \cdot \frac{(\sum y)^2}{n^2} \\
 s_Y^2 &= \frac{1}{n-1} \sum y^2 - n \cdot \frac{(\sum y)^2}{n^2} \\
 s_Y^2 &= \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1}
 \end{aligned}$$

ii. Tendenciosidade da estimativa da variância:

$$s^2 = \frac{\sum (y - m)^2}{n \text{ ou } n-1 ?}$$

$$\begin{aligned}\sum (y - m)^2 &= \sum (y - \mu + \mu - m)^2 \\ \sum (y - m)^2 &= \sum \{(y - \mu) - (m - \mu)\}^2 \\ \sum (y - m)^2 &= \sum \{(y - \mu)^2 - 2(y - \mu) \cdot (m - \mu) + (m - \mu)^2\} \\ \sum (y - m)^2 &= \sum (y - \mu)^2 - 2 \sum (y - \mu)(m - \mu) + \sum (m - \mu)^2\end{aligned}$$

$$\begin{aligned}\sum (y - \mu) &= \sum y_i - n \cdot \mu & \therefore \quad m = \frac{\sum y}{n} & \quad \therefore \quad \sum y = n \cdot m \\ \sum (y - \mu) &= n \cdot m - n \cdot \mu = n(m - \mu) \\ \sum (m - \mu)^2 &= n(m - \mu)^2 & \text{para uma determinada amostra } (m - \mu) = \text{constante}\end{aligned}$$

$$\begin{aligned}\sum (y - m)^2 &= \sum (y - \mu)^2 - 2n(m - \mu) \cdot (m - \mu) + n(m - \mu)^2 \\ \sum (y - m)^2 &= \sum (y - \mu)^2 - 2n(m - \mu)^2 + n(m - \mu)^2 & -2a + a = -a \\ \sum (y - m)^2 &= \sum (y - \mu)^2 - n(m - \mu)^2\end{aligned}$$

$$\text{Considerando } s^2 = \frac{\sum (y - m)^2}{n}$$

$$E(s^2) = \frac{1}{n} E\{\sum (y - \mu)^2 - n(m - \mu)^2\}$$

$$E(s^2) = \frac{1}{n} \{\sum E(y - \mu)^2 - n \cdot E(m - \mu)^2\}$$

$$E(s^2) = \frac{1}{n} \{n \cdot V(Y) - n \cdot V(m)\} \quad \therefore \quad V(m) = \frac{\sigma^2}{n}$$

$$E(s^2) = \frac{1}{n} \left\{ n \cdot \sigma^2 - n \frac{\sigma^2}{n} \right\}$$

$$E(s^2) = \frac{1}{n} \{n \cdot \sigma^2 - \sigma^2\} = \frac{1}{n} \{\sigma^2(n - 1)\} = \frac{(n - 1) \cdot \sigma^2}{n}$$

$$\text{Portanto, } s^2 = \frac{\sum (y - m)^2}{n}, \text{ é um estimador tendencioso (subestima) de } \sigma^2.$$

$$\text{Considerando } s^2 = \frac{\sum (y - m)^2}{n - 1}$$

$$E(s^2) = \frac{1}{n - 1} E\left\{\sum (y - \mu)^2 - n(m - \mu)^2\right\}$$

$$E(s^2) = \frac{1}{n - 1} \left\{\sum E(y - \mu)^2 - n \cdot E(m - \mu)^2\right\}$$

$$E(s^2) = \frac{1}{n - 1} \{n \cdot V(Y) - n \cdot V(m)\} \quad \therefore \quad V(m) = \frac{\sigma^2}{n}$$

$$E(s^2) = \frac{1}{n - 1} \left\{n \cdot \sigma^2 - n \frac{\sigma^2}{n}\right\}$$

$$E(s^2) = \frac{1}{n - 1} \{n \cdot \sigma^2 - \sigma^2\} = \frac{1}{n - 1} \{\sigma^2(n - 1)\} = \frac{(n - 1) \cdot \sigma^2}{n - 1} = \sigma^2$$

$$\text{Portanto, } s^2 = \frac{\sum (y - m)^2}{n - 1}, \text{ é um estimador não tendencioso de } \sigma^2.$$

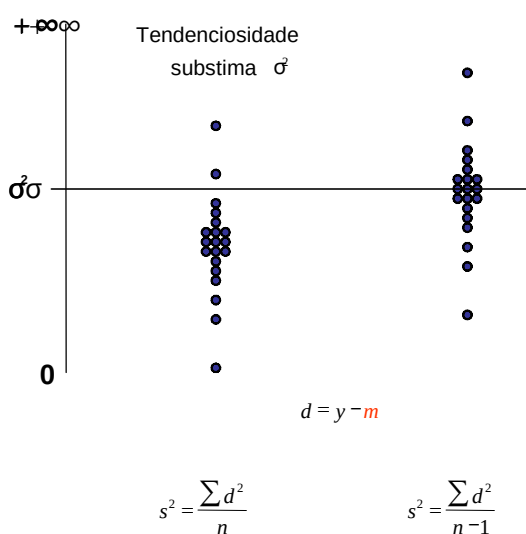


Figura 2.3 – Ilustração da tendenciosidade da estimativa de σ^2 se o somatório dos desvios em relação à média for dividido por n , ao invés de $n-1$.

2.6. Covariância

2.6.1. O que é

É uma medida estatística da associação linear entre duas variáveis aleatórias.

É definida como a esperança matemática do produto dos desvios, em relação às suas respectivas médias aritméticas.

2.6.2. O que quantifica

Quantifica o tipo e a magnitude da associação linear entre duas variáveis aleatórias.

Quanto ao tipo:

Positiva $\Rightarrow$ quando uma variável cresce a outra também cresce

Negativa $\Rightarrow$ quando uma variável cresce a outra diminui

Quanto ao grau:

Elevada $\Rightarrow$ as duas variáveis são estreitamente associadas, ou seja, o conhecimento de uma informa bastante sobre a outra.

Tendendo a zero $\Rightarrow$ as duas variáveis não são associadas, ou seja, o conhecimento de uma não informa nada sobre a outra. Neste caso as duas variáveis são consideradas independentes.

2.6.3. Simbologia e cálculo

É simbolizado por **COV** para populações e **cov** para amostras.

2.6.3.1. Cálculo

i. **Populações:**

$$COV(Y_1, Y_2) = E[(Y_1 - E(Y_1)) \cdot (Y_2 - E(Y_2))]$$

$$COV_{Pop}(Y_1, Y_2) = \frac{\sum [(Y_1 - \bar{Y}_1) \cdot (Y_2 - \bar{Y}_2)]}{N}$$

ii. **Amostras:**

a. σ é conhecido (caso raro):

$$cov_{Amo}(Y_1, Y_2) = \frac{\sum [(Y_1 - \bar{Y}_1) \cdot (Y_2 - \bar{Y}_2)]}{n}$$

b. σ é desconhecido (caso comum):

$$cov_{Amo}(Y_1, Y_2) = \frac{\sum [(Y_1 - m(Y_1)) \cdot (Y_2 - m(Y_2))]}{n - 1}$$

2.6.4. Unidade de expressão

A unidade de expressão é o produto das unidades de expressão das variáveis aleatórias em questão.

Vamos supor um exemplo em que se avalia o consumo de ração de aves de postura com a produção de ovos por semana:

$$COV_{ou\ cov} = \frac{(g\ dia^{-1} - g\ dia^{-1}) \cdot (un\ sem^{-1} - un\ sem^{-1})}{N\ ou\ n} = g\ dia^{-1} \cdot un\ sem^{-1}$$

2.6.4.1. Conceito

É muito comum a dificuldade de se compreender o significado da covariância, ou seja, compreender o conceito, o fundamento, antecedendo a qualquer cálculo.

A figura abaixo mostra com objetividade e clareza os fundamentos desta importante medida estatística, assim como fornece elementos para o entendimento da variação do grau de associação linear entre duas variáveis aleatórias quanto ao tipo (positiva ou negativa) e o grau (alta ou baixa):

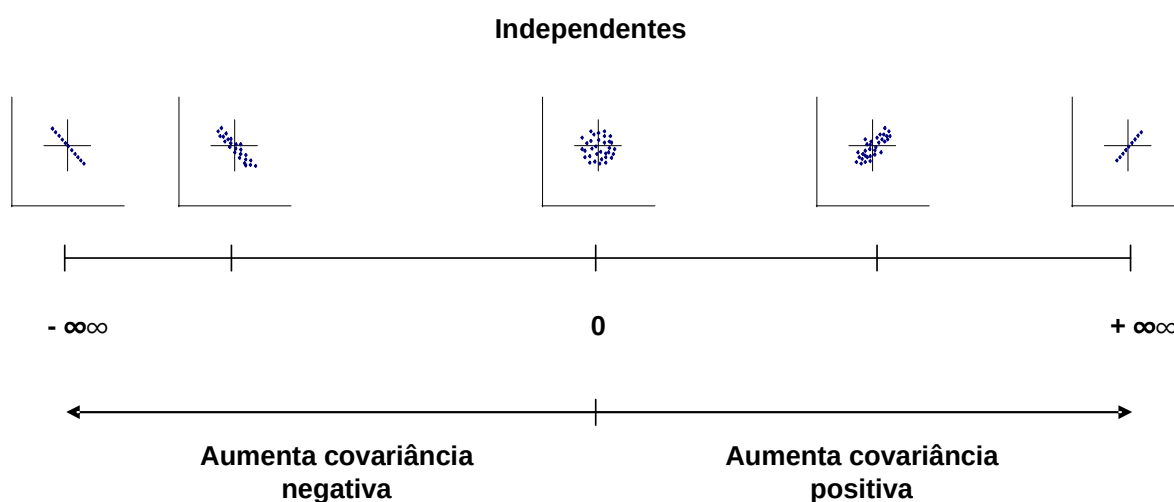
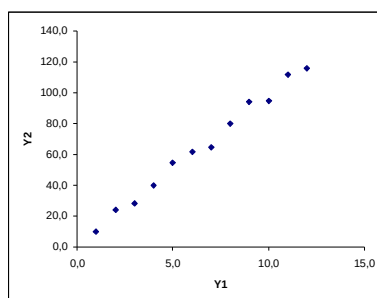


Figura 2.4 – Ilustração do significado da covariância.

2.6.5. Exemplos de cálculo e visualização das associações

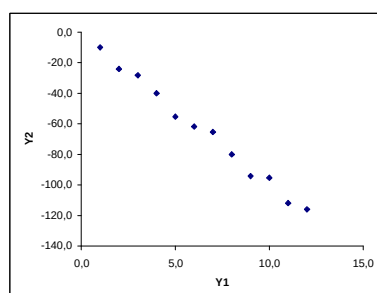
2.6.5.1. Variáveis com associação positiva e elevada

Obs	Y_1	Y_2	$Y_1 - m(Y_1)$	$Y_2 - m(Y_2)$	$Y_1 - m(Y_1) \cdot Y_2 - m(Y_2)$
1	1,00	10,00	-5,50	-55,08	302,96
2	2,00	24,00	-4,50	-41,08	184,88
3	3,00	28,00	-3,50	-37,08	129,79
4	4,00	40,00	-2,50	-25,08	62,71
5	5,00	55,00	-1,50	-10,08	15,13
6	6,00	62,00	-0,50	-3,08	1,54
7	7,00	65,00	0,50	-0,08	-0,04
8	8,00	80,00	1,50	14,92	22,38
9	9,00	94,00	2,50	28,92	72,29
10	10,00	95,00	3,50	29,92	104,71
11	11,00	112,00	4,50	46,92	211,13
12	12,00	116,00	5,50	50,92	280,04
$m(Y_1) = 6,50$		$m(Y_2) = 65,08$	$\Sigma[Y_1 - m(Y_1) \cdot Y_2 - m(Y_2)]/11 = \mathbf{126,14}$		



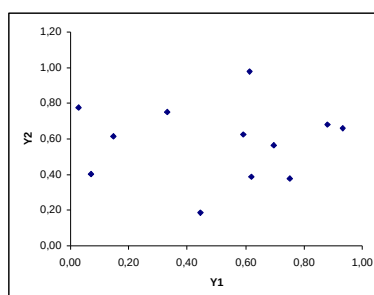
2.6.5.2. Variáveis com associação negativa e elevada

Obs	Y_1	Y_2	$Y_1 - m(Y_1)$	$Y_2 - m(Y_2)$	$Y_1 - m(Y_1) \cdot Y_2 - m(Y_2)$
1	1,00	-10,00	-5,50	-55,08	-302,96
2	2,00	-24,00	-4,50	-41,08	-184,88
3	3,00	-28,00	-3,50	-37,08	-129,79
4	4,00	-40,00	-2,50	-25,08	-62,71
5	5,00	-55,00	-1,50	-10,08	-15,13
6	6,00	-62,00	-0,50	-3,08	-1,54
7	7,00	-65,00	0,50	-0,08	0,04
8	8,00	-80,00	1,50	14,92	-22,38
9	9,00	-94,00	2,50	28,92	-72,29
10	10,00	-95,00	3,50	29,92	-104,71
11	11,00	-112,00	4,50	46,92	-211,13
12	12,00	-116,00	5,50	50,92	-280,04
$m(Y_1) = 6,50$		$m(Y_2) = -65,08$	$\Sigma[Y_1 - m(Y_1) \cdot Y_2 - m(Y_2)]/11 = \mathbf{-126,14}$		



2.6.5.3. Variáveis não associadas

Obs	Y_1	Y_2	$Y_1 - m(Y_1)$	$Y_2 - m(Y_2)$	$Y_1 - m(Y_1) \cdot Y_2 - m(Y_2)$
1	0,03	0,78	-0,48	0,19	-0,09
2	0,62	0,39	0,11	-0,19	-0,02
3	0,07	0,40	-0,44	-0,18	0,08
4	0,75	0,38	0,24	-0,21	-0,05
5	0,88	0,68	0,37	0,10	0,04
6	0,59	0,63	0,08	0,04	0,00
7	0,93	0,66	0,42	0,08	0,03
8	0,15	0,62	-0,36	0,03	-0,01
9	0,45	0,19	-0,06	-0,40	0,03
10	0,61	0,98	0,10	0,39	0,04
11	0,33	0,75	-0,18	0,17	-0,03
12	0,70	0,56	0,19	-0,02	0,00
$m(Y_1) = 0,56$		$m(Y_2) = 0,51$	$\Sigma[Y_1 - m(Y_1) \cdot Y_2 - m(Y_2)]/11 = 0,00$		



2.7. Teorema central do limite

2.7.1. O que é

Na medida em que aumenta o tamanho, n , a distribuição da média, m , de uma amostra aleatória, extraída de praticamente qualquer população, tende para a distribuição normal com média μ e desvio padrão $\sigma/\sqrt{n}$:

$$E(m) = \mu$$

$$DP(m) = \frac{\sigma}{\sqrt{n}}$$

$\therefore$

$$V(m) = \frac{\sigma^2}{n}$$

2.7.2. O que significa

Como a estimativa da média (média amostral) de uma variável aleatória é também uma variável aleatória, pode-se determinar sua esperança matemática (média) e sua dispersão (desvio padrão):

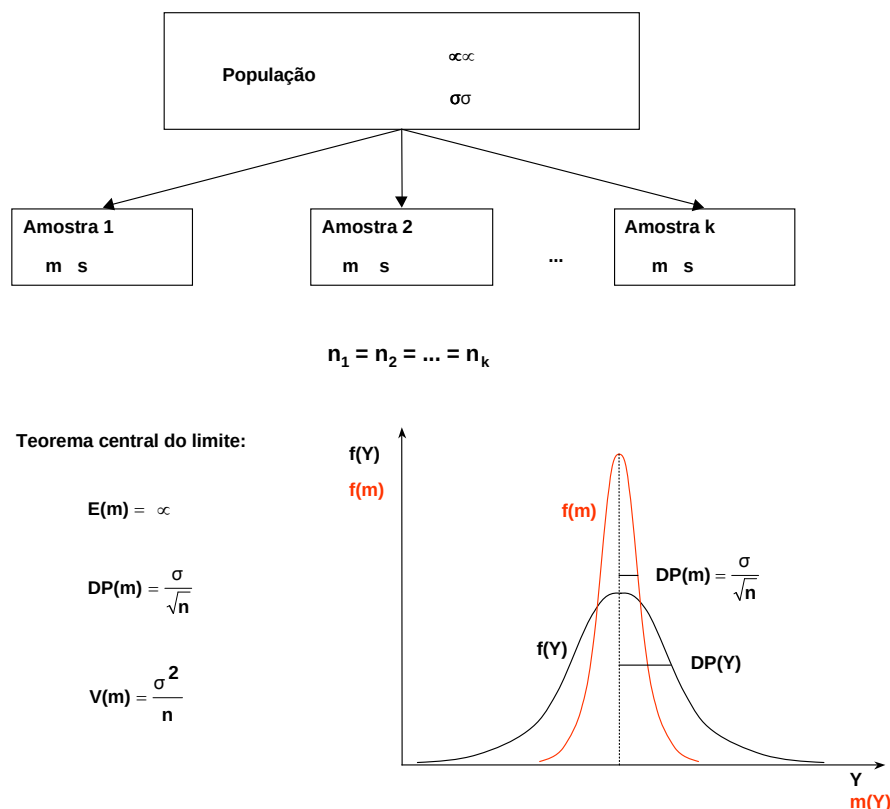


Figura 2.5 – Ilustração do teorema central do limite.

2.7.3. Como é usado

Na estatística experimental o caso mais comum de uso se dá quando é possível determinar a variância da média, $V(m)$, de um conjunto limitado de amostras (duas ou mais), não se conhece a variância populacional, e é necessário estimá-la:

$$V(m) = \frac{\sigma^2}{n} \quad \therefore \quad \sigma^2 = n \cdot V(m)$$

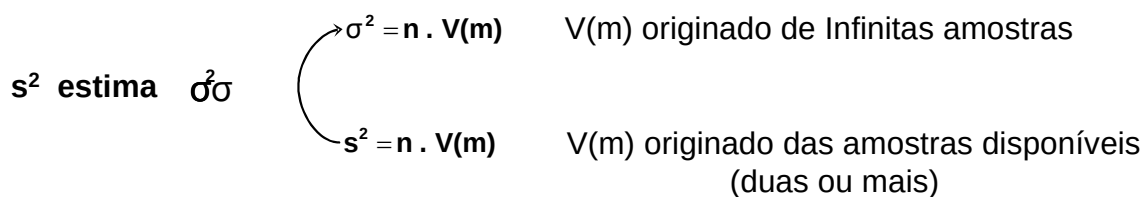


Figura 2.6 – Uso do teorema central do limite na estimação da variância

2.8. Teste de hipóteses

2.8.1. Hipótese: o que é

Trata-se de uma suposição sobre o valor de um parâmetro populacional ou quanto à natureza da distribuição de probabilidade populacional de uma variável aleatória.

Exemplos:

A precisão de dois métodos analíticos é igual fl ($\sigma_A^2 = \sigma_B^2$)
 As médias dos grupos são iguais fl ($\mu_A = \dots = \mu_K$)

2.8.2. Teste de hipóteses: o que é

É uma regra de decisão para aceitar, ou rejeitar, uma hipótese estatística com base nos elementos amostrais.

2.8.3. Tipos de hipóteses

H_0 : hipótese da igualdade : $\sigma_A^2 = \sigma_B^2$
 H_1 : hipóteses alternativas : $\sigma_A^2 > \sigma_B^2$; $\sigma_A^2 < \sigma_B^2$; $\sigma_A^2 \neq \sigma_B^2$

2.8.4. Tipos de erros

São os erros associados às decisões do teste de hipóteses:

		Realidade	
		H_0 verdadeira	H_0 falsa
Decisão	Aceitar H_0	Decisão correta ($1 - \alpha$)	Erro tipo II (β)
	Rejeitar H_0	Erro tipo I (α)	Decisão correta ($1 - \beta$)

O tomador da decisão (pesquisador) deseja, obviamente, reduzir ao mínimo as probabilidades dos dois tipos de erro na tomada de decisão, ou seja, na inferência estatística.

Infelizmente, esta é uma tarefa difícil, porque, para uma amostra de determinado tamanho, à medida que se diminui a probabilidade de incorrer em um erro do tipo I a probabilidade do erro tipo II aumenta, e vice-versa.

Estatisticamente a redução simultânea dos erros na inferência poderá ser alcançada apenas pelo aumento do tamanho da amostra.

2.9. Distribuição F

2.9.1. O que é

A definição mais comumente encontrada é a seguinte: a distribuição F é a razão entre duas variáveis aleatórias independentes com distribuição qui-quadrado, χ^2 .

Assim, uma distribuição F com ϕ_1 graus de liberdade no numerador, e ϕ_2 graus de liberdade no denominador é expressa por:

$$F(\phi_1, \phi_2) = \frac{\frac{\chi_{\phi_1}^2}{\phi_1}}{\frac{\chi_{\phi_2}^2}{\phi_2}}$$

Esta distribuição de probabilidade foi reduzida por Snedecor sendo sua denominação uma homenagem a Ronald Fisher. A função densidade de probabilidade é definida da seguinte forma:

$$f(F, \phi_1, \phi_2) = c \cdot \left(\frac{\phi_1}{\phi_2} \right)^{\frac{\phi_1}{2}} \cdot F^{\left(\frac{\phi_1}{2} - 1 \right)} \cdot \left(1 + \frac{\phi_1}{\phi_2} \cdot F \right)^{-\left(\frac{\phi_1 + \phi_2}{2} \right)} \quad c = \frac{\Gamma\left(\frac{\phi_1 + \phi_2}{2} \right)}{\Gamma\left(\frac{\phi_1}{2} \right) \cdot \Gamma\left(\frac{\phi_2}{2} \right)}$$

onde:

c: constante dependente de ϕ_1 e ϕ_2 determinada pela condição na qual a área sob a curva de probabilidade é igual a um.

ϕ_1 e ϕ_2 : graus de liberdade das amostras.

2.9.2. O que significa

Considerando que s^2 é um estimador não tendencioso de σ^2 :

$$E(F) = E\left(\frac{s_1^2}{s_2^2} \right) = \frac{E(s_1^2)}{E(s_2^2)} = \frac{\sigma^2}{\sigma^2} = 1$$

Ou seja, se infinitos pares de amostras aleatórias, cada amostra de tamanho fixo e constante, forem retirados de uma população normalmente distribuída, e a cada par a razão entre as estimativas da variância for calculada:

$$F = \frac{s_1^2}{s_2^2}$$

a média desses valores será igual a 1.

Entretanto, cada estimativa da variância está sujeita às variações normais decorrentes da amostragem aleatória dos indivíduos da população.

Assim, ao considerarmos um par qualquer, o valor F determinado poderá ser maior ou menor que 1.

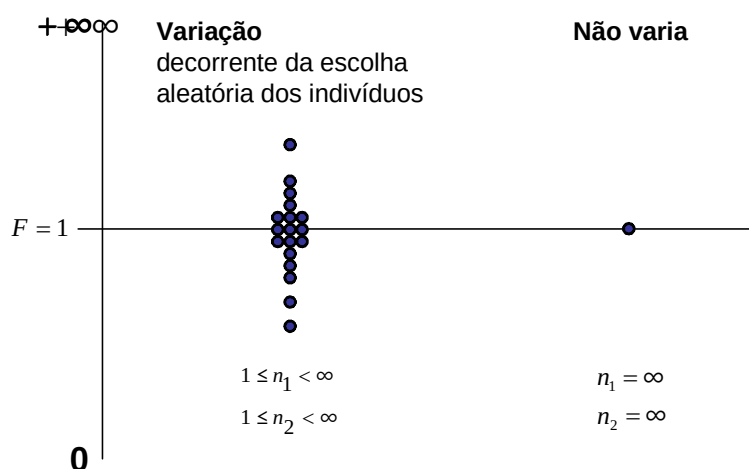


Figura 2.7 – Ilustração da variação de F decorrente da amostragem.

Uma curva específica da função densidade de probabilidade de F, que levará em consideração apenas o tamanho da amostra do par (ϕ_1 e ϕ_2), fornece a distribuição de probabilidades resultante de infinitas determinações do valor F.

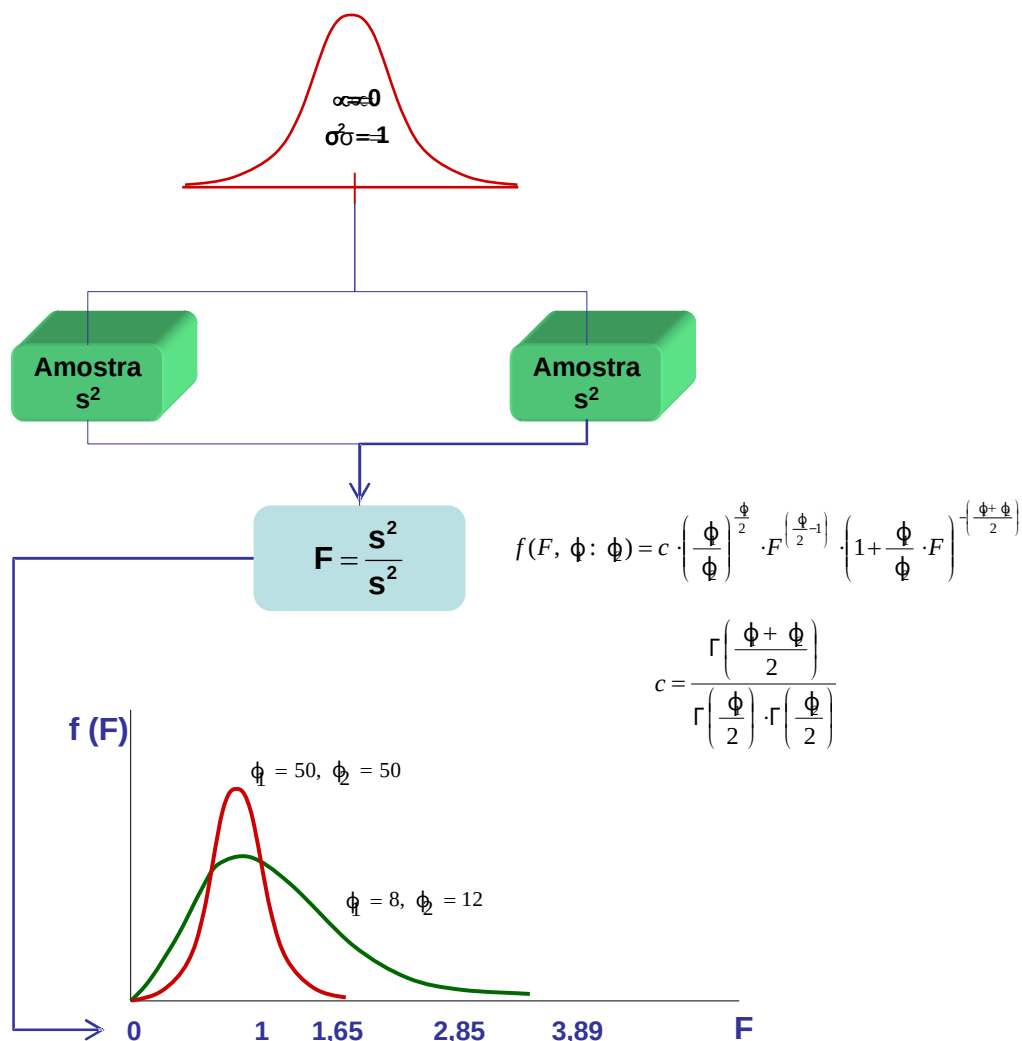
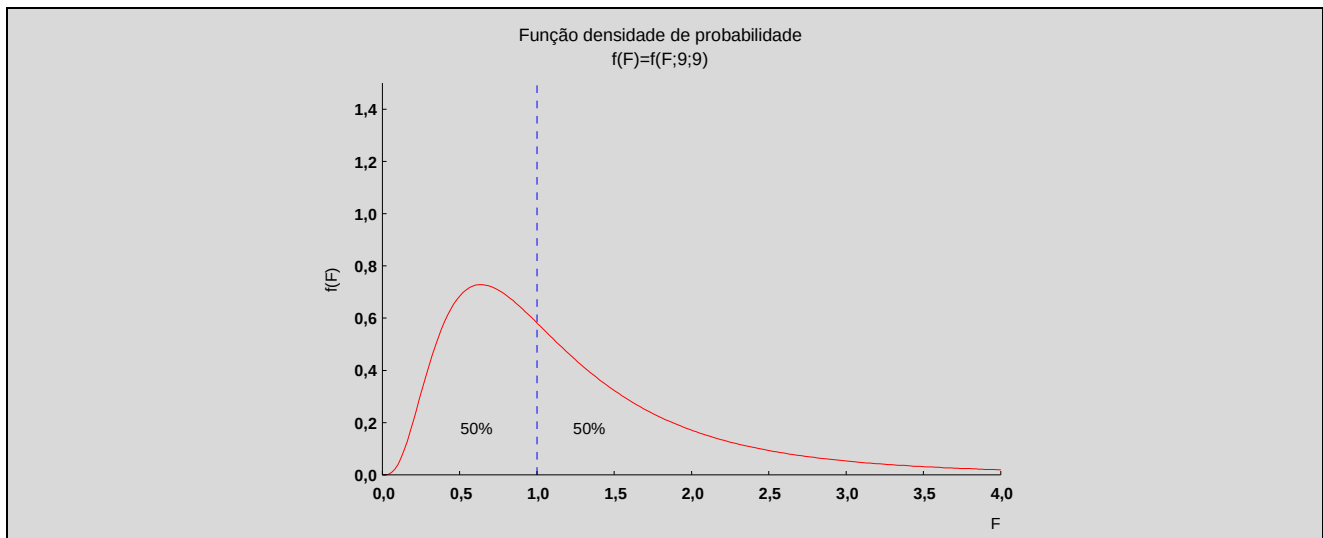


Figura 2.8 – Ilustração da origem da distribuição F.

A função densidade de probabilidade, $f(F)$, não é probabilidade. Somente quando integrada entre dois limites (a e b , com $a < b$), obtém-se a probabilidade do valor F encontrar-se situado entre os dois limites, ou seja:

$$P(a \leq F \leq b) = \int_a^b f(F) dF$$

Utilizando recursos computacionais o gráfico da distribuição F com tamanho das amostras igual a 10 ($\phi_1 = \phi_2 = 9$) foi gerado e encontra-se a seguir:



$$\int_0^1 f(F) dF = 0,50 = 50\%$$

$$\int_1^{\infty} f(F) dF = 0,50 = 50\%$$

2.9.3. Como é usada

A distribuição F é usada para se tomar decisões sobre as populações a partir de estimativas da variância (obtidas nas amostras) quando se testa hipóteses (inferências sobre as populações).

Um uso básico, por exemplo, permite a decisão se duas estimativas da variância podem, ou não, serem consideradas como provenientes de uma mesma população.

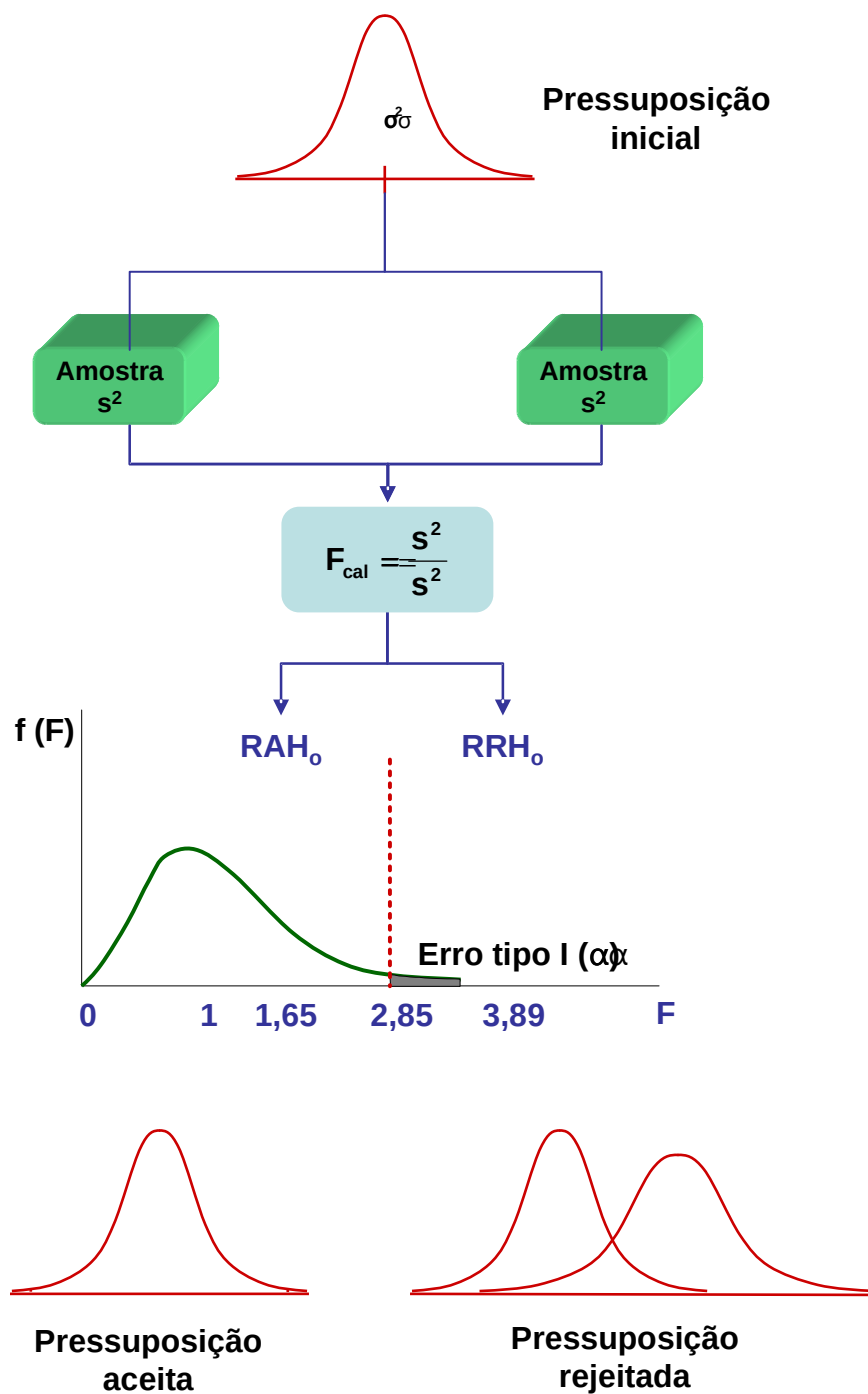


Figura 2.9 – Exemplo de uso da distribuição F.

2.9.4. Exatidão e precisão

Exatidão refere-se ao grau de aproximação do real, do objetivo ou do alvo.

Precisão refere-se ao grau de repetibilidade na aproximação do real, ou a proximidade de cada observação de sua própria média.

Exatidão	Fidelidade ao real ou certo
Precisão	Repetibilidade

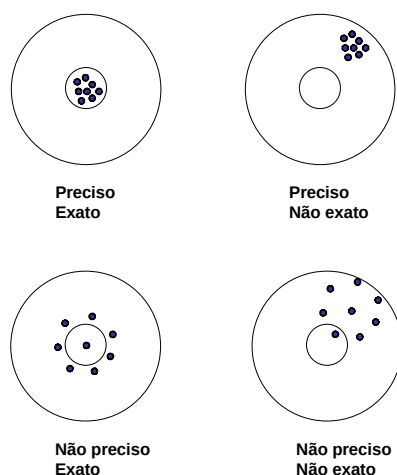


Figura 2.10 – Ilustração do conceito de precisão e exatidão.

Observações:

Os métodos analíticos padrões são exatos e precisos mas em geral são trabalhosos e caros.

Assim, em muitas situações eles são substituídos por métodos alternativos, mais rápidos e baratos, cuja principal característica desejável é a elevada precisão (repetibilidade), uma vez que a inexatidão (distanciamento do real), inerente ao método, pode ser corrigida por um fator de correção obtido entre o método padrão e o alternativo.

2.9.5. Exemplo básico de aplicação da distribuição F - comparação de precisão

Dois métodos de determinação da CTC do solo são usados em uma amostra de controle e fornecem os seguintes resultados em $\text{cmol}_c \text{ kg}^{-1}$:

	r_1	r_2	r_3	r_4	r_5	r_6	r_7	r_8	r_9	r_{10}	n	gl	m	s^2	s
A	10,2	8,7	9,5	12,0	9,0	11,2	12,5	10,9	8,9	10,6	10	9	10,35	1,76	1,33
B	9,9	9,2	10,4	10,5	11,0	11,3	9,6	9,4	10,0	10,4	10	9	10,17	0,46	0,68

A questão a ser investigada é se é possível, ou não, considerar as precisões dos dois métodos (população de resultados gerados por cada método) estatisticamente iguais, ou seja:

$$H_0 : \sigma_A^2 = \sigma_B^2$$

$$H_1 : \sigma_A^2 > \sigma_B^2$$

Caso de decida que os métodos apresentam igual precisão, $\sigma_A^2 = \sigma_B^2$, as diferenças entre os resultados obtidos serão atribuídas às flutuações estatísticas naturais e, neste caso, os métodos seriam similares e poderiam ser usados indiscriminadamente.

A estatística F pode ser usada para esta decisão.

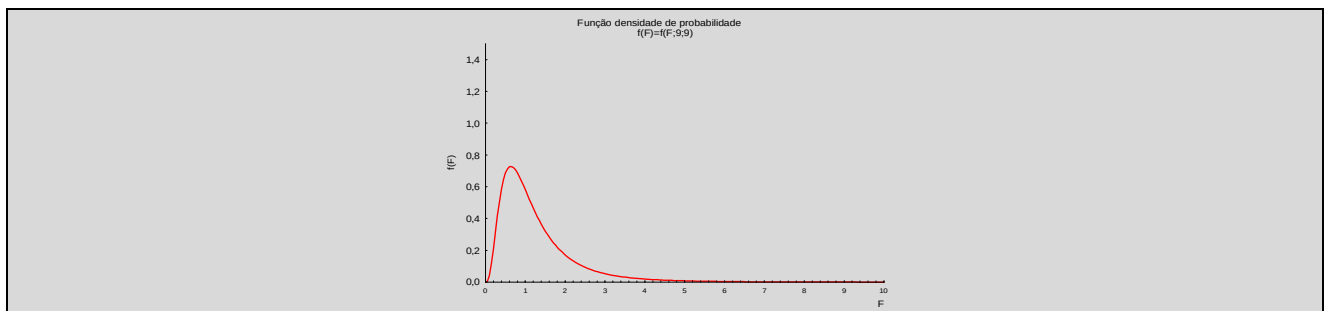
O teste faz uso da razão entre duas estimativa da variância, e como o teste é unilateral à direita, $\sigma_A^2 > \sigma_B^2$, o maior valor ocupa o numerador:

$$F_{cal} = \frac{s_A^2}{s_B^2} \text{ sendo } s_A^2 \geq s_B^2$$

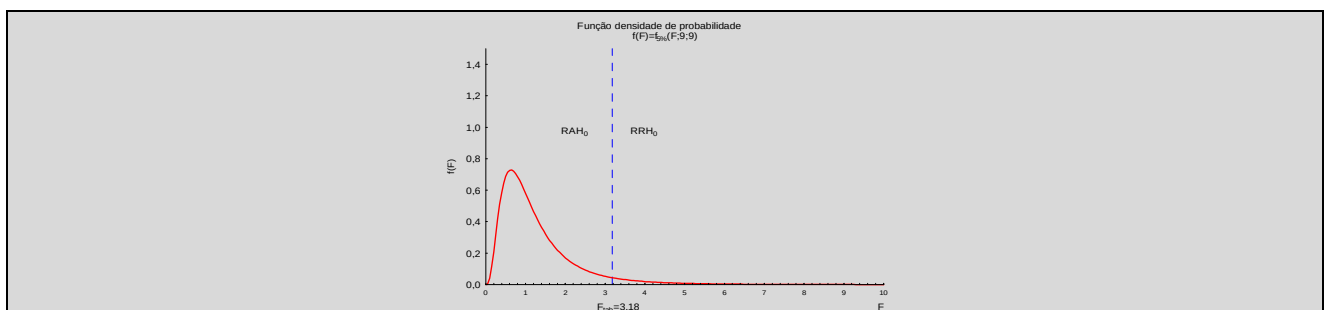
Esta decisão deve ser tomada adotando-se uma probabilidade de erro na decisão. Pode-se estabelecer, por exemplo, um erro máximo aceitável de 5%.

2.9.5.1. Mecanismo de decisão

Escolher a função densidade de probabilidades de F que apresente os graus de liberdade adequados (9:9).



O valor crítico, $F_{5\%}(9;9)$, pode ser obtido na tabela de F a 5% na interseção de 9 gl (numerador) na primeira linha com 9 gl (denominador) na primeira coluna.



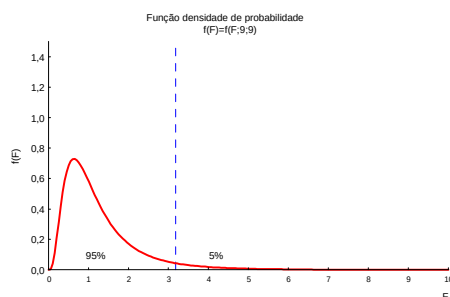
Considerar os resultados de cada um dos dois métodos como amostras (10 para cada método) aleatoriamente retiradas de uma mesma população normalmente distribuída:

	r_1	r_2	r_3	r_4	r_5	r_6	r_7	r_8	r_9	r_{10}	n	gl	m	s^2	s
A	10,2	8,7	9,5	12,0	9,0	11,2	12,5	10,9	8,9	10,6	10	9	10,35	1,76	1,33
B	9,9	9,2	10,4	10,5	11,0	11,3	9,6	9,4	10,0	10,4	10	9	10,17	0,46	0,68

Calcular o valor de prova (F_{cal}):

$$F_{cal} = \frac{s_A^2}{s_B^2} = 3,83$$

Caso se trate realmente de uma mesma população, o que implica em similaridade dos métodos, em 95% dos casos em que uma amostragem aleatória fosse realizada e o valor F_{cal} determinado ele seria igual ou estaria situado à esquerda da linha pontilhada.

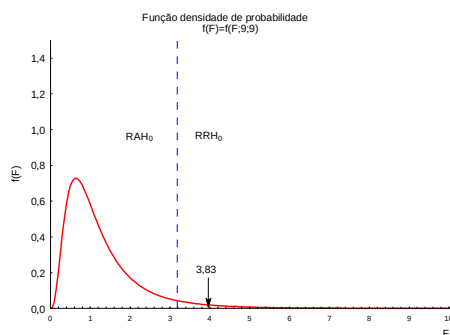


$$\int_0^{3,18} f(F) dF = 0,95 = 95\%$$

Nas mesmas condições anteriores (mesma população), em apenas 5% dos casos o valor F_{cal} assumiria valores iguais ou superiores a 3,18:

$$1 - \int_0^{3,18} f(F) dF = 1 - 0,95 = 0,05 = 5\%$$

Estes casos constituem o possível erro se decidirmos que os dados (resultados analíticos dos dois métodos) não podem ser considerados como provenientes de uma mesma população.



Portanto, como o valor de prova ($F_{\text{cal}} = 3,83$), e admitindo uma probabilidade de 5% de erro, deve-se decidir que os resultados produzidos pelos dois métodos não podem ser considerados como provenientes de uma mesma população.

A precisão dos métodos não pode ser considerada similar, significando que um método é mais preciso que o outro.

Implica dizer que o método (A: $s^2 = 1,76$) é menos preciso que o método (B: $s^2 = 0,46$), e que, para tomar esta decisão, admitiu-se um erro de 5%.

O significado do erro tipo I é muito claro:

- A razão entre duas estimativas da variância advindas de uma mesma população, oriundas de um par de amostras, cada uma com $n = 10$, pode assumir valores maiores ou iguais a 3,18 em 5% dos casos.
- Não se tem certeza absoluta se o caso analisado é, ou não, um desses possíveis casos.

Em síntese:

Consideraram-se os resultados das determinações dos dois métodos como sendo amostras aleatoriamente retiradas de uma mesma população básica, e admitiu-se que a variável aleatória, ou variável de resposta (determinação da CTC), apresenta distribuição normal.

A estatística F permitiu decidir, segundo uma determinada probabilidade de erro tipo I (em geral de 1 a 10%, o que implica em 99 a 90% de acerto, respectivamente), se a consideração inicial foi correta ou não, ou seja, se os resultados gerados pelos dois métodos podem ser considerados, ou não, como provenientes de uma mesma população básica:

Hipóteses:

$H_0 : \sigma_A^2 = \sigma_B^2$ (precisão igual $\Rightarrow$ população única)

$H_1 : \sigma_A^2 \gg \sigma_B^2$ (precisões distintas $\Rightarrow$ populações distintas)

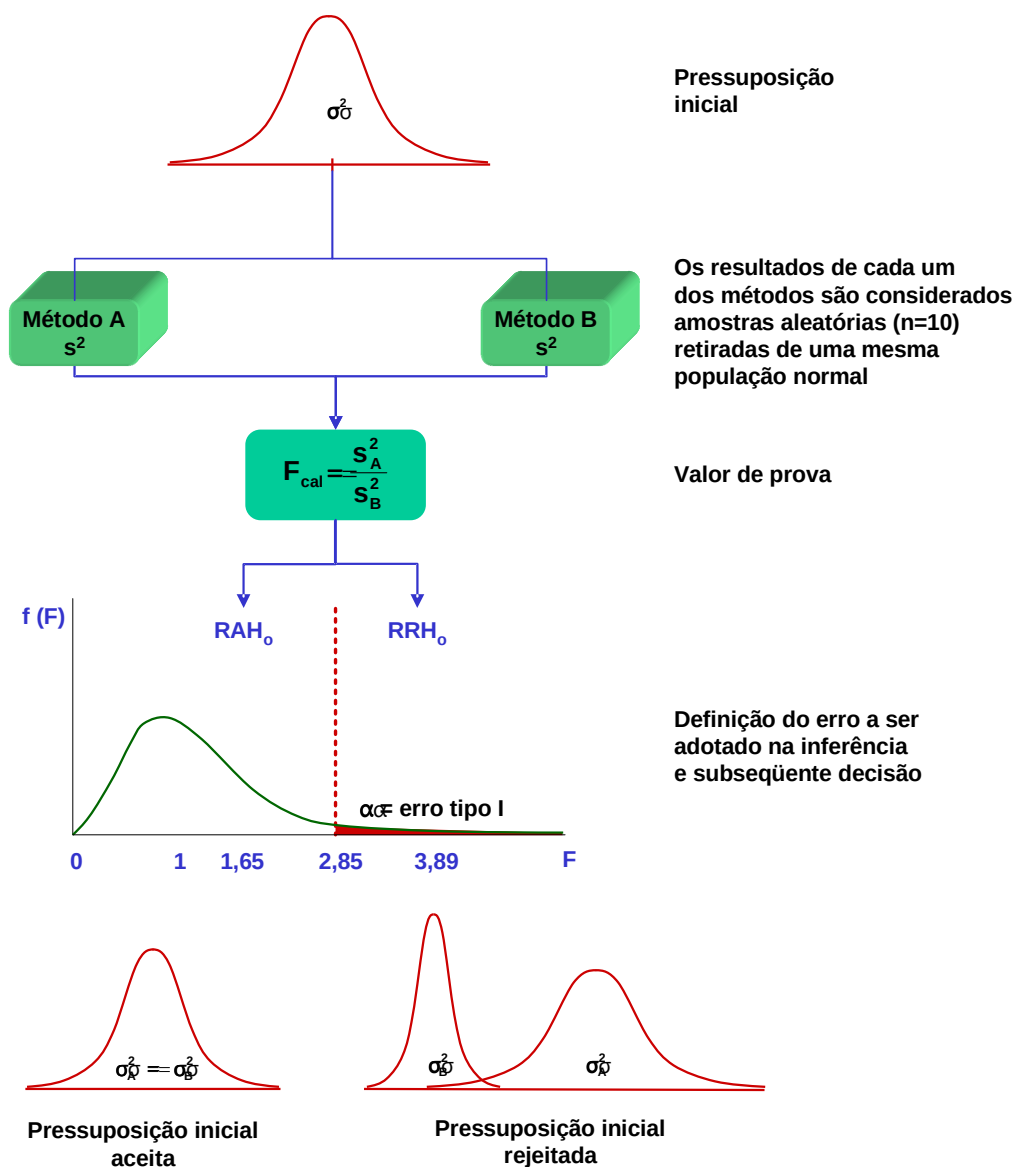


Figura 2.11 – Síntese do uso da distribuição F na inferência sobre precisão.

Denominando a linha pontilhada de F_{tab} :

$F_{cal} < F_{tab}$: aceita-se a igualdade

$F_{cal} \geq F_{tab}$: rejeita-se a igualdade

3. Análise de variância

3.1. Introdução

Análise de variância (ANOVA - ANalysis Of VAriance).

Alguns autores brasileiros preferem denominar ANAVA - ANálise de VAriância).

É uma técnica intensivamente utilizada pela estatística paramétrica para fazer inferências sobre médias populacionais através de suas estimativas, ou seja, das médias amostrais.

Nos experimentos agropecuários, em geral, o interesse é comparar:

Variedades

Manejo e alimentação de animais

Fontes e doses de fertilizantes

Preparos alternativos e métodos de conservação do solo

Formas de controle de pragas e doenças

Formas de controle de invasoras, etc.

A ANOVA é um procedimento básico para a tomada de decisão na avaliação de resultados experimentais.

3.2. Conceitos e uso

3.2.1. O que é?

A análise de variância de uma variável aleatória em estudo consiste na partição da soma de quadrados total dos desvios em relação à média em duas partes:

i. Uma parte associada às **fontes sistemáticas**, reconhecidas ou controladas de variação, ou seja, o que está em estudo: variedades, fertilizantes, rações, etc.

ii. Uma outra parte, de **natureza aleatória**, desconhecida ou não controlada, que constitui o erro experimental ou resíduo, medindo a influência dos erros: de mensuração e estocásticos.

3.2.2. Para que é usada?

Para fazer inferências sobre as médias populacionais pela comparação das médias amostrais.

3.2.3. Qual decisão é possível tomar?

Decidir, baseado na observação das amostras, segundo uma determinada probabilidade de erro, se as médias das populações dos tratamentos (o que está em estudo: variedades, fertilizantes, rações, etc) são estatisticamente iguais ou diferentes.

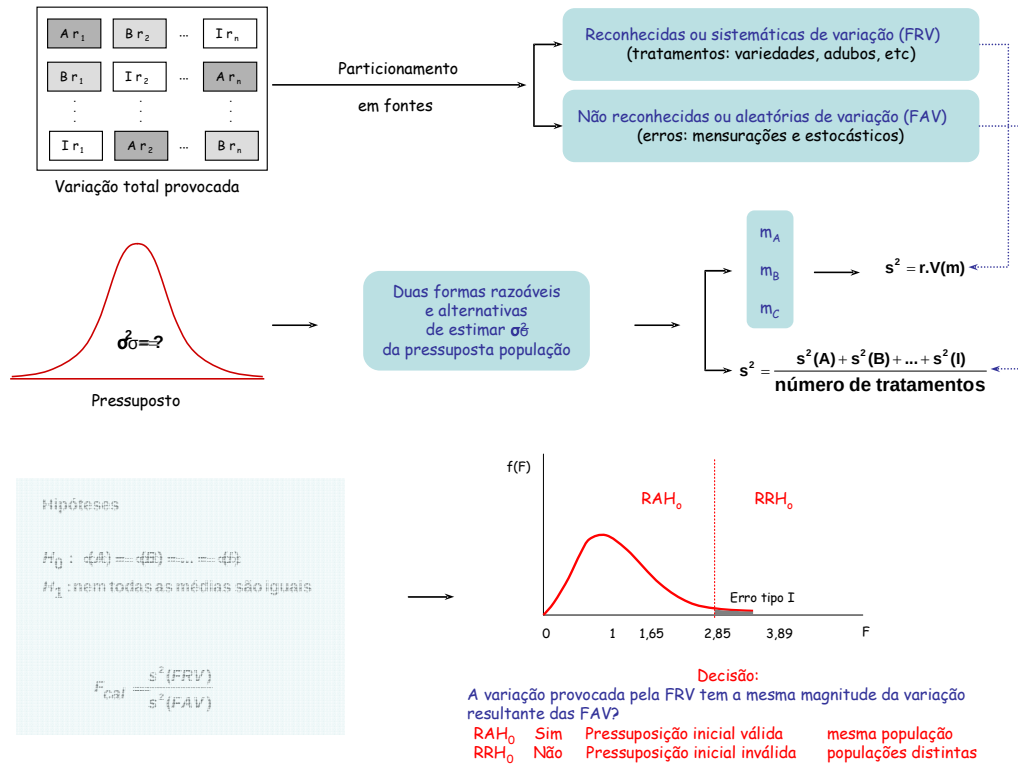


Figura 3.1 – Ilustração geral da análise de variância (modelo 1).

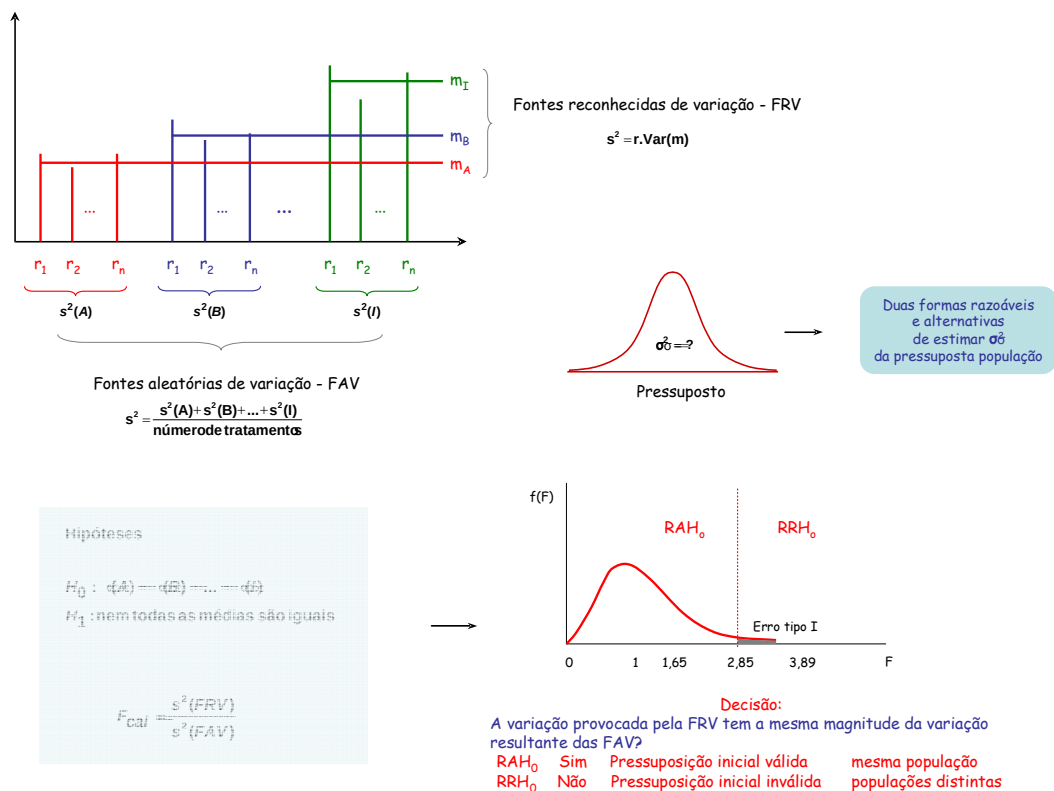


Figura 3.2 – Ilustração geral da análise de variância (modelo 2).

3.2.4. Exemplo

O desenvolvimento conceitual da análise de variância será feito a partir do resultado de um ensaio de produtividade de clones de cacau, abaixo transcrito, montado no delineamento inteiramente casualizado.

Produção de amêndoas (kg 10 plantas⁻¹ ano⁻¹) de cacau (5 anos)

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,0
B	60	55	66	61	54	61	357	6	59,5
C	59	47	44	49	62	60	321	6	53,5
D	45	33	34	48	42	44	246	6	41,0
	1.236							24	51,5

A questão a ser investigada (teste de hipóteses) é a seguinte: as produções dos clones de cacau são realmente diferentes?

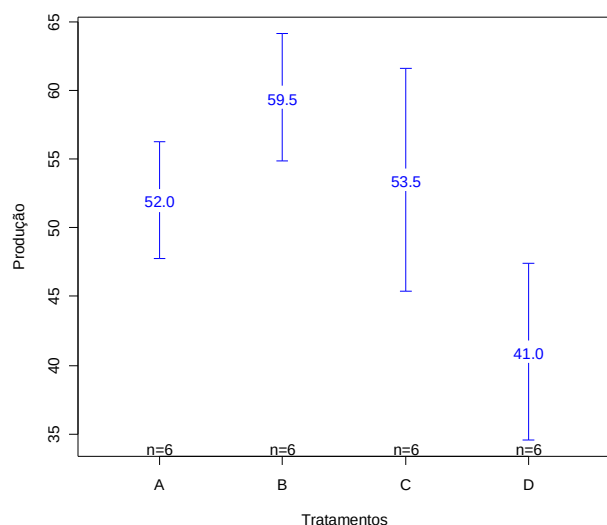


Figura 3.3 – Médias e dispersões dos tratamentos.

3.2.4.1. Teste de hipóteses

$$H_0: \alpha_A = \alpha_B = \alpha_C = \alpha_D$$

H_1 : Nem todas as médias são iguais

3.2.4.2. Procedimentos para a análise

a. Parte-se do pré-suposto de que cada tratamento é uma amostra – de tamanho igual ao número de repetições – retirada de uma mesma população, normalmente distribuída. Isto significa, a princípio, que as médias de todos os tratamentos são iguais, ou seja, iguais à média da pressuposta população.

b. Nestas condições, têm-se duas maneiras alternativas, e razoáveis, de estimar a variância da pressuposta população, σ^2 :

i. Tomar a média das variâncias de cada uma das amostras (ou tratamentos):

$$s^2 = \frac{\left(\frac{(58,0 - 52,0)^2 + \dots + (48,0 - 52,0)^2}{5} + \dots + \frac{(45,0 - 41,0)^2 + \dots + (44,0 - 41,0)^2}{5} \right)}{4} = 33,25$$

ii. Inferir σ^2 a partir da $V(m)$, isto é, a partir da variância da média amostral. Recordar que a variância da média amostral está relacionada com a variância da população da seguinte forma (teorema central do limite):

$$V(m) = \frac{\sigma^2}{n} \quad \therefore \quad \sigma^2 = n \cdot V(m)$$

Uma vez que n é conhecido, pois é o tamanho da amostra, ou melhor, o número de repetições do tratamento, é possível calcular $V(m)$:

$$V(m) = \frac{((52,0 - 51,5)^2 + (59,5 - 51,5)^2 + (53,5 - 51,5)^2 + (41,0 - 51,5)^2)}{3} = 59,5$$

$$s^2 = r \cdot V(m) = 59,5 \cdot 6 = 357,0$$

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,0
B	60	55	66	61	54	61	357	6	59,5
C	59	47	44	49	62	60	321	6	53,5
D	45	33	34	48	42	44	246	6	41,0
	1.236							24	51,5

c. Calcula-se o valor de prova, F_{cal} .

i. Foram obtidas duas estimativas da variância da pressuposta população básica (consideração inicial).

ii. Um teste estatístico – utilizando uma distribuição de probabilidades adequada – permitirá a conclusão se a consideração inicial é, ou não, válida.

iii. Como a distribuição de F fornece a distribuição de probabilidades do valor

F_{cal} :

$$F_{cal} = \frac{s^2}{s^2} = \frac{357,0}{33,25} = 10,74$$

pode-se usar esta distribuição e decidir se, de fato, a consideração inicial é, ou não, correta.

d. Estipulam-se as hipóteses

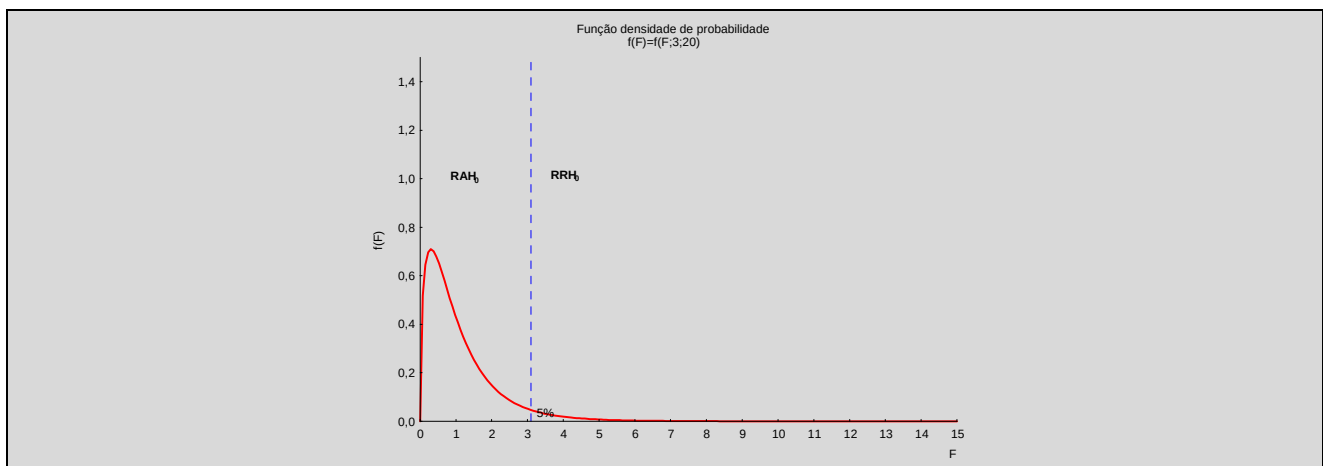
A partir do pré-suposto anteriormente estabelecido de que os tratamentos e suas repetições representam amostras feitas em uma mesma população básica, pode-se formular as seguintes hipóteses:

Hipóteses:

$H_0: \alpha_A = \alpha_B = \alpha_C = \alpha_D$	ou	$H_0: \text{Mesma população}$
$H_1: \text{Nem todas as médias são iguais}$	ou	$H_1: \text{Populações distintas}$

e. Adota-se um erro para a inferência

Para o exemplo será adotado um erro tipo I de 5%:



Se a consideração inicial for correta, ou seja, trata-se realmente de uma mesma população, em 95% das vezes, em média, que a razão entre duas estimativas da variância for calculada, F_{cal} , deveria ser encontrado um valor menor que 3,10, $P(F_{cal} < 3,10) = 95\%$. Neste caso a decisão seria aceitar H_0 .

Da mesma forma, em apenas 5% das vezes, também em média, que a relação fosse calculada, F_{cal} , seria encontrado um valor igual ou maior que 3,10, $P(F_{cal} \geq 3,10) = 5\%$. Neste caso a decisão seria rejeitar H_0 .

O erro tipo I (α) associado ao teste de hipóteses é muito claro: na situação “iii” seria rejeitada uma hipótese verdadeira. Isto é, os dados podem ser, de fato, provenientes

de uma mesma população básica, em outras palavras, valores F_{cal} iguais ou superiores a 3,10 podem efetivamente ocorrer, mas esses casos são muito raros, mais precisamente, ocorrem em média em apenas 5% dos casos.

A forma como se convencionou realizar o teste anterior é fornecida a seguir:

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,00
B	60	55	66	61	54	61	357	6	59,50
C	59	47	44	49	62	60	321	6	53,50
D	45	33	34	48	42	44	246	6	41,00
							1.236	24	51,50

$$C = (1.236)^2 / 24 = 63.654,00 \quad \therefore \quad \frac{(\sum y)^2}{n}$$

$$SQD_{\text{tot}} = [(58)^2 + (49)^2 + \dots + (44)^2] - C = 1.736,00$$

Observação:

Compare o cálculo efetuado acima, SQD_{tot} , e o cálculo posterior, que será efetuado no quadro da ANOVA, $SQD_{\text{tot}} / n-1$, com as duas fórmulas abaixo!

$$\sum y^2 - \frac{(\sum y)^2}{n}, \text{ é o numerador de uma fórmula muito conhecida: variância!}$$

$$s^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1}, \text{ o denominador, } n-1, \text{ são os graus de liberdade da ANOVA!}$$

$$SQD_{\text{tra}_m} = 6 [(52,00)^2 + (59,50)^2 + \dots + (41,00)^2] - C = 1.071,00$$

ou

$$SQD_{\text{tra}_t} = 1 / 6 [(312)^2 + (357)^2 + \dots + (246)^2] - C = 1.071,00$$

$$SQD_{\text{res}} = SQD_{\text{tot}} - SQD_{\text{tra}}$$

$$SQD_{\text{res}} = 1.736 - 1.071,00$$

$$SQD_{\text{res}} = 665,00$$

ANOVA

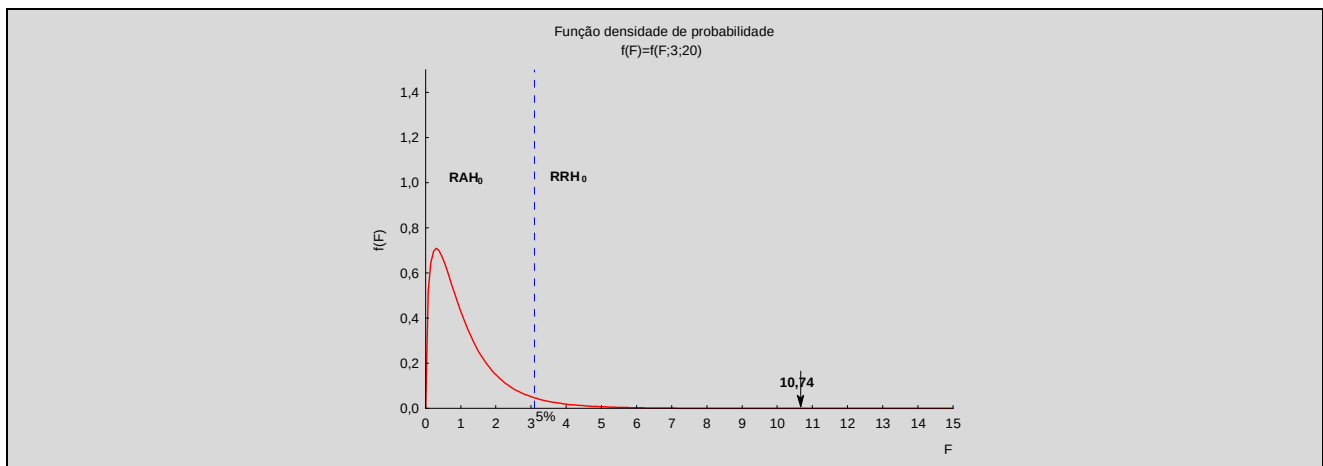
Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Tratamentos	3	1.071,00	357,00	10,74	0,0002
Resíduo	20	665,00	33,25		
Total	23	1.736,00			

Conclusão: Rejeita-se H_0 ao nível de 5% de probabilidade pelo teste F.

Observações:

A probabilidade do erro tipo I neste caso é de 0,02%.

Este valor (0,0002=0,02%) somente pode ser obtido com o uso de calculadoras adequadas ou via cálculo computacional.



ANOVA

Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Tratamentos	3	1.071,00	357,00	10,74	0,0002
Resíduo	20	665,00	33,25		
Total	23	1.736,00			

Observar que:

a. A soma de quadrados total dos desvios foi particionada em:

i. Uma parte associada à fonte reconhecida ou controlada de variação, ou seja, os tratamentos. Esta variação é denominada variação entre os tratamentos.

ii. Outra parte de natureza aleatória, não reconhecida ou não controlada, associada ao erro experimental ou resíduo. Esta variação é também denominada variação dentro dos tratamentos.

b. O erro experimental ou resíduo quantifica a variação observada dentro de cada tratamento, considerando todos os tratamentos. Possui duas causas:

- i. Erros de mensuração que ocorrem em todo o ciclo experimental (montagem, condução, coleta). Erros de medidas, pesagens, arredondamentos, etc.
- ii. Erros estocásticos, relacionados a irreprodutividade inerente os fenômenos biológicos.

Exemplos de alguns fatores relacionados a irreprodutividade:

As sementes ou mudas não são exatamente iguais.

As condições ambientais não são exatamente iguais para todas as unidades experimentais.

Enfim, não é possível garantir igualdade (material experimental e condições ambientais), para todos os fatores que podem influenciar a resposta da variável aleatória em estudo (produção dos clones de cacau).

3.2.5. Pressupostos da análise de variância

Para se usar a ANOVA na inferência estatística três pressuposições básicas devem ser atendidas:

 Para cada população, a variável de resposta é distribuída normalmente.

Implicação no exemplo: a produção de amêndoas de cacau precisa ser distribuída normalmente em cada clone.

 A variância da variável de resposta é a mesma para todas as populações.

Implicação no exemplo: as variâncias das produções de amêndoas de cacau precisam ser estatisticamente iguais (ou homogêneas) para todos os clones. Esta pressuposição recebe a denominação de invariância da variância ou homocedasticia.

 As observações precisam ser independentes.

Implicação no exemplo: a produção de amêndoas de cacau, para uma determinada repetição de um clone, precisa ser independente da produção de qualquer outra repetição do mesmo clone ou de clones diferentes. Em outras palavras, o erro de qualquer repetição não tem efeito sobre qualquer outra repetição do mesmo tratamento ou de tratamentos diferentes.

Em função da distribuição F ser considerada robusta, a inferência via ANOVA é ainda possível e eficiente, se os dados experimentais apresentarem ligeiros afastamentos (violações) das condições ideais (pressupostos).

Em casos de desvios acentuados das condições ideais, pode-se tentar o artifício, as vezes bem sucedido, da transformação dos dados. Por outro lado, os procedimentos da estatística não paramétrica (similares aos da paramétrica) devem ser usados nesses casos.

3.2.6. Demonstração da aplicação do teorema central do limite (TCL) na ANOVA

Em estatística experimental: $n = r$ (número de repetições)

$$s^2 = r \cdot V(m) \quad (1) \quad \text{TCL}$$

$$V(m) = \frac{SQD(m)}{n_m - 1} = \frac{\sum m^2 - \frac{(\sum m)^2}{n_m}}{n_m - 1} \quad (2) \quad n_m = \text{número de médias}$$

Substituindo (2) em (1)

$$s^2 = \frac{r \cdot \left[\sum m^2 - \frac{(\sum m)^2}{n_m} \right]}{n_m - 1} = \frac{\sum rm^2 - \frac{(\sum rm)^2}{n_m}}{n_m - 1}$$

$$\text{Pode se verificar que } C = \frac{(\sum y)^2}{n} = \frac{(\sum r_i \cdot m_i)^2}{\sum r_i}$$

Tratamentos com mesmo número de repetições: $r_1 = \dots = r_k = r$

$$\frac{(\sum r_i \cdot m_i)^2}{\sum r_i} = \frac{(\sum r \cdot m)^2}{r \cdot n_m} = \frac{(r \cdot \sum m)^2}{r \cdot n_m} = \frac{r^2 \cdot (\sum m)^2}{r \cdot n_m} = \frac{r \cdot (\sum m)^2}{n_m}$$

Assim

$$s^2 = \frac{r \cdot \left(\sum m^2 \right) - \frac{r \cdot (\sum m)^2}{n_m}}{n_m - 1} = \frac{r \cdot \left(\sum m^2 \right) - C}{n_m - 1} = \frac{6 \cdot \left[(52)^2 + \dots + (41)^2 \right] - C}{n_m - 1} = 357,0 = QMDtra$$

Exemplo ilustrativo da igualdade

$$\frac{(\sum y_i)^2}{n} = r \cdot \left[\frac{(\sum m_i)^2}{n_{mi}} \right] = C$$

considerando um mesmo número de repetições: $r_i = \dots = r_k = r$

	Repetições			r	m _i
	1	2	3		
A	1	2	3	3	2
B	4	5	6	3	5
Soma					15

$$\frac{(\sum y_i)^2}{n} = \frac{(21)^2}{6} = 73,5$$

$$r \left[\frac{(\sum m_i)^2}{n_{mi}} \right] = 3 \left[\frac{(7)^2}{2} \right] = 73,5$$

Cálculo da SQDtra utilizando médias e total de tratamentos com mesmo número de repetições aplicando o TCL:

Usando médias :

$$SQDtra_m = 6 \cdot \left[(52)^2 + \dots + (41)^2 \right] - C$$

Usando totais de tratamentos :

$$SQDtra_t = 6 \cdot \left[\left(\frac{312}{6} \right)^2 + \dots + \left(\frac{246}{6} \right)^2 \right] - C$$

$$SQDtra_t = 6 \cdot \left[\frac{(312)^2}{36} + \dots + \frac{(246)^2}{36} \right] - C$$

$$SQDtra_t = \frac{6}{36} \cdot \left[(312)^2 + \dots + (246)^2 \right] - C$$

$$SQDtra_t = \frac{1}{6} \cdot \left[(312)^2 + \dots + (246)^2 \right] - C$$

4. Noções básicas de experimentação

4.1. Introdução

Muito do que a humanidade adquiriu ao longo dos séculos foi através da experimentação.

A experimentação, entretanto, somente se definiu como técnica sistemática de pesquisa neste século, quando foi formalizada através da estatística.

Somente por meio da experimentação uma nova técnica poderá ser divulgada, com embasamento científico, sem desperdício de tempo e recursos financeiros, resguardando a credibilidade do pesquisador.

4.2. Público

Pesquisadores: necessitam de uma base sólida para planejar, executar, analisar e interpretar resultados de experimentos.

Extensionistas e técnicos: necessitam entender os experimentos e sua natureza, avaliar a confiabilidade dos resultados e trocar idéias com os pesquisadores pelo uso da linguagem técnica adequada.

4.3. Principais conceitos

Experimentação: é uma parte da estatística probabilística que estuda o planejamento, a execução, a coleta de dados, a análise e a interpretação dos resultados de experimentos.

Experimento: é um procedimento planejado, com base em hipóteses, com o objetivo de provocar variação em uma ou mais variáveis de resposta (variáveis aleatórias) no estudo de fenômenos ou processos, sob condições controladas.

Provocar variação: equivale a testar diferentes alternativas (tratamentos) no estudo dos fenômenos ou processos.

Exemplos:

Diferentes formas de:

Manejar ou alimentar um rebanho Combater doenças e pragas Adubar as culturas, etc.
--

Condições controladas: permite que o estudo seja repetido, o que é um fundamento do método científico.

Um experimento é constituído basicamente de um conjunto de unidades experimentais sobre as quais são aplicados os tratamentos, e das quais são obtidos os dados experimentais.

Parcela: termo de uso mais antigo para se referir a uma unidade de área do experimento e tem sido substituído por unidade experimental.

Unidade experimental (UE): trata-se de uma unidade de área, um conjunto de indivíduos ou uma parte de um indivíduo, sobre a qual um tratamento é aplicado e seus efeitos avaliados.

Unidade de observação (UO): trata-se da menor parte indivisa de uma unidade experimental.

Exemplos:

UNIDADE EXPERIMENTAL	UNIDADE DE OBSERVAÇÃO
GRUPO DE PLANTAS	UMA PLANTA
GRUPO DE ANIMAIS	UM ANIMAL
FOLHAS DE UMA PLANTA	CADA FOLHA DA PLANTA

Tratamentos: Identifica o que está em comparação e podem ser qualitativos ou quantitativos:

Qualitativos: diferenciam-se por suas qualidades, não podendo ser ordenados por algum critério numérico.
Exemplos: tipos, cultivares, métodos, espécies, marcas, etc.

Quantitativos: podem ser ordenados segundo algum critério numérico.
Exemplos: doses, idade, tempo, distâncias, densidade, etc.

Variáveis de resposta: são mensuradas nas unidades experimentais e estão sujeitas às variações provocadas pelas fontes reconhecidas (sob controle do pesquisador) e aleatórias ou não reconhecidas (fora de controle do pesquisador).

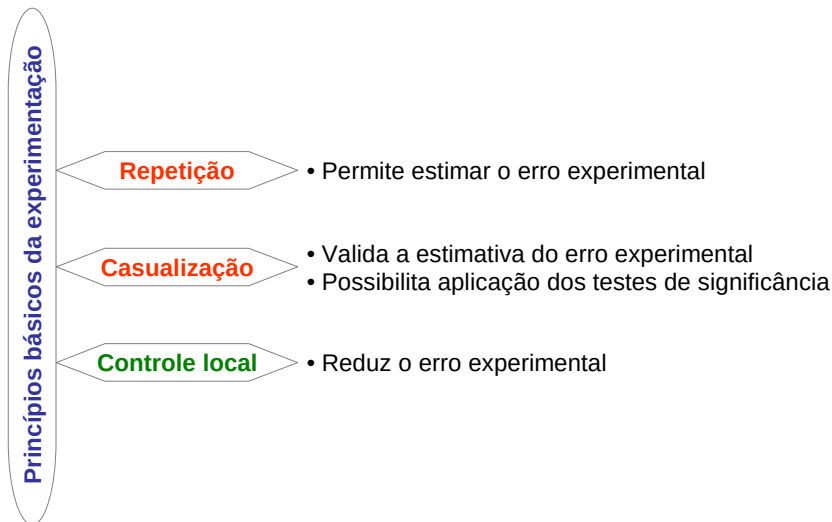
4.4. A origem agrícola

Boa parte da formalização que existe hoje em experimentação deve-se a Fisher (1890-1962), um estatístico que trabalhou na Estação Experimental de Agricultura de Rothamstead, na Inglaterra.



É a origem agrícola da experimentação que explica o uso de vários termos técnicos como parcela e tratamento

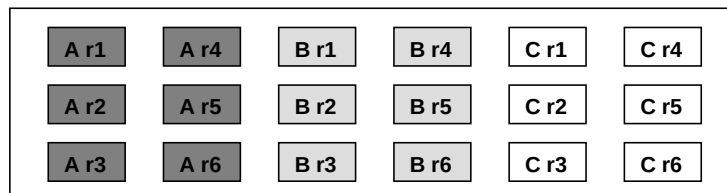
4.5. Princípios básicos da experimentação



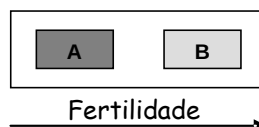
4.5.1. Repetição

A idéia em experimentação é comparar grupos, não apenas unidades experimentais.

As unidades experimentais de um mesmo grupo são consideradas repetições:



Se tivermos duas variedades de milho, A e B, plantadas em uma mesma área, o fato de A ter produzido mais do que B pouco significa, pois muitas explicações, além da variedade, por exemplo, podem justificar o resultado obtido:



Poderemos, porém, tentando contornar o problema, semear diversas parcelas com A e diversas parcelas com B e tomar a produção média de cada variedade: é onde intervém o princípio da repetição:

A r1	A r4	B r1	B r4	C r1	C r4
A r2	A r5	B r2	B r5	C r2	C r5
A r3	A r6	B r3	B r6	C r3	C r6

O número de repetições que devem ser utilizados em determinado experimento pode ser calculado através de fórmulas matemáticas. Estas fórmulas, entretanto, exigem que se tenham informações estatísticas anteriores sobre a variabilidade, o que, em geral não acontece.

O mais importante é a variabilidade do material experimental: quanto mais homogêneo menor o número de repetições necessárias para mostrar, com clareza, o efeito de um tratamento.

Do ponto de vista estatístico é sempre desejável que os experimentos tenham grande número de repetições, este número, entretanto, é limitado pelos recursos disponíveis (humanos, materiais, tempo, financeiros, etc).

Recomenda-se a adoção do que é usual na área de pesquisa, pois é através da repetição que se estima o erro experimental. Em geral planeja-se o experimento (tratamentos e repetições) de forma que se tenha, como recomendação prática geral, 12 ou mais gl associados ao resíduo.

Toda variação não explicada é tratada como variação casual (aleatória) e irá constituir o erro experimental.

4.5.2. Casualização

Foi formalmente proposta por Fischer na década de 1920.

Não casualizado:

A r1	A r4	B r1	B r4	C r1	C r4
A r2	A r5	B r2	B r5	C r2	C r5
A r3	A r6	B r3	B r6	C r3	C r6

Casualizado:

A r1	B r2	A r3	C r3	B r6	C r5
B r1	C r2	A r4	B r4	A r5	A r6
C r1	A r2	B r3	B r5	C r4	C r6

Vinte anos mais tarde esta técnica já estava definitivamente incorporada à experimentação agrícola.

Na área industrial passou a ser rotina após a II guerra mundial.

Na pesquisa médica, entretanto, só começou a ser aceita mais tarde (questões éticas e natureza do material experimental).

O princípio da casualização é uma das maiores contribuições dos estatísticos à ciência experimental.

Somente a casualização garante que as possíveis diferenças entre os tratamentos não sejam devidas ao favorecimento de um em detrimento aos demais (tendenciosidade).

Uma vez que tais diferenças existam, a utilização do princípio garante que elas não se devam a nenhum favorecimento.

É através da casualização que os erros experimentais tornam-se independentes, o que possibilitará os testes de significância.

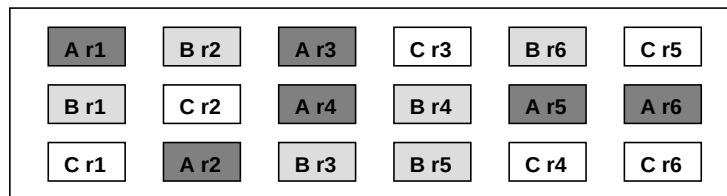
4.5.3. Controle local

É um princípio de uso muito freqüente, mas não obrigatório.

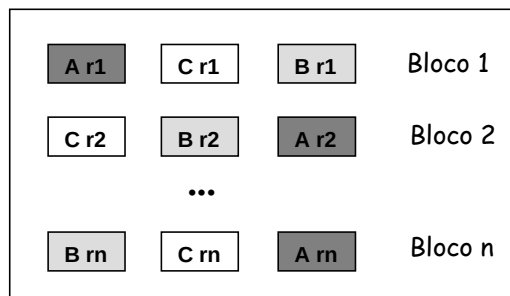
É uma forma de homogeneizar as condições experimentais.

Delineamentos mais usados:

Inteiramente casualizado (DIC):



Blocos casualizados (DBC):



Quadrado latino (DQL):

A 11	B 12	C 13	Linha 1
B 21	C 22	A 23	Linha 2
C 31	A 32	B 33	Linha 3
Coluna 1	Coluna 2	Coluna 3	

O controle local conduz sempre a uma diminuição do número de graus de liberdade associados ao erro experimental (ou resíduo), o que é, a princípio, indesejável.

Entretanto, quando ocorre uma diminuição considerável da variância residual, como em geral acontece quando o princípio é bem aplicado, o experimento apresenta maior precisão, melhorando, como consequência, a qualidade e a confiabilidade das inferências estatísticas.

4.6. Controle de qualidade de experimentos

Informações sobre qualidade orientam o pesquisador sobre os cuidados a serem tomados no planejamento, execução e análise dos resultados do experimento para manter o erro experimental em níveis aceitáveis.

A qualidade de um experimento pode ser avaliada, de forma comparativa, pela magnitude do erro experimental, que reflete a influência de todas as fontes não reconhecidas de variação sobre as variáveis de resposta.

A magnitude do erro experimental, por sua vez, pode ser avaliada pelo coeficiente de variação:

$$cv = \frac{s}{m} \cdot 100 \quad \therefore \quad cv = \frac{\sqrt{QMDres}}{m} \cdot 100$$

A precisão de um experimento pode ser considerada como alta, média ou baixa somente em relação a um grupo de experimentos semelhantes.

A título de ilustração são reproduzidas duas tabelas, ainda que genéricas, propondo classificações e apresentando informações estatísticas sobre qualidade de experimentos:

Tabela 4.1 – Classificação dos experimentos quanto aos coeficientes de variação

CLASSES DE CV	LIMITES DO CV, %	PRECISÃO
BAIXOS	≤ 10	ALTA
MÉDIOS	10-20	MÉDIA
ALTOS	20-30	BAIXA
MUITO ALTOS	≥ 30	MUITO BAIXA

Fonte: Gomes F.P. (1990)

Tabela 4.2 – Médias de coeficientes de variação (cv) e seu desvio padrão (s) sobre n experimentos, para algumas culturas e variáveis

CULTURA	VARIÁVEL	N	CV, %	S
ALGODÃO	RENDIMENTO	33	14,6	6,0
AMENDOIM	RENDIMENTO	10	13,6	7,4
CEREAIS DE INVERNO	RENDIMENTO	144	15,6	6,7
	DOENÇAS	7	18,4	10,8
MILHO	RENDIMENTO	205	14,7	8,9
	ALTURA-ESPIGA	24	12,8	4,0
	PESO- PARTE AÉREA	5	20,9	5,2
	PESO-RAIZ	5	33,1	18,0
	PESO-TOTAL	5	21,4	6,7
PLANTAS ARBÓREAS	PESO DE FRUTOS	62	40,6	26,7
	ALTURA	21	16,3	7,5
	NÚMERO DE FRUTOS	16	23,3	13,7

Fonte: Storck et all. (2000)

4.7. [Tipos de erros em experimentos](#)

Aleatório ou experimental: decorrente dos erros de mensuração e estocásticos, podendo ser reduzidos, mas nunca eliminados.

Sistemático: Tem origem no descuido ou na falta de equanimidade do experimentador ou de pessoas envolvidas. Dá-se quando determinado tratamento é favorecido (ou desfavorecido) em todas, ou na maioria, de suas repetições.

4.7.1. Principais fontes de erro e respectivos cuidados

4.7.1.1. Heterogeneidade das condições ambientais

Deve ser feito um ensaio em branco ou ensaio de uniformidade sem tratamentos para sua avaliação.

4.7.1.2. Heterogeneidade do material experimental

Realizar seleção rigorosa objetivando maximizar a padronização do material experimental ou adoção de controle local.

4.7.1.3. Condução diferenciada das unidades experimentais

Evitar tendenciosidade e manter um padrão equânime dos tratos necessários durante toda condução do experimento.

4.7.1.4. Competição intraparcelar

É muito difícil avaliar a influência da perda de uma unidade de observação devido à compensação do dossel pela menor competição além de provocar subestimação da variabilidade experimental.

Recomenda-se aumentar a densidade inicial e ir realizando periodicamente o descarte das unidades de observação pouco representativas, ou seja, as muito pouco desenvolvidas (irão subestimar o grupo ou tratamento) e as super desenvolvidas (irão superestimar o grupo ou tratamento), em relação às variáveis de resposta que se pretende avaliar.

4.7.1.5. Competição interparcelar

Descartar as unidades de observação que podem receber a influência dos tratamentos adjacentes (bordadura) e adotar como parcela útil às unidades de observação não influenciadas pelas adjacentes.

4.7.1.6. Pragas, doenças e acidentes

Deve-se realizar a avaliação do dano provocado e a influência da forma de controle sobre as variáveis de resposta, assim como, avaliação da possível repetição do experimento.

4.8. Planejamento de experimentos

O planejamento objetiva determinar, com antecedência, como será o experimento e como serão analisados os dados.

O projeto deve ser simples e suficientemente claro para que na falta de quem o planejou, outro pesquisador possa executá-lo, analisá-lo e obter conclusões.

Consultar STORK et al. (2000).

5. Delineamento inteiramente casualizado - DIC

5.1. Introdução

É o mais simples de todos os delineamentos experimentais. Os experimentos instalados de acordo com este delineamento são denominados experimentos inteiramente casualizados (DIC) ou experimentos ao acaso.

Para se utilizar este delineamento é necessário similaridade nas unidades experimentais. Como princípio norteador básico, a única diferença entre as unidades experimentais deve ser aquilo que está sendo testado, ou seja, os tratamentos, tudo o mais deve ser similar ou homogêneo.

Somente é eficiente nessas condições, ou seja, se for observada homogeneidade tanto das condições ambientais - que influenciam a manifestação do fenômeno, como do material experimental, anteriormente à aplicação dos tratamentos.

Devido a isto, seu uso mais comum se dá em condições controladas, ou seja, casas de vegetação, laboratórios, etc. Em condições de campo é necessário atenção em relação à(s) influência(s) das fontes de variação sistemáticas que podem reduzir a precisão do experimento, que em consequência, reduz as chances de se detectar diferenças entre os tratamentos, caso elas existam.

Os tratamentos são dispostos nas parcelas de forma inteiramente ao acaso, isto é, sem qualquer restrição do local que cada unidade experimental - associada a um tratamento, irá ocupar na área experimental.

5.2. Princípios utilizados

5.2.1. Repetição

Permite a estimativa do erro experimental ou resíduo.

Dependente da variabilidade do material experimental.

5.2.2. Casualização

Garante que as possíveis diferenças entre os tratamentos não sejam por favorecimento.

5.2.3. Vantagens e desvantagens

5.2.3.1. Vantagens

Flexibilidade quanto a número de tratamentos e repetições, embora um mesmo número de repetições seja desejável.

Análise de variância simples, mesmo se houver a perda de algumas unidades experimentais.

É o delineamento que apresenta o maior número de graus de liberdade associados ao resíduo.

5.2.3.2. Desvantagens

Muitas vezes é ineficiente, devido à presença de fontes de variação sistemáticas não controladas.

Pode ocorrer superestimação do erro experimental.

5.3. Modelo estatístico

$$Y_{ij} = \mu + t_i + e_{ij}$$

onde:

Y_{ij} = valor observado na parcela do tratamento i na repetição j

μ = média geral do experimento

t_i = efeito do tratamento i aplicado na parcela

e_{ij} = efeito dos fatores não controlados

5.4. Esquema de casualização dos tratamentos

Seja um experimento de comparação de produtividade de clones de cacau resistentes a vassoura de bruxa, envolvendo 4 tratamentos (A, B, C, D) em 6 repetições (24 unidades experimentais ou parcelas):

A (r_1)	B	D	B	A (r_5)	C
C	A (r_2)	B	D	C	B
A (r_6)	D	B	C	D	A (r_3)
C	D	A (r_4)	B	C	D

Figura 5.1 – Esquema da casualização das unidades experimentais.

5.5. Coleta de dados

Quadro 5.1 – Quadro para coleta de dados de experimentos no DIC

Tratamentos	Repetições			Totais	N.Repetições	Médias
	1	...	j			
A	y_{11}	...	y_{1j}	t_1	j	m_1
B	y_{21}	...	y_{2j}	t_2	j	m_2
.	.	.	.	.	.	.
.	.	.	.	.	.	.
.	.	.	.	.	.	.
i	y_{i1}	...	y_{ij}	t_i	j	m_i

Simbologia adotada: $y_{(tra,rep)}$

5.6. Análise de variância

5.6.1. Esquema da análise de variância

Quadro 5.2 – Quadro da análise de variância no DIC

Causa da variação	GL	SQD	QMD	F_{cal}
Tratamentos	i - 1	SQDtra	QMDtra	QMDtra/QMDres
Resíduo	i(j - 1)	SQDres	QMDres	
Total	ij - 1	SQDtot		

5.6.2. Teste de hipóteses

Em relação às médias populacionais

$H_0: \alpha_A = \alpha_B = \dots = \alpha_D$ $H_1: \text{Nem todas as } \alpha_I \text{ são iguais}$	ou	$H_0: \alpha_I = \alpha_K \text{ (para todo } I \neq K)$ $H_1: \text{Não } H_0$
--	----	--

Em relação ao modelo estatístico

$H_0: t_A = t_B = \dots = t_D = 0$ $H_1: \text{Nem todos os } t_I \text{ são iguais a zero}$	ou	$H_0: t_I = 0 \text{ (para todo } I)$ $H_1: \text{Não } H_0$
---	----	---

5.7. Exemplo com um mesmo número de repetiçõesQuadro 5.3 – Produção de amêndoas (kg 10 plantas⁻¹ ano⁻¹) de cacau aos 5 anos

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,00
B	60	55	66	61	54	61	357	6	59,50
C	59	47	44	49	62	60	321	6	53,50
D	45	33	34	48	42	44	246	6	41,00
							1.236	24	51,50

$$C = (1.236)^2 / 24 = 63.654,00$$

$$SQD_{tot} = [(58)^2 + (49)^2 + \dots + (44)^2] - C = 1.736,00$$

$$SQD_{tra_m} = 6 [(52,00)^2 + (59,50)^2 + \dots + (41,00)^2] - C = 1.071,00$$

ou

$$SQD_{tra_t} = 1 / 6 [(312)^2 + (357)^2 + \dots + (246)^2] - C = 1.071,00$$

$$SQD_{res} = SQD_{tot} - SQD_{tra} = 1.736 - 1.071,00 = 665,00$$

Hipóteses:

$$H_0: \alpha_I = \alpha_K \text{ (para todo } I \neq K)$$

$$H_1: \text{Nem todas as } \alpha_I \text{ são iguais}$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	3	1.071,00	357,00	10,74	0,0002
Resíduo	20	665,00	33,25		
Total	23	1.736,00			

$$cv = 100 \cdot (\sqrt{33,25/51,50}) = 11,20\%$$

Rejeita-se H_0 . Conclui-se que existe pelo menos um contraste entre médias de tratamentos estatisticamente diferente de zero, ao nível de 5% de probabilidade, pelo teste F.

5.7.1. Resíduo

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,00
B	60	55	66	61	54	61	357	6	59,50
C	59	47	44	49	62	60	321	6	53,50
D	45	33	34	48	42	44	246	6	41,00
							1.236	24	51,50

$$\begin{aligned} \text{Resíduo} = & [[(58 - 52,00)^2 + \dots + (48 - 52,00)^2] / 5 + \\ & [(60 - 59,50)^2 + \dots + (61 - 59,50)^2] / 5 + \\ & [(59 - 53,50)^2 + \dots + (60 - 53,50)^2] / 5 + \\ & [(45 - 41,00)^2 + \dots + (44 - 41,00)^2] / 5] / 4 = 33,25 \text{ (kg 10 plantas}^{-1} \text{ ano}^{-1})^2 \end{aligned}$$

O erro experimental (ou resíduo na ANOVA) é uma média aritmética das estimativas das variâncias dos tratamentos envolvidos na análise e quantifica a influência de todas as fontes de variação não controladas no experimento.

5.7.2. O coeficiente de variação e sua interpretação

$$cv = 100 \cdot \frac{\sqrt{33,25}}{51,50} = 11,20\%$$

O coeficiente de variação (cv) é uma medida relativa de dispersão, útil para a comparação, em termos relativos, do grau de concentração dos dados em torno da média.

É utilizado, muitas vezes, para comparar a variabilidade de diferentes experimentos, sobre um mesmo assunto, fornecendo uma idéia do quão preciso foi cada um dos experimentos.

Um mesmo experimento, conduzido de formas diferentes, pode originar resultados diferentes. A simples observação do cv pode informar o quão preciso foi cada um dos experimentos, complementando interpretação dos resultados.

É uma informação importante e deve ser apresentada após o quadro da ANOVA de todas as análises estatísticas de experimentos.

5.7.3. Testes de comparação de médias múltiplas

Quadro 5.4 – Comparação da sensibilidade dos diferentes testes de médias múltiplas

Clones	Média	Tukey	Duncan	SNK	t	Dunnett
B	59,50	a	a	a	a	Testemunha
C	53,50	a	a b	a	a	n.s
A	52,00	a	b	a	b	n.s
D	41,00	b	c	b	c	*

Obs: realizar os testes de Tukey, Duncan e SNK para aprendizagem e treinamento.

5.7.4. Hipóteses para os contrastes

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

5.7.5. Desdobramento dos gl associados a tratamentos em contrastes ortogonais

$$C_1 = (A, D) \text{ vs. } (B, C)$$

$$C_2 = A \text{ vs. } D$$

$$C_3 = B \text{ vs. } C$$

Estabelecendo os contrastes:

$$C_1 = 1A + 1D - 1B - 1C$$

$$C_2 = 1A - 1D$$

$$C_3 = 1B - 1C$$

Inicialmente calculamos as estimativas dos contrastes:

$$\hat{C}_1 = 1(312) + 1(246) - 1(357) - 1(321) = -120$$

$$\hat{C}_2 = 1(312) - 1(246) = 66$$

$$\hat{C}_3 = 1(357) - 1(321) = 36$$

Agora podemos calcular a soma de quadrados dos contrastes:

$$SQD(C_1) = (-120)^2 / 6 [(1)^2 + (1)^2 + (-1)^2 + (-1)^2] = 600,00$$

$$SQD(C_2) = (66)^2 / 6 [(1)^2 + (-1)^2] = 363,00$$

$$SQD(C_3) = (36)^2 / 6 [(1)^2 + (-1)^2] = 108,00$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(3)	(1.071,00)			
(A, D) vs (B, C)	1	600,00	600,00	18,05	0,0004
A vs D	1	363,00	363,00	10,92	0,0035
B vs C	1	108,00	108,00	3,25	0,0866
Resíduo	20	665,00	33,25		
Total	23	1.736,00			

Clones	Média
B	59,50
C	53,50
A	52,00
D	41,00

Conclusões com erro tipo I de 5%:

$C_1 = (A, D) \text{ vs } (B, C) : \text{Rejeita-se } H_0$
 $C_2 = A \text{ vs } D : \text{Rejeita-se } H_0$
 $C_3 = B \text{ vs } C : \text{Aceita-se } H_0$

5.8. Exemplo com número diferente de repetiçõesQuadro 5.5 – Produção de amêndoas (kg 10 plantas⁻¹ ano⁻¹) de cacau aos 5 anos

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	-	51	56	50	48	263	5	52,60
B	60	55	66	61	54	61	357	6	59,50
C	59	47	44	-	62	-	212	4	53,00
D	45	-	34	48	42	44	213	5	42,60
							1.045	20	52,25

$$C = (1.045)^2 / 20 = 54.601,25$$

$$SQD_{tot} = [(58)^2 + (51)^2 + \dots + (44)^2] - C = 1.297,75$$

$$SQD_{tra_t} = [(263^2/5) + (357^2/6) + \dots + (213^2/5)] - C = 783,85$$

$$SQD_{tra_m} = [5 (52,60)^2 + 6 (59,50)^2 + \dots + 5 (42,60)^2] - C = 783,85$$

$$SQD_{res} = SQD_{tot} - SQD_{tra} = 1.297,75 - 783,85 = 513,90$$

Hipóteses:

$$H_0: \alpha_I = \alpha_K \text{ (para todo } I \neq K)$$

$$H_1: \text{Nem todas as } \alpha_I \text{ são iguais}$$

ANOVA

Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Tratamentos	3	783,85	261,28	8,13	0,0016
Resíduo	16	513,90	32,12		
Total	19	1.297,75			

Rejeita-se H_0 . Conclui-se que existe pelo menos um contraste entre médias de tratamentos estatisticamente diferente de zero, ao nível de 5% de probabilidade, pelo teste F.

5.8.1. Desdobramento dos gl associados a tratamentos em contrastes ortogonais

Como temos três graus de liberdade associados a tratamentos podemos estabelecer até três contrastes ortogonais, mantendo os contrastes anteriores:

$$C_1 = (A, D) \text{ vs. } (B, C)$$

$$C_2 = A \text{ vs. } D$$

$$C_3 = B \text{ vs. } C$$

Uma forma prática para se estabelecer contrastes ortogonais entre totais de tratamentos de experimentos desbalanceados é a seguinte:

- Escrevem-se os totais de tratamentos envolvidos na comparação.
- Atribuí-se sinal positivo aos totais de um grupo e negativo aos totais do outro grupo.
- Verifica-se o número de repetições, n_1 , envolvidos no primeiro grupo e o número de repetições, n_2 , envolvidos no segundo grupo. Calcula-se o m.m.c. entre n_1 e n_2 .
- Divide-se o m.m.c. por n_1 ; o resultado será o coeficiente de cada total do primeiro grupo.
- Divide-se o m.m.c. por n_2 ; o resultado será o coeficiente de cada total do segundo grupo.

5.8.2. Estimação e teste de hipóteses para os contrastes

(5) (5) (6) (4)	Repetições envolvidas
$C_1 = 1A + 1D - 1B - 1C$	m.m.c.(10;10) = 10
(5) (5)	
$C_2 = 1A - 1D$	m.m.c.(5;5) = 5
(6) (4)	
$C_3 = 2B - 3C$	m.m.c.(6;4) = 12

Inicialmente calculamos as estimativas dos contrastes:

$$\hat{C}_1 = 1(263) + 1(213) - 1(357) - 1(212) = -93$$

$$\hat{C}_2 = 1(263) - 1(213) = 50$$

$$\hat{C}_3 = 2(357) - 3(212) = 78$$

Agora podemos calcular a soma de quadrados dos contrastes:

$$SQD(C_1) = (-93)^2 / [5(1)^2 + 5(1)^2 + 6(-1)^2 + 4(-1)^2] = 432,45$$

$$SQD(C_2) = (50)^2 / [5(1)^2 + 5(-1)^2] = 250,00$$

$$SQD(C_3) = (78)^2 / [6(2)^2 + 4(-3)^2] = 101,40$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(3)	(783,85)			
(A, D) vs (B, C)	1	432,45	432,45	13,46	0,0021
A vs D	1	250,00	250,00	7,78	0,0131
B vs C	1	101,40	101,40	3,16	0,0946
Resíduo	16	513,90	32,12		
Total	19	1297,75			

Clones	Média
B	52,60
C	59,50
A	53,00
D	42,60

Conclusões com erro tipo I de 5%:

$C_1 = (A, D) \text{ vs. } (B, C)$: Rejeita-se H_0 $C_2 = A \text{ vs. } D$: Rejeita-se H_0 $C_3 = B \text{ vs. } C$: Aceita-se H_0
--

5.9. Considerações finais

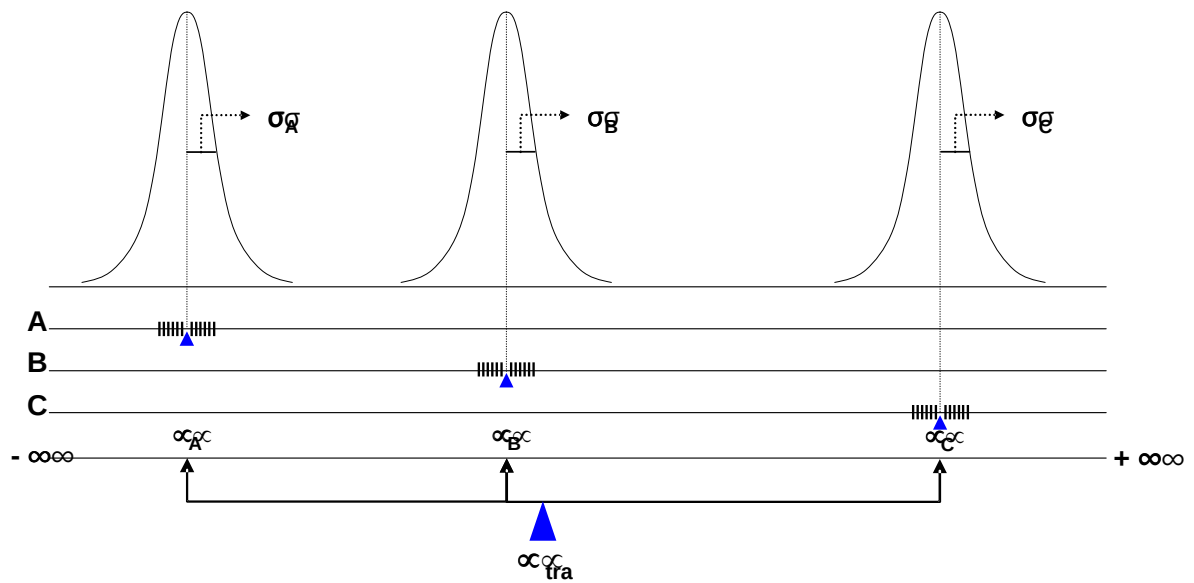
Embora seja simples, flexível e de fácil análise, no planejamento, na montagem, na condução e na coleta de dados nesse tipo de delineamento, é importante a presença e de um estatístico experimental experiente, assessorando todas as etapas do ciclo experimental.

O efeito de qualquer fonte de variação sistemática, além dos tratamentos, será atribuída ao erro experimental, reduzindo a precisão do experimento e, como consequência, diminuindo a probabilidade de se detectar diferenças entre tratamentos, caso elas existam.

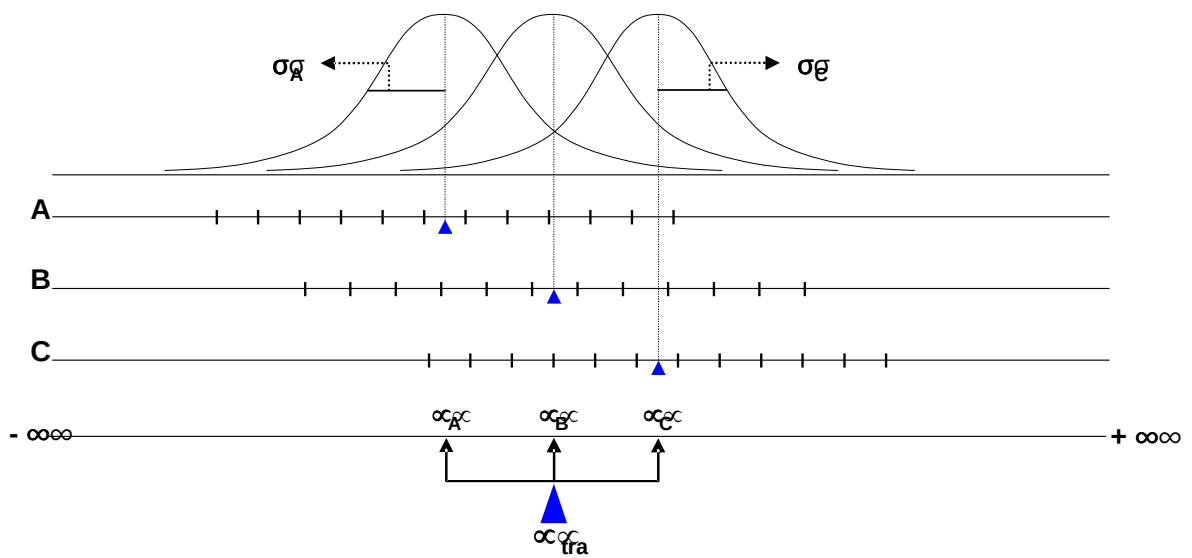
Nos exemplos apresentados procurou-se analisar o mesmo experimento, balanceado e desbalanceado, para que se consiga perceber a influência da perda de algumas unidades experimentais na análise.

Adicionalmente, com finalidades exclusivamente didáticas, foram apresentados os resultados de vários testes de comparação de médias múltiplas (tcm) além dos contrastes. Para as análises cotidianas, entretanto, deve-se optar por um dos métodos, preferencialmente na fase de planejamento do experimento.

Em razão dos argumentos apresentados e discutidos em sala de aula, recomenda-se a utilização preferencial pelos contrastes, dado a maior flexibilidade, abrangência e facilidade de cálculo.

5.10. Demonstrações e ilustrações

a) Médias de tratamentos distantes e erro experimental pequeno



b) Médias de tratamentos próximas e erro experimental grande

Figura 5.3 – Ilustração da ANOVA.

Demonstração da aplicação do teorema central do limite (TCL) na ANOVA

Como foi visto a origem conceitual do TCL, que nos informa sobre a distribuição da média amostral, foi feita concebendo-se infinitas repetições de uma amostra de tamanho n constante.

Em experimentos desbalanceados os tratamentos, considerados amostras de tamanho r , de uma pressuposta população, serão de tamanhos diferentes. Assim, a fórmula anterior (ver Análise de variância - ANOVA):

$$s^2 = \frac{r \cdot \left[\sum m^2 - \frac{(\sum m)^2}{n_m} \right]}{n_m - 1}$$

Fica assim:

$$s^2 = \frac{\sum r_i m_i^2 - \frac{(\sum r_i m_i)^2}{n_m}}{n_m - 1} = \frac{\sum r_i m_i - C}{n_m - 1}$$

$$s^2 = \frac{\left[\left(5 \cdot 52,60^2 \right) + \left(6 \cdot 59,50^2 \right) + \left(4 \cdot 53,00^2 \right) + \left(5 \cdot 42,60^2 \right) \right] - C}{4 - 1} = 261,28 = QMDtra$$

Exemplo ilustrativo da igualdade

$$\frac{(\sum y_i)^2}{n} = \frac{(\sum r_i \cdot m_i)^2}{\sum r_i} = C$$

considerando um número diferente de repetições: $r_i \neq \dots \neq r_k$

	Repetições							r_i	m_i	$r_i \cdot m_i$
	1	2	3	4	5	6	7			
A	1	2	3					3	2	6
B	4	5	6	7	8			5	6	30
C	9	10	11	12	13	14	15	7	12	84
Soma								15	20	120

$$\frac{(\sum y_i)^2}{n} = \frac{(120)^2}{15} = 225$$

$$\left[\frac{(\sum r_i \cdot m_i)^2}{\sum r_i} \right] = \frac{(120)^2}{15} = 225$$

Cálculo da SQDtra utilizando médias e totais de tratamentos com número diferente de repetições aplicando o TCL:

Usando médias :

$$SQDtra_m = \left[(r_1 \cdot m_1)^2 + \dots + (r_k \cdot m_k)^2 \right] - C$$

$$SQDtra_m = \left[(5 \cdot 52,60^2) + (6 \cdot 59,50^2) + (4 \cdot 53,00^2) + (5 \cdot 42,60^2) \right] - C$$

Usando totais de tratamentos :

$$SQDtra_t = \left[5 \cdot \left(\frac{263}{5} \right)^2 + 6 \cdot \left(\frac{357}{6} \right)^2 + \dots + 5 \cdot \left(\frac{213}{5} \right)^2 \right] - C$$

$$SQDtra_t = \left[5 \cdot \frac{(263)^2}{25} + 6 \cdot \frac{(357)^2}{36} + \dots + 5 \cdot \frac{(213)^2}{25} \right] - C$$

$$SQDtra_t = \left[\frac{(263)^2}{5} + \frac{(357)^2}{6} + \dots + \frac{(213)^2}{5} \right] - C$$

6. Testes de comparação de médias múltiplas

6.1. Introdução

Após a análise de variância (ANOVA) de um experimento, para comparar entre si as médias de tratamentos, uma das opções é o uso dos testes de comparação de médias múltiplas.

6.2. O fundamento dos testes

O fundamento consiste, para todos os testes, na obtenção do valor da diferença mínima significativa (dms), que permite a decisão dos testes de hipóteses, na comparação entre duas médias ou grupo de médias:

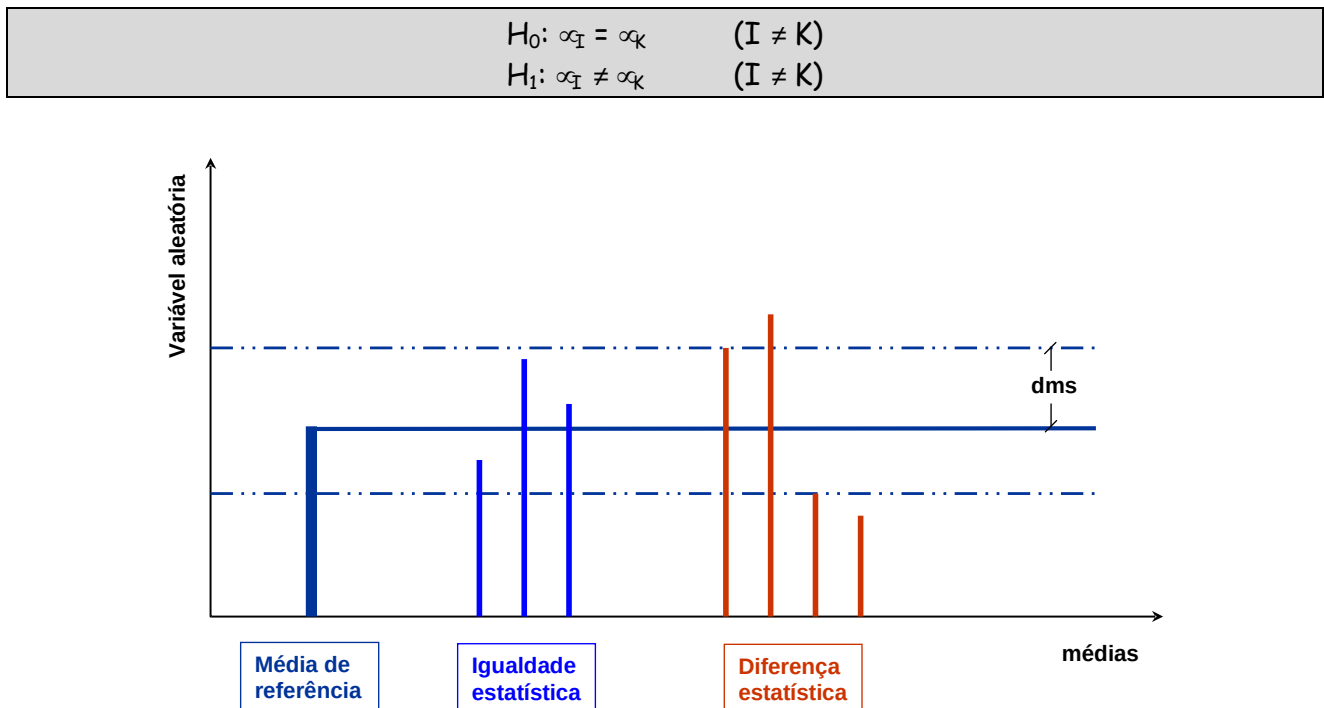


Figura 6.1 – Ilustração do fundamento dos testes de comparação entre médias.

Observação:

$$dms \propto QMDres \propto \text{Erro Experimental}$$

A diferença mínima significativa, para todos os testes, é diretamente proporcional ao quadrado médio do resíduo, que na ANOVA quantifica a influência de todas as fontes de variação não controladas.

Dessa forma, as inferências realizadas a partir dos testes aplicados a experimentos com elevado QMDres, e como consequência direta, com coeficiente de variação elevado, podem ser questionáveis.

6.3. Os testes

Para o estudo dos testes de médias será utilizado um exemplo em comum, conduzido no delineamento inteiramente casualizado (DIC) com 5 repetições, onde foram testadas quatro variedades (A, B, C e D) de milho:

Quadro 6.1 – Produção de milho em kg 100 m⁻²

Tra	Repetições					Totais	N.Repetições	Médias
	1	2	3	4	5			
A	25	26	20	23	21	115	5	23,00
B	31	25	28	27	24	135	5	27,00
C	22	26	28	25	29	130	5	26,00
D	33	29	31	34	28	155	5	31,00
						535	20	26,75

$$C = (535)^2 / 20 = 14.311,25$$

$$SQD_{tot} = [(25)^2 + (26)^2 + \dots + (28)^2] - C = 275,75$$

$$SQD_{trat} = 1 / 5 [(115)^2 + (135)^2 + \dots + (155)^2] - C = 163,75$$

$$SQD_{res} = SQD_{tot} - SQD_{tra} = 275,75 - 163,75 = 112,00$$

Hipóteses:

$$H_0: \alpha_I = \alpha_K \text{ (para todo } I \neq K)$$

$$H_1: \text{Nem todas as } \alpha_I \text{ são iguais}$$

ANOVA

FV	GL	SQD	QMD	Fcal	Pr
Tratamento	3	163,75	54,58	7,80	0,00197
Resíduo	16	112,00	7,00		
Total	19	275,75			

$$cv = 100 \cdot (\sqrt{7,00/26,75}) = 9,89\%$$

6.3.1. Teste de Duncan

É um dos teste que apresenta valores mais baixos da dms, implicando ser mais fácil detectar diferenças entre os tratamentos, caso elas existam.

6.3.1.1. Obtenção da dms

$$dms = Z_{\alpha} \sqrt{\frac{1}{2} \widehat{V}(\hat{C})}$$

$$|\hat{C}| \geq dms \Rightarrow *(\text{significativo})$$

$$|\hat{C}| < dms \Rightarrow ns(\text{não significativo})$$

$Z_{\alpha}(n_1; n_2)$: α = nível de significância do teste

n_1 = num. de médias envolvidas no teste

n_2 = num. gl. resíduo

6.3.1.2. Aplicação do teste

Inicialmente as médias devem ser ordenadas em ordem decrescente:

$$m_D = 31$$

$$m_B = 27$$

$$m_C = 26$$

$$m_A = 23$$

6.3.1.2.1. Para contrastes que abrangem 4 médias

$$\hat{C}_1 = m_D - m_A = 31 - 23 = 8^*$$

$$dms(4) = Z_{\alpha} \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} \quad \therefore \quad \widehat{V}(\hat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80$$

$$dms(4) = 3,235 \sqrt{\frac{1}{2} 2,80} = 3,83 \quad \therefore \quad Z_{5\%}(4; 16) = 3,235$$

6.3.1.2.2. Para contrastes que abrangem 3 médias

$$\hat{C}_2 = m_D - m_C = 31 - 26 = 5^*$$

$$\hat{C}_3 = m_B - m_A = 27 - 23 = 4^*$$

$$dms(3) = Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80$$

$$dms(3) = 3,144 \cdot \sqrt{\frac{1}{2} 2,80} = 3,72 \quad \therefore \quad Z_{5\%}(3; 16) = 3,144$$

6.3.1.2.3. Para testar contrastes que abrangem 2 médias

$$\widehat{C}_4 = m_D - m_B = 31 - 27 = 4^*$$

$$\widehat{C}_5 = m_B - m_C = 27 - 26 = 1^{ns}$$

$$\widehat{C}_6 = m_C - m_A = 26 - 23 = 3^{ns}$$

$$dms(2) = Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80$$

$$dms(2) = 2,998 \cdot \sqrt{\frac{1}{2} 2,80} = 3,55 \quad \therefore \quad Z_{5\%}(2; 16) = 2,998$$

Quadro 6.2 – Diferenças mínimas significativas usadas nas comparações

	m _D	m _B	m _C	m _A
m _D	-	dms(2)	dms(3)	dms(4)
m _B		-	dms(2)	dms(3)
m _C			-	dms(2)
m _A				-

Quadro 6.3 – Resultado das comparações

	m _D	m _B	m _C	m _A
m _D	-	*	*	*
m _B		-	ns	*
m _C			-	ns
m _A				-

6.3.1.3. Apresentação dos resultados e conclusão

A partir do Quadro 6.3 elaborase o resultado final que pode ser dado utilizando barras ou letras:

m _D = 31		m _D = 31	a	
m _B = 27		m _B = 27	b	
m _C = 26	ou	m _C = 26	b	c
m _A = 23		m _A = 23	c	

Utilizando barras:

As médias de tratamentos ligadas por uma mesma barra, não diferem entre si, pelo teste de Duncan a 5% de probabilidade.

Utilizando letras:

As médias de tratamentos que apresentam pelo menos uma mesma letra em comum, não diferem entre si, pelo teste de Duncan a 5% de probabilidade.

6.3.2. Teste de Dunnett

Usado quando as comparações que interessam ao pesquisador são entre um tratamento padrão (quase sempre a testemunha) e cada um dos demais tratamentos:

6.3.2.1. Obtenção da dms

$$dms = t_d \cdot \sqrt{\widehat{V}(\widehat{C})}$$

$$|\widehat{C}| \geq dms \Rightarrow *(\text{significativo})$$

$$|\widehat{C}| < dms \Rightarrow ns(\text{não significativo})$$

$$t_{d\alpha}(n_1; n_2): \alpha = \text{nível de significância do teste}$$

$$n_1 = \text{num. total de tratamentos}$$

$$n_2 = \text{num. gl. resíduo}$$

6.3.2.2. Aplicação do teste

Considerando o tratamento A como padrão ou testemunha, os contrastes a serem testados são:

$$\widehat{C}_1 = m_B - m_A = 27 - 23 = 4^{ns}$$

$$\widehat{C}_2 = m_C - m_A = 26 - 23 = 3^{ns}$$

$$\widehat{C}_3 = m_D - m_A = 31 - 23 = 8^*$$

$$dms = t_d \cdot \sqrt{\widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80$$

$$dms = 2,71 \cdot \sqrt{2,80} = 4,53 \quad \therefore \quad t_{d5\%}(4; 16) = 2,71$$

Quadro 6.4 – Resultado das comparações

	m_A
m_D	*
m_B	ns
m_C	ns

6.3.2.3. Apresentação dos resultados e conclusão

$m_D = 31$	a
$m_B = 27$	b
$m_C = 26$	b
$m_A = 23$	b (Testemunha)

As médias de tratamentos que apresentam pelo menos uma mesma letra em comum, não diferem entre si, pelo teste de Dunnett a 5% de probabilidade.

6.3.3. Teste de Tukey

Usado para contraste entre duas médias, é simples e de fácil aplicação.

É um dos testes que apresentam maior valor da dms, o que implica em maiores dificuldades em detectar diferenças entre as médias, caso elas existam.

6.3.3.1. Obtenção da dms

$$dms = q \cdot \sqrt{\frac{1}{2} V(\hat{C})}$$

$$|\hat{C}| \geq dms \Rightarrow *(\text{significativo})$$

$$|\hat{C}| < dms \Rightarrow ns(\text{não significativo})$$

$q_\alpha(n_1; n_2)$: α = nível de significância do teste

n_1 = num. total de tratamentos

n_2 = num. gl. resíduo

6.3.3.2. Aplicação do teste

Médias ordenadas:

$$m_D = 31$$

$$m_B = 27$$

$$m_C = 26$$

$$m_A = 23$$

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80$$

$$dms = 4,05 \cdot \sqrt{\frac{1}{2} 2,80} = 4,79 \quad \therefore \quad q_{5\%}(4; 16) = 4,05$$

$$\widehat{C}_1 = m_D - m_B = 4^{ns}$$

$$\widehat{C}_2 = m_D - m_C = 5^*$$

$$\widehat{C}_3 = m_D - m_A = 8^*$$

$$\widehat{C}_4 = m_B - m_C = 1^{ns}$$

$$\widehat{C}_5 = m_B - m_A = 4^{ns}$$

$$\widehat{C}_7 = m_C - m_A = 3^{ns}$$

$$dms = 4,79$$

Quadro 6.5 – Diferenças mínimas significativas usadas nas comparações

	m_D	m_B	m_C	m_A
m_D	-	dms	dms	dms
m_B		-	dms	dms
m_C			-	dms
m_A				-

Quadro 6.6 – Resultado das comparações

	m_D	m_B	m_C	m_A
m_D	-	ns	*	*
m_B		-	ns	ns
m_C			-	ns
m_A				-

6.3.3.3. Apresentação dos resultados e conclusão

A partir do Quadro 6.6 elabora-se o resultado final:

$m_D = 31$	a	
$m_B = 27$	a	b
$m_C = 26$		b
$m_A = 23$		b

As médias de tratamentos seguidas de pelo menos uma letra em comum não diferem entre si pelo teste de Tukey ao nível de 5% de probabilidade.

6.3.4. Teste de Student – Newman – Keuls (SNK)

Usa a metodologia do teste de Duncan e a tabela do teste de Tukey, sendo de rigor intermediário entre os dois.

6.3.4.1. Obtenção da dms

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})}$$

$$\begin{aligned} |\hat{C}| &\geq dms \Rightarrow *(\text{significativo}) \\ |\hat{C}| &< dms \Rightarrow ns(\text{não significativo}) \end{aligned}$$

6.3.4.2. Aplicação do teste

6.3.4.2.1. Para contrastes que abrangem 4 médias

$$\hat{C}_1 = m_D - m_A = 31 - 23 = 8 *$$

$$\begin{aligned} dms(4) &= q \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore \quad \widehat{V}(\hat{C}) &= QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80 \\ dms(4) &= 4,05 \cdot \sqrt{\frac{1}{2} 2,80} = 4,79 & \therefore \quad q_{5\%}(4;16) &= 4,05 \end{aligned}$$

6.3.4.2.2. Para contrastes que abrangem 3 médias

$$\hat{C}_2 = m_D - m_C = 31 - 26 = 5^*$$

$$\hat{C}_3 = m_B - m_A = 27 - 23 = 4^{ns}$$

$$\begin{aligned} dms(3) &= q \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore \quad \widehat{V}(\hat{C}) &= QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80 \\ dms(3) &= 3,65 \cdot \sqrt{\frac{1}{2} 2,80} = 4,32 & \therefore \quad q_{5\%}(3; 16) &= 3,65 \end{aligned}$$

6.3.4.2.3. Para contrastes que abrangem 2 médias

$$\hat{C}_4 = m_D - m_B = 31 - 27 = 4^*$$

$$\hat{C}_5 = m_B - m_C = 27 - 26 = 1^{ns}$$

$$\hat{C}_6 = m_C - m_A = 26 - 23 = 3^{ns}$$

$$\begin{aligned} dms(2) &= q \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore \quad \widehat{V}(\hat{C}) &= QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80 \\ dms(2) &= 3,00 \cdot \sqrt{\frac{1}{2} 2,80} = 3,55 & \therefore \quad q_{5\%}(2; 16) &= 3,00 \end{aligned}$$

Quadro 6.7 – Diferenças mínimas significativas usadas nas comparações

	m _D	m _B	m _C	m _A
m _D	-	dms(2)	dms(3)	dms(4)
m _B		-	dms(2)	dms(3)
m _C			-	dms(2)
m _A				-

Quadro 6.8 – Resultado das comparações

	m _D	m _B	m _C	m _A
m _D	-	*	*	*
m _B		-	ns	ns
m _C			-	ns
m _A				-

6.3.4.3. Apresentação dos resultados e conclusão

A partir do Quadro 6.8 elabora-se o resultado final:

$m_D = 31$	a
$m_B = 27$	b
$m_C = 26$	b
$m_A = 23$	b

As médias de tratamentos seguidas de pelo menos uma mesma letra em comum não diferem entre si, pelo teste de SNK, a 5% de probabilidade.

6.3.5. Teste de Scheffé

Usado para testar todo e qualquer contraste, sendo considerado um teste bastante rigoroso:

6.3.5.1. Obtenção da dms

$$dms = \sqrt{(I - 1) \cdot F \cdot \widehat{V}(\widehat{C})}$$

$$|\widehat{C}| \geq dms \Rightarrow *(\text{significativo})$$

$$|\widehat{C}| < dms \Rightarrow ns(\text{não significativo})$$

I = num. de tratamentos

$F_{\alpha}(n_1; n_2)$: α = nível de significância do teste

n_1 = num. gl. tratamento

n_2 = num. gl. resíduo

6.3.5.2. Teste de Scheffé - médias de tratamentos

Aplicar o teste de Scheffé para comparar o seguinte contraste

$$C = A \text{ vs. } D$$

$$\widehat{C} = m_A - m_D = 23 - 31 = -8^*$$

$$dms = \sqrt{(I - 1) \cdot F \cdot \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} (1^2 + (-1)^2) = 2,80$$

$$dms = \sqrt{(4 - 1) \cdot 3,24 \cdot 2,80} = 5,22 \quad \therefore \quad F_{5\%}(3; 16) = 3,24$$

$|\hat{C}| \geq dms$, o contraste é significativo, ou seja, existe diferença entre a produção das variedades pelo teste de Scheffé ao nível de 5% de probabilidade.

6.3.5.3. Teste de Scheffé - grupos de médias de tratamentos

Supondo que neste exemplo, as variedades A e B sejam de porte normal e as variedades C e D de porte baixo, a produção desses dois grupos pode ser comparada pelo teste de Scheffé:

$$Y = (A, B) \text{ vs. } (C, D)$$

$$\hat{C} = m_A + m_B - m_C - m_D$$

$$\hat{C} = 23 + 27 - 26 - 31 = -7^{ns}$$

$$dms = \sqrt{(I-1) \cdot F \cdot \widehat{V}(\hat{C})}$$

$$\widehat{V}(\hat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{7}{5} [(1)^2 + (1)^2 + (-1)^2 + (-1)^2] = 5,60$$

$$dms = \sqrt{(4-1) \cdot 3,24 \cdot 5,60} = 7,38 \quad \therefore \quad F_{5\%}(3; 16) = 3,24$$

$|\hat{C}| < dms$, o contraste não é significativo, ou seja, não há diferença entre as médias de produção entre as variedades de porte normal e porte baixo.

6.4. Exemplo de aplicação em experimentos desbalanceados

Será utilizado o mesmo experimento anterior, porém, considerando a perda de algumas unidades experimentais:

Quadro 6.2 – Produção de milho em kg 100 m⁻²

Tra	Repetições					Totais	N.Repetições	Médias
	1	2	3	4	5			
A	-	26	20	23	21	90	4	22,50
B	31	25	28	27	24	135	5	27,00
C	22	26	-	25	29	102	4	25,50
D	33	29	31	34	-	127	4	31,75
						454	17	26,69

$$\begin{aligned}
C &= (454)^2 / 17 = 12.124,47 \\
SQD_{tot} &= [(26)^2 + (20)^2 + \dots + (34)^2] - C = 269,53 \\
SQD_{tra_t} &= [1/4(90)^2 + 1/5(135)^2 + \dots + 1/4(127)^2] - C = 178,78 \\
SQD_{res} &= SQD_{tot} - SQD_{tra} = 269,53 - 178,78 = 90,75
\end{aligned}$$

Hipóteses:

$$\begin{aligned}
H_0: \alpha_I &= \alpha_K \text{ (para todo } I \neq K) \\
H_1: &\text{Nem todas as } \alpha_I \text{ são iguais}
\end{aligned}$$

ANOVA

FV	GL	SQD	QMD	Fcal	Pr
Tratamento	3	178,78	59,59	8,54	0,00216
Resíduo	13	90,75	6,98		
Total	16	269,53			

$$cv = 100 \cdot (\sqrt{6,98/26,69}) = 9,90\%$$

6.4.1. Teste de Duncan

$$m_D = 31,75 \text{ (4)}$$

$$m_B = 27,00 \text{ (5)}$$

$$m_C = 25,50 \text{ (4)}$$

$$m_A = 22,50 \text{ (4)}$$

6.4.1.1. Para contrastes que abrangem 4 médias: 4 vs. 4 repetições

$$\hat{C}_1 = m_D - m_A = 31,75 - 22,50 = 9,25 *$$

$$\begin{aligned}
dms(4) &= Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore & \quad \widehat{V}(\hat{C}) = QMD_{res} \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{6,98}{4} (1^2 + (-1)^2) = 3,49 \\
dms(4) &= 3,29 \cdot \sqrt{\frac{1}{2} 3,49} = 4,35 & \therefore & \quad Z_{5\%}(4; 13) = 3,29
\end{aligned}$$

6.4.1.2. Para contrastes que abrangem 3 médias: 4 vs. 4 repetições

$$\hat{C}_2 = m_D - m_C = 31,75 - 25,50 = 6,25^*$$

$$\begin{aligned} dms(3) &= Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore & \quad \widehat{V}(\hat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{6,98}{4} (1^2 + (-1)^2) = 3,49 \\ dms(3) &= 3,29 \cdot \sqrt{\frac{1}{2} 3,49} = 4,35 & \therefore & \quad Z_{5\%}(3; 13) = 3,29 \end{aligned}$$

6.4.1.3. Para contrastes que abrangem 3 médias: 4 vs. 5 repetições

$$\hat{C}_3 = m_B - m_A = 27,00 - 22,50 = 4,50^*$$

$$\begin{aligned} dms(3) &= Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore & \quad \widehat{V}(\hat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = 6,98 \left(\frac{1^2}{5} + \frac{(-1)^2}{4} \right) = 3,14 \\ dms(3) &= 3,29 \cdot \sqrt{\frac{1}{2} 3,14} = 4,12 & \therefore & \quad Z_{5\%}(3; 13) = 3,29 \end{aligned}$$

6.4.1.4. Para testar contrastes que abrangem 2 médias: 4 vs. 5 repetições

$$\hat{C}_4 = m_D - m_B = 31,75 - 27,00 = 4,75^*$$

$$\hat{C}_5 = m_B - m_C = 27,00 - 25,50 = 1,50^{ns}$$

$$\begin{aligned} dms(2) &= Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} & \therefore & \quad \widehat{V}(\hat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = 6,98 \left(\frac{1^2}{5} + \frac{(-1)^2}{4} \right) = 3,14 \\ dms(2) &= 3,055 \cdot \sqrt{\frac{1}{2} 3,14} = 3,83 & \therefore & \quad Z_{5\%}(2; 13) = 3,055 \end{aligned}$$

6.4.1.5. Para testar contrastes que abrangem 2 médias: 4 vs. 4 repetições

$$\hat{C}_6 = m_C - m_A = 25,50 - 22,50 = 3,00^{ns}$$

$$dms(2) = Z \cdot \sqrt{\frac{1}{2} \widehat{V}(\hat{C})} \quad \therefore \quad \widehat{V}(\hat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{6,98}{4} (1^2 + (-1)^2) = 3,49$$

$$dms(2) = 3,055 \cdot \sqrt{\frac{1}{2} 3,49} = 4,04 \quad \therefore \quad Z_{5\%}(2; 13) = 3,055$$

Quadro 6.9 – Diferenças mínimas significativas usadas nas comparações

	m_D	m_B	m_C	m_A
m_D	-	$dms(2) \text{ 4r vs. 5r} = 4,80$	$dms(3) \text{ 4r vs. 4r} = 4,35$	$dms(4) \text{ 4r vs. 4r} = 3,29$
m_B		-	$dms(2) \text{ 4r vs. 5r} = 4,80$	$dms(3) \text{ 4r vs. 5r} = 4,12$
m_C			-	$dms(2) \text{ 4r vs. 4r} = 4,04$
m_A				-

Quadro 6.10 – Resultado das comparações

	m_D	m_B	m_C	m_A
m_D	-	4,75*	6,25*	9,25*
m_B		-	1,50ns	4,50*
m_C			-	3,00ns
m_A				-

A partir do Quadro 6.10 elabora-se o resultado final:

$m_D = 31,75$	a		
$m_B = 27,00$		b	
$m_C = 26,50$	b		c
$m_A = 22,50$		c	

As médias que apresentam pelo menos uma mesma letra em comum, não diferem entre si, pelo teste de Duncan a 5% de probabilidade.

6.4.2. Teste de Tukey

$$m_D = 31,75 \text{ (4)}$$

$$m_B = 27,00 \text{ (5)}$$

$$m_C = 25,50 \text{ (4)}$$

$$m_A = 22,50 \text{ (4)}$$

6.4.2.1. Para testar contrastes que abrangem 2 médias: 5 vs. 4 repetições

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = 6,98 \left(\frac{1^2}{5} + \frac{(-1)^2}{4} \right) = 3,14$$

$$dms = 4,15 \cdot \sqrt{\frac{1}{2} 3,14} = 5,20 \quad \therefore \quad q_{5\%}(4; 13) = 4,15$$

6.4.2.2. Para testar contrastes que abrangem 2 médias: 4 vs. 4 repetições

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{6,98}{4} (1^2 + (-1)^2) = 3,49$$

$$dms = 4,15 \cdot \sqrt{\frac{1}{2} 3,49} = 5,48 \quad \therefore \quad q_{5\%}(4; 13) = 4,15$$

$$\hat{C}_1 = m_D - m_B = 4,75^{ns} \quad (4 \text{ vs. } 5r)$$

$$\hat{C}_2 = m_D - m_C = 6,25^* \quad (4 \text{ vs. } 4r)$$

$$\hat{C}_3 = m_D - m_A = 9,25^* \quad (4 \text{ vs. } 4r)$$

$$\hat{C}_4 = m_B - m_C = 1,50^{ns} \quad (4 \text{ vs. } 5r)$$

$$\hat{C}_5 = m_B - m_A = 4,50^{ns} \quad (4 \text{ vs. } 5r)$$

$$\hat{C}_7 = m_C - m_A = 3,00^{ns} \quad (4 \text{ vs. } 4r)$$

$$4 \text{ vs. } 4 \text{ repetições} \rightarrow dms = 4,15$$

$$4 \text{ vs. } 5 \text{ repetições} \rightarrow dms = 5,48$$

Quadro 6.11 – Resultado das comparações

	m_D	m_B	m_C	m_A
m_D	-	4,75ns	6,25*	9,25*
m_B		-	1,50ns	4,50ns
m_C			-	3,00ns
m_A				-

$m_D = 31,75$	a
$m_B = 27,00$	a
$m_C = 25,50$	b
$m_A = 22,50$	b

As médias seguidas de pelo menos uma letra em comum não diferem entre si pelo teste de Tukey ao nível de 5% de probabilidade.

7. Estudo e aplicação de contrastes

7.1. Introdução

Muitas vezes é mais eficiente, e até mesmo mais informativo, proceder ao desdobramento do número de graus de liberdade associados a tratamentos dentro da própria análise de variância, ao invés de utilizar os métodos de comparação de médias múltiplas.

Neste caso o pesquisador está interessado em algumas comparações, em alguns contrastes apenas.

O pesquisador estará testando hipóteses formuladas nas fases de planejamento do experimento, antecedendo a qualquer observação ou análise de seus dados.

Embora a não observação destas sugestões, de boa conduta experimental, não inviabilize a aplicação dos contrastes.

As informações possíveis de serem obtidas pela aplicação e teste dos contrastes, em geral, são de maior eficiência e abrangência que a simples comparações de médias.

Adicionalmente, a aplicação de contrastes é mais fácil e rápida que os testes de comparação de médias.

7.2. Definição

Normalmente, se trabalha com contrastes entre totais de tratamentos.

O caso mais comum é aquele em que os tratamentos possuem o mesmo número de repetições.

Nestas condições, uma função linear do tipo:

$$C = a_1T_1 + \dots + a_iT_i$$

é denominada contraste de totais de tratamentos se:

$$a_1 + \dots + a_i = 0 \quad \therefore \quad \sum a_i = 0$$

onde $a_1 + \dots + a_i$, são os coeficientes dos totais dos tratamentos $T_1, \dots, T_i$, respectivamente.

Assim, por exemplo:

$$C_1 = T_1 - T_2$$

$$C_2 = T_1 + T_2 - 2T_3$$

são contrastes entre totais de tratamentos, pois a soma dos coeficientes, de cada um, individualmente, é zero. Ou seja:

$$\sum a_i = 0$$

Quando os totais de tratamentos (T_i) são obtidos com número diferente de repetições (r_i) a função linear do tipo:

$$C = a_1 T_1 + \dots + a_i T_i$$

será um contraste entre totais de tratamentos se:

$$r_1 a_1 + \dots + r_i a_i = 0 \quad \therefore \quad \sum r_i a_i = 0$$

7.3. Contrastes entre totais de tratamentos com um mesmo número de repetições

7.3.1. Cálculo da soma de quadrados dos desvios

A soma de quadrados de um contraste C , a partir de totais de tratamentos, T oriundos de um mesmo número de repetições, é dada por:

$$SQD(C) = \frac{\widehat{C}^2}{r_1 a_1^2 + \dots + r_i a_i^2} = \frac{\widehat{C}^2}{r(a_1^2 + \dots + a_i^2)} = \frac{\widehat{C}^2}{r \sum a_i^2}$$

onde:

$\widehat{C}$: é a estimativa do contraste

r : o número de repetições dos tratamentos

Esta soma de quadrados é parte da soma de quadrados para tratamentos e a ela se atribui um (1) grau de liberdade.

7.3.2. Ortogonalidade

A orthogonalidade entre contrastes significa que as comparações são independentes.

Em outras palavras, a variação de um contraste é totalmente independente da variação de outro qualquer que lhe seja ortogonal, indicando uma independência entre as comparações.

Dois contrastes entre totais de tratamentos

$$C_1 = a_1T_1 + \dots + a_iT_i$$

$$C_2 = b_1T_1 + \dots + b_iT_i$$

são ortogonais se:

$$a_1b_1 + \dots + a_ib_i = 0 \quad \therefore \quad \sum a_ib_i = 0$$

Ou seja, o somatório dos produtos dos coeficientes é igual a zero.

7.4. Contrastes entre totais de tratamentos com número diferentes de repetições

7.4.1. Cálculo da soma de quadrados dos desvios

Neste caso, a soma de quadrados do contraste é dada por:

$$SQD(C) = \frac{\widehat{C}^2}{r_1a_1^2 + \dots + r_ia_i^2} = \frac{\widehat{C}^2}{\sum r_ia_i^2}$$

7.4.2. Ortogonalidade

Os contrastes entre totais de tratamentos:

$$C_1 = a_1T_1 + \dots + a_iT_i$$

$$C_2 = b_1T_1 + \dots + b_iT_i$$

com número diferente de repetições são ortogonais se:

$$r_1a_1b_1 + \dots + r_ia_ib_i = 0 \quad \therefore \quad \sum r_ia_ib_i = 0$$

Uma maneira prática, que garante a obtenção de comparações independentes (ortogonais) entre si é a seguinte:

- Dividem-se os tratamentos em dois grupos, para estabelecer a primeira comparação.
- Para estabelecer as novas comparações, não se pode mais comparar tratamentos de um grupo com tratamentos do outro grupo. Somente se podem comparar os tratamentos remanescentes dentro de cada grupo original obtidos em "a".
- Dividem-se os grupos em subgrupos e somente se compara dentro de cada subgrupo.

Exemplos:

$$C_1 = T_1 \text{ vs. } (T_2, T_3, T_4, T_5)$$

$$C_2 = (T_2, T_3) \text{ vs. } (T_4, T_5)$$

$$C_3 = T_2 \text{ vs. } T_3$$

$$C_4 = T_4 \text{ vs. } T_5$$

$$C_1 = (T_1, T_2, T_3) \text{ vs. } (T_4, T_5, T_6)$$

$$C_2 = (T_1) \text{ vs. } (T_2, T_3)$$

$$C_3 = T_2 \text{ vs. } T_3$$

$$C_4 = (T_4) \text{ vs. } (T_5, T_6)$$

$$C_5 = T_5 \text{ vs. } T_6$$

Observações:

Comparando n tratamentos, pode-se obter n-1 contrastes ortogonais.

Não existe uma regra fixa para o estabelecimento dos contrastes, desde que sejam satisfeitas as condições de contraste e de ortogonalidade.

Os contrastes devem ser estabelecidos de forma a possibilitarem ao pesquisador testar as hipóteses estatísticas estabelecidas. Para o primeiro exemplo as seguintes perguntas estarão sendo formuladas para serem testadas:

$$C_1 = T_1 \text{ vs. } (T_2, T_3, T_4, T_5)$$

$$C_2 = (T_2, T_3) \text{ vs. } (T_4, T_5)$$

$$C_3 = T_2 \text{ vs. } T_3$$

$$C_4 = T_4 \text{ vs. } T_5$$

T_1 difere estatisticamente da média conjunta de (T_2, T_3, T_4, T_5) ?

A média conjunta $(T_2 \text{ e } T_3)$ difere estatisticamente da média conjunta de $(T_4 \text{ e } T_5)$?

T_2 difere de T_3 ?

T_4 difere de T_5 ?

7.5. Regras para obtenção de contrastes ortogonais7.5.1. Contrastes com um mesmo número de repetições

- Escreve-se os totais de tratamentos envolvidos na comparação.
- Atribue-se sinal positivo aos totais de um grupo e negativo aos totais do outro grupo.
- Verifica-se o número de tratamentos (n_1) envolvidos no primeiro grupo, e o número de tratamentos (n_2) envolvidos no segundo grupo. Em seguida calcula-se o mmc entre n_1 e n_2 .
- Divide-se o mmc por n_1 . O resultado será o coeficiente de cada total do primeiro grupo.
- Divide-se o mmc por n_2 . O resultado será o coeficiente de cada total do segundo grupo.

Exemplo:

$C_1 = T_1 \text{ vs. } (T_2, T_3, T_4, T_5)$	$C_1 = 4T_1 - T_2 - T_3 - T_4 - T_5$	$(1;4) : mmc = 4$
$C_2 = (T_2, T_3) \text{ vs. } (T_4, T_5)$	$C_2 = T_2 + T_3 - T_4 - T_5$	$(2;2) : mmc = 2$
$C_3 = T_2 \text{ vs. } T_3$	$C_3 = T_2 - T_3$	$(1;1) : mmc = 1$
$C_4 = T_4 \text{ vs. } T_5$	$C_4 = T_4 - T_5$	$(1;1) : mmc = 1$

7.5.2. Contrastes com número diferente de repetições

- Escreve-se os totais de tratamentos envolvidos na comparação.
- Atribui-se sinal positivo aos totais de um grupo e negativo aos totais do outro grupo.
- Verifica-se o número de repetições (r_1) envolvidos no primeiro grupo, e o número de repetições (r_2) envolvidos no segundo grupo. Em seguida calcula-se o mmc entre r_1 e r_2 .
- Divide-se o mmc por r_1 . O resultado será o coeficiente de cada total do primeiro grupo.
- Divide-se o mmc por r_2 . O resultado será o coeficiente de cada total do segundo grupo.

Exemplo:

	$r :$	6	6	4	5	6	
$C_1 = (T_1, T_2, T_3, T_4) \text{ vs. } T_5$	$C_1 = 2T_1 + 2T_2 + 2T_3 + 2T_4 - 7T_5$						$(21;6) : mmc = 42$
$C_2 = (T_1, T_2, T_3) \text{ vs. } T_4$	$C_2 = 5T_1 + 5T_2 + 5T_3 - 16T_4$						$(16;5) : mmc = 80$
$C_3 = (T_1, T_2) \text{ vs. } T_3$	$C_3 = T_1 + T_2 - 3T_3$						$(12;4) : mmc = 12$
$C_4 = T_1 \text{ vs. } T_2$	$C_4 = T_1 - T_2$						$(6;6) : mmc = 6$

Observações:

Considere que os números de repetições iniciais, r , para cada tratamento foram 6.

Foram perdidas 2 parcelas no tratamento T_3 .

Foi perdida uma parcela no tratamento T_4 .

7.6. Variância de contrastes

Variância de um contraste:

$$C = a_1 \varphi + \dots + a_k \varphi$$

$$V(C) = V(a_1 \varphi + \dots + a_k \varphi)$$

$$V(C) = a_1^2 V(\varphi) + \dots + a_k^2 V(\varphi) \quad \therefore \quad V(\varphi) = 0 \quad \text{considerando } i = 1 \dots k$$

$$V(C) = a_1^2 0 + \dots + a_k^2 0$$

$$V(C) = 0$$

Variância da estimativa de um contraste:

$$\widehat{C} = a_1 m_1 + \dots + a_k m_k$$

$$V(\widehat{C}) = V(a_1 m_1 + \dots + a_k m_k) \quad \therefore \quad \text{Admitindo as médias independentes}$$

$$V(\widehat{C}) = a_1^2 V(m_1) + \dots + a_k^2 V(m_k) \quad \therefore \quad \text{Admitindo que } m_i \text{ vem de } r_i \text{ repetições}$$

$$V(\widehat{C}) = a_1^2 \frac{\sigma^2}{r_1} + \dots + a_k^2 \frac{\sigma^2}{r_k}$$

Pode-se usar s^2 como estimativa de σ^2 , neste caso será determinada a estimativa da variância da estimativa de um contraste:

$$\widehat{V}(\widehat{C}) = a_1^2 \frac{s^2}{r_1} + \dots + a_k^2 \frac{s^2}{r_k}$$

$$\widehat{V}(\widehat{C}) = s^2 \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) \quad \therefore \quad \text{Como } s^2 = \text{QMDres}$$

$$\widehat{V}(\widehat{C}) = \text{QMDres} \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right)$$

Esta fórmula será intensamente utilizada nos testes de comparação de médias múltiplas (Tukey, Duncan, SNK, etc).

7.7. Compreensão do cálculo as soma de quadrados dos desvios de contrastes

7.7.1. Com médias de tratamentos

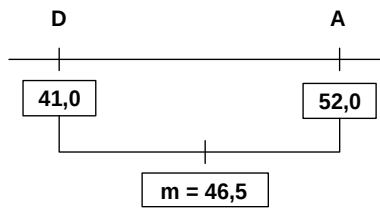
$$m_B = 59,50$$

$$m_C = 53,50$$

$$m_A = 52,00$$

$$m_D = 41,00$$

$$C_2 = A - D$$

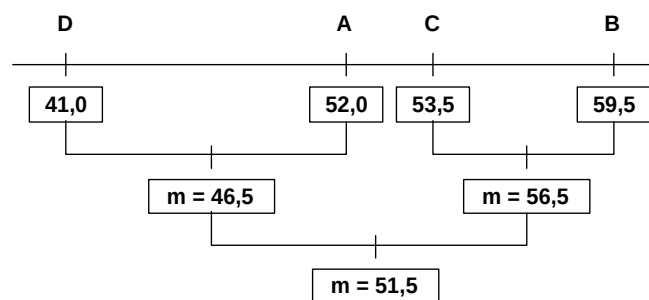


$$SQD_{C_2} = 6 \cdot \sum d^2$$

$$SQD_{C_2} = 6[(41,0 - 46,5)^2 + (52,0 - 46,5)^2]$$

$$SQD_{C_2} = 363,0$$

$$C_1 = (A, D) \text{ vs } (B, C)$$



$$SQD_{C_1} = 12 \cdot \sum d^2$$

$$SQD_{C_1} = 12[(46,5 - 51,5)^2 + (56,5 - 51,5)^2]$$

$$SQD_{C_1} = 600,0$$

7.7.2. Com os totais de tratamentos

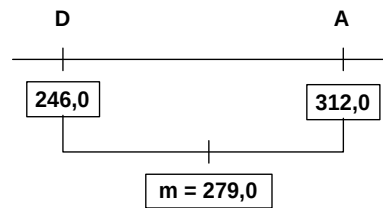
$$T_B = 357,0$$

$$T_C = 321,0$$

$$T_A = 312,0$$

$$T_D = 246,0$$

$$C_2 = A - D$$

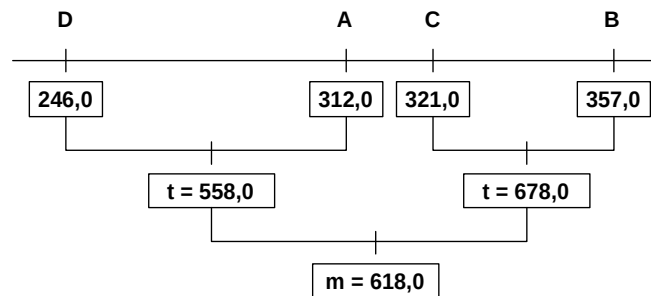


$$SQD_{C_2} = \frac{1}{6} \cdot \sum d^2$$

$$SQD_{C_2} = \frac{1}{6} [(246,0 - 279,0)^2 + (312,0 - 279,0)^2]$$

$$SQD_{C_2} = 363,0$$

$$C_1 = (A, D) \text{ vs } (B, C)$$



$$SQD_{C_1} = \frac{1}{12} \cdot \sum d^2$$

$$SQD_{C_1} = \frac{1}{12} [(558,0 - 618,0)^2 + (678,0 - 618,0)^2]$$

$$SQD_{C_1} = 600,0$$

8. [Reflexões sobre a análise de variância](#)

8.1. [Introdução](#)

A Análise de variância (ANOVA - ANalysis Of VAriance, que alguns autores brasileiros preferem denominar ANAVA - ANÁLise de VAriância) é uma técnica intensivamente utilizada na estatística paramétrica para fazer inferências sobre as médias populacionais a partir de suas estimativas (médias amostrais).

Nos experimentos agropecuários, em geral, o interesse é comparar diferentes variedades, fertilizantes, rações, formas de controle de pragas e doenças, controle de invasoras, etc.

Quando se ensina esta técnica matemática, utilizada para a partição da soma de quadrados dos desvios total de uma variável aleatória, em uma parte atribuída às fontes reconhecidas, sistemáticas ou controladas de variação, e uma outra parte, atribuída aos efeitos aleatórios ou não controlados, habitualmente, muita ênfase é dada à parte puramente algébrica da técnica. Por outro lado, muita pouca atenção é dedicada à compreensão e ao significado destes procedimentos. A consequência desse hábito é que o estudante memoriza as fórmulas e os procedimentos, torna-se capaz de montar o quadro da ANOVA, realizar os testes estatísticos e retirar conclusões sem, no entanto, entender muito bem o que está se passando.

Algumas pessoas, entretanto, não se dão por satisfeitas apenas com a parte algébrica e mecânica deste procedimento estatístico, ou seja, de serem capazes apenas de analisar e interpretar dados experimentais. Querem entender mais. Para estas pessoas é que este texto foi escrito e tem sido aperfeiçoado continuamente.

Ao entender, com conhecimento de causa, o significado menos aparente e evidente de uma análise de variância o usuário pode perceber, por exemplo, o porque de em algumas situações experimentais não encontrar diferenças significativas entre os tratamentos, assim como, pode avaliar se o delineamento adotado, a montagem e a condução do experimento foram adequados aos propósitos. A análise de variância pode fornecer informações valiosas a este respeito.

Não bastassem os argumentos apresentados, a ANOVA é um procedimento básico para a tomada de decisão na avaliação de resultados experimentais. Entender realmente o que se passa por trás da parte puramente algébrica, nunca será um conhecimento desnecessário, podendo trazer clareza de idéias e conceitos para quem a utiliza.

8.2. [Reflexões](#)

As reflexões desenvolvidas utilizam um exemplo numérico já analisado, originalmente apresentado na apostila sobre delineamento inteiramente casualizado (DIC), do curso de Metodologia e Estatística Experimental da Universidade Estadual de Santa Cruz.

Trata-se de um experimento montado no delineamento inteiramente casualizado completo, com 6 repetições, onde foram avaliadas a produção de amêndoas ($\text{kg } 10 \text{ plantas}^{-1} \text{ ano}^{-1}$) de 4 clones de cacau tolerantes a vassoura de bruxa. Os resultados experimentais são representados no Quadro 8.1 a seguir:

Quadro 8.1 – Produção de amêndoas (kg 10 plantas¹ ano⁻¹) de cacau aos 5 anos de idade

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,00
B	60	55	66	61	54	61	357	6	59,50
C	59	47	44	49	62	60	321	6	53,50
D	45	33	34	48	42	44	246	6	41,00
							1.236	24	51,50

Hipóteses:

$$H_0: \alpha_A = \alpha_B = \alpha_C = \alpha_D$$

H_1 : Nem todas as médias são iguais

A questão a ser investigada (teste de hipóteses) é a seguinte: Os clones de cacau são realmente diferentes? Ou seja, as diferenças entre as estimativas das médias de cada clone, m_i , são devidas a diferenças nas médias, α_i , das populações básicas, onde α_i representa o rendimento médio do clone i ? Ou tais diferenças entre as m_i podem ser atribuídas apenas às flutuações aleatórias?

Para ilustrar, suponhamos que solicitássemos a três pessoas que cada uma retirasse uma amostra (de 6 plantas) da população de plantas de apenas uma dos clones, o A por exemplo, calculasse a estimativa da média, e os resultados obtidos fossem os apresentados no Quadro 8.2.

Quadro 8.2 – Amostras da produção de amêndoas (kg 10 plantas⁻¹ ano⁻¹) de cacau aos 5 anos do clone A obtidas por cada uma das três pessoas

Amostra	Média amostral (m_A)
1	51,85
2	52,63
3	53,00

Observa-se que a estimativa da média, m , do clone A (m_A), obtida por cada pessoa (Quadro 8.2), foi diferente da anteriormente obtida, 52,00 (Quadro 8.1), além de diferirem entre si. Ocorreu algum erro?

Não ocorreu nenhum erro! Naturalmente, é de se esperar que cada pessoa selecione uma amostra diferente, obtendo assim diferentes estimativas da média, m_A . Ou seja, são estimativas da média, m , do clone A, obtidas a partir de diferentes amostras, e não a verdadeira média, α_A , da população básica do clone A. Esta sim, α_A , não varia, e é em geral desconhecida (α é um parâmetro da população).

Como era de se esperar, as flutuações amostrais naturais refletem-se em pequenas diferenças nas m_i , mesmo que as α_i sejam idênticas. Podemos, então, reformular a pergunta de forma mais objetiva: As diferenças nas m_i do Quadro 8.1, são da mesma magnitude que as do Quadro 8.2, e assim atribuíveis a flutuações aleatórias da

estimativa da média, ou são suficientemente grandes para sugerir uma diferença nas básicas? Em outras palavras, as estimativas das médias caracterizam, ou refletem, populações diferentes dos clones, representadas pelos diferentes tratamentos, ou, na realidade, as diferenças são devidas a flutuações aleatórias na amostragem e, neste caso, os diferentes clones podem ser considerados, de fato, quanto à produção, uma mesma população, não apresentando diferenças entre si.

No presente caso a primeira explicação parece a mais plausível. Mas como elaborar um teste formal para demonstrar isto? O teste formal é obtido através da técnica matemática da análise da variância (ANOVA).

A análise de variância de uma variável aleatória em estudo (produção no presente caso) consiste na partição da soma de quadrados dos desvios total em componentes associados às fontes sistemáticas, reconhecidas ou controladas de variação, neste caso os clones, e uma outra parte, de natureza aleatória, desconhecida ou não controlada que constitui o erro experimental ou resíduo.

Para se proceder a análise de variância dos dados experimentais do Quadro 8.1, os procedimentos são listados a seguir:

Parte-se do pré-suposto de que cada tratamento é uma amostra – de tamanho igual ao número de repetições – retirada de uma mesma população básica, normalmente distribuída. Isto significa, a princípio, que todos os tratamentos são iguais;

Nestas condições, têm-se duas maneiras alternativas, e razoáveis, de estimar a variância da população básica, σ^2 :

i. Tomar a média das variâncias de cada uma das amostras:

$$s^2 = \frac{\left(\frac{(58,0 - 52,0)^2 + \dots + (48,0 - 52,0)^2}{5} + \dots + \frac{(45,0 - 41,0)^2 + \dots + (44,0 - 41,0)^2}{5} \right)}{4} = 33,25$$

ii. Inferir σ^2 a partir da $V(m)$, isto é, a partir da variância das médias amostrais. Recordar que a variância da média amostral está relacionada com a variância da população, σ^2 , da seguinte forma:

$$\text{se, } V(m) = \frac{\sigma^2}{n} \quad \therefore \quad \sigma^2 = n \cdot V(m)$$

$$\text{então, } V(m) = \frac{s^2}{n} \quad \therefore \quad s^2 = n \cdot V(m)$$

Uma vez que n é conhecido, pois é o tamanho da amostra, ou melhor, o número de repetições do tratamento é possível calcular $V(m)$:

$$V(m) = \frac{((52,0 - 51,5)^2 + (59,5 - 51,5)^2 + (53,5 - 51,5)^2 + (41,0 - 51,5)^2)}{3} = 59,5$$

$$s^2 = V(m) \cdot n = 59,5 \cdot 6 = 357,0$$

Como foram obtidas duas estimativas da variância, σ^2 , da pré-suposta população básica (lembrar da consideração inicial), é possível formular hipóteses e realizar um teste estatístico utilizando uma distribuição de probabilidades adequada para a conclusão se a consideração inicial é, ou não, válida.

Como a distribuição de F fornece a distribuição de probabilidades do valor F_{cal} :

$$F_{cal} = \frac{s^2}{s^2} = \frac{357,0}{33,25} = 10,74$$

pode-se utilizar esta distribuição e decidir se, de fato, a consideração inicial é, ou não, correta. Em outras palavras, se as produções dos clones de cacau podem, ou não, ser consideradas como provenientes de uma mesma população básica. Posto de outra forma, se as produções dos clones são estatisticamente iguais ou diferentes.

A partir do pré-suposto anteriormente estabelecido de que os tratamentos e suas repetições representam amostras feitas em uma mesma população básica, pode-se formular as seguintes hipóteses:

Hipóteses:

$$H_0: \alpha_A = \alpha_B = \alpha_C = \alpha_D$$

H_1 : Nem todas as médias são iguais

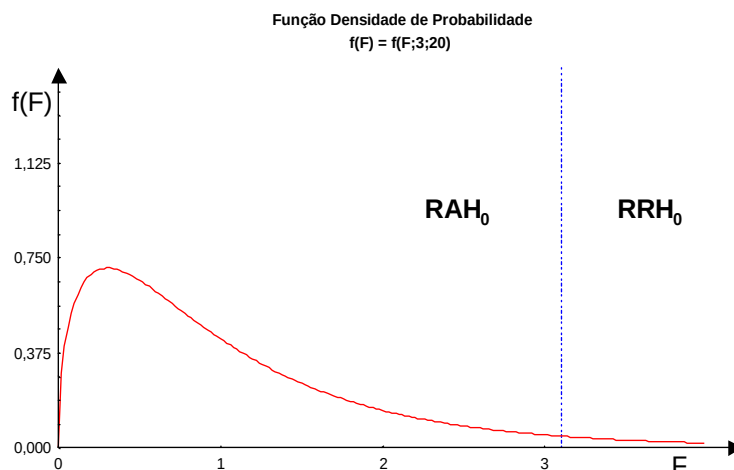
cujo significado é:

H_0 : mesma população

H_1 : populações distintas

Para testar estas hipóteses utiliza-se a estatística F:

a. A primeira providência é estipular o erro tipo I. Para o exemplo será adotado 5%:



b. Se a consideração inicial for correta, ou seja, trata-se realmente de uma mesma população, em 95% das vezes, em média, que a relação entre duas estimativas da variância for calculada, F_{cal} , deveria ser encontrado um valor menor que 3,10, $P(F_{cal} < 3,10) = 95\%$. Neste caso a decisão seria aceitar H_0 .

c. Da mesma forma, em apenas 5% das vezes, também em média, que a relação fosse calculada, F_{cal} , seria encontrado um valor igual ou maior que 3,10, $P(F_{cal} \geq 3,10) = 5\%$. Neste caso a decisão seria rejeitar H_0 .

O erro tipo I (α) associado ao teste de hipóteses é muito claro: na situação “c” seria rejeitada uma hipótese verdadeira. Isto é, os dados podem ser, de fato, provenientes de uma mesma população básica, em outras palavras, valores F_{cal} iguais ou superiores a 3,10 podem efetivamente ocorrer, mas estes casos são muito raros, mais precisamente, em apenas 5% dos casos.

Estes mesmos cálculos são convencionalmente feitos da seguinte forma:

Tra	Repetições						Totais	N.Repetições	Médias
	1	2	3	4	5	6			
A	58	49	51	56	50	48	312	6	52,00
B	60	55	66	61	54	61	357	6	59,50
C	59	47	44	49	62	60	321	6	53,50
D	45	33	34	48	42	44	246	6	41,00
							1.236	24	51,50

$$C = (1.236)^2 / 24 = 63.654,00$$

$$SQD_{tot} = [(58)^2 + (49)^2 + \dots + (44)^2] - C = 1.736,00$$

$$SQD_{trat} = 1 / 6 [(312)^2 + (357)^2 + \dots + (246)^2] - C = 1.071,00$$

$$SQD_{res} = SQD_{tot} - SQD_{tra} = 1.736 - 1.071,00 = 665,00$$

ANOVA

Causa da variação	GL	SQD	QMD	Fcal
Tratamentos	3	1.071,00	357,00	10,74*
Resíduo	20	665,00	33,25	
Total	23	1.736,00		

* Significativo ao nível de 5% de probabilidade.

É comum não se conseguir visualizar que cada quadrado médio dos desvios do quadro da ANOVA é, na realidade, o resultado da aplicação da conhecida fórmula para calcular a variância amostral:

$$s^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1}$$

o denominador, $n-1$, são os graus de liberdade da ANOVA;

$$\frac{(\sum y)^2}{n}$$

é o valor C;

$$\sum y^2 - \frac{(\sum y)^2}{n}$$

é o numerador da fórmula

$$s^2 = \frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n-1}$$

Concluindo a análise:

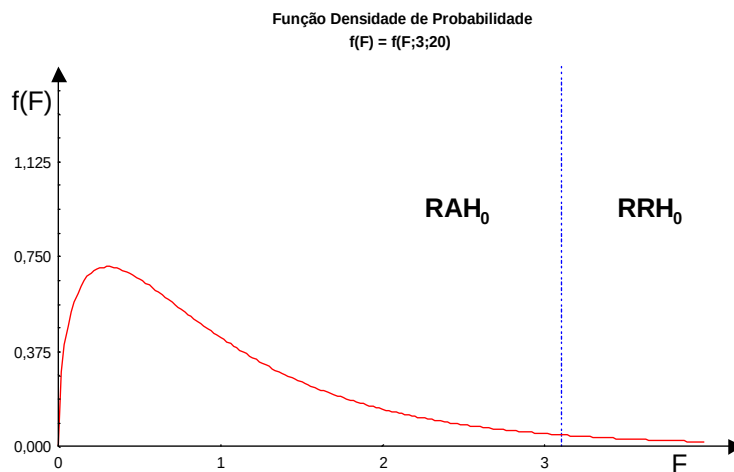


Figura 8.1 – Distribuição F mostrando RAH_0 : região de aceitação de H_0 e RRH_0 : região de rejeição de H_0 .

No presente caso o que está em comparação é uma amostra de tamanho 4 (3 gl) e uma amostra de tamanho 21 (20 gl).

$$F_{cal} = \frac{s^2(3\text{ gl})}{s^2(20\text{ gl})}$$

O valor $F = 3,10$ marca o limite do valor F onde, em média, em apenas 5% dos casos em que comparássemos as variâncias de duas amostras advindas de uma mesma população, obteríamos valores superiores a este.

O valor obtido ($F_{cal} = 10,74$), portanto, é um valor que ocorreria em muitos poucos casos se realmente as produções dos clones fossem iguais, ou seja, provenientes de uma mesma população básica, conforme a consideração inicial. E estes poucos casos

constituem-se nos possíveis valores associados aos erros de decisão neste teste de hipóteses.

$F_{5\%}(3;20) = 3,10$. Como $F_{\text{cal}}(10,74) \geq F_{\text{tab}}(3,10)$ Rejeita-se H_0 . Conclui-se que existe pelo menos um contraste entre as médias de tratamento estatisticamente diferente de zero, ao nível de 5% de probabilidade, pelo teste F.

Pronto! Está realizada a análise de variância e concluiu-se que, pelo menos uma média, é estatisticamente diferente das demais. Para saber quais são os melhores clones, procede-se, preferencialmente, ao desdobramento dos graus de liberdade devidos a tratamento em contrastes ortogonais, no próprio quadro da análise de variância:

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	Fcal
Tratamentos	(3)	(1.071,00)		
(B, C) vs. (A, D)	1	600,00	600,00	18,05*
B vs. C	1	108,00	108,00	3,25ns
A vs. D	1	363,00	363,00	10,92*
Resíduo	20	665,00	33,25	
Total	23	1.736,00		

* Significativo ao nível de 5% de probabilidade.

(B, C) vs. (A, D)	Rejeita-se H_0
B vs. C	Aceita-se H_0
A vs. D	Rejeita-se H_0

ou realiza-se um dos testes de comparação de médias múltiplas:

Quadro 8.3 – Comparação dos diferentes clones por vários testes estatísticos

Clones	Média	Tukey	Duncan	S-N-K	t	Dunnnett
B	59,50	a	a	a	a	Testemunha
C	53,50	a	a b	a	a	n.s
A	52,00	a	b	a	b	n.s
D	41,00	b	c	b	c	*

n.s., *: não significativo e significativo ao nível de 5% de probabilidade, respectivamente.

Neste último caso conclui-se: os clones seguidos de uma mesma letra não diferem estatisticamente entre si ao nível de (...) de probabilidade pelo teste (...).

8.3. Bloco de perguntas 1

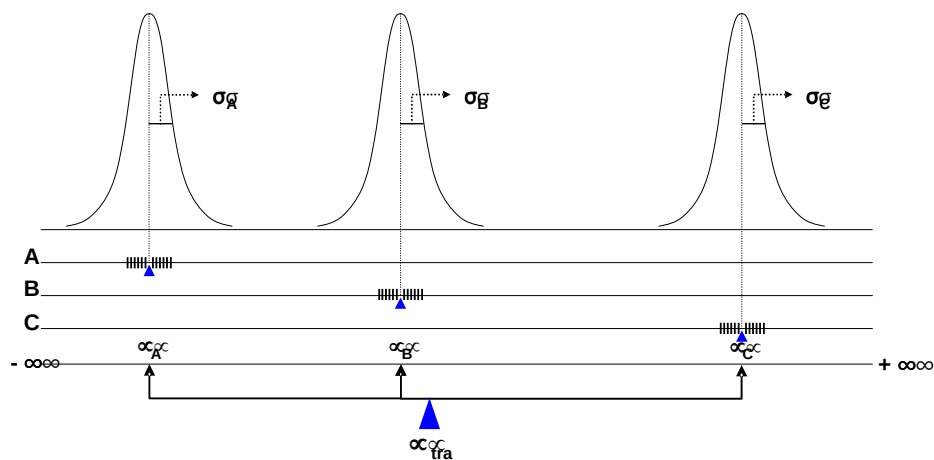
Perguntas de um produtor rural – leigo em estatística, mas que se interessa pelos resultados de seus trabalhos - ao observar os resultados analisados:

1. Qual o significado de se dizer: significativo ao nível de 5% de probabilidade pelo teste F na ANOVA?
2. Se ao invés de 5% de probabilidade fosse utilizado 1 ou 10% de probabilidade, poderia haver alguma diferença nos resultados encontrados?
3. Em caso afirmativo, qual a consequência, em termos de risco caso eu acatasse os clones superiores de seu experimento, em cada caso (1 ou 10%)?
4. Para reduzir ao máximo a probabilidade do “erro” na tomada de decisão, não seria interessante trabalhar com valores mais baixos, por exemplo, 0,1 ou 0,01%? (Obs: o produtor não entende o relacionamento dos erros, tipo I e II, envolvidos na tomada de decisão de um teste de hipóteses. Portanto, explique de forma clara e objetiva a consequência da redução proposta na tomada de decisão em termos dos clones serem consideradas iguais ou diferentes).
5. Estou observando seu quadro de comparação de médias múltiplas e vejo que os resultados obtidos pelos diferentes métodos não são iguais! Ocorreu algum erro, ou esses testes possuem sensibilidade diferenciada para a detecção de possíveis diferenças entre médias de tratamentos?
6. Sendo verdade que existe sensibilidade diferenciada, quais os testes de comparação de médias múltiplas são mais sensíveis (a diferença mínima significativa, dms, é reduzida) na detecção de possíveis diferenças entre médias de tratamentos? Quais os pouco sensíveis (a diferença mínima significativa, dms, é elevada)? Quais os de sensibilidade intermediária?
7. Se eu desejar maior segurança na comparação entre as médias, ou seja, uma vez que o método detecta diferenças entre as médias populacionais estas são realmente diferentes, qual, entre os métodos apresentados, seria o mais recomendado? Justifique.
8. É possível classificar um experimento em relação à qualidade dos procedimentos adotados, ou seja, se este experimento foi bem planejado e bem conduzido? Em caso afirmativo, como seria classificado este experimento.
9. O clone D é o que tenho plantado. Baseado em fundamentos estatísticos, haveria algum ganho de produtividade se fossem plantados os clones C ou A? Que decisão tomar?
10. Para o contexto atual da cacauicultura, supondo os clones como igualmente resistentes a vassoura-de-bruxa, com fundamentos estatísticos, quais clones seriam mais recomendados para a propagação e plantio?

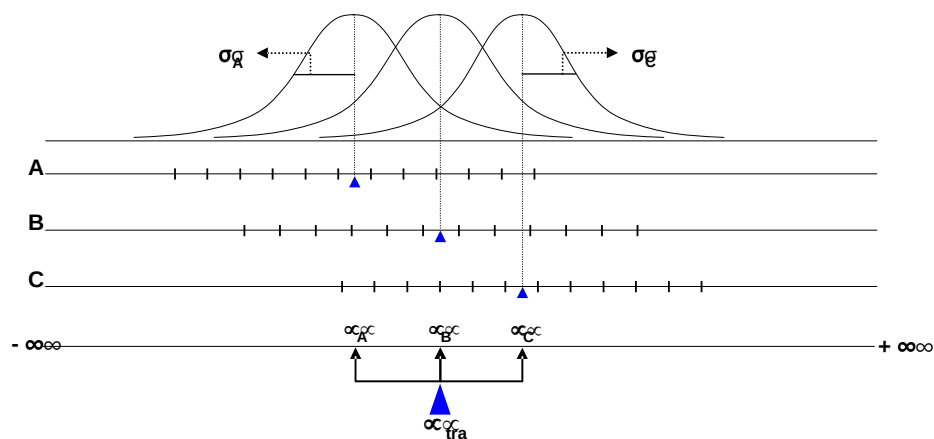
Por tudo o quanto tem sido discutido, você teria condições de apresentar respostas claras e objetivas para estas questões. Afinal, são perguntas de um produtor rural, leigo em estatística experimental.

Considerando a média dos cursos de graduação lecionados na formação acadêmica dos profissionais das ciências da terra, se você possui conceitos e idéias claras sobre estas questões, já é um bom começo. Entretanto, se você quer entender um pouco mais, e talvez até pense em fazer um curso de mestrado, seria desejável ir um pouco mais longe.

Imagine o planejamento, a montagem e a condução de um mesmo experimento, semelhante ao do experimento analisado, realizado de duas formas distintas, cujos resultados são ilustrados na Figura 8.2.



a) Médias de tratamentos distantes e erro experimental pequeno



b) Médias de tratamentos próximas e erro experimental grande

Figura 8.2 – Resultados experimentais hipotéticos para a comparação de três tratamentos dispostos no DIC.

Na situação “a” observa-se que existe uma elevada concentração das repetições de cada tratamento individual em relação à sua média. Ou seja, o desvio padrão, s , de cada tratamento individual apresenta um valor reduzido em relação aos da situação “b”. Em outras palavras, a dispersão das repetições em relação às suas respectivas médias é reduzida. Observa-se, também, que na situação “a” as médias encontram-se mais espaçadas uma das outras que na situação “b”.

Qual seria sua expectativa ao realizar uma análise de variância, seguida de um dos métodos apresentados para comparação dos tratamentos: contrastes ou testes de comparação de médias múltiplas? Em outras palavras, em que situação você esperaria encontrar diferenças significativas entre os tratamentos, na situação “a” ou na “b”?

Para detalhar nossas reflexões, vamos iniciar pelo teste básico que é realizado na ANOVA, o teste F. O teste F é o resultado da razão entre duas estimativas da variância, s^2 . Uma de natureza reconhecida (FRV), controlada ou sistemática no numerador, também denominada variação entre os grupos, e uma outra, de natureza aleatória (FAV), desconhecida ou não controlada no denominador, também denominada variação dentro dos grupos. Assim:

$$F_{cal} = \frac{s^2(FRV)}{s^2(FAV)}$$

Lembrar que o quadrado médio dos desvios do erro experimental ou resíduo (QMDres) representa a variação aleatória, e que somente é possível obtê-la pela análise das repetições de cada tratamento, individualmente. Conforme já discutido, o erro experimental ou resíduo, nada mais é que a média aritmética das variâncias de todos os tratamentos envolvidos na análise:

$$QMDres = \frac{s^2(A) + s^2(B) + s^2(C) + s^2(D)}{4}$$

Para o exemplo numérico fornecido:

$$\begin{aligned} \text{Resíduo} = & [[(58 - 52,00)^2 + \dots + (48 - 52,00)^2] / 5 + \\ & [(60 - 59,50)^2 + \dots + (61 - 59,50)^2] / 5 + \\ & [(59 - 53,50)^2 + \dots + (48 - 53,50)^2] / 5 + \\ & [(45 - 41,00)^2 + \dots + (44 - 41,00)^2] / 5] / 4 = 33,25 \end{aligned}$$

Sempre comparando uma situação em relação à outra (a vs. b), vamos analisar as possibilidades:

No caso “a” esperaríamos um elevado valor do numerador de F, uma vez que as estimativas das médias, m_i , dos “diferentes” tratamentos, encontram-se bastante dispersas em torno da média geral dos tratamentos ($\bar{x}$ tratamentos). Esperaríamos, também, um reduzido valor no denominador de F_{cal} , pois o valor do QMDres seria reduzido, uma vez que as repetições de cada tratamento individual apresentam reduzida dispersão em relação às suas respectivas médias.

Desta forma, o valor de F_{cal} deveria ser elevado. Assim sendo, a chance (probabilidade) do valor obtido, F_{cal} , ser superior a 1 (um) e de se encontrar na região de rejeição de H_0 , RRH_0 , seria elevada (Figura 8.1). Neste caso rejeitaria-se H_0 em um

determinado nível de probabilidade (ou probabilidade de erro), em prol de sua hipótese alternativa H_1 .

Ou seja, o teste F da análise de variância estaria indicando que nem todas as médias de tratamentos, $\bar{y}_i$, seriam estatisticamente iguais. Não se esqueça que as hipóteses são sempre realizadas considerando-se as médias das populações básicas, μ_i , e que para isto utiliza-se as estimativas das médias, $\bar{y}_i$, e suas respectivas estimativas das variâncias, s_i^2 : o que está sendo feito é inferência estatística.

No caso “b” esperaríamos um reduzido valor do numerador de F, uma vez que as estimativas das médias, $\bar{y}_i$, dos diferentes tratamentos, encontram-se pouco dispersas em torno da média geral dos tratamentos ($\bar{y}_{\text{tratamentos}}$). Esperaríamos, também, um elevado valor no denominador de F_{cal} , pois o valor do QMDres seria elevado, uma vez que as repetições de cada tratamento individual apresentam elevada dispersão em relação às suas correspondentes médias.

Desta forma, o valor de F_{cal} deveria ser reduzido. Assim sendo, a chance (probabilidade) do valor obtido, F_{cal} , ser superior a 1 (um) e de se encontrar na região de rejeição de H_0 , RRH_0 , seria reduzida (Figura 8.1). Neste caso, aceitar-se-ia H_0 em um determinado nível de probabilidade (ou probabilidade de erro). Ou seja, o teste F da análise de variância estaria indicando que todas as médias de tratamentos, $\bar{y}_i$, seriam estatisticamente iguais.

Observe também que neste caso, “b”, a partir dos dados apresentados poderíamos confeccionar uma única curva de densidade de probabilidade normal. Ou seja, é como se os “diferentes clones” formassem uma única população, tal é a proximidade de suas médias em relação a média geral, e tal a magnitude da dispersão dos dados de produtividade de amêndoas ($\text{kg } 10 \text{ plantas}^{-1} \text{ ano}^{-1}$), dos “diferentes” tratamentos, em relação às suas respectivas médias, ou seja, as repetições.

Agora reflita: a compreensão espacial do significado da análise de variância, vista até aqui, em comparação apenas com os procedimentos apenas algébricos usuais, pode auxiliar, ou não, na compreensão do significado da ANOVA?

Vamos ainda um pouco mais longe?

8.4. [Bloco de perguntas 2](#)

Você é interrogado por um outro colega profissional, que não teve a oportunidade de compreender muito bem os fundamentos da estatística experimental:

11. Detalhe o que pode ter influenciado, nas fases de planejamento, condução e colheita do experimento para um reduzido valor do resíduo no caso “a”?
12. Detalhe o que pode ter influenciado, nas fases de planejamento, condução e colheita do experimento para um elevado valor do resíduo no caso “b”?
13. No caso “b” se as médias dos tratamentos tivessem se apresentado mais dispersas em relação à média geral dos tratamentos, mantida as mesmas variâncias individuais de cada tratamento, isto aumentaria ou reduziria a chance dos tratamentos serem estatisticamente diferentes? Justifique.
14. Ainda no caso “b”, se a dispersão das repetições em relação a média de cada tratamento individual, fosse reduzida, e fossem mantidas as mesmas médias dos tratamentos, isto aumentaria ou reduziria a chance dos tratamentos mostrarem-se estatisticamente diferentes? Justifique.

15. O aumento do número de repetições do experimento aumentaria ou reduziria a probabilidade de acerto na tomada de decisão das hipóteses? Justifique.

Vamos caminhar ainda um pouco mais? Agora na direção de como as análises estatísticas são feitas utilizando-se computadores pessoais e programas estatísticos. Fica aqui, de antemão, a seguinte mensagem: embora sejam ferramentas de extrema importância para a análise rápida de experimentos, de pouca utilidade são estes programas se o usuário não possuir idéias e conceitos claros sobre o que são e como interpretar os resultados experimentais obtidos com o auxílio computacional. Dito de outra forma, os programas computacionais são apenas ferramentas que realizam cálculos rapidamente, possibilitam o armazenamento e a recuperação rápida das informações e dos dados, permitem visualizações gráficas - o que auxilia bastante a compreensão dos resultados; e nada mais que isto. Conceitos estatísticos simples e claros são fundamentais para sua utilização.

A seguir é apresentado o programa, feito para executar a análise estatística do exemplo, assim como os resultados fornecidos. A análise estatística completa foi obtida utilizando-se o programa SAS ("Statistical Analysis System"). Na atualidade, este é um dos mais completos, confiáveis e utilizados programas para análises estatísticas em computadores em todo o mundo. Cabe ressaltar, entretanto, que existem muitos outros bons programas em universidades, empresas e no mercado.

8.5. Análise computacional de um experimento

8.5.1. Programa para a análise

```
/* Informo um nome (apelido) do arquivo de dados para o SAS..*/
DATA DIC;

/* Informo para que não seja apresentado data e número da página no relatório..*/
OPTIONS LS = 80 NODATE NONUMBER;

/* Informo a ordem das variáveis e que os dados estão em linhas (@@)..*/
INPUT TRA$ REP PROD @@;

/* entre os dois pontos e vírgulas abaixo são fornecidos os dados..*/
CARDS
;
A 1 58 A 2 49 A 3 51 A 4 56 A 5 50 A 6 48
B 1 60 B 2 55 B 3 66 B 4 61 B 5 54 B 6 61
C 1 59 C 2 47 C 3 44 C 4 49 C 5 62 C 6 60
D 1 45 D 2 33 D 3 34 D 4 48 D 5 42 D 6 44
;
PROC GLM DATA=DIC; /* Tipo de análise a ser executada e o nome do arquivo de dados..*/
  CLASS TRA; /* Declarei a classe TRA */
  MODEL PROD = TRA; /* Informa-se que a produção é função dos tratamentos (TRA)..*/

  CONTRAST ' (B, C) vs. (A, D)' TRA -1 1 1 -1; /* Plano de contrastes
  CONTRAST 'B vs. C' TRA 0 1 -1 0;
  CONTRAST 'A vs. D' TRA 1 0 0 -1;

  TITLE 'ANOVA - DELINEAMENTO INTEIRAMENTE CASUALIZADO'; /* Título para o relatório..*/
  MEANS TRA/TUKEY; /* Informo os tipos de testes de médias a serem executados..*/
  MEANS TRA/DUNCAN;
  MEANS TRA/DUNNETT("B"); /* Informo qual é o tratamento testemunha..*/
  MEANS TRA/SNK;
RUN; /* Informo ao programa para executar os comandos listados acima..*/
```

Obs: as palavras entre /* */ não são interpretadas pelo programa, ou seja, são apenas comentários para documentar o programa.

8.5.2. Resultados fornecidos

8.5.2.1. Análise de variância

General Linear Models Procedure

Class Level Information

Class	Levels	Values
TRA	4	A B C D

Number of observations in data set = 24

ANOVA - DELINEAMENTO INTEIRAMENTE CASUALIZADO

General Linear Models Procedure

Dependent Variable: PROD

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	1071.000000	357.000000	10.74	0.0002
Error	20	665.000000	33.250000		
Corrected Total	23	1736.000000			

R-Square	C.V.	Root MSE	PROD Mean
0.616935	11.19666	5.7662813	51.500000

Source	DF	Type I SS	Mean Square	F Value	Pr > F
TRA	3	1071.000000	357.000000	10.74	0.0002

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
(B, C) vs. (A, D)	1	600.000000	600.000000	18.05	0.0004
B vs. C	1	108.000000	108.000000	3.25	0.0866
A vs. D	1	363.000000	363.000000	10.92	0.0035
Error	20	665.000000	33.250000		

8.5.2.2. Testes de comparação de médias

8.5.2.2.1. Teste de Tukey

General Linear Models Procedure

Tukey's Studentized Range (HSD) Test for variable: PROD

NOTE: This test controls the type I experimentwise error rate, but generally has a higher type II error rate than REGWQ.

Alpha= 0.05 df= 20 MSE= 33.25

Critical Value of Studentized Range= 3.958

Minimum Significant Difference= 9.3181

Means with the same letter are not significantly different.

Tukey Grouping	Mean	N	TRA
A	59.500	6	B
A	53.500	6	C
A	52.000	6	A
B	41.000	6	D

8.5.2.2.2. Teste de Duncan

General Linear Models Procedure

Duncan's Multiple Range Test for variable: PROD

NOTE: This test controls the type I comparisonwise error rate, not the experimentwise error rate

Alpha= 0.05 df= 20 MSE= 33.25
 Number of Means 2 3 4
 Critical Range 6.945 7.289 7.509

Means with the same letter are not significantly different.

Duncan Grouping		Mean	N	TRA
	A	59.500	6	B
B	A	53.500	6	C
B		52.000	6	A
	C	41.000	6	D

8.5.2.2.3. Teste de Dunnett

General Linear Models Procedure

Dunnett's T tests for variable: PROD

NOTE: This tests controls the type I experimentwise error for comparisons of all treatments against a control.

Alpha= 0.05 Confidence= 0.95 df= 20 MSE= 33.25
 Critical Value of Dunnett's T= 2.540
 Minimum Significant Difference= 8.4575

Comparisons significant at the 0.05 level are indicated by '***'.

TRA Comparison	Simultaneous Lower Confidence Limit	Difference Between Means	Simultaneous Upper Confidence Limit	
C - B	-14.458	-6.000	2.458	
A - B	-15.958	-7.500	0.958	
D - B	-26.958	-18.500	-10.042	***

8.5.2.2.4. Teste de Student – Newman – Keuls

General Linear Models Procedure

Student-Newman-Keuls test for variable: PROD

NOTE: This test controls the type I experimentwise error rate under the complete null hypothesis but not under partial null hypotheses.

Alpha= 0.05 df= 20 MSE= 33.25
 Number of Means 2 3 4
 Critical Range 6.9445189 8.422726 9.318121

Means with the same letter are not significantly different.

SNK Grouping	Mean	N	TRA
A	59.500	6	B
A	53.500	6	C
A	52.000	6	A
B	41.000	6	D

Muito prático, não? Observa-se que no caso da análise realizada com o auxílio computacional não aparecem no quadro da ANOVA os conhecidos asteriscos (*, ** e ***) indicativos da significância de cada valor F calculado (F_{cal}). Ao invés disso, o programa apresenta o valor da probabilidade do erro tipo I, α , ou seja, a probabilidade de rejeitarmos a hipótese H_0 , sendo esta, de fato, verdadeira. Ou seja, decidir que os clones são diferentes quando na verdade são iguais. No caso da ANOVA realizada, o valor desta probabilidade foi 0,0002, ou seja, a probabilidade de estarmos errados ao rejeitarmos H_0 é de apenas 0,02%, e a de estarmos corretos em nossa decisão é de 0,98% ($1 - \alpha$).

Se o valor F calculado (F_{cal}) fosse, por exemplo 3,10, mantidos os mesmos graus de liberdade para a fonte de variação reconhecida em teste, tratamentos, e também para o resíduo, o valor que apareceria na coluna da probabilidade do programa, à frente do valor F_{cal} seria exatamente 0,050. Neste caso ao rejeitarmos H_0 , teríamos 5% de probabilidade de estarmos errados e 95% ($1 - \alpha$) de estarmos corretos. Observe a coincidência destes valores com os assinalados na Figura 8.1. Você não acha que a forma como o programa emite o relatório da ANOVA é muito mais informativa que utilizando apenas a tabela?

Seu raciocínio deve estar ficando ainda mais claro não? Vamos ainda um pouco mais longe?

8.6. Bloco de perguntas 3

Você agora é interrogado por um estatístico experimental:

16. O que é quantificado na ANOVA pelo erro experimental ou resíduo? Em outras palavras, ele reflete a influência de quais fontes de variação?
17. No exemplo analisado o que é quantificado na ANOVA pelo efeito de tratamento? Em outras palavras, ele reflete a influência de qual fonte de variação?
18. Faça uma análise comparativa qualitativa entre os testes de comparação de médias múltiplas apresentados (Tukey, Duncan, etc.) em relação à análise de contrastes. Ou seja, compare os métodos em conjunto com os contrastes. De sua opinião em relação à flexibilidade (comparações possíveis de serem obtidas) e facilidade de cálculos.
19. Se a probabilidade apresentada no teste F da ANOVA para a fonte de variação “tratamento” fosse 0,062 (6,2%), neste caso não significativo a 5%, você ainda assim continuaria a análise estatística e realizaria um dos métodos de comparação de médias (contrastos ou testes de comparação de médias múltiplas) ou não? Justifique sua decisão. Observação: Visualize a possibilidade de um conjunto de médias de tratamentos se apresentar muito próximas entre si, e apenas uma das médias se distanciar do restante do grupo. Lembre-se que a variância devida ao efeito dos tratamentos é uma medida aproximada da dispersão média de cada tratamento em torno da média geral do experimento.

20. Um dos pressupostos básicos para a realização de uma ANOVA é que exista homocedasticidade (invariância da variância) entre os “diferentes” tratamentos. O que isto significa?
21. No quadro da ANOVA onde se realizou o desdobramento dos graus de liberdade em contrastes ortogonais, qual é a conclusão quando os clones comparados são B vs. C? Você recomendaria os dois indistintamente ou preferiria recomendar o B? Justifique.
22. Considerando a análise realizada, utilize o teste de comparação de médias adequado para testar adicionalmente o contraste (B, C, A) vs. (D) e conclua ao nível de 5% de probabilidade.
23. Qual a seria a forma mais eficiente (e a única) de aumentarmos a confiabilidade de nossas decisões, ou seja, afirmar que existem diferenças estatísticas em relação às fontes de variação controladas quando, de fato, elas existem, e que não existem diferenças quando, também de fato, elas não existem?

9. Delineamento em blocos casualizados - DBC

9.1. Introdução

É o mais conhecido e utilizado entre os delineamentos experimentais. Os experimentos montados segundo este delineamento são denominados experimentos em blocos casualizados ou blocos ao acaso (DBC).

Consiste em dividir o material experimental em grupos homogêneos ou blocos, pois pressupõe a existência de similaridade dentro de cada bloco individual. Entre blocos, entretanto, pode haver variação, à vontade.

Compõe-se de tantos blocos quantas forem as repetições dos tratamentos.

Os tratamentos são designados às parcelas, dentro de cada bloco, de forma inteiramente aleatória ou casual.

A confecção dos blocos, no entanto, não é fruto do acaso, mas da intervenção direta do pesquisador, que deverá decidir onde e como serão estabelecidos, segundo as necessidades do experimento e de seus propósitos.

O DBC é utilizado quando se deseja controlar pelo menos uma causa ou fonte de variação adicional reconhecida, além do efeito dos tratamentos.

9.2. Princípios utilizados

9.2.1. Repetição

Permite a estimativa do erro experimental ou resíduo, sendo seu número dependente da variabilidade do material experimental.

9.2.2. Casualização

Garante que as possíveis diferenças entre tratamentos não seja por favorecimento.

9.2.3. Controle local

É feito através do uso de blocos homogêneos.

Garante que as possíveis variações entre as repetições, devido à heterogeneidade das condições experimentais, e ou, do material experimental, não seja atribuída ao erro experimental ou resíduo.

9.2.4. Exemplos de controle local

Falta de uniformidade do terreno (gradientes de fertilidade, umidade, etc).

Gradientes de luminosidade, e ou, temperatura no interior de casas de vegetação.

Animais com peso inicial, Sexo, idade, condições, etc, diferente ao se estudar ganho de peso.

Idade de lactação diferente dos animais ao se estudar a influência de diferentes manejos.

Uso de mais de uma pessoa para se avaliar certos caracteres, mais de um equipamento, etc.

Deve ficar claro que entre blocos pode haver grande variação, pois esta variação não refletirá, apenas por si, em um elevado valor do quadrado médio do resíduo. No entanto, no interior de cada bloco, as condições experimentais, e ou, o material experimental, devem ser tão homogêneos quanto possível.

9.3. Vantagens e desvantagens

9.3.1. Vantagens

As unidades experimentais são agrupadas em blocos homogêneos, permitindo, em geral, maior precisão que no DIC.

Não há restrições no número de tratamentos ou blocos.

A análise estatística é simples.

9.3.2. Desvantagens

Quando a variação entre as unidades experimentais dentro dos blocos é grande, resulta em um grande erro experimental.

Isto geralmente ocorre quando o número de tratamentos é grande e não é possível assegurar uniformidade entre as unidades experimentais dentro dos blocos.

9.4. Modelo estatístico

$$y_{ij} = \mu + t_i + b_j + e_{ij}$$

onde,

y_{ij}	= Valor observado na parcela do tratamento, i , no bloco, j
μ	= Média geral do experimento
t_i	= Efeito do tratamento, i , aplicado na parcela
b_j	= Efeito do bloco, j
e_{ij}	= Efeito dos fatores não controlados

9.5. Esquema de casualização dos tratamentos

Seja um experimento envolvendo 5 tratamentos (A, B, C, D, E) em 4 repetições (20 unidades experimentais ou parcelas):

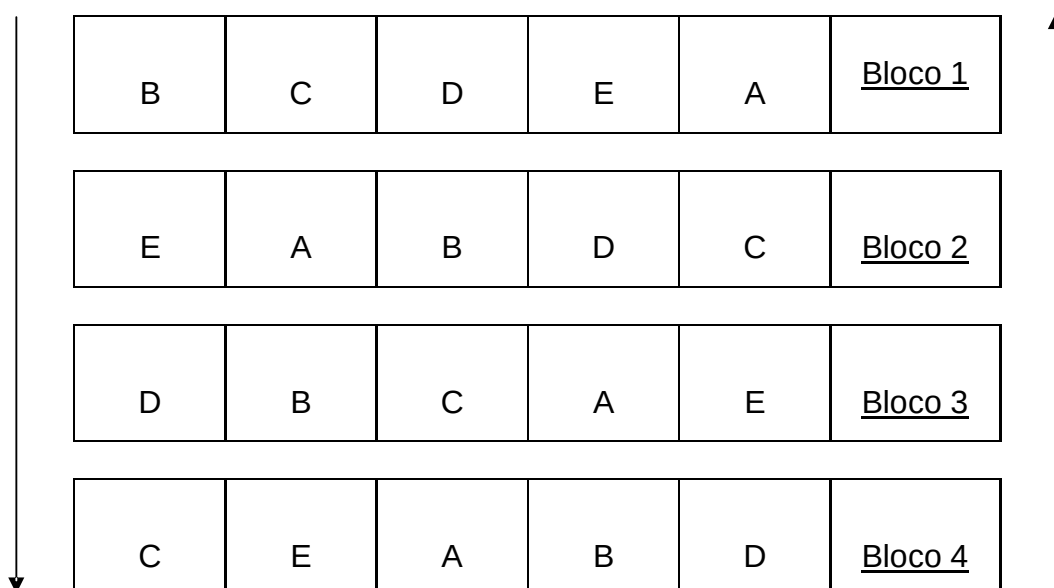


Figura 9.1 – Esquema da casualização das unidades experimentais. As setas à esquerda da figura estão indicando os sentidos dos possíveis gradientes.

9.6. Coleta de dados

Quadro 9.1 – Quadro para coleta de dados de experimentos no DBC

Tratamentos	Blocos			Totais	Médias
	1	...	j		
A	y_{11}	...	y_{1j}	t_1	m_1
B	y_{21}	...	y_{2j}	t_2	m_2
.	.	.	.	.	.
.	.	.	.	.	.
.	.	.	.	.	.
I	y_{i1}	...	y_{ij}	t_i	m_i

9.7. Análise de variância

9.7.1. Esquema da análise de variância

Quadro 9.2 – Quadro da análise de variância no DBC

Causa da variação	GL	SQD	QMD	F_{cal}
Blocos	$j - 1$	SQDblo	QMDblo	QMDblo/QMDres
Tratamentos	$i - 1$	SQDtra	QMDtra	QMDtra/QMDres
Resíduo	$(i - 1)(j - 1)$	SQDres	QMDres	
Total	$n - 1$	SQDtot		

9.7.2. Teste de hipóteses

Relativas aos tratamentos

$$H_0: \alpha_A = \alpha_B = \dots = \alpha_I$$

$$H_1: \text{Nem todas as médias de tratamentos são iguais}$$

Relativas aos blocos

$$H_0: \alpha_{B1} = \alpha_{B2} = \dots = \alpha_{Bj}$$

$$H_1: \text{Nem todas as médias de blocos são iguais}$$

9.8. Exemplo com um mesmo número de repetições

Os dados abaixo foram obtidos de um experimento no DBC com 4 repetições. Os tratamentos constaram de 5 variedades de macieira e o peso médio dos frutos, em gramas, de cada variedade é dado a seguir:

Quadro 9.3 – Peso dos frutos, em gramas, das variedades de macieira

Tra	Repetições				Totais
	1	2	3	4	
A	142,36	144,78	145,19	138,88	571,21
B	139,28	137,77	144,44	130,61	552,10
C	140,73	134,06	136,07	144,11	554,97
D	150,88	135,83	136,97	136,36	560,04
E	153,49	165,02	151,75	150,22	620,48
Totais	726,74	717,46	714,42	700,18	2.858,80

$$C = (2.858,80)^2 / 20 = 408.636,87$$

$$SQD_{tot} = [(142,36)^2 + (144,78)^2 + \dots + (150,22)^2] - C = 1.273,95$$

$$SQD_{tra} = 1 / 4 [(571,21)^2 + (552,10)^2 + \dots + (620,48)^2] - C = 794,93$$

$$SQD_{blo} = 1 / 5 [(726,74)^2 + (717,46)^2 + \dots + (700,18)^2] - C = 72,70$$

$$SQD_{res} = SQD_{tot} - SQD_{tra} - SQD_{blo} = 1.273,95 - 794,93 - 72,70 = 406,35$$

Hipóteses relativas aos tratamentos:

$$H_0: \alpha_I = \alpha_K \text{ (para todo } I \neq K)$$

$$H_1: \text{Nem todas as } \alpha_I \text{ são iguais}$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Blocos	3	72,70	24,23	0,72	0,5614
Tratamentos	4	794,93	198,73	5,87	0,0074
Resíduo	12	406,35	33,86		
Total	19	1.273,95			

$$cv = 100 \cdot (\sqrt{33,86/142,94}) = 4,07\%$$

Rejeita-se H_0 . Conclui-se que existe pelo menos um contraste entre médias de tratamentos estatisticamente diferente de zero, ao nível de 5% de probabilidade, pelo teste F.

9.8.1. Testes de comparação de médias múltiplas

Quadro 9.4 – Comparação da sensibilidade dos diferentes testes de médias múltiplas

Variedades	Média	Tukey	Duncan	SNK	t	Dunnnett
E	155,12	a	a	a	a	*
A	142,80	a b	b	b	b	n.s
D	140,01	b	b	b	b	Testemunha
C	138,74	b	b	b	b	n.s
B	138,03	b	b	b	b	n.s

Obs: realizar os testes de Tukey, Duncan e SNK para treinamento.

9.8.2. Desdobramento dos gl associados a tratamentos em contrastes ortogonais

Como temos quatro graus de liberdade associados a tratamentos podemos estabelecer até quatro contrastes ortogonais, por exemplo:

$$C_1 = D \text{ vs}(A, B, C, E)$$

$$C_2 = (A, E) \text{ vs}(B, C)$$

$$C_3 = A \text{ vs} E$$

$$C_4 = B \text{ vs} C$$

Estabelecendo os contrastes ortogonais:

$$C_1 = 4D - 1A - 1B - 1C - 1E$$

$$C_2 = 1A + 1E - 1B - 1C$$

$$C_3 = 1A - 1E$$

$$C_4 = 1B - 1C$$

Inicialmente calculamos as estimativas dos contrastes:

$$\hat{C}_1 = 4(560,04) - 1(571,21) - 1(552,10) - 1(554,97) - 1(620,48) = - 58,60$$

$$\hat{C}_2 = 1(571,21) + 1(620,48) - 1(552,10) - 1(554,97) = 84,62$$

$$\hat{C}_3 = 1(571,21) - 1(620,48) = - 49,27$$

$$\hat{C}_4 = 1(552,10) - 1(554,97) = - 2,87$$

Agora podemos calcular a soma de quadrados dos contrastes:

$$SQD(C_1) = (- 58,60)^2 / 4 [(4)^2 + (-1)^2 + (-1)^2 + (-1)^2 + (-1)^2] = 42,92$$

$$SQD(C_2) = (84,62)^2 / 4 [(1)^2 + (1)^2 + (-1)^2 + (-1)^2] = 447,53$$

$$SQD(C_3) = (- 49,27)^2 / 4 [(1)^2 + (-1)^2] = 303,44$$

$$SQD(C_4) = (- 2,87)^2 / 4 [(1)^2 + (-1)^2] = 1,03$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Blocos	3	72,70			
Tratamentos	(4)	(794,93)			
D vs. (A,B,C,E)	1	42,92	42,92	1,27	0,2822
(A,E) vs. (B,C)	1	447,53	447,53	13,22	0,0034
A vs. E	1	303,44	303,44	8,96	0,0112
B vs. C	1	1,03	1,03	0,03	0,8645
Resíduo	12	406,35	33,86		
Total	19	1.273,95			

9.9. Considerações finais

Embora o delineamento em blocos casualizados seja simples, flexível e de fácil análise, no planejamento, na montagem, na condução e na coleta de dados nesse tipo de delineamento, é importante a presença e de um estatístico experimental experiente, assessorando todas as etapas do ciclo experimental.

As etapas cruciais são: a identificação das fontes de variação intervenientes, a forma de distribuir os blocos e a definição do número de blocos necessários.

A análise de experimentos onde foram perdidas algumas unidades experimentais implica na adoção de procedimentos adequados para a análise, que envolvem, em geral, a estimação da parcela perdida utilizando critérios estatísticos.

10. Delineamento em quadrado latino - DQL

10.1. Introdução

Utiliza-se este delineamento quando é possível reconhecer duas fontes de variação antes da aplicação dos tratamentos.

Cada uma dessas fontes de variação (linhas e colunas) deve ter o mesmo número de níveis, n , que o número de tratamentos, n^2 .

As unidades experimentais são arranjadas em um quadrado, $n \times n$, e os i tratamentos são aplicados ao acaso, de tal forma que cada tratamento aparece, exatamente, uma vez em cada linha e em cada coluna.

O número de tratamentos é igual ao número repetições. Dentro das linhas, e dentro das colunas, deve-se ter a maior uniformidade possível.

Os quadrados latinos constituem um bom tipo de delineamento, mas sua flexibilidade é muito menor em relação ao DBC.

10.2. Princípios utilizados

10.2.1. Repetição

Permite a estimativa do erro experimental ou resíduo, sendo seu número dependente da variabilidade do material experimental.

10.2.2. Casualização

Garante que as possíveis diferenças entre tratamentos não seja por favorecimento.

10.2.3. Controle local

É feito através do uso de linhas e colunas homogêneas.

Garante que as possíveis variações entre as repetições, devido à heterogeneidade das condições experimentais, e ou, do material experimental, não seja atribuída ao erro experimental ou resíduo.

10.2.4. Exemplos de causas de variação controladas por este delineamento

Gradientes de fertilidade e umidade, perpendiculares entre si, no solo e no interior de casas de vegetação.

Animais de mesma idade nas linhas e de mesmo peso inicial nas colunas ao se estudar ganho de peso, etc.

Aplicador e máquinas diferentes ao se estudar controles alternativos de invasoras, pragas e doenças.

Heterogeneidade em áreas experimentais de uso intensivo.

10.3. Vantagens e desvantagens

10.3.1. Vantagens

Possibilidade de se controlar, simultaneamente, duas fontes de variação sistemáticas em adição aos tratamentos.

10.3.2. Desvantagens

Pouca flexibilidade.

Redução no número de graus de liberdade associados ao resíduo.

Excessivo número de unidades experimentais necessárias quando o número de tratamentos é grande.

10.4. Modelo estatístico

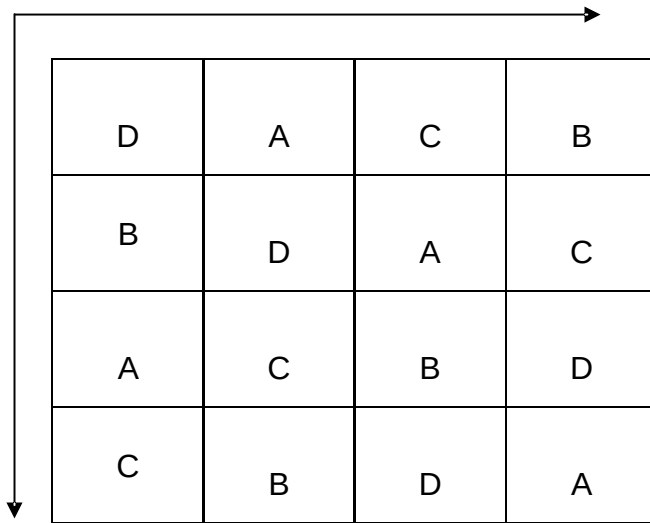
$$y_{ijk} = \mu + I_i + c_j + (t_k)_{ij} + e_{ijk}$$

onde,

y_{ijk}	= Valor observado na parcela do tratamento k na linha i e na coluna j
μ	= Média geral do experimento
I_i	= Efeito da linha i
c_j	= Efeito da coluna j
$(t_k)_{ij}$	= Efeito do tratamento k na linha i e na coluna j
e_{ijk}	= Efeito dos fatores não controlados

10.5. Esquema de casualização dos tratamentos

Seja um experimento envolvendo 4 tratamentos (A, B, C, D):



D	A	C	B
B	D	A	C
A	C	B	D
C	B	D	A

Figura 9.1 – Esquema da casualização das unidades experimentais. As setas à esquerda da figura estão indicando as direções dos possíveis gradientes.

Observa-se que cada tratamento é casualizado, tendo que estar presente uma única vez em cada linha e uma única vez em cada coluna.

10.6. Coleta de dados

Quadro 9.1 – Quadro para coleta de dados de experimentos no DQL

Linha	Coluna			Totais de linhas
	1	...	j	
1	Y_{11k}	...	Y_{1jk}	l_1
2	Y_{21k}	...	Y_{2jk}	l_2
.	.	.	.	.
.	.	.	.	.
.	.	.	.	.
i				l_i
Totais de colunas	c_1	...	c_j	

10.7. Análise de variância

10.7.1. Esquema da análise de variância

Causa da variação	GL	SQD	QMD	F _{cal}
Linhas	k - 1	SQDlin	QMDlin	QMDlin/QMDres
Colunas	k - 1	SQDcol	QMDcol	QMDcol/QMDres
Tratamentos	k - 1	SQDtra	QMDtra	QMDtra/QMDres
Resíduo	(k - 2) (k - 1)	SQDres	QMDres	
Total	k ² - 1	SQDtot		

10.7.2. Teste de hipóteses relativas aos tratamentos

$$H_0: \alpha_A = \alpha_B = \dots = \alpha_K$$

H_1 : Nem todas as médias são iguais

Caso haja interesse em testar as fontes de variação que foram alocadas nas linhas e colunas, hipótese semelhantes aos dos tratamentos devem ser formuladas para ambas.

10.8. Exemplo com um mesmo número de repetições

Os dados abaixo foram obtidos de um experimento de competição de cana-de-açúcar. Foram utilizadas cinco variedades (A, B, C, D e E) dispostas no delineamento em quadrado latino. As produções de cana-planta, em kg parcela⁻¹, são dadas a seguir:

Quadro 9.2 – Peso de cana-planta, em kg parcela⁻¹

						T. linhas				
D	432	A	518	B	458	C	583	E	331	2.322
C	724	E	478	A	524	B	550	D	400	2.676
E	489	B	384	C	556	D	297	A	420	2.146
B	494	D	500	E	313	A	486	C	501	2.294
A	515	C	660	D	438	E	394	B	318	2.325
T. colunas	2.654	2.540	2.289	2.310	1.970	11.763				

$$C = (11.763)^2 / 25 = 5.634.726,76$$

$$SQD_{tot} = [(432)^2 + (518)^2 + \dots + (318)^2] - C = 257.724,24$$

$$SQD_{lin} = 1 / 5 [(2.322)^2 + (2.676)^2 + \dots + (2.325)^2] - C = 30.480,64$$

$$SQD_{col} = 1 / 5 [(2.654)^2 + (2.540)^2 + \dots + (1.970)^2] - C = 55.640,64$$

Quadro auxiliar

Totais de tratamentos	N.Repetições
A = 2.463	5
B = 2.204	5
C = 3.024	5
D = 2.067	5
E = 2.005	5

$$SQD_{tra} = 1 / 5 [(2.463)^2 + (2.204)^2 + \dots + (2.005)^2] - C = 137.488,24$$

$$SQD_{res} = SQD_{tot} - SQD_{lin} - SQD_{col} - SQD_{tra}$$

$$SQD_{res} = 257.724,24 - 30.480,64 - 55.640,64 - 137.488,24$$

$$SQD_{res} = 34.114,72$$

Hipóteses:

$$H_0: \sigma_I = \sigma_K \text{ (para todo } I \neq K)$$

$$H_1: \text{Nem todas as } \sigma_I \text{ são iguais}$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Linhas	4	30.480,64			
Colunas	4	55.640,64			
Tratamentos	4	137.488,24	34.372,06	12,09	0,0004
Resíduo	12	34.114,72	2.842,89		
Total	24	257.724,24			

$$cv = 11,33 \%$$

Rejeita-se H_0 . Conclui-se que existe pelo menos um contraste entre médias de tratamentos estatisticamente diferente de zero, ao nível de 5% de probabilidade, pelo teste F.

10.8.1. Testes de comparação de médias múltiplas

Quadro 9.3 – Comparação da sensibilidade dos diferentes testes de médias múltiplas

Variedades	Média	Tukey	Duncan	SNK	t	Dunnett
C	604,80	a	a	a	a	*
A	492,60	b	b	b	b	n.s
B	440,80	b	b c	b	b c	n.s
D	413,40	b	c	b	c	Testemunha
E	401,00	b	c	b	c	n.s

Obs: realizar os testes de Tukey, Duncan e SNK para treinamento.

10.8.2. Desdobramento dos gl de tratamentos em contrastes ortogonais

$$C_1 = D \text{ vs}(A, B, C, E)$$

$$C_2 = (A, E) \text{ vs}(B, C)$$

$$C_3 = A \text{ vs} E$$

$$C_4 = B \text{ vs} C$$

$$C_1 = 4D - 1A - 1B - 1C - 1E$$

$$C_2 = 1A + 1E - 1B - 1C$$

$$C_3 = 1A - 1E$$

$$C_4 = 1B - 1C$$

$$\hat{C}_1 = 4(2.067) - 1(2.463) - 1(2.204) - 1(3.024) - 1(2.005) = - 1.428,00$$

$$\hat{C}_2 = 1(2.463) + 1(2.005) - 1(2.204) - 1(3.024) = - 760,00$$

$$\hat{C}_3 = 1(2.463) - 1(2.005) = 458,00$$

$$\hat{C}_4 = 1(2.204) - 1(3.024) = - 820,00$$

$$SQD(C_1) = (- 1.428)^2 / 5 [(4)^2 + (-1)^2 + (-1)^2 + (-1)^2 + (-1)^2] = 20.391,84$$

$$SQD(C_2) = (- 760)^2 / 5 [(1)^2 + (1)^2 + (-1)^2 + (-1)^2] = 28.880,00$$

$$SQD(C_3) = (458)^2 / 5 [(1)^2 + (-1)^2] = 20.976,40$$

$$SQD(C_4) = (- 820)^2 / 5 [(1)^2 + (-1)^2] = 67.240,00$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Linhas	4	30.480,64			
Colunas	4	55.640,64			
Tratamentos	(4)	(137.488,24)			
D vs. (A,B,C,E)	1	20.391,84	20.391,84	7,17	0,0201
(A,E) vs. (B,C)	1	28.880,00	28.880,00	10,16	0,0078
A vs. E	1	20.976,40	20.976,40	7,38	0,0187
B vs. C	1	67.240,00	67.240,00	23,65	0,0004
Resíduo	12	34.114,72	2.842,89		
Total	24	257.724,24			

Variedades	Média
C	604,80
A	492,60
B	440,80
D	413,40
E	401,00

10.9. Considerações finais

As etapas cruciais para a utilização deste delineamento são: a identificação das fontes de variação intervenientes, a forma de distribuir as linhas e as colunas, assim como a definição do número de repetições necessárias.

A análise de experimentos onde foram perdidas algumas unidades experimentais implica na adoção de procedimentos adequados para a análise, que envolvem, em geral, a estimação da parcela perdida utilizando critérios estatísticos.

O efeito de qualquer possível fonte de variação sistemática dentro das linhas, e ou, colunas, além dos tratamentos, será atribuída ao erro experimental, diminuindo a probabilidade de se detectar possíveis diferenças entre tratamentos, caso existam.

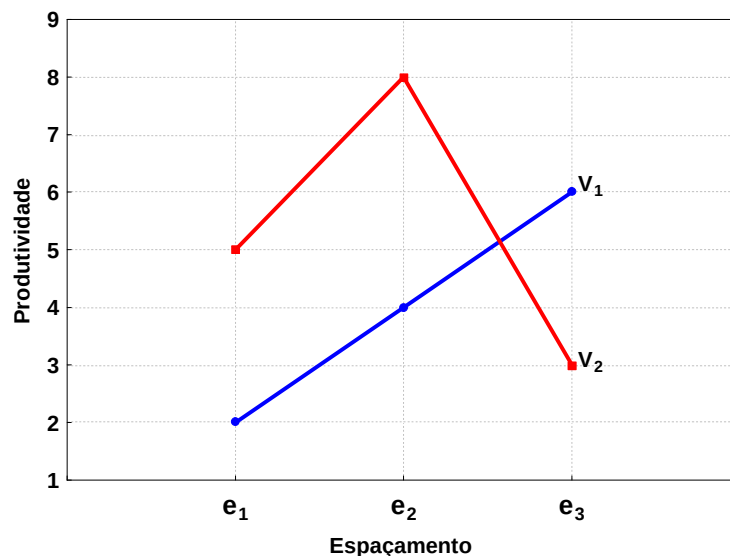
11. Experimentos fatoriais

11.1. Introdução

Os experimentos fatoriais não constituem um delineamento, são formas de montar e analisar experimentos.

Podem ser executados em qualquer um dos delineamentos (DIC, DBC, DQL, etc) onde se estudam simultaneamente dois ou mais fatores.

São mais eficientes do que os experimentos simples, com um só conjunto de tratamentos, permitindo retirar conclusões mais abrangentes.



Cada nível de um fator se combina com cada um dos níveis dos outros fatores, constituindo um tratamento.

Assim, em um experimento com dois fatores A e B, onde o fator A tem 4 níveis ($a_1, \dots, a_4$) e o fator B tem 3 níveis ($b_1, \dots, b_3$), teremos, então, um fatorial 4×3 e os tratamentos, resultantes de todas as combinações possíveis, são:

a_1b_1	a_2b_1	a_3b_1	a_4b_1
a_1b_2	a_2b_2	a_3b_2	a_4b_2
a_1b_3	a_2b_3	a_3b_3	a_4b_3

Um fatorial 3^3 se caracteriza pela combinação de 3 fatores (expoente), cada um com 3 níveis (base), resultando assim em 27 diferentes combinações, constituindo os tratamentos. Assim, poderíamos combinar:

3 doses de N
 3 doses de P
 3 doses de K

Um fatorial $3^1 \times 2^2$ se caracteriza pela combinação de 3 fatores (soma dos expoentes), sendo um fator com 3 níveis e os outros dois fatores com 2 níveis, resultando assim em 12 combinações que constituem os tratamentos. Assim, poderíamos combinar:

3 doses de N

2 doses de P

2 doses de K

A notação genérica destes experimentos é dada por: $(\text{Níveis})^{\text{Fatores}}$

Exemplos de notação:

$3^1 \times 2^2$: 3 fatores: 3 níveis de um fator
2 níveis de dois fatores / 12 tratamentos.

$4^2 \times 3^2$: 4 fatores: 4 níveis de dois fatores
3 níveis de dois fatores / 144 tratamentos

$4^1 \times 2^4$: 5 fatores: 4 níveis de um fator
2 níveis de quatro fatores / 64 tratamentos

11.2. Classificação dos efeitos

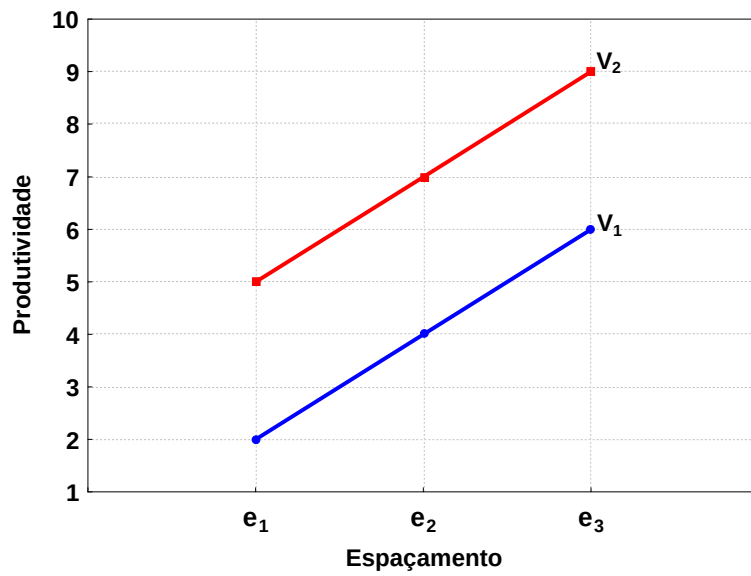
11.2.1. Efeito principal

É o efeito de cada fator independentemente da influência de outros fatores.

11.2.2. Efeito da interação

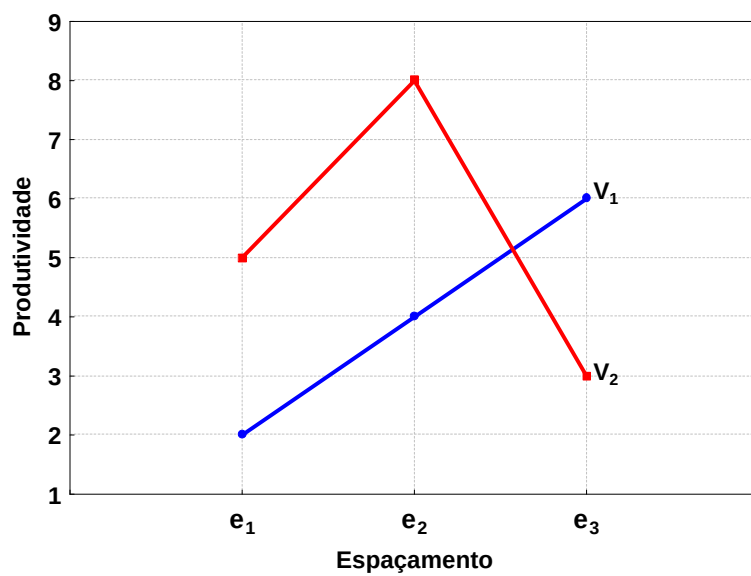
É a resposta diferencial da combinação de tratamentos que não se deve aos efeitos principais. Ocorre interação quando a resposta, ou efeitos, dos níveis de um fator são modificados pelos níveis do(s) outro(s) fator(es).

	E	e_1	e_2	e_3
V				
v_1		2	4	6
v_2		5	7	9



Não há interação

	E	e_1	e_2	e_3
V				
v_1		2	4	6
v_2		5	8	3



Há interação

11.3. Vantagens e desvantagens

11.3.1. Vantagens

A grande vantagem dos experimentos fatoriais é a possibilidade do estudo das interações e sua grande versatilidade, uma vez que pode ser utilizado em vários delineamentos experimentais.

11.3.2. Desvantagens

A principal desvantagem é o rápido crescimento das unidades experimentais com o aumento dos fatores e seu níveis, podendo, contudo, ser contornado por técnicas de confundimento e o uso de matrizes experimentais.

11.4. Modelo estatístico

$$y_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + e_{ijk}$$

$$i = 1, \dots, I$$

$$j = 1, \dots, J$$

$$k = 1, \dots, K$$

onde,

y_{ijk} = observação relativa ao i -ésimo nível do fator A e ao j -ésimo nível do fator B na repetição k
 μ = média geral
 α_i = efeito do i -ésimo nível do fator A, definido por: $\alpha_i = \alpha_{.i} - \mu$
 β_j = efeito do j -ésimo nível do fator B, definido por: $\beta_j = \alpha_{.j} - \mu$
 $\alpha\beta_{ij}$ = efeito da interação entre o i -ésimo nível do fator A e o j -ésimo nível do fator B, definido por: $\alpha\beta_{ij} = \alpha_{ij} - (\mu + \alpha_i + \beta_j)$
 e_{ijk} = erro aleatório associado à observação y_{ijk}

11.5. Coleta de dados

Quadro 11.1 - Coleta de dados de experimentos fatoriais

a_1			a_2			...			a_i		
b_1	...	b_j	b_1	...	b_j	b_1	...	b_j	b_1	...	b_j
y_{111}	...	y_{1j1}	y_{211}	...	y_{2j1}		...		y_{i11}	...	y_{ij1}
.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.
y_{11k}	...	y_{1jk}	y_{21k}	...	y_{2jk}		...		y_{i1k}	...	y_{ijk}

11.6. Análise de variância

11.6.1. Esquema da análise de variância

O esquema da análise de variância será dependente do delineamento adotado na montagem do experimento.

Para um experimento montado no DBC, teríamos:

Quadro 11.2 – Quadro da análise de variância do experimento fatorial no DBC

Causa da variação	GL	SQD	QMD	F _{cal}
Blocos	k - 1	SQDblo		
Tratamentos	(IJ - 1)	(SQDtra)		
A	I - 1	SQD(A)	QMD(A)	QMD(A)/QMDres
B	J - 1	SQD(B)	QMD(B)	QMD(B)/QMDres
A x B	(I - 1)(J - 1)	SQD(AxB)	QMD(AxB)	QMD(AxB)/QMDres
Resíduo	IJ(k - 1)	SQDres	QMDres	
Total	IJK - 1	SQDtot		

11.6.2. Testes de hipóteses

$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$ $H_1: \text{Não } H_0$	$H_0: \alpha_1 = \dots = \alpha_I = 0$ $H_1: \text{Não } H_0$	$H_0: \beta_1 = \dots = \beta_J = 0$ $H_1: \text{Não } H_0$
--	--	--

11.7. Exemplos

11.7.1. Experimento montado no DIC com interação não significativa

Seja um experimento realizado para se estudar variedade de milho, fator A, e espaçamento, fator B, sendo variedade com 3 níveis e espaçamento com 4 níveis, totalizando 12 tratamentos ($3^1 \times 4^1$), dispostos no delineamento inteiramente casualizado com 3 repetições. Os totais de tratamentos constam no quadro a seguir:

Quadro 11.3 - Totais de tratamentos da produção de milho em kg parcela⁻¹

A	B					Totais de A
		b ₁	b ₂	b ₃	b ₄	
a ₁		(3) 120	(3) 132	(3) 150	(3) 162	(12) 564
a ₂		(3) 126	(3) 141	(3) 162	(3) 171	(12) 600
a ₃		(3) 144	(3) 150	(3) 171	(3) 186	(12) 651
Totais de B		(9) 390	(9) 423	(9) 483	(9) 519	(36) 1.815

$$SQD_{tot} = 1.489,67 \text{ (fornecido)}$$

$$C = (1.815)^2 / 36 = 91.506,25$$

$$SQD_{tra} = 1/3 [(120)^2 + (132)^2 + \dots + (186)^2] - C = 1.454,75$$

$$SQD(A) = 1/12 [(564)^2 + (600)^2 + (651)^2] - C = 318,50$$

$$SQD(B) = 1/9 [(390)^2 + (423)^2 + (483)^2 + (519)^2] - C = 1.124,75$$

$$SQD(A \times B) = SQD_{tra} - SQD(A) - SQD(B)$$

$$SQD(A \times B) = 1.454,75 - 318,50 - 1.124,75$$

$$SQD(A \times B) = 11,50$$

$$SQ_{Res} = SQ_{tot} - SQ_{tra}$$

$$SQ_{Res} = 1.498,67 - 1.454,75$$

$$SQ_{Res} = 43,92$$

Hipóteses:

$$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \alpha_1 = \dots = \alpha_I = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \beta_1 = \dots = \beta_J = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(11)	(1.454,75)			
A (variedade)	2	318,50	159,25	87,02	< 0,0001
B (espaçamento)	3	1.124,75	374,21	204,49	< 0,0001
A x B	6	11,50	1,92	1,05	0,4193
Resíduo	24	43,92	1,83		
Total	35	1.489,67			

Conclusões:

Não existe interação entre variedade e espaçamento. Isto significa que o comportamento de um fator não depende, ou não é influenciado, pelos níveis do outro fator, sendo portanto, independentes. Neste caso os fatores podem ser estudados isoladamente.

Existe pelo menos um contraste entre médias de variedades, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

Existe pelo menos um contraste entre médias de espaçamentos, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

Observações:

Devemos ser cautelosos em relação à primeira conclusão.

Quando o número de graus de liberdade associados a uma fonte de variação em teste pelo teste F, é elevado, pode ocorrer o efeito de diluição.

Para maior segurança nesta afirmativa, é recomendado o estudo da interação, como será visto em outros exemplos.

Assumindo que realmente não existe interação, para comparar as médias dos efeitos principais podemos desdobrar os graus de liberdade associados a cada um dos fatores em contrastes ortogonais, ou aplicar um dos testes de comparação de médias múltiplas.

Teste de Tukey aplicado nos fatores A (variedade) e B (espaçamento):

Fator A:	$m(a_i) = TA_i / 12$ (observações)
$m(a_3) = 54,25$ a	(651 12)
$m(a_2) = 50,00$ b	(600 12)
$m(a_1) = 47,00$ c	(564 12)

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{1,83}{12} (1^2 + (-1)^2) = 0,31$$

$$dms = 3,53 \cdot \sqrt{\frac{1}{2} 0,31} = 1,38 \quad \therefore \quad q_{5\%}(3, 24) = 3,53$$

As médias das variedades seguidas de pelo menos uma mesma letra, não diferem entre si, ao nível de 5% de probabilidade, pelo teste de Tukey.

Fator B:	$m(b_i) = TB_i / 9$
$m(b_4) = 57,66$ a	(519 9)
$m(b_3) = 53,66$ b	(483 9)
$m(b_2) = 47,00$ c	(423 9)
$m(b_1) = 43,33$ d	(390 9)

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{1,83}{9} (1^2 + (-1)^2) = 0,41$$

$$dms = 3,90 \cdot \sqrt{\frac{1}{2} 0,41} = 1,76 \quad \therefore \quad q_{5\%}(4, 24) = 3,90$$

As médias dos espaçamentos seguidas de pelo menos uma mesma letra, não diferem entre si, ao nível de 5% de probabilidade, pelo teste de Tukey.

11.7.2. Experimento montado no DIC com interação significativaQuadro 11.4 – Produção de batatas em kg parcela⁻¹

Irrigação		Calagem							
		Com				Sem			
		1	2	3	4	1	2	3	4
Com		32,70	30,50	31,55	28,00	28,40	28,50	25,86	29,68
Sem		18,05	18,10	20,72	19,80	18,13	21,00	19,50	20,50

Quadro 11.5 - Totais de tratamentos da produção de batatas em kg parcela⁻¹

Irrigação	Calagem		Totais
	Com	Sem	
Com	(4) 122,75	(4) 112,44	(8) 235,19
Sem	(4) 76,67	(4) 79,13	(8) 155,80
Totais	(8) 199,42	(8) 191,57	(16) 390,99

$$C = (390,99)^2 / 16 = 9.554,57$$

$$SQD_{tot} = [(32,70)^2 + (30,50)^2 + \dots + (20,50)^2] - C = 437,79$$

$$SQD_{tra} = 1/4 [(122,75)^2 + \dots + (79,13)^2] - C = 407,97$$

$$SQD_{irr} = 1/8 [(235,19)^2 + (155,80)^2] - C = 393,92$$

$$SQD_{cal} = 1/8 [(199,42)^2 + (191,57)^2] - C = 3,85$$

$$SQD(irr \times cal) = SQD_{tra} - SQD_{irr} - SQD_{cal}$$

$$SQD(irr \times cal) = 407,97 - 393,92 - 3,85$$

$$SQD(irr \times cal) = 10,19$$

$$SQD_{res} = SQD_{tot} - SQD_{tra}$$

$$SQD_{res} = 437,79 - 407,97$$

$$SQD_{res} = 29,82$$

Hipóteses:

$$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \alpha_1 = \dots = \alpha_I = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \beta_1 = \dots = \beta_J = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(3)	(407,97)			
Irrigação (irr)	1	393,92	393,92	158,51	0,0001
Calagem (cal)	1	3,85	3,85	1,55	0,2369
irr x cal	1	10,19	10,19	4,10	0,0657
Resíduo	12	29,82	2,49		
Total	15	437,78			

cv =

Conclusões:

Baseados na ANOVA anterior poderia-se concluir que não existe interação entre os fatores Irrigação e Calagem ao nível de 5% de probabilidade.

Isto significaria que o comportamento de um fator não depende, ou não é influenciado, pelos níveis do outro fator, sendo portanto, independentes.

Entretanto, o aprofundamento da análise irá mostrar que a interação é significativa ao nível de 5% de probabilidade.

Estudo da interação via contrastes:

Irrigação	Calagem		Totais
	Com	Sem	
Com	(4) 122,75	(4) 112,44	(8) 235,19
Sem	(4) 76,67	(4) 79,13	(8) 155,80
Totais	(8) 199,42	(8) 191,57	(16) 390,99

$$C_1 = C_{cal} \text{ vs. } S_{cal} / C_{irr} = 1C_{cal} - 1S_{cal}$$

$$C_2 = C_{cal} \text{ vs. } S_{cal} / S_{irr} = 1C_{cal} - 1S_{cal}$$

$$C_3 = C_{irr} \text{ vs. } S_{irr} / C_{cal} = 1C_{irr} - 1S_{irr}$$

$$C_4 = C_{irr} \text{ vs. } S_{irr} / S_{cal} = 1C_{irr} - 1S_{irr}$$

$$\hat{C}_1 = 1C_{cal} - 1S_{cal} = 122,75 - 112,44 = 10,31$$

$$\hat{C}_2 = 1C_{cal} - 1S_{cal} = 76,67 - 79,13 = -2,46$$

$$\hat{C}_3 = 1C_{irr} - 1S_{irr} = 122,75 - 76,67 = 46,08$$

$$\hat{C}_4 = 1C_{irr} - 1S_{irr} = 112,44 - 79,13 = 33,31$$

$$\begin{aligned} \text{SQD}(C_1) &= (10,31)^2 / 4[(1)^2 + (-1)^2] = 13,29 \\ \text{SQD}(C_2) &= (-2,46)^2 / 4[(1)^2 + (-1)^2] = 0,76 \\ \text{SQD}(C_3) &= (46,08)^2 / 4[(1)^2 + (-1)^2] = 263,42 \\ \text{SQD}(C_4) &= (33,31)^2 / 4[(1)^2 + (-1)^2] = 138,69 \end{aligned}$$

Hipóteses:

$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$ $H_1: \text{Não } H_0$	$H_0: \alpha_1 = \dots = \alpha_I = 0$ $H_1: \text{Não } H_0$	$H_0: \beta_1 = \dots = \beta_J = 0$ $H_1: \text{Não } H_0$
--	--	--

ANOVA

Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Tratamentos	(3)	(407,97)			
Irrigação (irr)	1	393,92	393,92	158,51	0,0001
Ccal vs. Scal / Cirr	1	13,29	13,29	5,35	0,0393
Ccal vs. Scal / Sirr	1	0,76	0,76	0,30	0,5913
Calagem (cal)	1	3,85	3,85	1,55	0,2369
Cirr vs. Sirr / Ccal	1	265,42	265,42	106,80	0,0001
Cirr vs. Sirr / Scal	1	138,69	138,69	55,81	0,0001
Resíduo	12	29,82	2,49		
Total	15	437,78			

Formas como são apresentadas as análises estatísticas:

i. Sem interpretação

Causa da variação	GL	QMD	Pr
Tratamentos	(3)		
Irrigação (irr)	1	393,92	0,0001
Ccal vs. Scal / Cirr	1	13,29	0,0393
Ccal vs. Scal / Sirr	1	0,76	0,5913
Calagem (cal)	1	3,85	0,2369
Cirr vs. Sirr / Ccal	1	265,42	0,0001
Cirr vs. Sirr / Scal	1	138,69	0,0001
Resíduo	12	2,49	
Total	15		

ii. Com interpretação

Causa da variação	GL	QMD
Tratamentos	(3)	
Irrigação (irr)	1	393,92 ***
Ccal vs. Scal / Cirr	1	13,29 *
Ccal vs. Scal / Sirr	1	0,76 ns
Calagem (cal)	1	3,85 ns
Cirr vs. Sirr / Ccal	1	265,42 ***
Cirr vs. Sirr / Scal	1	138,69 ***
Resíduo	12	2,49
Total	15	

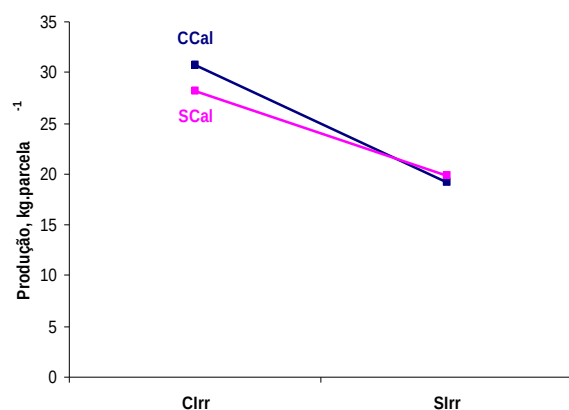
*, **, *** e ns = significativo a 5, 1 e 0,1 % de probabilidade e não significativo respectivamente pelo teste F.

ANOVA conclusiva

Causa da variação	GL	QMD	Pr
Tratamentos	(3)		
Irrigação (irr)	1	393,92	0,0001
Ccal vs. Scal / Cirr	1	13,29	0,0393
Ccal vs. Scal / Sirr	1	0,76	0,5913
Calagem (cal)	1	3,85	0,2369
Cirr vs. Sirr / Ccal	1	265,42	0,0001
Cirr vs. Sirr / Scal	1	138,69	0,0001
Resíduo	12	2,49	
Total	15		

Quadro 11.6 - Médias da produção de batatas em kg parcela⁻¹

Irrigação	Calagem	
	Com	Sem
Com	30,69	28,11
Sem	19,17	19,78

Figura 11.1 – Médias da produção de batatas em kg parcela⁻¹.

Observações:

Deve ser sempre considerado que os resultados de um experimento são válidos apenas para as condições em que foi realizado o experimento.

Extrapolações somente podem ser realizadas, cercadas dos devidos cuidados, apenas para condições muito similares as predominantes durante a condução do experimento.

11.7.3. Experimento montado no DBC com interação significativa

Em um experimento fatorial 3 x 4 no DBC com 3 repetições, são dados:

A	B	b ₁	b ₂	b ₃	b ₄	Totais de A
a ₁		(3) 69,40	(3) 74,50	(3) 78,40	(3) 82,60	(12) 304,90
a ₂		(3) 74,50	(3) 79,40	(3) 84,80	(3) 71,50	(12) 310,20
a ₃		(3) 64,50	(3) 63,50	(3) 65,20	(3) 62,80	(12) 256,00
Totais de B		(9) 208,40	(9) 217,40	(9) 228,40	(9) 216,90	(36) 871,10

$$SQD_{res} = 24,64 \text{ (fornecido)}$$

$$C = (871,10)^2 / 36 = 21.078,20$$

$$SQD_{tra} = 1/3 [(69,40)^2 + (74,50)^2 + \dots + (62,80)^2] - C = 215,54$$

$$SQD(A) = 1/12 [(304,90)^2 + (310,20)^2 + (256,00)^2] - C = 148,80$$

$$SQD(B) = 1/9 [(208,40)^2 + (217,40)^2 + \dots + (216,90)^2] - C = 22,41$$

$$SQD(A \times B) = SQD_{tra} - SQD(A) - SQD(B)$$

$$SQD(A \times B) = 215,54 - 148,80 - 22,41$$

$$SQD(A \times B) = 44,32$$

Hipóteses:

$$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \alpha_1 = \dots = \alpha_I = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \beta_1 = \dots = \beta_J = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Blocos	2				
Tratamentos	(11)	(215,54)			
A	2	148,80	74,40	66,43	< 0,0001
B	3	22,41	7,47	6,67	0,0023
A x B	6	44,32	7,39	6,59	0,0004
Resíduo	22	24,64	1,12		
Total	35				

Conclusões:

Existe interação entre os fatores A e B ao nível de 5% de probabilidade. Isto significa que o comportamento de um fator depende, ou é influenciado, pelos níveis do outro fator, sendo portanto, dependentes.

Neste caso, não estudamos os fatores isoladamente e sim, modificamos a análise anterior desdobrando a interação e avaliando o comportamento de um fator em cada nível do outro fator.

i. Estudo do fator A dentro dos níveis do fator B:

B A	b ₁	b ₂	b ₃	b ₄	Totais de A
a ₁	(3) 69,40	(3) 74,50	(3) 78,40	(3) 82,60	(12) 304,90
a ₂	(3) 74,50	(3) 79,40	(3) 84,80	(3) 71,50	(12) 310,20
a ₃	(3) 64,50	(3) 63,50	(3) 65,20	(3) 62,80	(12) 256,00
Totais de B	(9) 208,40	(9) 217,40	(9) 228,40	(9) 216,90	(36) 871,10

$$SQD(A / b_1) = 1/3 [(69,40)^2 + \dots + (64,50)^2] - [(208,40)^2 / 9] = 16,67$$

$$SQD(A / b_2) = 1/3 [(74,50)^2 + \dots + (63,50)^2] - [(217,40)^2 / 9] = 44,20$$

$$SQD(A / b_3) = 1/3 [(78,40)^2 + \dots + (65,20)^2] - [(228,40)^2 / 9] = 66,60$$

$$SQD(A / b_4) = 1/3 [(82,60)^2 + \dots + (62,80)^2] - [(216,90)^2 / 9] = 65,67$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(11)	(215,54)			
Fator B	3	22,41			
A/b ₁	2	16,67	8,33	7,44	0,0034
A/b ₂	2	44,20	22,10	19,73	< 0,0001
A/b ₃	2	66,60	33,30	29,73	< 0,0001
A/b ₄	2	65,67	32,83	29,31	< 0,0001
Resíduo	22	24,64	1,12		

Conclusão:

Dentro de cada nível de B, existe pelo menos um contraste entre médias dos níveis do fator A, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

ii. Estudo dos níveis de A dentro de cada nível de B via contrastes ortogonais:

Os contrastes de interesse são:

a_2 vs. (a_1 , a_3)

a_1 vs. a_3

A/b₁

$$C_1 = 2a_2 - a_1 - a_3$$

$$C_2 = a_1 - a_3$$

$$\widehat{C}_1 = 2(74,50) - 69,40 - 64,50 = 15,10$$

$$\widehat{C}_2 = 69,40 - 64,50 = 4,90$$

$$SQD(C_1) = (15,10)^2 / 3[(2)^2 + (-1)^2 + (-1)^2] = 12,67$$

$$SQD(C_2) = (4,90)^2 / 3[(1)^2 + (-1)^2] = 4,00$$

A/b₂

$$\widehat{C}_1 = 2(79,40) - 74,50 - 63,50 = 20,80$$

$$\widehat{C}_2 = 74,50 - 63,50 = 11,00$$

$$SQD(C_1) = (20,80)^2 / 3[(2)^2 + (-1)^2 + (-1)^2] = 24,04$$

$$SQD(C_2) = (11,00)^2 / 3[(1)^2 + (-1)^2] = 20,17$$

A/b₃

$$\widehat{C}_1 = 2(84,80) - 78,40 - 65,20 = 26,00$$

$$\widehat{C}_2 = 78,40 - 65,20 = 13,20$$

$$SQD(C_1) = (26,00)^2 / 3[(2)^2 + (-1)^2 + (-1)^2] = 37,56$$

$$SQD(C_2) = (13,20)^2 / 3[(1)^2 + (-1)^2] = 29,04$$

A/b₄

$$\widehat{C}_1 = 2(71,50) - 82,60 - 62,80 = -2,40$$

$$\widehat{C}_2 = 82,60 - 62,80 = 19,80$$

$$SQD(C_1) = (-2,40)^2 / 3[(2)^2 + (-1)^2 + (-1)^2] = 0,32$$

$$SQD(C_2) = (19,80)^2 / 3[(1)^2 + (-1)^2] = 65,34$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA conclusiva

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(11)	(215,54)			
Fator B	3	22,41			
A/b ₁	(2)	(16,67)	8,33	7,44	0,0034
a ₂ vs. (a ₁ , a ₃)	1	12,67	12,67	11,31	0,0028
a ₁ vs. a ₃	1	4,00	4,00	3,57	0,0720
A/b ₂	(2)	(44,20)	22,10	19,73	< 0,0001
a ₂ vs. (a ₁ , a ₃)	1	24,04	24,04	21,46	0,0001
a ₁ vs. a ₃	1	20,17	20,17	18,01	0,0003
A/b ₃	(2)	(66,60)	33,30	29,73	< 0,0001
a ₂ vs. (a ₁ , a ₃)	1	37,56	37,56	33,54	< 0,0001
a ₁ vs. a ₃	1	29,04	29,04	25,93	< 0,0001
A/b ₄	(2)	(65,67)	32,83	29,31	< 0,0001
a ₂ vs. (a ₁ , a ₃)	1	0,32	0,32	0,29	0,5983
a ₁ vs. a ₃	1	65,34	65,34	58,34	< 0,0001
Resíduo	22	24,64	1,12		

Quadro 11.7 – Médias de tratamentos

A	B			
	b ₁	b ₂	b ₃	b ₄
a ₁	23,13	24,83	26,13	27,53
a ₂	24,83	26,47	28,27	23,83
a ₃	21,50	21,17	21,73	20,93

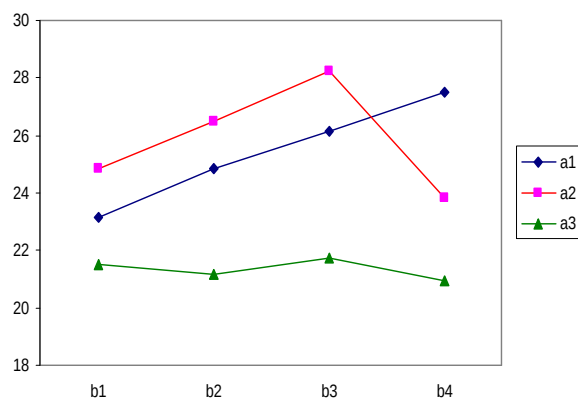


Figura 11.2 – Médias de tratamentos.

iii. Estudo do fator B dentro dos níveis do fator A:

A	B				Totais de A
	b ₁	b ₂	b ₃	b ₄	
a ₁	(3) 69,40	(3) 74,50	(3) 78,40	(3) 82,60	(12) 304,90
a ₂	(3) 74,50	(3) 79,40	(3) 84,80	(3) 71,50	(12) 310,20
a ₃	(3) 64,50	(3) 63,50	(3) 65,20	(3) 62,80	(12) 256,00
Totais de B	(9) 208,40	(9) 217,40	(9) 228,40	(9) 216,90	(36) 871,10

$$SQD(B / a_1) = 1/3 [(69,40)^2 + \dots + (82,60)^2] - [(304,90)^2 / 12] = 31,64$$

$$SQD(B / a_2) = 1/3 [(74,50)^2 + \dots + (71,50)^2] - [(310,20)^2 / 12] = 33,96$$

$$SQD(B / a_3) = 1/3 [(64,50)^2 + \dots + (62,80)^2] - [(256,00)^2 / 12] = 1,13$$

Hipóteses:

$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$	$H_0: \alpha_1 = \dots = \alpha_I = 0$	$H_0: \beta_1 = \dots = \beta_J = 0$
$H_1: \text{Não } H_0$	$H_1: \text{Não } H_0$	$H_1: \text{Não } H_0$

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(11)	(215,54)			
Fator A	2	148,80			
B/a ₁	3	31,64	10,55	9,42	0,0003
B/a ₂	3	33,96	11,32	10,11	0,0002
B/a ₃	3	1,13	0,38	0,33	0,8037
Resíduo	22	24,64	1,12		

Conclusões:

Dentro de cada nível de a_1 e a_2 , existe pelo menos um contraste entre médias dos níveis do fator B, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

Todos os contrastes entre médias dos níveis de B dentro de a_3 , são estatisticamente nulos, ao nível de 5% de significância.

iv. Estudo dos níveis de B dentro de cada nível de A via contrastes ortogonais:

Os contrastes de interesse são:

b_1 vs. (b_2, b_3, b_4)

b_2 vs. (b_3, b_4)

b_3 vs. b_4

B/a₁

$$C_1 = 3b_1 - b_2 - b_3 - b_4$$

$$C_2 = 2b_2 - b_3 - b_4$$

$$C_3 = b_3 - b_4$$

$$\widehat{C}_1 = 3(69,40) - 74,50 - 78,40 - 82,60 = -27,30$$

$$\widehat{C}_2 = 2(74,50) - 78,40 - 82,60 = -12,00$$

$$\widehat{C}_3 = 78,40 - 82,60 = -4,20$$

$$SQD(C_1) = (-27,30)^2 / 3[(3)^2 + (-1)^2 + (-1)^2 + (-1)^2] = 20,70$$

$$SQD(C_2) = (-12,00)^2 / 3[(2)^2 + (1)^2 + (-1)^2] = 8,00$$

$$SQD(C_3) = (-4,20)^2 / 3[(1)^2 + (-1)^2] = 2,94$$

B/a₂

$$\widehat{C}_1 = 3(74,50) - 79,40 - 84,80 - 71,50 = -12,20$$

$$\widehat{C}_2 = 2(79,40) - 84,80 - 71,50 = 2,50$$

$$\widehat{C}_3 = 84,80 - 71,50 = 13,30$$

$$SQD(C_1) = (-12,20)^2 / 3[(3)^2 + (-1)^2 + (-1)^2 + (-1)^2] = 4,13$$

$$SQD(C_2) = (2,50)^2 / 3[(2)^2 + (1)^2 + (-1)^2] = 0,35$$

$$SQD(C_3) = (13,30)^2 / 3[(1)^2 + (-1)^2] = 29,48$$

B/a₃

$$\widehat{C}_1 = 3(64,50) - 63,50 - 65,20 - 62,80 = 2,00$$

$$\widehat{C}_2 = 2(63,50) - 65,20 - 62,80 = -1,00$$

$$\widehat{C}_3 = 65,20 - 62,80 = 2,40$$

$$SQD(C_1) = (2,00)^2 / 3[(3)^2 + (-1)^2 + (-1)^2] = 0,11$$

$$SQD(C_2) = (-1,00)^2 / 3[(2)^2 + (1)^2 + (-1)^2] = 0,06$$

$$SQD(C_3) = (2,40)^2 / 3[(1)^2 + (-1)^2] = 0,96$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA conclusiva

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(11)	(215,54)			
Fator A	2	148,80			
B/a ₁	3	31,64	10,55	9,42	0,0003
b ₁ vs. (b ₂ , b ₃ , b ₄)	1	20,70	20,70	18,48	0,0003
b ₂ vs. (b ₃ , b ₄)	1	8,00	8,00	7,14	0,0139
b ₃ vs. b ₄	1	2,94	2,94	2,63	0,1194
B/a ₂	3	33,96	11,32	10,11	0,0002
b ₁ vs. (b ₂ , b ₃ , b ₄)	1	4,13	4,13	3,69	0,0679
b ₂ vs. (b ₃ , b ₄)	1	0,35	0,35	0,31	0,5818
b ₃ vs. b ₄	1	29,48	29,48	26,32	<0,0001
B/a ₃	3	1,13	0,38	0,33	0,8037
b ₁ vs. (b ₂ , b ₃ , b ₄)	1	0,11	0,11	0,10	0,7569
b ₂ vs. (b ₃ , b ₄)	1	0,06	0,06	0,60	0,8191
b ₃ vs. b ₄	1	0,96	0,96	0,86	0,3646
Resíduo	22	24,64	1,12		

Quadro 11.8 – Médias de tratamentos

A	B				
		b ₁	b ₂	b ₃	b ₄
a ₁		23,13	24,83	26,13	27,53
a ₂		24,83	26,47	28,27	23,83
a ₃		21,50	21,17	21,73	20,93

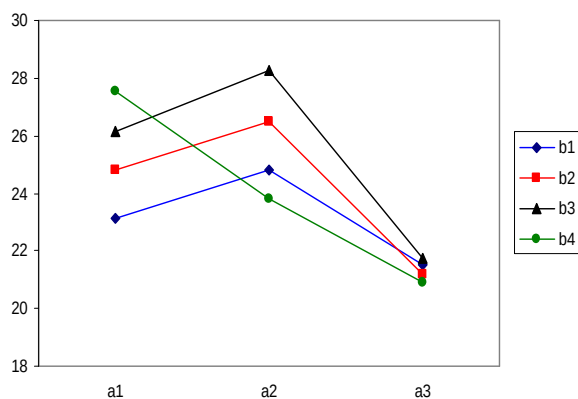


Figura 11.3 – Médias de tratamentos.

11.7.4. Experimento montado no DIC com interação significativa

Quadro 11.9 – Qualidade de mudas em função do recipiente e da espécie

		Espécie							
		e ₁				e ₂			
		1	2	3	4	1	2	3	4
Recipiente	r ₁	26,2	26,0	25,0	25,4	24,8	24,6	26,7	25,2
	r ₂	25,7	26,3	25,1	26,4	19,6	21,1	19,0	18,6
	r ₃	22,8	19,4	18,8	19,2	19,8	21,4	22,8	21,3

Quadro 11.10 - Totais de tratamentos

	e ₁	e ₂	Totais (r)
r ₁	(4) 102,60	(4) 101,30	(8) 203,90
r ₂	(4) 103,50	(4) 78,30	(8) 181,80
r ₃	(4) 80,20	(4) 85,30	(8) 165,50
Totais (e)	(12) 286,30	(12) 264,90	(24) 551,20

$$C = (551,20)^2 / 24 = 12.659,23$$

$$SQD_{tot} = [(26,2)^2 + \dots + (21,3)^2] - C = 198,79$$

$$SQD_{tra} = 1/4 [(102,60)^2 + \dots + (85,30)^2] - C = 175,70$$

$$SQD_{rec} = 1/8 [(203,90)^2 + \dots + (165,50)^2] - C = 92,86$$

$$SQD_{esp} = 1/12 [(286,30)^2 + (264,90)^2] - C = 19,08$$

$$SQDtra = SQDr + SQDe + SQD(r \times e)$$

$$SQD(r \times e) = SQDtra - SQDr - SQDe$$

$$SQD(r \times e) = 175,70 - 92,86 - 19,08$$

$$SQD(r \times e) = 63,76$$

$$SQDtot = SQDtra + SQDres$$

$$SQDres = SQDtot - SQDtra$$

$$SQDres = 198,79 - 175,70$$

$$SQDres = 19,08$$

Hipóteses:

$$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \alpha_1 = \dots = \alpha_I = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \beta_1 = \dots = \beta_J = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(5)	(175,70)			
Recipiente (r)	2	92,86	46,43	36,20	< 0,0001
Espécie (e)	1	19,08	19,08	14,88	0,0012
r x e	2	63,76	31,88	24,85	< 0,0001
Resíduo	18	23,09	1,28		
Total	23	198,79			

cv =

Conclusões:

Existe interação entre recipiente e espécie. Isto significa que o comportamento de um fator depende, ou é influenciado, pelos níveis do outro fator, sendo portanto, dependentes. Neste caso os fatores não podem ser estudados isoladamente.

Existe pelo menos um contraste entre médias de recipientes, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

Existe pelo menos um contraste entre médias de espécies, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

i. Estudo dos níveis de espécie nos níveis de recipiente:

	e_1	e_2	Totais (r)
r_1	(4) 102,60	(4) 101,30	(8) 203,90
r_2	(4) 103,50	(4) 78,30	(8) 181,80
r_3	(4) 80,20	(4) 85,30	(8) 165,50
Totais (e)	(12) 286,30	(12) 264,90	(24) 551,20

$$SQD(e / r_1) = 1/4 [(102,60)^2 + (101,30)^2] - [(203,90)^2/8] = 0,21$$

$$SQD(e / r_2) = 1/4 [(103,50)^2 + (78,30)^2] - [(181,80)^2/8] = 79,38$$

$$SQD(e / r_3) = 1/4 [(80,20)^2 + (85,30)^2] - [(165,50)^2/8] = 3,25$$

Obs: os mesmos resultados calculados via contrastes:

$$C_1 = e_1 \text{ vs. } e_2 / r_1$$

$$C_2 = e_1 \text{ vs. } e_2 / r_2$$

$$C_3 = e_1 \text{ vs. } e_2 / r_3$$

$$\hat{C}_1 = e_1 - e_2 = 102,60 - 101,3 = 1,30$$

$$\hat{C}_2 = e_1 - e_2 = 103,50 - 78,30 = 25,20$$

$$\hat{C}_3 = e_1 - e_2 = 80,20 - 85,30 = -5,10$$

$$SQD(C_1) = (1,30)^2 / 4[(1)^2 + (-1)^2] = 0,21$$

$$SQD(C_2) = (25,20)^2 / 4[(1)^2 + (-1)^2] = 79,38$$

$$SQD(C_3) = (-5,10)^2 / 4[(1)^2 + (-1)^2] = 3,25$$

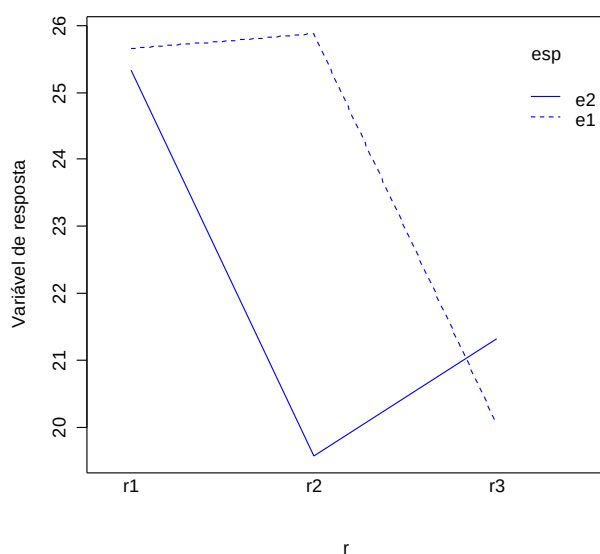
Hipóteses:

$$H_0: |C_i| = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(5)	(175,70)			
Recipiente (r)	2	91,86			
e ₁ vs e ₂ / r ₁	1	0,21	0,21	0,16	0,6897
e ₁ vs e ₂ / r ₂	1	79,38	79,38	61,88	< 0,0001
e ₁ vs e ₂ / r ₃	2	3,25	3,25	2,53	< 0,1288
Resíduo	18	23,09	1,28		
Total	23	198,79			



	e ₁	e ₂
r ₁	25,65	25,33
r ₂	25,88	19,58
r ₃	20,50	21,33

Figura 11.4 – Médias de tratamentos.

ii. Estudo dos níveis de recipiente nos níveis de espécie:

	e ₁	e ₂	Totais (r)
r ₁	(4) 102,60	(4) 101,30	(8) 203,90
r ₂	(4) 103,50	(4) 78,30	(8) 181,80
r ₃	(4) 80,20	(4) 85,30	(8) 165,50
Totais (e)	(12) 286,30	(12) 264,90	(24) 551,20

$$SQD(r / e_1) = 1/4 [(102,60)^2 + \dots + (80,20)^2] - [(286,30)^2/12] = 87,12$$

$$SQD(r / e_2) = 1/4 [(101,30)^2 + \dots + (85,30)^2] - [(264,90)^2/12] = 69,50$$

Hipóteses:

$$H_0: |C_i| = |C_j| = 0 \quad (\text{para } i \neq j)$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(5)	(175,70)			
Espécie (e)	1	19,08			
r / e ₁	2	87,12	43,56	33,96	< 0,0001
r / e ₂	2	69,50	34,75	27,09	< 0,0001
Resíduo	24	23,09	1,28		
Total	23	198,79			

iii. Estudo da interação via contrastes:

	e ₁	e ₂	Totais (r)
r ₁	(4) 102,60	(4) 101,30	(8) 203,90
r ₂	(4) 103,50	(4) 78,30	(8) 181,80
r ₃	(4) 80,20	(4) 85,30	(8) 165,50
Totais (e)	(12) 286,30	(12) 264,90	(24) 551,20

a. Estudo dos níveis de recipiente no nível e₁ de espécie:

$$C_1 = r_1 \text{ vs. } (r_2, r_3)$$

$$C_2 = r_2 \text{ vs. } r_3$$

$$\hat{C}_1 = 2r_1 - r_2 - r_3 = 2(102,60) - 103,50 - 80,20 = 21,50$$

$$\hat{C}_2 = r_2 - r_3 = 103,50 - 80,20 = 23,30$$

$$SQD(C_1) = (21,50)^2 / 4[(2)^2 + (1)^2 + (-1)^2] = 19,26$$

$$SQD(C_2) = (23,30)^2 / 4[(1)^2 + (-1)^2] = 67,86$$

b. Estudo dos níveis de recipiente no nível e₂ de espécie:

$$C_1 = r_1 \text{ vs. } (r_2, r_3)$$

$$C_2 = r_2 \text{ vs. } r_3$$

$$\hat{C}_1 = 2r_1 - r_2 - r_3 = 2(101,30) - 78,30 - 85,30 = 39,00$$

$$\hat{C}_2 = r_2 - r_3 = 76,67 - 79,13 = -7,00$$

$$SQD(C_1) = (39,00)^2 / 4[(2)^2 + (1)^2 + (-1)^2] = 63,38$$

$$SQD(C_2) = (-7,00)^2 / 4[(1)^2 + (-1)^2] = 6,13$$

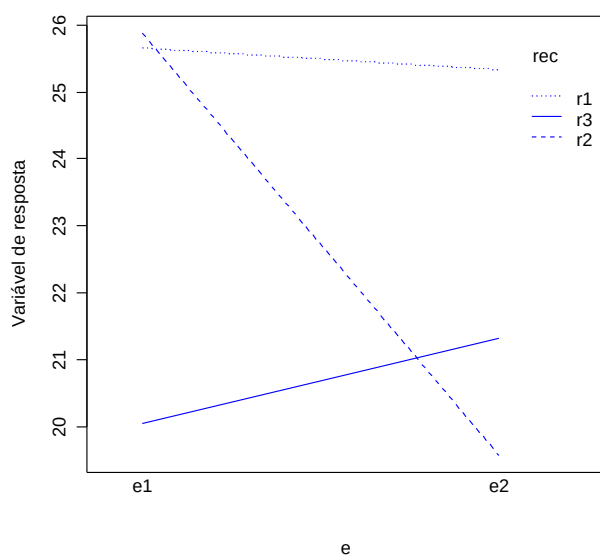
Hipóteses:

$$H_0: |C_i| = |C_j| = 0 \quad (\text{para } i \neq j)$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Tratamentos	(5)	(175,70)			
Espécie (e)	1	19,08			
r / e ₁	(2)	(87,12)			
r ₁ vs (r ₂ , r ₃)	1	19,26	19,26	15,01	0,0011
r ₂ vs r ₃	1	67,86	67,86	52,89	< 0,0001
r / e ₂	(2)	(69,50)			
r ₁ vs (r ₂ , r ₃)	1	63,38	63,38	49,40	< 0,0001
r ₂ vs r ₃	1	6,12	6,12	4,77	0,0424
Resíduo	18	23,09	1,28		
Total	23	198,79			



	e ₁	e ₂
r ₁	25,65	25,33
r ₂	25,88	19,58
r ₃	20,50	21,33

Figura 11.5 – Médias de tratamentos.

12. Experimentos em parcelas subdivididas

12.1. Introdução

Alguns autores consideram que os experimentos em parcelas subdivididas (“split plot”) não constituem um delineamento, mas um esquema de análise. Assim, podem ser utilizados em qualquer um dos delineamentos como: DIC, DBC, DQL, entre outros. Entretanto, é comum encontrar autores que os consideram como delineamentos.

Esses experimentos se caracterizam pela sua estruturação através de tratamentos principais ou primários nas parcelas, e estas, por sua vez, são constituídas de tratamentos secundários, que são as subparcelas.

Pode-se distinguir dois tipos, em conformidade com a estruturação das subparcelas:

Subdivididas no espaço Subdivididas no tempo

As parcelas poderão estar dispostas em qualquer tipo de delineamento. Os mais usuais, entretanto, são o inteiramente casualizado ou em blocos casualizados.

Tem-se dois resíduos distintos: o resíduo (a) referente às parcelas e o resíduo (b), correspondente às subparcelas dentro das parcelas. Em decorrência disso existem dois tipos de tratamentos em comparação: os principais e os secundários.

12.2. Fatorial vs. parcela subdividida

Deve ser feito um experimento em parcelas subdivididas toda vez que:

A parcela é uma unidade física, ou seja, um vaso, um animal, uma pessoa que pode receber vários tratamentos secundários.

O tratamento principal exige grandes parcelas, como é o caso da irrigação e de alguns processos industriais.

O pesquisador quer comparar tratamentos secundários com maior precisão.

Os experimentos em parcelas subdivididas são freqüentemente usados para tratamentos fatoriais, onde a natureza do material experimental, ou as operações envolvidas, tornam difícil o manuseio de todas as combinações dos fatores de uma mesma maneira.

O erro experimental das parcelas é geralmente maior que o erro experimental das subparcelas. Ou seja, em geral, o erro da subparcela é menor que aquele que seria observado se todas as combinações de tratamentos fossem arranjadas aleatoriamente dentro do delineamento escolhido, como no fatorial normal.

É importante, então, alocar os fatores de forma a obter maior precisão na comparação das interações e efeitos médios dos tratamentos de maior interesse, alocando-os nas subparcelas, uma vez que a sensibilidade em detectar diferenças significativas, caso elas existam, é maior nos tratamentos alocados nas subparcelas que nas parcelas.

12.3. Classificação dos efeitos

12.3.1. Efeito principal

É o efeito de cada fator independentemente da influência dos outros fatores.

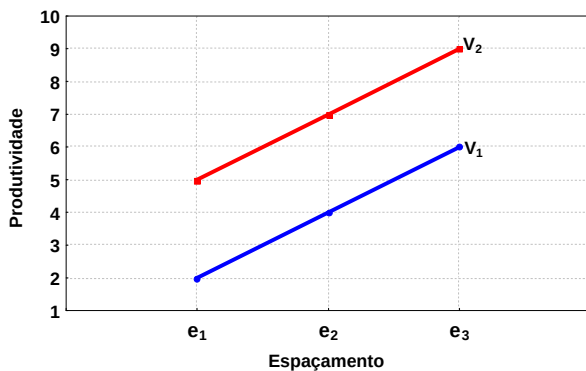
12.3.2. Efeito da interação

É a resposta diferencial da combinação de tratamentos que não se deve a efeitos principais. Ocorre interação quando os efeitos dos níveis de um fator são modificados por níveis do outro fator.

Assim temos:

Caso A

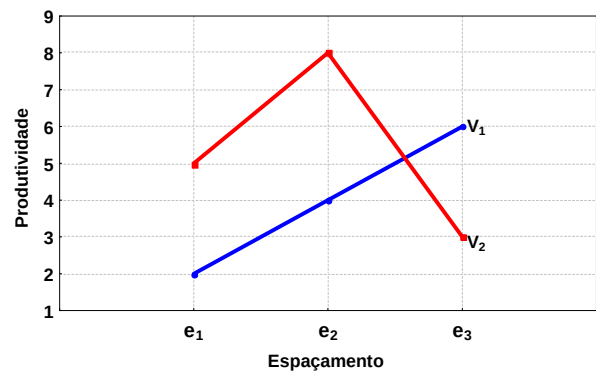
	E	e ₁	e ₂	e ₃
V				
v ₁		2	4	6
v ₂		5	7	9



Não há interação

Caso B

	E	e ₁	e ₂	e ₃
V				
v ₁		2	4	6
v ₂		5	8	3



Há interação

12.4. Vantagens e desvantagens

12.4.1. Vantagens

As grandes vantagens dos experimentos em parcelas subdivididas é a possibilidade do estudo das interações e sua grande versatilidade, uma vez que pode ser utilizado em vários delineamentos experimentais.

Em relação aos experimentos fatoriais pode, em determinadas situações, apresentar maiores facilidades operacionais.

12.4.2. Desvantagens

A principal desvantagem, similarmente ao experimentos fatoriais, é o rápido crescimento das unidades experimentais com o aumento dos fatores e seu níveis, podendo, contudo, ser contornado por técnicas de confundimento e o uso de matrizes experimentais.

Outra desvantagem é a diferença de sensibilidade do teste F entre o fator que está alocado na parcela e o fator alocado na subparcela.

Adicionalmente, a análise estatística é mais difícil que nos blocos casualizados ou nos quadrados latinos e que algumas comparações entre combinações de tratamentos se tornam relativamente complicadas.

12.5. Modelo estatístico

$$y_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + e_{ijk}$$

$$i = 1, \dots, I$$

$$j = 1, \dots, J$$

$$k = 1, \dots, K$$

onde,

y_{ijk} = observação relativa ao i-ésimo nível do fator A e ao j-ésimo nível do fator B na repetição k

μ = média geral

α_i = efeito do i-ésimo nível do fator A, definido por: $\alpha_i = \mu_{i.} - \mu$

β_j = efeito do j-ésimo nível do fator B, definido por: $\beta_j = \mu_{.j} - \mu$

$\alpha\beta_{ij}$ = efeito da interação entre o i-ésimo nível do fator A e o j-ésimo nível do fator B, definido por: $\alpha\beta_{ij} = \mu_{ij.} - (\mu + \alpha_i + \beta_j)$

e_{ijk} = erro aleatório associado à observação y_{ijk}

12.6. Coleta de dados

Quadro 12.1 - Coleta de dados de experimentos em parcelas subdivididas

a_1			a_2			...			a_i		
b_1	...	b_j	b_1	...	b_j	b_1	...	b_j	b_1	...	b_j
y_{111}	...	y_{1j1}	y_{211}	...	y_{2j1}		...		y_{i11}	...	y_{ij1}
.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.
y_{11k}	...	y_{1jk}	y_{21k}	...	y_{2jk}		...		y_{i1k}	...	y_{ijk}

Para a análise de variância manual, dependendo do delineamento adotado, é necessário a elaboração de quadros auxiliares.

12.7. Análise de variância

12.7.1. Teste de hipóteses

$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$ $H_1: \text{Não } H_0$	$H_0: \alpha_1 = \dots = \alpha_I = 0$ $H_1: \text{Não } H_0$	$H_0: \beta_1 = \dots = \beta_J = 0$ $H_1: \text{Não } H_0$
--	--	--

Quadro 12.2 – Quadro da análise de variância de experimentos em parcelas subdivididas no DIC

Causa da variação	GL	SQD	QMD	F_{cal}
Fator na parcela (A)	$I - 1$	SQD(A)	QMD(A)	QMD(A)/QMDres(a)
Resíduo (a)	$I(k-1)$	SQDres(a)	QMDres(a)	
Parcelas	$(IK - 1)$	SQDpar		
Fator na subparcela (B)	$(J-1)$	SQD(B)	QMD(B)	QMD(B)/QMDres(b)
A x B	$(I - 1)(J - 1)$	SQD(AxB)	QMD(AxB)	QMD(AxB)/QMDres(b)
Resíduo (b)	$IJ(k - 1)$	SQDres(b)	QMDres(b)	
Total	$IJK - 1$	SQDtot		

Quadro 12.3 – Quadro da análise de variância de experimentos em parcelas subdivididas no DBC

Causa da variação	GL	SQD	QMD	F _{cal}
Blocos	k - 1	SQDblo		
Fator na parcela (A)	I - 1	SQD(A)	QMD(A)	QMD(A)/QMDres(a)
Resíduo (a)	(I - 1)(k - 1)	SQDres(a)	QMDres(a)	
Parcelas	(Ik-1)	SQDpar		
Fator na subparcela (B)	J - 1	SQD(B)	QMD(B)	QMD(B)/QMDres(b)
A x B	(I - 1)(J - 1)	SQD(AxB)	QMD(AxB)	QMD(AxB)/QMDres(b)
Resíduo (b)	I(J - 1)(k - 1)	SQDres(b)	QMDres(b)	
Total	IJK - 1	SQDtot		

12.8. Exemplo: parcela subdividida no espaço

Os dados a seguir referem-se ao brix de frutos de 5 variedades de mangueira, colhidos de 3 pés por variedade. De cada pé foram colhidos 4 frutos, um de cada um dos pontos cardeais. O experimento foi montado no delineamento inteiramente casualizado.

Quadro 12.4 - Brix dos frutos

Variedade	Norte	Sul	Leste	Oeste	Totais	Totais
1	(1) 18,0	(1) 17,1	(1) 17,6	(1) 17,6	(4) 70,3	(12) 210,7
	(1) 17,5	18,8	18,1	17,2	(4) 71,6	
	(1) 17,8	16,9	17,6	16,5	(4) 68,8	
2	(1) 16,3	15,9	16,5	18,3	(4) 67,0	(12) 191,8
	(1) 16,6	14,3	16,3	17,5	(4) 64,7	
	(1) 15,0	14,0	15,9	15,2	(4) 60,1	
3	(1) 16,0	16,2	17,9	16,1	(4) 66,2	(12) 196,0
	(1) 19,5	14,9	15,0	15,3	(4) 64,7	
	(1) 16,3	16,4	16,0	16,4	(4) 65,1	
4	(1) 16,6	15,2	14,2	15,5	(4) 61,5	(12) 194,3
	(1) 15,9	13,2	18,0	17,3	(4) 64,4	
	(1) 17,5	15,8	16,7	18,4	(4) 68,4	
5	(1) 18,9	18,6	15,3	17,0	(4) 69,8	(12) 211,3
	(1) 18,5	13,7	18,2	18,3	(4) 68,7	
	(1) 21,5	16,4	18,3	16,6	(4) 72,8	
Totais	(15) 261,9	(15) 237,4	(15) 251,6	(15) 253,2	(60) 1.004,1	(60) 1.004,1

Fonte: Gomes, F.P. (1990).

$$C = (1.004,1)^2 / 60 = 16.803,61$$

$$SQD_{tot} = [(18,0)^2 + (17,1)^2 + \dots + (16,6)^2] - C = 137,58$$

$$SQD_{var} = 1 / 12 [(210,7)^2 + (191,8)^2 + \dots + (211,3)^2] - C = 29,55$$

$$SQD_{par} = 1 / 4 [(70,3)^2 + (71,6)^2 + \dots + (72,8)^2] - C = 45,26$$

$$SQD_{res(a)} = SQD_{par} - SQD_{var}$$

$$SQD_{res(a)} = 45,26 - 29,55 = 15,71$$

$$SQD_{pca} = 1 / 15 [(261,9)^2 + (237,4)^2 + \dots + (253,2)^2] - C = 20,60$$

	pca	Norte	Sul	Leste	Oeste	Médias
var						
1		(3) 53,3	(3) 52,8	(3) 53,3	(3) 51,3	(12) 17,56
2		(3) 47,9	(3) 44,2	(3) 48,7	(3) 51,0	(12) 15,98
3		(3) 51,8	(3) 47,5	(3) 48,9	(3) 47,8	(12) 16,33
4		(3) 50,0	(3) 44,2	(3) 48,9	(3) 51,2	(12) 16,19
5		(3) 58,9	(3) 48,7	(3) 51,8	(3) 51,9	(12) 17,61
Médias		(15) 17,46	(15) 15,83	(15) 16,77	(15) 16,88	

$$SQD(var, pca) = 1 / 3 [(53,3)^2 + (52,8)^2 + \dots + (51,9)^2] - C = 70,27$$

$$SQD(var, pca) = SQD(var, pca) - SQD_{var} - SQD_{pca}$$

$$SQD(var, pca) = 70,27 - 29,55 - 20,60$$

$$SQD(var, pca) = 20,12$$

$$SQD_{res(b)} = SQD_{tot} - SQD_{par} - SQD_{pca} - SQD(var, pca)$$

$$SQD_{res(b)} = 137,58 - 45,26 - 20,60 - 20,12$$

$$SQD_{res(b)} = 51,60$$

Hipóteses:

$$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \alpha_1 = \dots = \alpha_I = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \beta_1 = \dots = \beta_J = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Variedade (var)	4	29,55	7,39	4,71	0,0214
Resíduo (a)	10	15,71	1,57		
Parcelas	(14)	(45,26)			
Pontos cardeais (pca)	3	20,60	6,87	3,99	0,0167
var x pca	12	20,12	1,68	0,97	0,4970
Resíduo (b)	30	51,60	1,72		
Total	59	137,58			

Conclusões:

Não existe interação entre os fatores variedades e pontos cardeais. Isto significa que o comportamento de um fator não depende, ou não é influenciado, pelos níveis do outro fator, sendo portanto, independentes. Neste caso os fatores podem ser estudados isoladamente.

Existe pelo menos um contraste entre médias de variedades, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

Existe pelo menos um contraste entre médias de pontos cardeais, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

Observações:

Devemos ser cautelosos em relação a esta última conclusão.

Como temos discutido em sala de aula, quando o número de graus de liberdade associados a uma fonte de variação em teste pelo teste F, pode ocorrer o efeito de diluição. Para maior segurança nesta afirmativa, é recomendado o estudo da interação.

Assumindo que realmente não existe interação, para comparar as médias dos efeitos principais podemos desdobrar os graus de liberdade associados a cada um dos fatores em contrastes ortogonais, ou aplicar um dos testes de comparação de médias múltiplas.

12.8.1. Teste de Tukey aplicado aos efeitos principais

i. Teste de Tukey aplicado nas variedades

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{1,57}{12} (1^2 + (-1)^2) = 0,26$$

$$dms = 4,65 \cdot \sqrt{\frac{1}{2} 0,26} = 1,68 \quad \therefore \quad q_{5\%}(5; 10) = 4,65$$

	$m_5 = 17,61$	$m_1 = 17,56$	$m_3 = 16,33$	$m_4 = 16,19$	$m_2 = 15,98$
$m_5 = 17,61$	-	0,05 ^{ns}	1,28 ^{ns}	1,42 ^{ns}	1,63 ^{ns}
$m_1 = 17,56$		-	1,23 ^{ns}	1,37 ^{ns}	1,58 ^{ns}
$m_3 = 16,33$			-	0,14 ^{ns}	0,35 ^{ns}
$m_4 = 16,19$				-	0,21 ^{ns}
$m_2 = 15,98$					-

$m_5 = 17,61$ a
 $m_1 = 17,56$ a
 $m_3 = 16,33$ a
 $m_4 = 16,19$ a

$$m_2 = 15,98 \text{ a}$$

ii. Teste de Tukey aplicado nos pontos cardeais

$$dms = q \cdot \sqrt{\frac{1}{2} \widehat{V}(\widehat{C})} \quad \therefore \quad \widehat{V}(\widehat{C}) = QMDres \left(\frac{a_1^2}{r_1} + \dots + \frac{a_k^2}{r_k} \right) = \frac{1,72}{15} (1^2 + (-1)^2) = 0,23$$

$$dms = 3,85 \cdot \sqrt{\frac{1}{2} 0,23} = 1,30 \quad \therefore \quad q_{5\%}(4;30) = 3,85$$

	$m_N = 17,46$	$m_O = 16,88$	$m_L = 16,77$	$m_S = 15,83$
$m_N = 17,46$	-	$0,58^{ns}$	$0,69^{ns}$	$1,63^*$
$m_O = 16,88$		-	$0,11^{ns}$	$1,05^{ns}$
$m_L = 16,77$			-	$0,94^{ns}$
$m_S = 15,83$				-

m_N	=	17,61	a
m_O	=	17,56	a b
m_L	=	16,33	a b
m_S	=	16,19	b

Para os efeitos principais, as médias seguidas de pelo menos uma letra em comum não diferem entre si pelo teste de Tukey ao nível de 5% de probabilidade.

12.9. Exemplo: parcela subdividida no tempo

Os dados a seguir referem-se a contagem da colonização de um antagonista (*trichoderma* – TVC) aplicado sobre as vassouras-de-bruxa de uma cultura de cacau no município de Itabuna- BA em 2000. Na aplicação apl_1 o antagonista foi aplicado de 15 em 15 dias (0, 15, 30, 45 e 60), na apl_2 de 30 em 30 dias (0, 30 e 60) e apl_3 não recebeu aplicação do antagonista (testemunha). As avaliações foram feitas aos 15, 45 e 75 dias após o início das aplicações. O experimento foi montado no delineamento em blocos casualizados com 3 repetições.

Quadro 12.5 – Colonização do TVC em vassouras-de-bruxa, %

A p l i c a ç ã o	Blocos				Totais		
	apl ₁		blo ₁	blo ₂	blo ₃		
		15	(1) 18,75	43,75	18,75		
		45	56,25	75,00	75,00		
		75	68,75	93,75	87,50		
				(3) 143,75	(3) 212,50	(3) 181,25	(9) 537,50
	apl ₂						
		15	37,50	43,75	68,75		
		45	50,00	75,00	93,75		
75		62,50	75,00	100,00			
			(3) 150,00	(3) 193,75	(3) 262,50	(9) 606,25	
apl ₃							
	15	0	0	0			
	45	0	0	0			
	75	0	0	0			
			(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00	
Totais		(9) 293,75	(9) 406,25	(9) 443,75	(27) 1.143,75		

$$C = (1.143,75)^2 / 27 = 48.450,52$$

$$SQD_{tot} = [(18,75)^2 + (43,75)^2 + \dots + (0)^2] - C = 34.244,79$$

$$SQD_{blo} = 1 / 9 [(293,75)^2 + (406,25)^2 + (443,75)^2] - C = 1.354,17$$

$$SQD_{apl} = 1 / 9 [(537,50)^2 + (606,25)^2 + (0)^2] - C = 24.487,85$$

$$SQD_{par} = 1 / 3 [(143,75)^2 + (212,50)^2 + \dots + (0)^2] - C = 24.487,85$$

$$SQD_{par} = SQD_{apl} + SQD_{blo} + SQD_{res(a)}$$

$$SQD_{res(a)} = SQD_{par} - SQD_{apl} - SQD_{blo}$$

$$SQD_{res(a)} = 27.421,88 - 24.487,85 - 1.354,17$$

$$SQD_{res(a)} = 1.579,86$$

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$SQD_{tem} = 1 / 9 [(231,25)^2 + (425,00)^2 + (487,50)^2] - C = 3.967,01$$

$$SQD(apl, tem) = 1 / 3 [(81,25)^2 + \dots + (0)^2] - C = 31.015,63$$

$$SQD(apl, tem) = SQD_{apl} + SQD_{tem} + SQD(apl \times tem)$$

$$SQD(apl \times tem) = SQD(apl, tem) - SQD_{apl} - SQD_{tem}$$

$$SQD(apl \times tem) = 31.015,63 - 24.487,85 - 3.967,01$$

$$SQD(apl \times tem) = 2.560,77$$

$$SQD_{tot} = SQD_{par} + SQD_{tem} + SQD(apl \times tem) + SQD_{res(b)}$$

$$SQD_{res(b)} = SQD_{tot} - SQD_{par} - SQD_{tem} - SQD(apl \times tem)$$

$$SQD_{res(b)} = 34.244,79 - 27.421,88 - 3.967,01 - 2.560,77$$

$$SQD_{res(b)} = 295,13$$

Hipóteses:

$$H_0: \alpha\beta_{11} = \dots = \alpha\beta_{IJ} = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \alpha_1 = \dots = \alpha_I = 0$$

$$H_1: \text{Não } H_0$$

$$H_0: \beta_1 = \dots = \beta_J = 0$$

$$H_1: \text{Não } H_0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Bloco	2	1.354,17	677,09	1,71	0,29
Aplicação (apl)	2	24.487,85	12.243,93	31,00	0,0037
Resíduo (a)	4	1.579,86	394,97		
Parcelas	(8)	(27.421,88)			
Tempo (tem)	2	3.967,01	1.983,51	80,65	< 0,0001
apl x tem	4	2.560,77	640,19	26,03	< 0,0001
Resíduo (b)	12	295,13	24,59		
Total	26	34.244,79			

Conclusões:

Existe interação entre os fatores Aplicação e Tempo. Isto significa que o comportamento de um fator depende, ou é influenciado, pelos níveis do outro fator, sendo portanto, dependentes. Neste caso os fatores não podem ser estudados isoladamente.

Existe pelo menos um contraste entre médias de Aplicação, estatisticamente diferente de zero, ao nível de 5% de probabilidade;

Existe pelo menos um contraste entre médias de Tempo, estatisticamente diferente de zero, ao nível de 5% de probabilidade.

12.9.1. Desdobramento da interação

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(3) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(3) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

Estudo da Aplicação dentro dos níveis de Tempo:

$$SQD(apl/tem_{15}) = 1 / 3 [(81,25)^2 + (150,00)^2 + (0)^2 - (231,25)^2 / 9] = 3.758,68$$

$$SQD(apl/tem_{45}) = 1 / 3 [(206,25)^2 + (218,75)^2 + (0)^2 - (425,00)^2 / 9] = 10.060,75$$

$$SQD(apl/tem_{75}) = 1 / 3 [(250,00)^2 + (237,50)^2 + (0)^2 - (487,50)^2 / 9] = 13.229,17$$

Estudo do Tempo dentro dos níveis de Aplicação:

$$SQD(tem/apl_1) = 1 / 3 [(81,25)^2 + (206,25)^2 + (250,00)^2 - (537,50)^2 / 9] = 5.112,85$$

$$SQD(tem/apl_2) = 1 / 3 [(150,00)^2 + (218,75)^2 + (237,50)^2 - (606,25)^2 / 9] = 1.414,93$$

$$SQD(tem/apl_3) = 1 / 3 [(0,00)^2 + (0,00)^2 + (0,00)^2 - (0,00)^2 / 9] = 0,00$$

Desdobramento do efeito de apl/tem₁₅ em contrastes ortogonais:

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$C_1 = (apl_1, apl_2) \text{ vs } apl_3$$

$$C_2 = apl_1 \text{ vs } apl_2$$

$$C_1 = apl_1 + apl_2 - 2 \text{ } apl_3$$

$$C_2 = apl_1 - apl_2$$

$$\hat{C}_1 = 81,25 + 150,00 - 2(0,00) = 231,25$$

$$\hat{C}_2 = 81,25 - 150,00 = -68,75$$

$$SQD(C_1) = (231,25)^2 / 3 [(1)^2 + (1)^2 + (-2)^2] = 2.970,92$$

$$SQD(C_2) = (-68,75)^2 / 3 [(1)^2 + (-1)^2] = 787,76$$

Desdobramento do efeito de apl/tem₄₅ em contrastes ortogonais:

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$C_1 = (apl_1, apl_2) \text{ vs } apl_3$$

$$C_2 = apl_1 \text{ vs } apl_2$$

$$C_1 = apl_1 + apl_2 - 2apl_3$$

$$C_2 = apl_1 - apl_2$$

$$\hat{C}_1 = 206,25 + 218,75 - 2(0,00) = 425,00$$

$$\hat{C}_2 = 206,25 - 218,75 = -12,50$$

$$SQD(C_1) = (425,00)^2 / 3 [(1)^2 + (1)^2 + (-2)^2] = 10.034,72$$

$$SQD(C_2) = (-12,50)^2 / 3 [(1)^2 + (-1)^2] = 26,04$$

Desdobramento do efeito de apl/tem₇₅ em contrastes ortogonais:

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$C_1 = (apl_1, apl_2) \text{ vs. } apl_3$$

$$C_2 = apl_1 \text{ vs. } apl_2$$

$$C_1 = apl_1 + apl_2 - 2apl_3$$

$$C_2 = apl_1 - apl_2$$

$$\hat{C}_1 = 250,00 + 237,50 - 2(0,00) = 487,50$$

$$\hat{C}_2 = 250,00 - 237,50 = 12,50$$

$$SQD(C_1) = (487,50)^2 / 3 [(1)^2 + (1)^2 + (-2)^2] = 13.203,13$$

$$SQD(C_2) = (12,50)^2 / 3 [(1)^2 + (-1)^2] = 26,04$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Bloco	2	1.354,17	677,09	1,71	0,29
Aplicação (apl)	2	24.487,85	12.243,93	31,00	0,0037
Resíduo (a)	4	1.579,86	394,97		
Parcelas	(8)	(27.421,88)			
Tempo	2	3.967,01			
apl/tem ₁₅	2	3.758,68	1.879,34	76,43	< 0,0001
(apl ₁ , apl ₂) vs. apl ₃	1	2.970,92	2.970,92	120,82	< 0,0001
apl ₁ vs. apl ₂	1	787,76	787,76	32,04	< 0,0001
apl/tem ₄₅	2	10.060,75	5.030,38	204,57	< 0,0001
(apl ₁ , apl ₂) vs. apl ₃	1	10.034,72	10.034,72	408,08	< 0,0001
apl ₁ vs. apl ₂	1	26,04	26,04	1,06	0,32
apl/tem ₇₅	2	13.229,17	6.614,59	268,99	< 0,0001
(apl ₁ , apl ₂) vs. apl ₃	1	13.203,13	13.203,13	536,93	< 0,0001
apl ₁ vs. apl ₂	1	26,04	26,04	1,06	0,32
Resíduo (b)	12	295,13	24,59		
Total	26	34.244,79			

	tem ₁₅	tem ₄₅	tem ₇₅
apl ₁	27,08	68,75	83,33
apl ₂	50,00	72,92	79,17
apl ₃	0,00	0,00	0,00

Em todos os tempos (tem₁₅ a tem₇₅) a percentagem de colonização do TVC nas aplicações apl₁ e apl₂ foram estatisticamente superiores a apl₃ ao nível de 5% de significância pelo teste F.

Na avaliação tem₁₅ a aplicação apl₂ foi estatisticamente superior a apl₁, não tendo sido detectadas diferenças significativas para os demais tempos ao nível de 5% de significância pelo teste F.

Desdobramento do efeito de tem/apl₁ em contrastes ortogonais:

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$C_1 = (\text{tem}_{45}, \text{tem}_{75}) \text{ vs. } \text{tem}_{15}$$

$$C_2 = \text{tem}_{45} \text{ vs. } \text{tem}_{75}$$

$$C_1 = \text{tem}_{45} + \text{tem}_{75} - 2\text{tem}_{15}$$

$$C_2 = \text{tem}_{45} - \text{tem}_{75}$$

$$\hat{C}_1 = 206,25 + 250,00 - 2(81,25) = 293,75$$

$$\hat{C}_2 = 206,25 - 250,00 = -43,75$$

$$\text{SQD}(C_1) = (293,75)^2 / 3 [(1)^2 + (1)^2 + (-2)^2] = 4.793,84$$

$$\text{SQD}(C_2) = (-43,75)^2 / 3 [(1)^2 + (-1)^2] = 319,01$$

Desdobramento do efeito de tem/apl₂ em contrastes ortogonais:

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$C_1 = (\text{tem}_{45}, \text{tem}_{75}) \text{ vs. } \text{tem}_{15}$$

$$C_2 = \text{tem}_{45} \text{ vs. } \text{tem}_{75}$$

$$C_1 = \text{tem}_{45} + \text{tem}_{75} - 2\text{tem}_{15}$$

$$C_2 = \text{tem}_{45} - \text{tem}_{75}$$

$$\hat{C}_1 = 218,75 + 237,50 - 2(150,00) = 156,25$$

$$\hat{C}_2 = 218,75 - 237,50 = -18,75$$

$$SQD(C_1) = (156,25)^2 / 3 [(1)^2 + (1)^2 + (-2)^2] = 1.356,34$$

$$SQD(C_2) = (-18,75)^2 / 3 [(1)^2 + (-1)^2] = 58,59$$

Desdobramento do efeito de tem/apl₃ em contrastes ortogonais:

	tem ₁₅	tem ₄₅	tem ₇₅	Totais
apl ₁	(3) 81,25	(3) 206,25	(3) 250,00	(9) 537,50
apl ₂	(3) 150,00	(3) 218,75	(3) 237,50	(9) 606,25
apl ₃	(3) 0,00	(3) 0,00	(3) 0,00	(9) 0,00
Totais	(9) 231,25	(9) 425,00	(9) 487,50	(27) 1.143,75

$$C_1 = (\text{tem}_{45}, \text{tem}_{75}) \text{ vs. } \text{tem}_{15}$$

$$C_2 = \text{tem}_{45} \text{ vs. } \text{tem}_{75}$$

$$C_1 = \text{tem}_{45} + \text{tem}_{75} - 2\text{tem}_{15}$$

$$C_2 = \text{tem}_{45} - \text{tem}_{75}$$

$$\hat{C}_1 = 0,00 + 0,00 - 2(0,00) = 0,00$$

$$\hat{C}_2 = 0,00 - 0,00 = 0,00$$

$$SQD(C_1) = (487,50)^2 / 3 [(1)^2 + (1)^2 + (-2)^2] = 0,00$$

$$SQD(C_2) = (12,50)^2 / 3 [(1)^2 + (-1)^2] = 0,00$$

Hipóteses:

$$H_0: |C_i| = 0 \quad i = 1 \dots n$$

$$H_1: |C_i| > 0$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Bloco	2	1.354,17	677,09	1,71	0,29
Aplicação (apl)	2	24.487,85	12.243,93	31,00	0,0037
Resíduo (a)	4	1.579,86	394,97		
Parcelas	(8)	(27.421,88)			
Aplicação (apl)	2	24.487,85			
tem/apl ₁	2	5.112,85	2.556,43	103,96	< 0,0001
(tem ₄₅ , tem ₇₅) vs. tem ₁₅	1	4.793,84	4.793,84	194,95	< 0,0001
tem ₄₅ vs. tem ₇₅	1	319,01	319,01	12,97	< 0,0001
tem/apl ₂	2	1.414,83	707,42	28,77	< 0,0001
(tem ₄₅ , tem ₇₅) vs. tem ₁₅	1	1.356,34	1.356,34	55,16	< 0,0001
tem ₄₅ vs. tem ₇₅	1	58,59	58,59	2,38	0,15
tem/apl ₃	2	0,00	0,00	0,00	1,00
(tem ₄₅ , tem ₇₅) vs. tem ₁₅	1	0,00	0,00	0,00	1,00
tem ₄₅ vs. tem ₇₅	1	0,00	0,00	0,00	1,00
Resíduo (b)	12	295,13	24,59		
Total	26	34.244,79			

	tem ₁₅	tem ₄₅	tem ₇₅
apl ₁	27,08	68,75	83,33
apl ₂	50,00	72,92	79,17
apl ₃	0,00	0,00	0,00

Na aplicação apl₁ o tempo tem₁₅ é estatisticamente inferior a média de tem₄₅ e tem₇₅, e entre estas, tem₄₅ é estatisticamente inferior a tem₇₅ ao nível de 5% de significância pelo teste F.

Na aplicação apl₂ o tempo tem₁₅ é estatisticamente inferior a média de tem₄₅ e tem₇₅, e entre estas, não foi detectada diferença ao nível de 5% de significância pelo teste F.

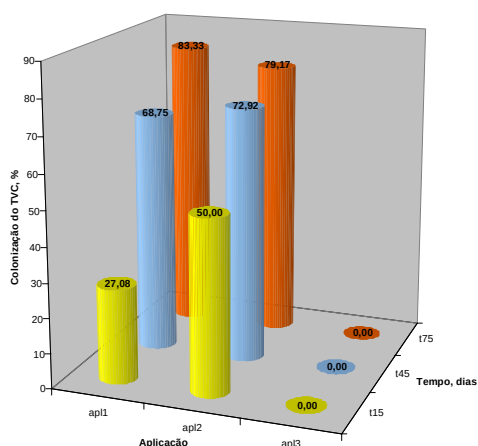


Figura 12.1 – Colonização do TVC em função da forma (apl₁: 15 x 15 dias, apl₂: 30 x 30 dias e apl₃: testemunha) de aplicação e do tempo.

13. Correlação linear simples

13.1. Introdução

A análise de correlação linear simples (Pearson, 1896) , outros tipos de análise de correlação (parcial, múltipla, canônica) e a análise de regressão, são técnicas estatísticas utilizadas no estudo quantitativo de experimentos.

Enquanto a análise de regressão linear simples nos mostra como duas variáveis se relacionam linearmente, a análise de correlação linear simples nos mostra apenas o grau da associação, ou de proporcionalidade, entre estas duas variáveis.

Conquanto a correlação seja uma técnica menos potente que a regressão, as duas se acham tão intimamente ligadas que a correlação freqüentemente é útil na interpretação da regressão.

Muitas das técnicas de análise multivariada tem na correlação a medida estatística básica para estudar a associação entre variáveis aleatórias.

13.2. Definição

ρ : Correlação populacional

r : Estimativa da correlação ou correlação amostral

$$\rho = \frac{COV_{Pop}(Y_1, Y_2)}{\sigma(Y_1) \cdot \sigma(Y_2)}$$

$$r = \frac{cov_{Amo}(Y_1, Y_2)}{s(Y_1) \cdot s(Y_2)}$$

$$COV(Y_1, Y_2) = E[(Y_1 - E(Y_1))(Y_2 - E(Y_2))]$$

$$COV_{Pop}(Y_1, Y_2) = \frac{\sum[(Y_1 - \sigma(Y_1))(Y_2 - \sigma(Y_2))]}{N}$$

$$cov_{Amo}(Y_1, Y_2) = \frac{\sum[(Y_1 - \sigma(Y_1))(Y_2 - \sigma(Y_2))]}{n}$$

$$cov_{Amo}(Y_1, Y_2) = \frac{\sum[(Y_1 - m(Y_1))(Y_2 - m(Y_2))]}{n - 1}$$

13.3. Conceitos e compreensão a partir de um exemplo

Consideremos duas variáveis aleatórias:

M : rendimento acadêmico em matemática

L : rendimento acadêmico em línguas

Quadro 13.1 - Rendimento acadêmico

Obs	01	02	03	04	05	06	07	08
M	36	80	50	58	72	60	56	68
L	35	65	60	39	48	44	48	61

$$\sum M = 480$$

$$m(M) = 60$$

$$s(M) = 13,65$$

$$\sum L = 400$$

$$m(L) = 50$$

$$s(L) = 10,93$$

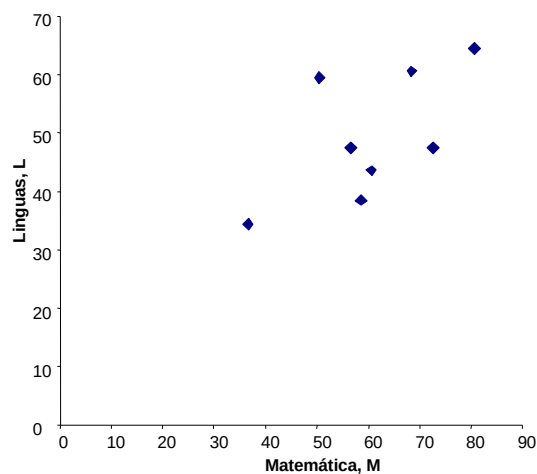
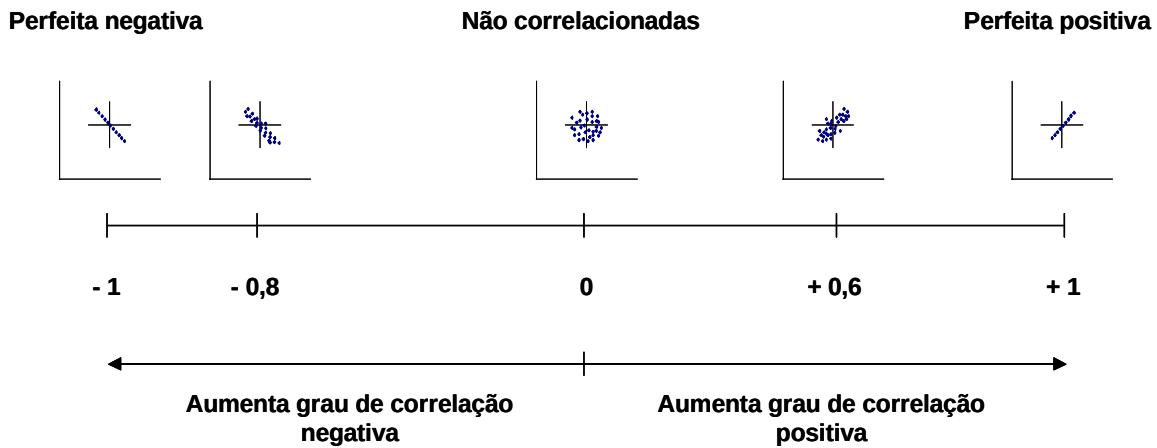


Figura 13.1 - Gráfico da dispersão entre M e L.

Necessita-se de um índice que forneça o grau de associação, ou de proporcionalidade, linear entre as duas variáveis aleatórias (M e L).



Para testar Σm_l como este índice:

$$m = m_i - m(M)$$

$$\text{cov}(Y_1, Y_2) = \frac{1}{n-1} \cdot \Sigma [(Y_1 - m(Y_1)) (Y_2 - m(Y_2))]$$

$$l = l_i - m(L)$$

deve-se sobrepôr aos pontos dispersos nos eixos cartesianos, os eixos das médias de matemática e línguas (M e L):

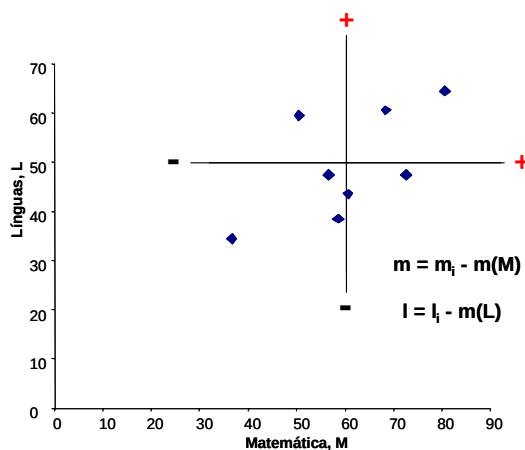
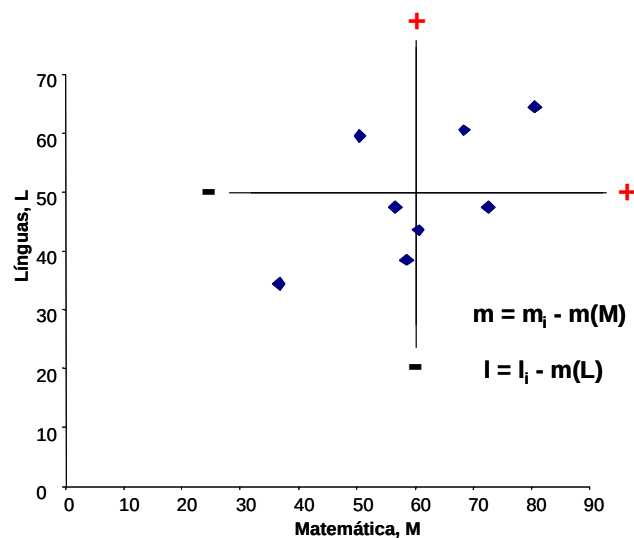


Figura 13.2 - Gráfico da dispersão entre M e L com as médias transladadas.

Quadro 13.2 – Cálculo do índice Σml

Obs	M	L	$m = (M_i - m(M))$	$l = (L_i - m(L))$	$m.l$
1	36	35	- 24	- 15	360
2	80	65	20	15	300
3	50	60	- 10	10	- 100
4	58	39	- 2	- 11	22
5	72	48	12	- 2	- 24
6	60	44	0	- 6	0
7	56	48	- 4	- 2	8
8	68	61	8	11	88
$m(M) = 60$ $m(L) = 50$ $s(M) = 13,65$ $s(L) = 10,93$					$\Sigma ml = 654$



Se M e L caminharem juntas, isto é, enquanto uma aumenta a outra também aumenta, e enquanto uma diminui a outra também diminui, a maior parte das observações recairão nos 1º e 3º quadrantes.

Conseqüentemente, a maior parte dos produtos ($m.l$) serão positivos, bem como sua soma (Σml), demonstrando um relacionamento positivo entre M e L.

Mas se M e L estão relacionadas negativamente, isto é, uma aumenta enquanto a outra diminui, a maior parte das observações recairão nos 2º e 4º quadrantes, dando um valor negativo para o índice Σml .

Concluí-se, então, que como índice do grau de associação, ou proporcionalidade, entre as duas variáveis, Σml , pelo menos, tem sinal correto.

Além disso, quando não houver relação entre M e L as observações tenderão a serem distribuídas igualmente pelos quatro quadrantes, os termos positivos e negativos se cancelarão e Σml tenderá para zero.

Há apenas duas maneiras de melhorar Σml como medida do grau de associação, ou proporcionalidade, linear entre duas variáveis aleatórias:

- i. Primeiro: Σml é dependente do tamanho da amostra:

Suponha que tivéssemos observado o mesmo diagrama de dispersão para uma amostra com o dobro do tamanho.

Então, $\sum ml$ também seria o dobro, muito embora a configuração da tendência das variáveis permaneça a mesma.

Para evitar este problema dividimos $\sum ml$ pelo tamanho da amostra:

$$\frac{\sum ml}{n-1} = \frac{1}{n-1} \left[\sum (M_i - m(M)) \cdot (L_i - m(L)) \right]$$

Ao ser eliminada a influência do tamanho da amostra, nesta medida do grau de associação, ou proporcionalidade, linear entre duas variáveis aleatórias, obtém-se uma medida bastante útil em estatística denominada covariância, neste caso representada por $\text{COV}(M, L)$:

$$\text{cov}(M, L) = \frac{\sum ml}{n-1} = \frac{\sum (M_i - m(M)) \cdot (L_i - m(L))}{n-1}$$

ii. Segundo: pode-se perceber que a covariância tem um ponto fraco: é influenciada pelas unidades de medida das variáveis envolvidas.

Suponha que o teste de matemática tenha valor 50 ao invés de 100.

Os valores relacionados aos desvios de matemática, m , serão apenas a metade, e isto irá influenciar o valor da covariância - muito embora, em essência, o grau da associação, ou proporcionalidade, linear entre matemática e línguas não tenha se modificado.

Em outras palavras, a covariância depende das unidades de medida das variáveis.

Esta dificuldade pode ser contornada se medirmos ambas as variáveis em termos de uma unidade padronizada.

Ou seja, dividindo-se m e l pelos seus respectivos desvios padrões:

$$\frac{1}{n-1} \sum \left(\frac{m}{s(M)} \right) \left(\frac{l}{s(L)} \right) = \frac{1}{n-1} \left[\sum \left(\frac{M_i - m(M)}{s(M)} \right) \cdot \left(\frac{L_i - m(L)}{s(L)} \right) \right]$$

Ao eliminar a influência do tamanho da amostra (i), obtém-se a covariância; e ao eliminar a influência das unidades de medida das variáveis (ii) define-se, finalmente, o que é denominado correlação linear simples entre M e L , $r(M, L)$, por vezes chamada de correlação de Pearson:

$$r(M, L) = \frac{\text{cov}(M, L)}{s(M) \cdot s(L)}$$

Assim, para calcularmos a correlação entre M e L:

$$\text{cov}(M, L) = \frac{\sum (M_i - m(M)) \cdot (L_i - m(L))}{n - 1} = \frac{654}{7} = 93,43$$

$$r(M, L) = \frac{\text{cov}(M, L)}{s(M) \cdot s(L)} = \frac{93,43}{13,65 \cdot 10,93} = 0,63$$

Observações:

Limites da correlação: $-1 \leq (\rho \text{ ou } r) \leq +1$

13.4. Pressuposições da correlação

O relacionamento entre as variáveis tem forma linear.

As duas variáveis são aleatórias por natureza e medidas em escalas intervalares ou proporcionais, não podendo ser categóricas ou nominais.

As variáveis apresentam distribuição normal bivariada.

Enquanto medida do grau de associação, ou proporcionalidade, entre duas variáveis aleatórias a covariância possui uma vantagem: não é influenciada pelo tamanho da amostra; e uma desvantagem: é influenciada pela unidade de medida das variáveis.

Ao dividi-la pelos respectivos desvios padrões das variáveis aleatórias obtém-se o coeficiente de correlação linear, $r(M, L)$, que não é influenciado nem pelo tamanho da amostra e nem pelas unidades de medida das variáveis.

O quadrado do coeficiente de correlação indica a proporção da variação em uma variável explicada ou predita pela variação na outra variável:

$$r = 0,63 \quad r^2 = 0,3922$$

39,22% da variação observada em M é explicada pela variação em L, e vice-versa.

Uma fórmula prática para cálculo da correlação linear simples é apresentada abaixo:

$$r(M, L) = \frac{\text{cov}(M, L)}{s(M) \cdot s(L)} = \frac{\sum (M_i - m(M)) \cdot (L_i - m(L))}{n - 1} \cdot \frac{1}{s(M) \cdot s(L)}$$

Pode-se calcular a correlação linear na ausência do conhecimento das médias das duas variáveis. A equação acima, retrabalhada, origina:

$$r(M, L) = \frac{n \cdot \sum ML - \sum M \cdot \sum L}{\sqrt{n \sum M^2 - (\sum M)^2} \cdot \sqrt{n \sum L^2 - (\sum L)^2}}$$

Que é a fórmula mais conhecida e utilizada para o cálculo do coeficiente de correlação linear simples.

Quadro 13.3 – Cálculo do coeficiente de correlação para o exemplo dado

Obs	M	L	ML
1	36	35	1.260
2	80	65	5.200
3	50	60	3.000
4	58	39	2.262
5	72	48	3.456
6	60	44	2.640
7	56	48	2.688
8	68	61	4.148
	$\sum M = 480$	$\sum L = 400$	
n=8	$\sum M^2 = 30.104$	$\sum L^2 = 20.836$	$\sum ML = 24.654$
	$(\sum M)^2 = 230.400$	$(\sum L)^2 = 160.000$	

$$r(M, L) = \frac{n \cdot \sum ML - \sum M \cdot \sum L}{\sqrt{n \sum M^2 - (\sum M)^2} \cdot \sqrt{n \sum L^2 - (\sum L)^2}}$$

$$r(M, L) = \frac{8 \cdot 24.654 - 480 \cdot 400}{\sqrt{8 \cdot 30.104 - 230.400} \cdot \sqrt{8 \cdot 20.836 - 160.000}} = 0,63$$

Considerações finais:

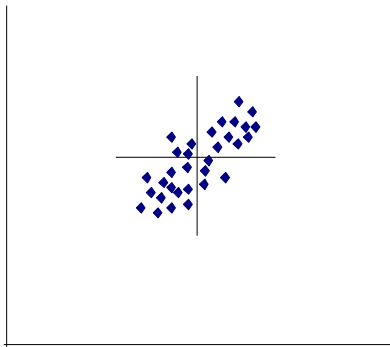
A existência de correlação entre duas variáveis aleatórias não implica em casualidade. Ou seja, não implica que a variação de uma provoca variação na outra. Para esta afirmativa é necessário variar os níveis de uma das variáveis (preditora), mantendo-se fixos todos os outros fatores que podem influenciar, e observar o que ocorre com a variável de resposta.

O montante da variação em uma variável é explicada pela variação da outra pode ser medido elevando-se o coeficiente de correlação linear, r , ao quadrado: r^2 .

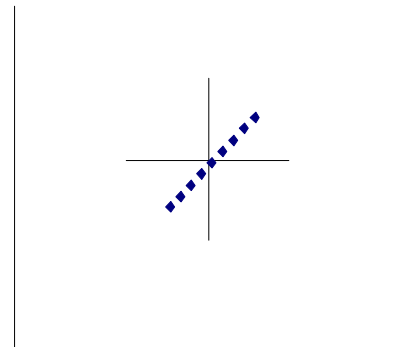
As utilidades básicas da medida são:

Análise exploratória

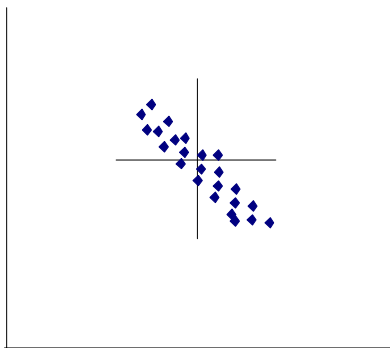
Predição.



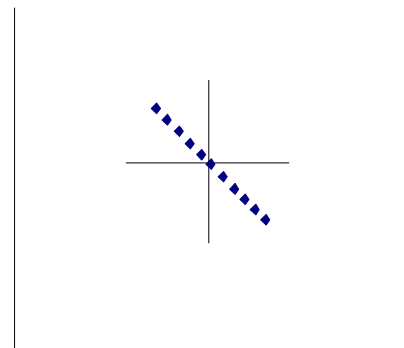
a. $r = 0,6$



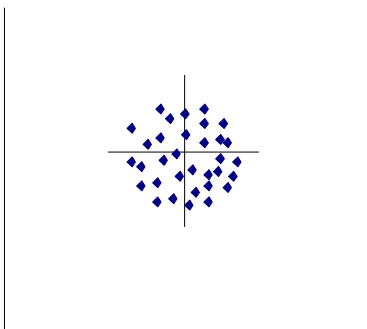
b. $r = 1$



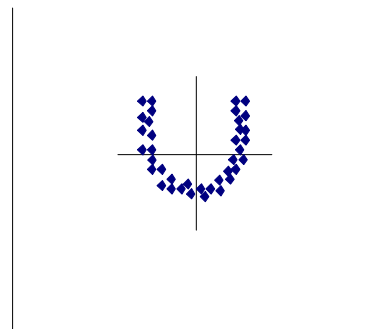
c. $r = -0,8$



d. $r = -1$



e. $r = 0$



f. $r = 0$

Figura 13.3 - Diagramas ilustrativos dos possíveis valores de r .

Observar que em f, muito embora seja possível identificar um tipo de associação entre as duas variáveis aleatórias, esta associação não é do tipo linear.

14. Introdução ao estudo de regressão linear simples

14.1. Introdução

$$IS = 78,9103007 - 0,3418326^{**}.T + 0,7287253^{**}.C - 0,0027154^{**}.T^2 - 0,0041295^{**}.C^2 + 0,0017052^{**}.T.C$$

$$R^2 = 77,17\%$$

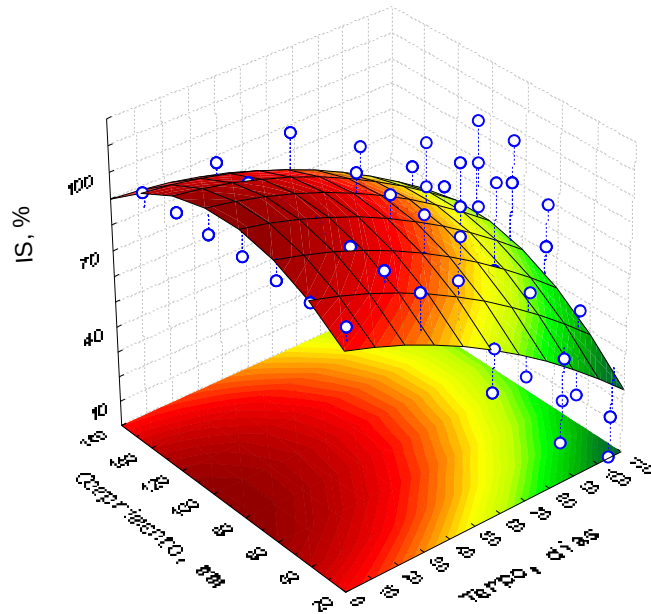


Figura 14.1 – Exemplo ilustrativo de regressão linear múltipla. O índice de sobrevivência (IS) do clone TSH 565 em função do comprimento remanescente foliar e do tempo, após preparo para propagação massal.

Nos experimentos em que os tratamentos são níveis crescentes de pelo menos um fator quantitativo, como por exemplo: adubo, herbicida, irrigação; é estritamente incorreto a utilização dos testes de comparação de médias múltiplas (TCMM), ou análise de contrastes (AC), para estudar seus efeitos sobre as variáveis aleatórias mensuradas.

Essas técnicas, TCMM e AC, são utilizadas na análise qualitativa de experimentos.

Quando os tratamentos são níveis crescentes de pelo menos um fator quantitativo, os ensaios devem ser analisados por intermédio da análise quantitativa de experimentos, isto é, regressão, e ou, correlação.

Embora as técnicas e princípios sejam comuns a ambos os métodos (regressão e correlação), existem diferenças conceituais que devem ser consideradas.

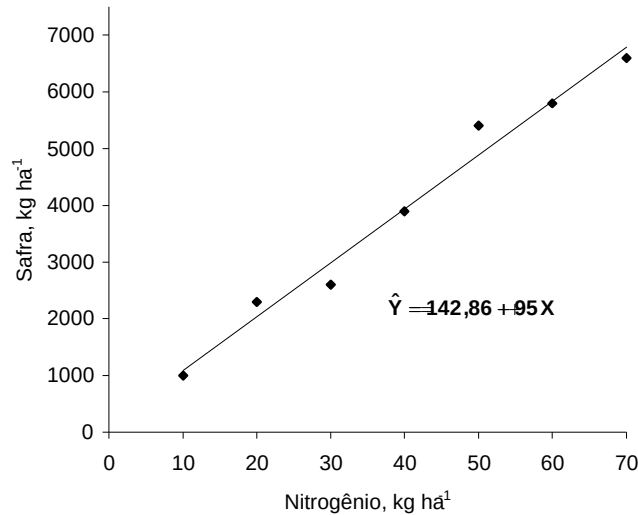
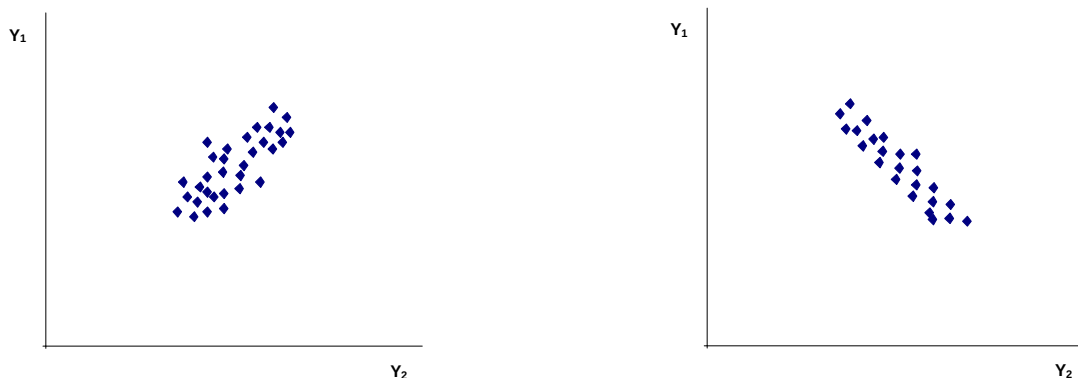


Figura 14.2 – Exemplo ilustrativo de regressão linear simples. A safra do milho em função de doses crescentes de adubo nitrogenado aplicado em cobertura.

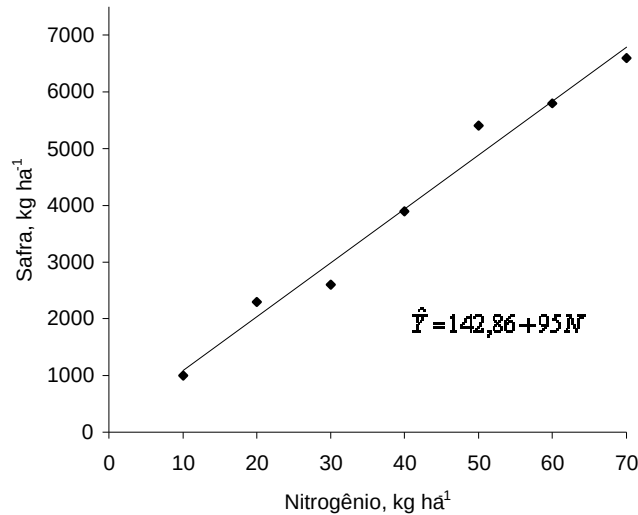
A análise de correlação é indicada para estudar o grau de associação linear entre variáveis aleatórias. Ou seja, essa técnica é empregada, especificamente, para se avaliar o grau de covariação entre duas variáveis aleatórias: se uma variável aleatória Y_1 aumenta, o que acontece com uma outra variável aleatória Y_2 : aumenta, diminui ou não altera?



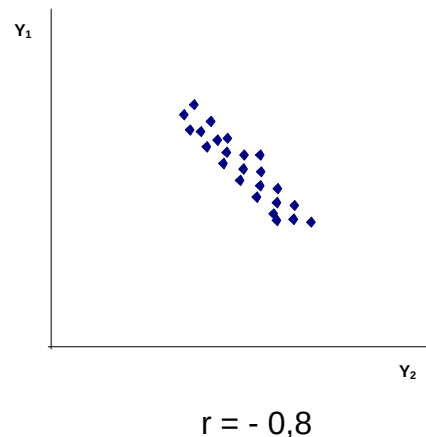
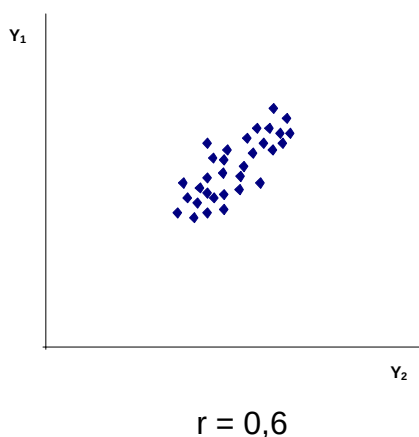
Na análise de regressão uma resposta unilateral é esperada: alterações em X (fator quantitativo) podem implicar em alterações em Y , mas alterações em Y não resultam em alterações em X .

Enquanto a análise de regressão linear nos mostra como as variáveis se relacionam linearmente, a análise de correlação vai nos mostrar apenas o grau desse mesmo relacionamento.

Na análise de regressão estimamos toda uma função $Y = f(X)$, a equação de regressão:



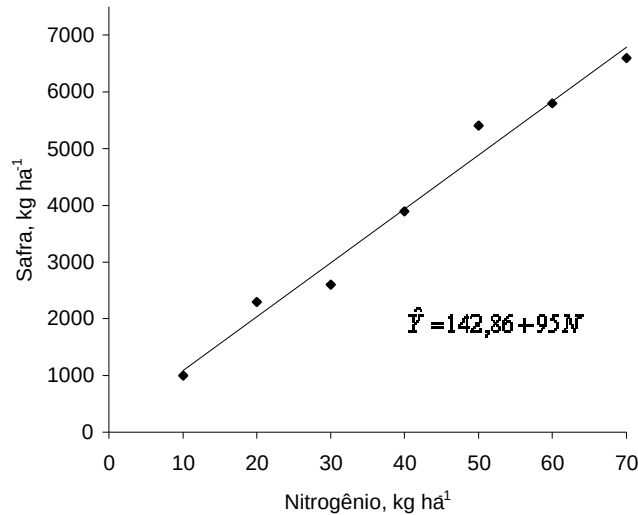
A análise de correlação, por sua vez, nos fornece apenas um número, um índice, que quantifica o grau da associação linear entre duas variáveis aleatórias:



Quando se deseja verificar a existência de alguma relação estatística entre uma ou mais variáveis fixas, independentes, sobre uma variável aleatória, denominada dependente, utiliza-se a análise de regressão (embora essa análise possa, também, ser utilizada para estabelecer a relação funcional entre duas ou mais variáveis aleatórias).

Para exemplificar, vamos considerar que conduzimos um experimento submetendo plantas de milho a doses crescentes de nitrogênio.

Naturalmente, a produção será dependente da quantidade aplicada desse fertilizante, X:



Assim, o fertilizante nitrogenado aplicado é a variável independente, e cada uma das quantidades aplicadas são seus níveis, x_i (10 ... 70 kg ha⁻¹).

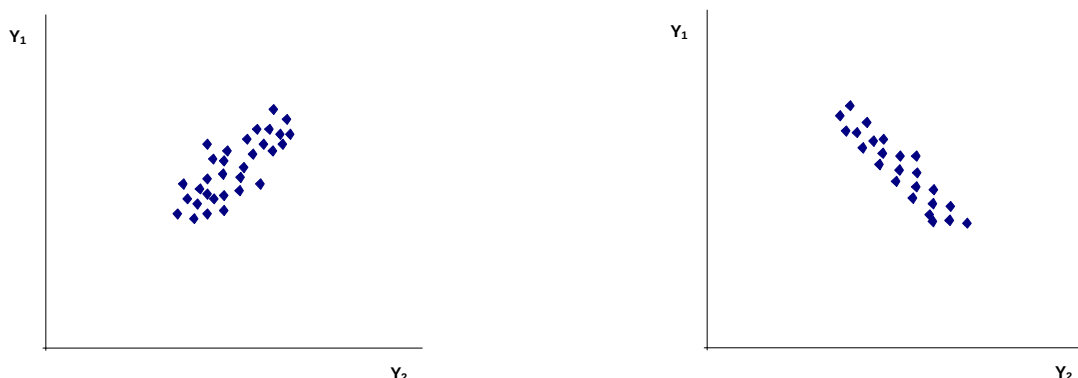
Cada variável aleatória mensurada na cultura do milho, sujeita a influência dos níveis x_i da variável independente, ou seja, das doses de nitrogênio, é chamada “variável dependente” ou “fator resposta”.

Poderia-se medir, por exemplo, o número de espigas por planta (Y_1), a altura média das plantas (Y_2), o peso de 1.000 grãos (Y_3), o teor de proteínas dos grãos (Y_4), o teor de gordura dos grãos (Y_5), etc.

Como a aplicação do fertilizante não depende da safra, sendo, ao contrário, determinada independentemente pelo pesquisador, designamo-la “variável independente” ou “regressor”.

Podemos estudar via análise de regressão o efeito da variável, neste caso, fixa, independente, X (dose de nitrogênio), sobre as variáveis aleatórias, ou dependentes, Y_i (produção de matéria seca, teor de proteínas dos grãos, teor de gordura dos grãos, etc.). Diz-se regressão de Y sobre X.

Posteriormente, caso seja de interesse, podemos utilizar a análise de correlação para estudar o grau de associação linear, por exemplo, entre o teor de proteínas e o teor de gordura dos grãos, sendo ambas variáveis aleatórias:



Ou seja, poderemos estudar via correlação linear simples o grau de associação entre um par qualquer (Y_i, Y_i) . Por exemplo, se o teor de proteínas aumenta, o que acontece com o teor de gordura (aumenta, diminui ou não altera). Estaremos, então, interessados em averiguar a covariação entre estas duas variáveis aleatórias.

Nada impede, entretanto, que o estudo entre o teor de proteínas e teor de gordura seja feito, por meio da análise de regressão. Nesses casos, seria indiferente a posição ocupada por cada uma das variáveis aleatórias, ou seja, a posição Y_i (dependente) ou X_i (independente).

O incorreto seria estudar via análise de correlação o efeito do nitrogênio (variável fixa) sobre a produção de matéria seca dos grãos de milho (variável aleatória), ou sobre os teores de proteína, gordura, etc.

Em síntese, o método da análise de regressão pode ser utilizado sempre que existir uma relação funcional entre uma variável chamada dependente e uma outra chamada independente (regressão linear simples) ou entre uma variável dependente e duas ou mais variáveis independentes (regressão linear múltipla).

Ajustamento

Se precisarmos considerar como a safra depende de diferentes quantidades de nitrogênio, deveremos definir a aplicação do nitrogênio segundo uma escala numérica.

Se grafarmos a safra, Y , decorrente das diversas aplicações, X , de nitrogênio, poderemos observar uma dispersão análoga a Figura 14.3:

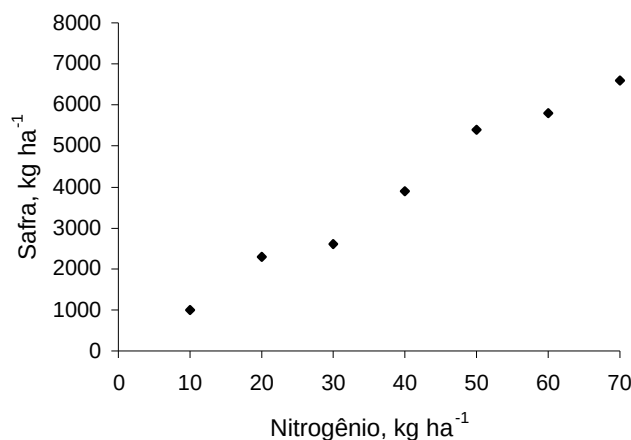


Figura 14.3 - Relação observada entre a safra e a aplicação de nitrogênio.

A aplicação de nitrogênio afeta a safra.

Podemos, por meio de uma equação, relacionando X e Y , descrever como afeta.

Estimar uma equação é geometricamente equivalente a ajustar uma curva àqueles dados dispersos, isto é, a "regressão de Y sobre X ".

Esta equação será útil como descrição breve e precisa de prever a safra Y para qualquer quantidade X de nitrogênio.

Como safr depende do nitrogênio, a safra é chamada "variável dependente" ou "fator resposta", Y .

A aplicação do nitrogênio não depende da safra, sendo, ao contrário, determinada independentemente pelo pesquisador, é chamada a “variável independente” ou “regressor”, X .

Vamos considerar um estudo sobre a influência do N (nitrogênio) aplicado em cobertura sobre a safra do milho.

Suponhamos que só dispomos de recursos para fazer sete observações experimentais.

O pesquisador fixa então sete valores de X (sete níveis do regressor), fazendo apenas uma observação Y (fator resposta), em cada caso, tal como se vê na Figura 14.4:

X Nitrogênio kg ha^{-1}	Y Safr kg ha^{-1}
10	1.000
20	2.300
30	2.600
40	3.900
50	5.400
60	5.800
70	6.600

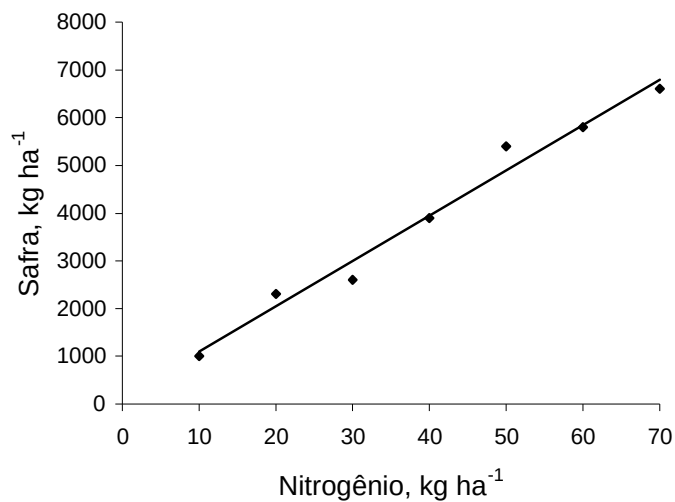
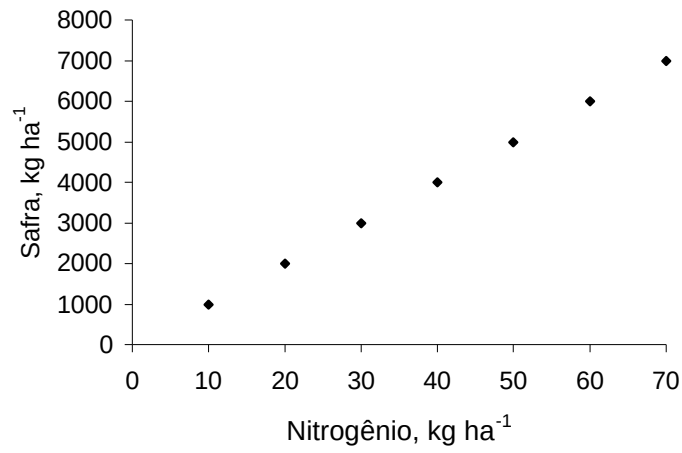


Figura 14.4 - Dados e reta ajustada a olho aos dados apresentados.

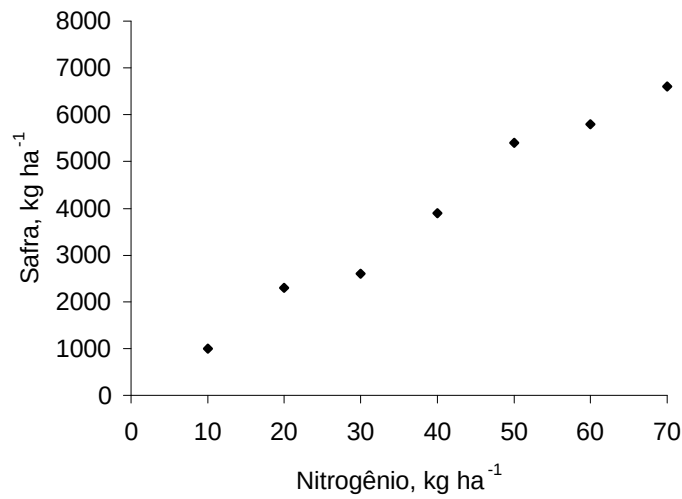
Até onde é bom um ajustamento feito a olho, tal como o da Figura 14.4?

Verificar a ilustração de vários graus de dispersão (Figura 14.5).

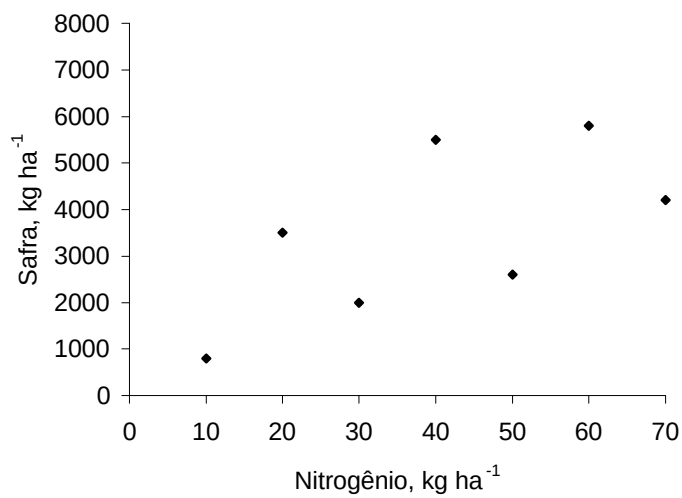
Necessitamos então de um método objetivo, que possa ser estendido ao maior número de situações, onde o ajustamento a olho esteja fora de questão.



a.



b.



c.

Figura 14.5 - Ilustração de diversos graus de dispersão.

14.1.1. Critérios para se ajustar uma reta

Precisamente, o que é um bom ajustamento?

A resposta óbvia seria: um ajustamento que acusa pequeno erro total.

A Figura 14.6 ilustra um erro típico (desvio).

O erro ou a falta de ajustamento é definido como a distância vertical entre o valor observado Y_i e o valor ajustado $\hat{Y}_i$ na reta, isto é, $(Y_i - \hat{Y}_i)$:

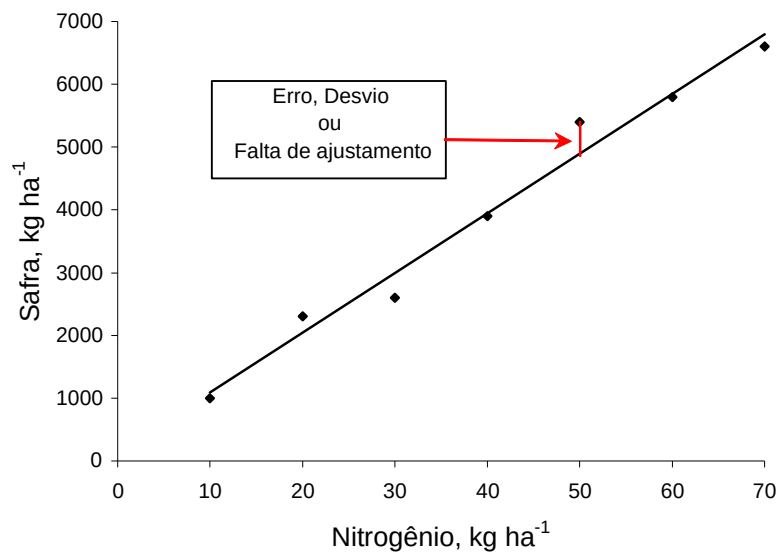


Figura 14.6 - Erro típico no ajustamento de uma reta.

O método mais comumente utilizado para se ajustar uma reta aos pontos dispersos é o que minimiza a soma de quadrados dos erros:

$$\sum_{i=1}^n (Y_i - \hat{Y}_i)^2$$

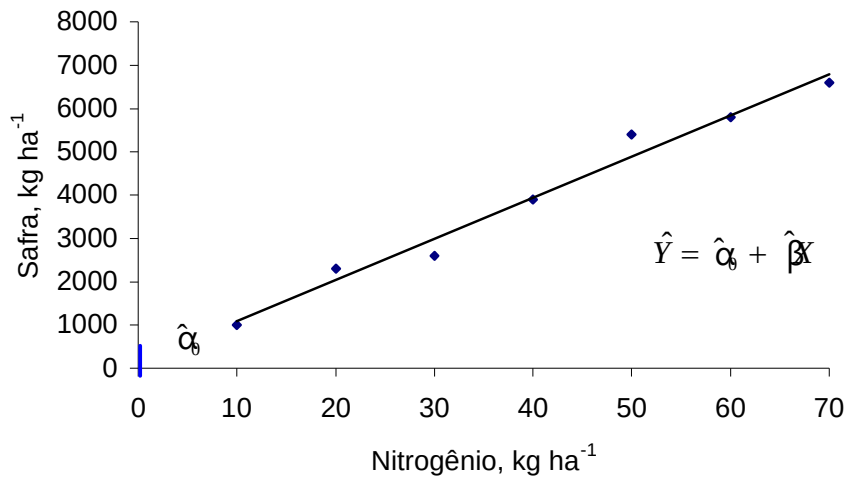
conhecido como critério dos “mínimos quadrados” ou “mínimos quadrados dos erros”. Sua justificativa inclui as seguintes observações:

O quadrado elimina o problema do sinal, pois torna positivos todos os erros.

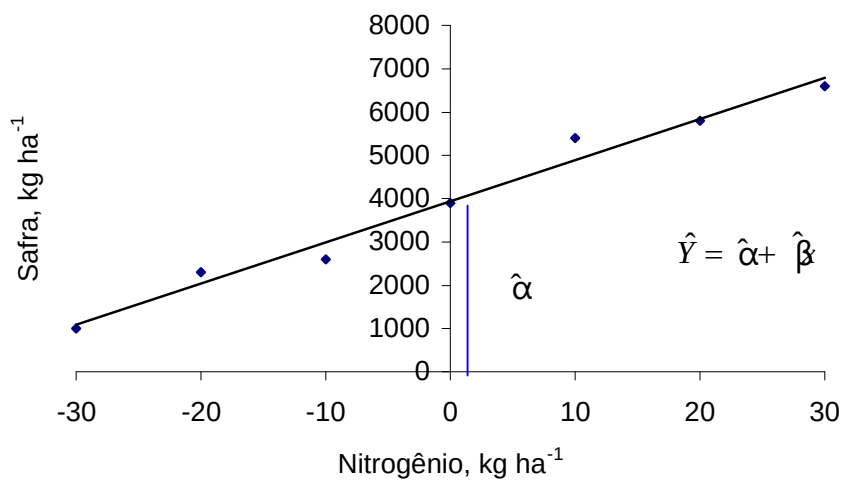
A álgebra dos mínimos quadrados é de manejo relativamente fácil.

14.1.2. Ajustando uma reta

O conjunto de valores X e Y observados na Figura 14.4 é grafado novamente na Figura 14.7(a):



a.



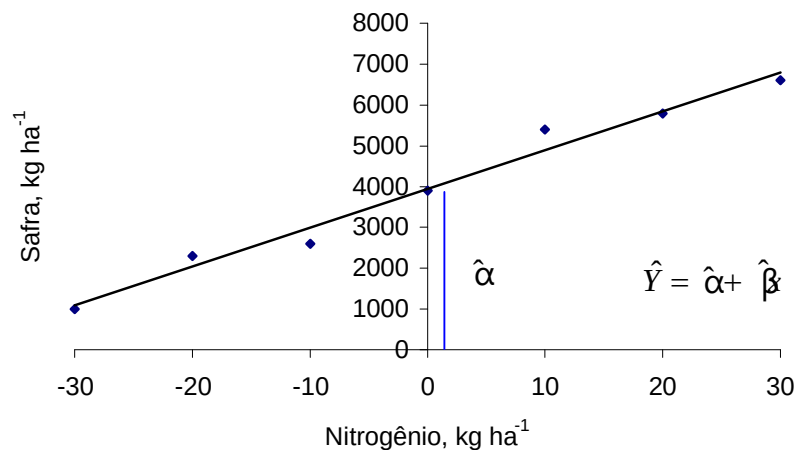
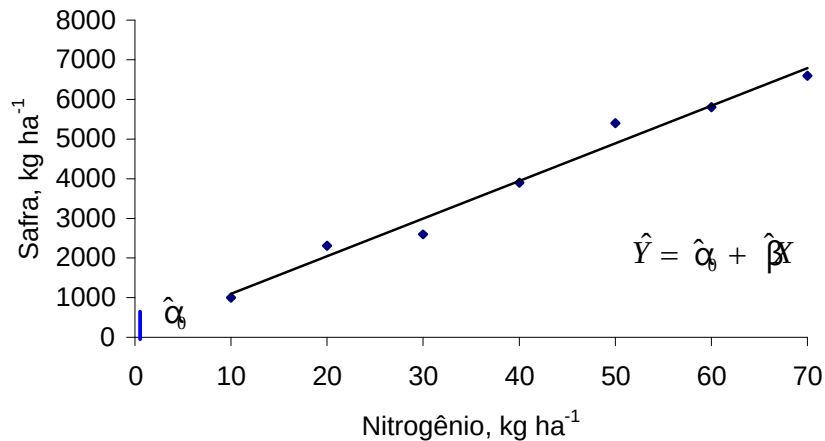
b.

Figura 14.7 - Translação de eixos. (a) Regressão utilizando os valores originais. (b) Regressão após transladar Y.

Estágio 1: Expressar X em termos de desvios a contar de sua média, isto é, definir uma nova variável x (minúsculo), tal que:

$$x = X - \bar{X}$$

Isto equivale a uma translação geométrica de eixos:



Observa-se que o eixo Y foi deslocado para a direita, de 0 a $\bar{X}$.

O novo valor x torna-se positivo, ou negativo, conforme X esteja a direita ou a esquerda de $\bar{X}$.

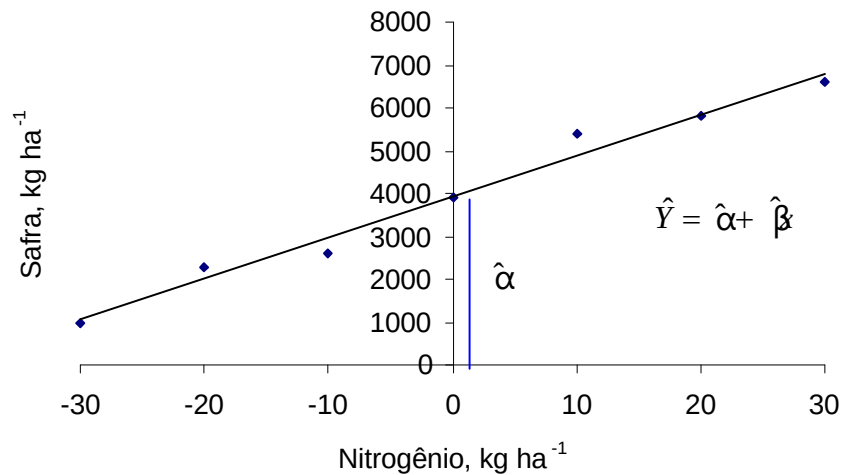
Não há modificação nos valores de Y .

O intercepto $\hat{\alpha}$ difere do intercepto original, $\hat{\alpha}_0$, mas o coeficiente angular, $\hat{\beta}$, permanece o mesmo.

Medir X como desvio a contar de $\bar{X}$ simplifica os cálculos porque a soma dos novos valores x é igual a zero, isto é:

$$\sum x_i = 0 \quad \therefore \quad \sum x_i = \sum (X_i - \bar{X}) = \sum X_i - n\bar{X} = n\bar{X} - n\bar{X} = 0$$

Estágio 2: Ajustar a reta da Figura 14.7(b), isto é, a reta: $\hat{Y} = \hat{\alpha} + \hat{\beta}x$



Devemos ajustar a reta aos dados, escolhendo valores para $\hat{\alpha}$ e $\hat{\beta}$, que satisfaçam o critério dos mínimos quadrados. Ou seja, escolher valores de $\hat{\alpha}$ e $\hat{\beta}$ que minimizem

$$\sum_{i=1}^n (Y_i - \hat{Y}_i)^2 \quad \text{Equação 01}$$

Cada valor ajustado $\hat{Y}_i$ estará sobre a reta estimada:

$$\hat{Y}_i = \hat{\alpha} + \hat{\beta}_i \quad \text{Equação 02}$$

Assim, estamos diante da seguinte situação: devemos encontrar os valores $\hat{\alpha}$ e $\hat{\beta}$ de modo a minimizar a soma de quadrados dos erros.

Considerando as Equações 01 e 02, isto pode ser expresso algebricamente como:

$$\sum_{i=1}^n (Y_i - \hat{Y}_i)^2 \quad \therefore \quad \hat{Y}_i = \hat{\alpha} + \hat{\beta}_i$$

$$S(\hat{\alpha}, \hat{\beta}) = \sum (Y_i - (\hat{\alpha} + \hat{\beta}_i))^2 = \sum (Y_i - \hat{\alpha} - \hat{\beta}_i)^2$$

Utilizou-se $S(\hat{\alpha}, \hat{\beta})$ para enfatizar que esta expressão depende de $\hat{\alpha}$ e $\hat{\beta}$. Ao variarem $\hat{\alpha}$ e $\hat{\beta}$ (quando se tentam várias retas), $S(\hat{\alpha}, \hat{\beta})$ variará também.

Pergunta-se então, para que valores de $\hat{\alpha}$ e $\hat{\beta}$ haverá um mínimo de erros?

A resposta a esta pergunta nos fornecerá a reta “ótima” (de mínimos quadrados dos erros).

A técnica de minimização mais simples é fornecida pelo cálculo. A minimização de $S(\hat{\alpha}, \hat{\beta})$ exige o anulamento simultâneo de suas derivadas parciais:

Igualando a zero a derivada parcial em relação a $\hat{\alpha}$:

$$\frac{\partial}{\partial \hat{\alpha}} \sum (Y_i - \hat{\alpha} - \hat{\beta}_i)^2 = \sum 2(-1)(Y_i - \hat{\alpha} - \hat{\beta}_i) = 0$$

Dividindo ambos os termos por (-2) e reagrupando:

$$\sum Y_i - n \hat{\alpha} - \hat{\beta} \sum x_i = 0 \quad \therefore \quad \sum x_i = 0$$

$$\sum Y_i - n \hat{\alpha} - 0 = 0$$

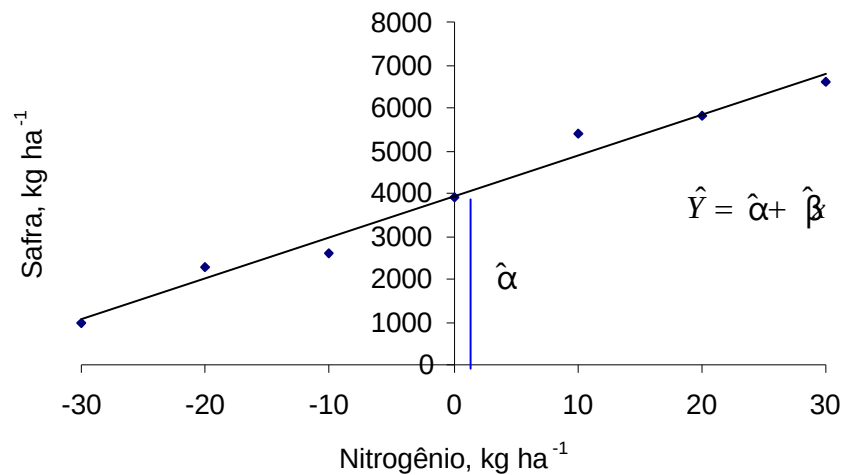
$$\sum Y_i - n \hat{\alpha} = 0$$

$$n \hat{\alpha} = \sum Y_i$$

$$\hat{\alpha} = \frac{\sum Y_i}{n} = \bar{Y}$$

Assim, a estimativa de mínimos quadrados para $\hat{\alpha}$ é simplesmente o valor médio de Y.

Verifica-se que isto assegura que a reta de regressão ajustada deve passar pelo ponto $(\bar{x}, \bar{Y})$, que pode ser interpretado como o centro de gravidade da amostra de n pontos:



É preciso também anular a derivada parcial em relação a $\hat{\beta}$:

$$\frac{\partial}{\partial \hat{\beta}} \sum (Y_i - \hat{\alpha} - \hat{\beta}x_i)^2 = \sum 2(-x_i)(Y_i - \hat{\alpha} - \hat{\beta}x_i) = 0$$

Dividindo ambos os termos por (-2):

$$\sum x_i (Y_i - \hat{\alpha} - \hat{\beta}x_i) = 0$$

Reagrupando:

$$\sum x_i Y_i - \hat{\alpha} \sum x_i - \hat{\beta} \sum x_i^2 = 0 \quad \therefore \quad \sum x_i = 0$$

$$\sum x_i Y_i - 0 - \hat{\beta} \sum x_i^2 = 0$$

$$\sum x_i Y_i - \hat{\beta} \sum x_i^2 = 0$$

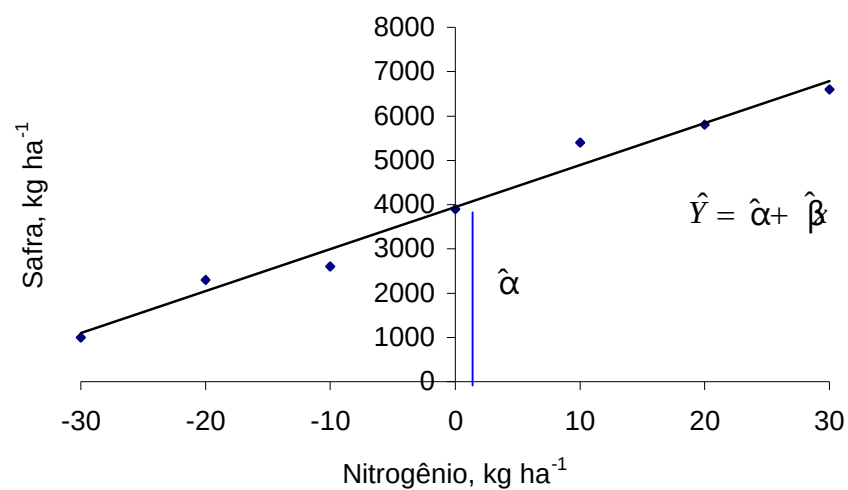
$$\hat{\beta} \sum x_i^2 = \sum x_i Y_i$$

$$\hat{\beta} = \frac{\sum x_i Y_i}{\sum x_i^2}$$

Podemos sintetizar da seguinte forma:

Com os valores x medidos como desvios a contar de sua média, os valores $\hat{\alpha}$ e $\hat{\beta}$ de mínimos quadrados dos erros são:

$$\hat{\alpha} = \frac{\sum Y_i}{n} = \bar{Y}$$



Para os dados da Figura 14.4, $\hat{\alpha}$ e $\hat{\beta}$ acham-se calculados no Quadro 14.1.

Quadro 14.1 - Cálculos dos valores necessários

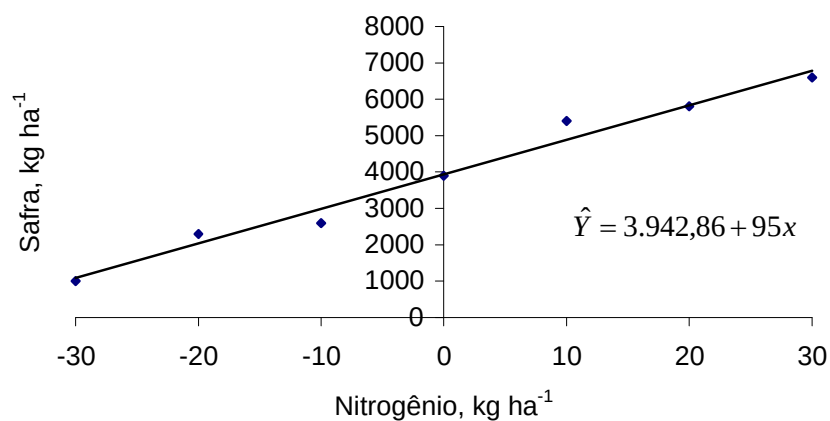
X	$x = X - \bar{X}$ $x = X - 40$	Y	xY	x^2
10	-30	1.000	-30.000	900
20	-20	2.300	-46.000	400
30	-10	2.600	-26.000	100
40	0	3.900	0	0
50	10	5.400	54.000	100
60	20	5.800	116.000	400
70	30	6.600	198.000	900
$\sum X = 280$		$\sum Y = 27.600$		
$\bar{X} = \frac{1}{N} \sum X$		$\bar{Y} = \frac{1}{N} \sum Y$		
$\bar{X} = \frac{280}{7} = 40$		$\bar{Y} = \frac{27.600}{7}$		
$\sum x = 0$		$\bar{Y} = 3.942,86$		
		$\sum xY = 266.000$		$\sum x^2 = 2.800$

$$\hat{\alpha} = \frac{\sum Y_i}{n} = \bar{Y} \therefore \hat{\alpha} = \frac{27.600}{7} = 3.942,86$$

$$\hat{\beta} = \frac{\sum x_i Y_i}{\sum x_i^2} \therefore \hat{\beta} = \frac{266.000}{2.800} = 95,00$$

$$\hat{Y} = 3.942,86 + 95x$$

Equação 03



Estágio 3: A regressão pode agora ser transformada para o sistema original de referência:

$$\hat{Y} = 3.942,86 + 95x \quad \therefore \quad x = (X - \bar{X})$$

$$\hat{Y} = 3.942,86 + 95(X - \bar{X})$$

$$\hat{Y} = 3.942,86 + 95(X - 40)$$

$$\hat{Y} = 3.942,86 + 95X - 3.800$$

$$\hat{Y} = 142,86 + 95X$$

Equação 04

$$\hat{Y} = 3.942,86 + 95x$$

Equação 03

Comparando as Equações 03 e 04, observa-se que:

O coeficiente angular da reta de regressão ajustada ($\hat{\beta} = 95X$) permanece inalterado.

A única diferença é o intercepto, $\hat{\alpha}$, onde a reta tangencia o eixo Y.

O intercepto original foi facilmente reobtido.

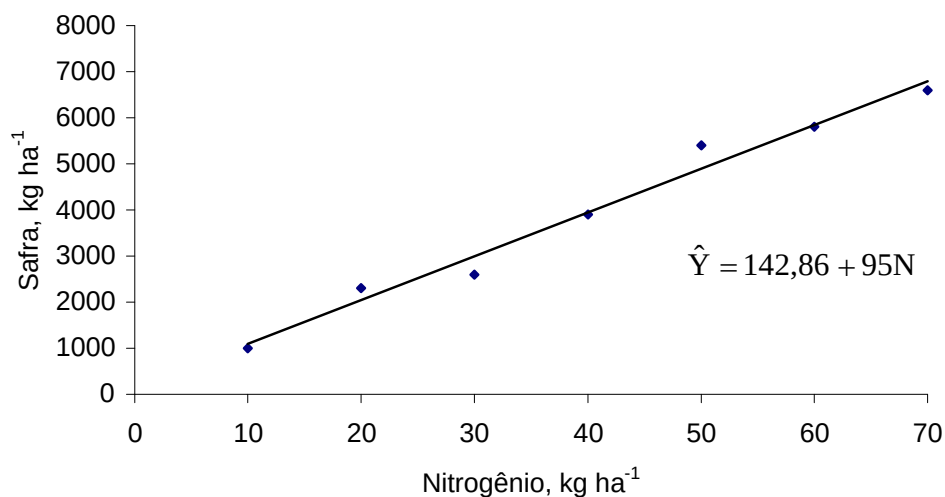


Figura 14.8 - Gráfico dos pontos dispersos com a reta ajustada.

Esta equação é útil como descrição breve e precisa de predizer a safra, em kg ha^{-1} , para qualquer quantidade de nitrogênio, também em kg ha^{-1} , aplicada.

Observar que:

Se nenhum nitrogênio for aplicado à cultura, a safra estimada será de 142,86 kg.

Esta safra se deve a absorção pela cultura do N disponível no solo, possivelmente associado ao ciclo orgânico.

No intervalo das doses aplicadas (10 a 70 kg), considerando-se um hectare, para cada kg de nitrogênio aplicado, a cultura responde com 95 kg de grãos.

14.2. Análise de variância da regressão

Para se decidir quão bem o modelo ajustado é adequado à natureza dos dados experimentais, pode-se lançar mão da análise de variância da regressão (ANOVAR).

Para o caso em estudo, a ANOVAR irá particionar a variação total (SQDtot) da variável dependente - ou fator resposta - em função das variações nos níveis da variável independente - ou regressor, em duas partes:

Uma parte associada ao modelo ajustado (SQDDreg): soma de quadrados dos desvios devido à regressão, que quantifica o quanto da variação total da safra, provocada pela variação das doses de nitrogênio, é explicada pelo modelo ajustado.

Uma outra parte associada à falta de ajuste (SQDDerr): soma de quadrados dos desvios devido ao erro, que quantifica o montante da variação total da safra, provocada pela variação da dose de nitrogênio, que não é explicada pelo modelo ajustado.

Para o exemplo em análise a ANOVAR teria a seguinte estrutura:

Hipóteses:

$H_0: \beta_i = 0$	ou	$H_0: Y \neq \alpha_0 + \beta X$
$H_1: \beta_i > 0$	ou	$H_1: Y = \alpha_0 + \beta X$

Significado de H_0 : A equação de regressão não explica a variação da variável dependente Y, em decorrência da variação da variável independente X, ao nível de ...% de probabilidade.

Significado de H_1 : A equação de regressão explica a variação da variável dependente Y, em decorrência da variação da variável independente X, ao nível de ...% de probabilidade.

ANOVAR

Causa da variação	GL
Regressão	1
Erro	5
Total	6

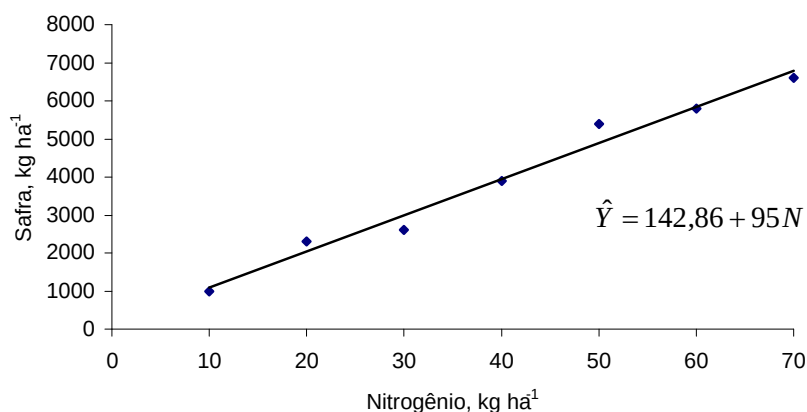
Existem várias formas de realizar estes cálculos.

Objetivando clareza de idéias e conceitos, a forma que será empregada utilizará o conceito mais elementar da estatística, ou seja, a variância:

$$\text{Quadrado médio dos desvios} = s^2 = \frac{SQD}{n-1} \therefore SQD = \sum (Y_i - m)^2$$

Vejamos¹:

N , kg ha ⁻¹	Safra_Obs	Safra_Est
10	1.000	1092,86
20	2.300	2042,86
30	2.600	2992,86
40	3.900	3942,86
50	5.400	4892,86
60	5.800	5842,86
70	6.600	6792,86



¹ Obs = Observado: valores observados de Y

Est = Estimado: valores estimados para Y a partir da equação de regressão.

SQDtot

Obs	$m_{(Obs)}$	$Obs - m_{(Obs)}$	$[Obs - m_{(Obs)}]^2$
1.000	3.942,86	-2.942,86	8.660.408,16
2.300	3.942,86	-1.642,86	2.698.979,59
2.600	3.942,86	-1.342,86	1.803.265,31
3.900	3.942,86	-42,86	1.836,73
5.400	3.942,86	1.457,14	2.123.265,31
5.800	3.942,86	1.857,14	3.448.979,59
6.600	3.942,86	2.657,14	7.060.408,16
			<u>25.797.142,86</u>

SQDreg

Est	$m_{(Est)}$	$Est - m_{(Est)}$	$[Est - m_{(Est)}]^2$
1.093	3.942,86	-2.850,00	8.122.500,00
2.043	3.942,86	-1.900,00	3.610.000,00
2.993	3.942,86	-950,00	902.500,00
3.943	3.942,86	0,00	0,00
4.893	3.942,86	950,00	902.500,00
5.843	3.942,86	1.900,00	3.610.000,00
6.793	3.942,86	2.850,00	8.122.500,00
			<u>25.270.000,00</u>

SQDerr

Obs	Est	Erro($Obs - Est$)	$m_{(Erro)}$	$Erro - m_{(Erro)}$	$[Erro - m_{(Erro)}]^2$
1.000	1.092,86	-92,86	0,00	-92,86	8.622,45
2.300	2.042,86	257,14	0,00	257,14	66.122,45
2.600	2.992,86	-392,86	0,00	-392,86	154.336,73
3.900	3.942,86	-42,86	0,00	-42,86	1.836,73
5.400	4.892,86	507,14	0,00	507,14	257.193,88
5.800	5.842,86	-42,86	0,00	-42,86	1.836,73
6.600	6.792,86	-192,86	0,00	-192,86	37.193,88
					<u>527.142,86</u>

ANOVAR

Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Regressão	1	25.270.000,00	25.270.000,00	239,69	< 0,0001
Erro	5	527.142,86	105.428,57		
Total	6	25.797.142,86			

Conclusão: rejeita-se H_0 ao nível de 5% de probabilidade pelo teste F.

Ou seja, a equação de regressão ajustada explica a variação da safra, em decorrência da variação das doses de nitrogênio, ao nível de 5% de probabilidade pelo teste F.

14.2.1. Cálculos alternativos da soma de quadrados dos desvios

É possível demonstrar algebricamente que:

$$SQD_{tot} = \sum Y_i^2 - \frac{(\sum Y_i)^2}{n}$$

$$SQD_{reg} = \hat{\alpha} \sum Y_i + \hat{\beta} \sum X_i Y_i - \frac{(\sum Y_i)^2}{n}$$

$$SQD_{err} = SQD_{tot} - SQD_{reg}$$

Esta forma de realizar os cálculos da soma de quadrados dos desvios, embora menos compreensível a primeira vista, é a mais prática e deve ser preferencialmente utilizada.

X	Y	Y ²	XY
10	1.000	1.000.000	10.000
20	2.300	5.290.000	46.000
30	2.600	6.760.000	78.000
40	3.900	15.210.000	156.000
50	5.400	29.160.000	270.000
60	5.800	33.640.000	348.000
70	6.600	43.560.000	462.000
	27.600	134.620.000	1.370.000

$$SQD_{tot} = \sum Y_i^2 - \frac{(\sum Y_i)^2}{n} = 134.620.000 - \frac{(27.600)^2}{7} = 25.797.142,86$$

$$SQD_{reg} = \hat{\alpha}_0 \sum Y_i + \hat{\beta} \sum X_i Y_i - \frac{(\sum Y_i)^2}{n}$$

$$SQD_{reg} = 142,85714286 \cdot 27.600 + 95 \cdot 1.370.000 - \frac{(27.600)^2}{7}$$

$$SQD_{reg} = 25.270.000$$

$$SQD_{err} = SQD_{tot} - SQD_{reg}$$

$$SQD_{err} = 25.797.142,86 - 25.270.000$$

$$SQD_{err} = 527.142,86$$

ANOVAR

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Regressão	1	25.270.000,00	25.270.000,00	239,69	< 0,0001
Erro	5	527.142,86	105.428,57		
Total	6	25.797.142,86			

14.2.2. Coefficiente de determinação da regressão

O coeficiente de determinação do modelo de regressão, r^2 , é uma medida do grau de ajuste do modelo aos dados experimentais:

$$r^2 = \frac{SQD_{reg}}{SQD_{tot}} \therefore 0 \leq r^2 \leq 1$$

Este coeficiente, nos dá uma informação do quão bem, ou não, o modelo utilizado se ajusta a natureza dos dados experimentais. Para o exemplo em análise:

$$r^2 = \frac{25.270.000,00}{25.797.142,86} = 0,9796 = 97,96\%$$

Interpretação: 97,96% da variação total da safra, em decorrência da variação da dose de nitrogênio, é explicada pelo modelo de regressão ($\hat{Y} = 142,86 + 95N$) ajustado.

14.2.3. Relação entre o coeficiente de determinação e o coeficiente de correlação

Se análise de regressão linear simples for realizada entre duas variáveis aleatórias, a relação existente entre o coeficiente de determinação da regressão, r^2 , e o coeficiente de correlação, r , é a seguinte:

$$r = \sqrt{r^2}$$

Nos casos da regressão ter sido realizada entre uma variável aleatória e uma variável fixa, esta relação não possui significado estatístico.

14.2.4. Observações a respeito da regressão

Quando os dados não provêm de um delineamento experimental, como no exemplo analisado, a ANOVA pode ser realizada da forma apresentada, e se terá chegado ao fim da análise.

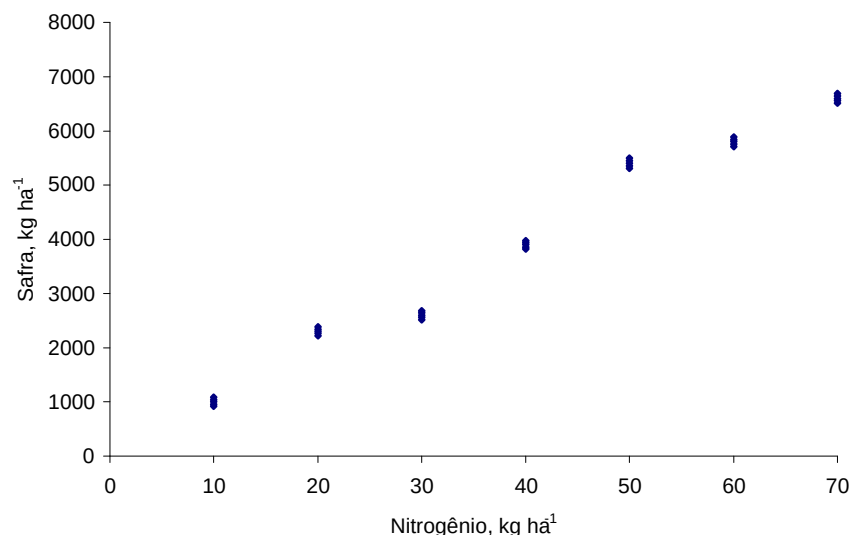
Entretanto, quando os dados provêm de um delineamento experimental, onde são observadas repetições, e por conseguinte existe um erro experimental, além do erro devido a falta de ajuste do modelo:

O ajustamento segue os mesmos princípios, ou seja, geralmente, é realizado observando-se as médias de cada tratamento.

A análise de variância sofre ligeiras alterações, como será visto no exemplo a seguir.

14.2.5. Análise de regressão de dados provenientes de delineamentos experimentais

Ao aplicar-se o princípio da repetição, cada nível de nitrogênio terá que ser repetido um certo número de vezes.



Considerando mais detalhadamente, a parte “puramente aleatória” de Y_i . O termo erro ou perturbação, de onde provém? Por que não obtemos um valor preciso e exato da safra (Y_i) em cada repetição, já que cada dose de nitrogênio (X_i) é fixa?

O erro pode ser encarado como a soma de duas componentes:

Erro de mensuração.

Erro estocástico. Ocorre em consequência da irreprodutividade inerente aos fenômenos biológicos, podendo ser reduzido mediante um controle experimental rígido.

O erro estocástico pode ser encarado como a influência sobre a safra de muitas variáveis omissas, ou não controladas, cada uma com um pequeno efeito individual.

Exemplo:

Os dados abaixo são provenientes de um ensaio experimental em que foram utilizadas sete doses de nitrogênio aplicado em cobertura sobre a produtividade de milho. O Experimento foi montado no delineamento inteiramente casualizado, DIC, com cinco repetições. Os dados são fornecidos abaixo:

Quadro 14.2 – Produção de milho, kg ha⁻¹

N kg.ha ⁻¹	Repetições					Totais	Rep.	Médias
	1	2	3	4	5			
10	1.000	916	958	1.084	1.042	5.000	5	1.000
20	2.340	2.220	2.300	2.260	2.380	11.500	5	2.300
30	2.559	2.518	2.682	2.641	2.600	13.000	5	2.600
40	3.976	3.900	3.862	3.938	3.824	19.500	5	3.900
50	5.448	5.304	5.352	5.400	5.496	27.000	5	5.400
60	5.843	5.886	5.800	5.714	5.757	29.000	5	5.800
70	6.600	6.555	6.690	6.510	6.645	33.000	5	6.600
						138.000	35	3.942,86

$$C = (138.000)^2 / 35 = 544.114.285,71$$

$$SQD_{tot} = [(1.000)^2 + (916)^2 + \dots + (6.645)^2] - C = 129.112.384,29$$

$$SQD_{tra} = 1/5 [(5.000)^2 + (11.510)^2 + \dots + (33.000)^2] - C = 128.985.714,29$$

$$SQD_{res} = SQD_{tot} - SQD_{tra} = 129.112.384,29 - 128.985.714,29 = 126.670,00$$

Hipóteses:

$$H_0: \alpha_{10} = \dots = \alpha_{70}$$

H_1 : Nem todas as médias são iguais

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	6	128.985.714,29	21.497.619,05	4.751,98	< 0,0001
Resíduo	28	126.670,00	4.523,93		
Total	34	129.112.384,29			

Conclusão: rejeita-se H_0 ao nível de significância de 5% pelo teste F.

Como as médias dos tratamentos deste experimento em análise foram utilizadas na parte referente a ajustamento, anteriormente visto, podemos, então, montar o quadro de análise de variância completo do experimento:

Hipóteses:

$H_0: \beta_i = 0$	ou	$H_0: Y \neq \alpha_0 + \beta X$
$H_1: \beta_i > 0$	ou	$H_1: Y = \alpha_0 + \beta X$

ANOVA

Causa da variação	GL	SQD	QMD	F_{cal}	Pr
Tratamentos	(6)	(128.985.714,29)			
Dev. regressão	1	126.350.000,00	126.350.000,00	27.929,26	< 0,0001
Ind. regressão	5	2.635.714,29	527.142,86	116,52	< 0,0001
Resíduo	28	126.670,00	4.523,93		
Total	34	129.112.384,29			

Observa-se que a soma de quadrados, e os respectivos graus de liberdade, associados a tratamentos foram desdobrados em duas partes:

Uma parte associada ao modelo de regressão utilizado ($\hat{Y} = 142,86 + 95N$).

Uma parte associada à falta de ajuste ou erro de ajustamento:

Para a obtenção da soma de quadrados do devido à regressão e ao independente da regressão tem-se duas opções:

- Realizar todos os cálculos das somas de quadrados dos desvios considerando agora todas as repetições, o que embora possa ser feito, é um processo mais trabalhoso.
- Utilizar o teorema do limite central (que facilita bastante os cálculos):

$$Var(m) = \frac{\sigma^2}{n} \therefore \sigma^2 = Var(m) \cdot n$$

$$SQD(m) = \frac{SQD}{n} \therefore SQD = SQD(m) \cdot n \therefore \text{Como } n = r$$

$$SQDDreg = 25.270.000,00 \cdot 5 = 126.350.000,00$$

$$SQDDireg = 527.142,86 \cdot 5 = 2.635.714,29$$

14.3. Critérios para decisão de um modelo ajustado e considerações finais

Para se chegar a uma conclusão final sobre um modelo de regressão ajustado aos dados experimentais deve-se considerar o seguinte conjunto de observações:

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Tratamentos	(6)	(128.985.714,29)			
Dev. regressão	1	126.350.000,00	126.350.000,00	27.929,26	< 0,0001
Ind. regressão	5	2.635.714,29	527.142,86	116,52	< 0,0001
Resíduo	28	126.670,00	4.523,93		
Total	34	129.112.384,29			

O modelo é adequado à natureza do fenômeno em estudo, ou adequado ao que se sabe sobre o fenômeno?

O coeficiente de determinação (r^2) é elevado?

No quadro final da análise de variância o efeito do devido a regressão é significativo?

No quadro final da análise de variância o efeito do devido ao independente da regressão é não significativo?

Informações adicionais:

Nem sempre se consegue respostas favoráveis a todo o conjunto destes pontos (a ... d).

Quanto mais próximo da situação ideal: melhor o modelo ajustado.

É necessário bom censo e muita prática para se realizar bons ajustes de modelos de regressão aos dados experimentais.

Individualmente, a análise de regressão é um dos mais amplos tópicos da estatística e da estatística experimental.

A abordagem utilizada, embora não seja a usual para trabalhos do dia a dia, é a mais simples, prática e objetiva para um estudo introdutório, possibilitando um entendimento inicial claro aos modelos de regressão linear.

14.4. Exemplo de análise completa de um experimento

Os dados abaixo são provenientes de um ensaio experimental realizado em casa de vegetação, montado no delineamento em blocos casualizados, com cinco repetições, para avaliar o efeito de doses de fósforo na produção de matéria seca da parte aérea do milho:

Quadro 14.3 – Matéria seca da parte aérea das plantas de milho, g vaso⁻¹

P mg.kg ⁻¹	Blocos					Totais	Rep.	Médias
	1	2	3	4	5			
0,0	6,73	6,93	6,65	6,78	6,61	33,70	5	6,74
32,5	8,72	8,65	8,74	8,56	8,98	43,65	5	8,73
65,0	11,12	10,88	11,02	10,65	10,78	54,45	5	10,89
97,5	12,36	12,51	12,61	12,84	12,48	62,80	5	12,56
130,0	14,23	14,09	14,13	14,04	14,06	70,55	5	14,11
	53,16	53,06	53,15	52,87	52,91	265,15	25	

$$\begin{aligned}
 C &= (265,15)^2 / 25 = 2.812,181 \\
 SQD_{tot} &= [(6,73)^2 + (6,93)^2 + \dots + (14,06)^2] - C = 173,663 \\
 SQD_{blo} &= 1/5 [(53,16)^2 + (53,06)^2 + \dots + (52,91)^2] - C = 0,014 \\
 SQD_{tra} &= 1/5 [(33,70) + (43,65) + \dots + (70,55)] - C = 173,211 \\
 SQD_{res} &= SQD_{tot} - SQD_{blo} - SQD_{tra} = 0,438
 \end{aligned}$$

Hipóteses:

$$H_0: \alpha_0 = \dots = \alpha_{130}$$

$$H_1: \alpha_a > \alpha_b, \text{ para } a \neq b$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Bloco	4	0,014	0,004	0,132	0,969
Tratamentos	4	173,211	43,303	1.580,533	< 0,0001
Resíduo	16	0,438	0,027		
Total	24	173,663			

Conclusão: rejeita-se H_0 ao nível de significância de 5% pelo teste F.

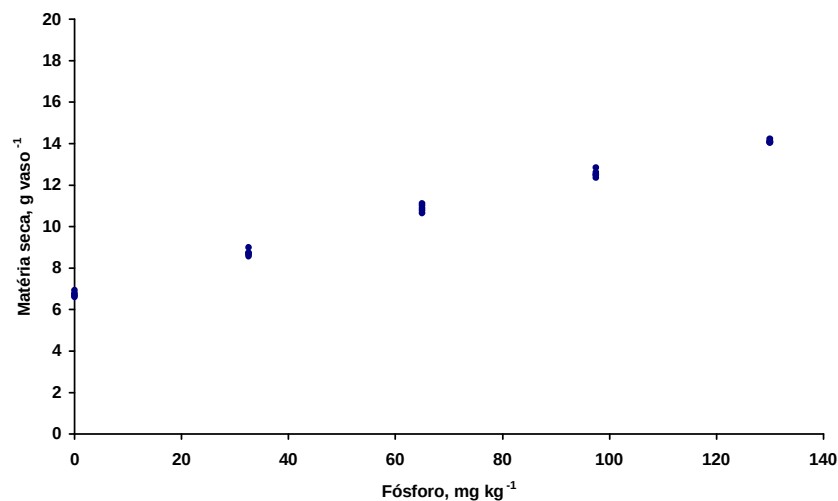


Figura 14.9 - A visualização dos dados experimentais em um gráfico de dispersão auxilia na escolha do modelo a ser ajustado.

Ao se tentar ajustar um modelo de regressão aos dados experimentais a ANOVA permitirá a decisão se a equação obtida é adequada, ou não, como forma de prever a matéria seca da parte aérea produzida pelas plantas de milho, em g vaso⁻¹, para qualquer quantidade de fósforo aplicado no intervalo estudado, em mg kg⁻¹.

Para isto, a soma de quadrados de tratamentos (SQDtra) deverá ser particionada em:

Uma parte explicada ou devida à equação de regressão a ser ajustada.

Uma outra parte que não é explicada por esta equação de regressão, ou seja, independe da regressão ajustada:

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Bloco	4	0,014	0,004	0,132	0,969
Tratamentos	(4)	(173,211)	43,303	1.580,533	< 0,0001
Dev. regressão	?	?	?	?	?
Ind. regressão	?	?	?	?	?
Resíduo	16	0,438	0,027		
Total	24	173,663			

Ajustando um modelo linear: $\hat{Y} = \alpha + \beta X$

Quadro 14.4 - Valores necessários para o ajustamento do modelo linear¹

X	$x = X - \bar{X}$ $x = X - 65$	Y	xY	x ²
0,0	-65,0	6,74	-438,10	4.225,00
32,5	-32,5	8,73	-283,73	1.056,25
65,0	0,0	10,89	- 0,00	0,00
97,5	32,5	12,56	408,20	1.056,25
130,0	65,0	14,11	917,15	4.225,00

$$\begin{aligned}
 \sum X &= 325 & \sum Y &= 53,03 \\
 \bar{X} &= \frac{1}{N} \sum X & \bar{Y} &= \frac{1}{N} \sum Y \\
 \bar{X} &= \frac{325}{5} = 65,0 & \bar{Y} &= \frac{53,03}{5} = 10,606 \\
 \sum x &= 0 & \sum xY &= 603,53 & \sum x^2 &= 10.562,50
 \end{aligned}$$

$$\hat{\alpha} = \frac{\sum Y_i}{n} = \bar{Y} \therefore \hat{\alpha} = \frac{53,03}{5} = 10,606$$

¹ Se o leitor realizar os cálculos utilizando apenas o número de casas decimais apresentadas encontrará diferenças de resultados ao longo deste tópico da apostila. Estas diferenças devem-se às aproximações. Nos cálculos estatísticos intermediários recomenda-se trabalhar com o máximo possível de casas decimais (utilizou-se 17 casas decimais).

$$\hat{\beta} = \frac{\sum x_i Y_i}{\sum x_i^2} \therefore \hat{\beta} = \frac{603,53}{10.562,50} = 0,0571$$

$$\hat{Y} = 10,606 + 0,0571.x \quad \therefore \quad x = (X - \bar{X})$$

$$\hat{Y} = 10,606 + 0,0571.(X - \bar{X})$$

$$\hat{Y} = 10,606 + 0,0571.(X - 65)$$

$$\hat{Y} = 10,606 + 0,0571.X - 3,714$$

$$\hat{Y} = 6,892 + 0,0571.X$$

Equação ajustada:

$$\hat{Y} = 6,892 + 0,0571.P$$

Quadro 14.5 - Valores necessários para a análise de variância da regressão

X	Y	Y ²	X.Y
0,0	6,74	45,4276	0,000
32,5	8,73	76,2129	283,725
65,0	10,89	118,5921	707,850
97,5	12,56	157,7536	1.224,600
130,0	14,11	199,0921	1.834,300

$$\sum Y = 53,03 \quad \sum Y^2 = 597,0783 \quad \sum XY = 4.050,475$$

$$\hat{Y} = 6,892 + 0,0571.P$$

$$SQD_{tot} = \sum Y_i^2 - \frac{(\sum Y_i)^2}{n}$$

$$SQD_{tot} = 597,0783 - \frac{(53,03)^2}{5} = 34,642$$

$$SQDreg = \hat{\alpha} \sum Y_i + \hat{\beta} \sum X_i Y_i - \frac{(\sum Y_i)^2}{n}$$

$$SQDreg = 6,892 \cdot 53,03 + 0,0571 \cdot 4.050,475 - \frac{(53,03)^2}{5}$$

$$SQDreg = 34,484$$

$$SQDerr = SQDtot - SQDreg$$

$$SQDerr = 34,642 - 34,484 = 0,158$$

Ilustração da ANOVAR apenas para efeito de comparação com a ANOVA:

ANOVAR

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Regressão	1	34,484	34,484	656,31	0,0001
Erro	3	0,158	0,053		
Total	4	34,642			

Coeficiente de determinação:

$$r^2 = \frac{SQDreg}{SQDtot} = \frac{34,484}{34,642} = 0,995 = 99,54\%$$

$$V(m) = \frac{\sigma^2}{n} \quad \therefore \quad \sigma^2 = V(m) \cdot n$$

$$V(m) = \frac{s^2}{n} \quad \therefore \quad s^2 = V(m) \cdot n \quad \therefore \quad (s^2 \text{ estima } \sigma^2)$$

$$SQD(m) = \frac{SQD}{n} \quad \therefore \quad SQD = SQD(m) \cdot n \quad \therefore \quad (n = r)$$

$$SQDDreg = 34,484 \cdot 5$$

$$SQDDreg = 172,422$$

$$SQDDireg = 0,158 \cdot 5$$

$$SQDDireg = 0,788$$

ou

$$SQDDireg = SQDtra - SQDDreg$$

$$SQDDireg = 173,211 - 172,422 = 0,788$$

Coeficiente de determinação:

$$r^2 = \frac{SQDDreg}{SQDtra} = \frac{172,422}{173,211} = 0,995 = 99,54\%$$

Hipóteses:

$$H_0: |\beta_i| = 0$$

ou

$$H_0: Y \neq \alpha_0 + \beta_1 X$$

$$H_1: |\beta_i| > 0$$

ou

$$H_1: Y = \alpha_0 + \beta_1 X$$

ANOVA

Causa da variação	GL	SQD	QMD	F _{cal}	Pr
Bloco	4	0,014	0,004	0,132	0,9685
Tratamentos	(4)	(173,211)	43,303	1.580,533	< 0,0001
Dev. regressão	1	172,422	172,422	6.293,348	< 0,0001
Ind. regressão	3	0,788	0,263	9,599	0,0010
Resíduo	16	0,438	0,027		
Total	24	173,663			

Conclusão: rejeita-se H_0 ao nível de significância de 5% pelo teste F.Interpretação:

A equação ajustada explica significativamente as variações na matéria seca da parte aérea das plantas de milho, decorrentes das variações nas doses de fósforo, a 5% de probabilidade.

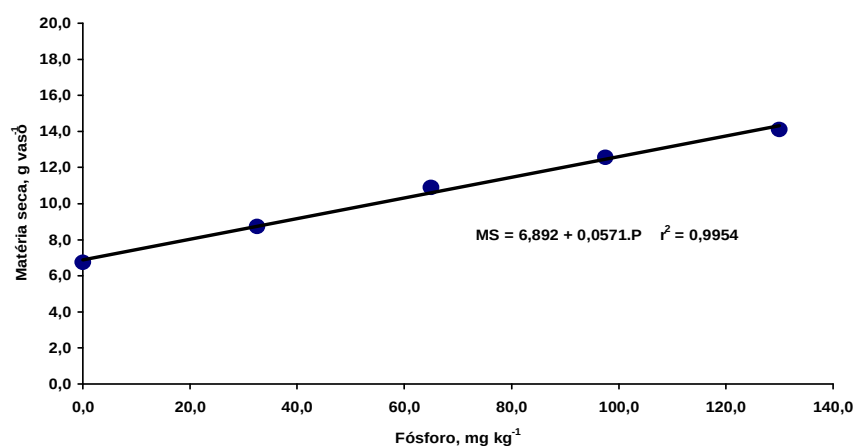


Figura 14.10 – Matéria seca da parte aérea das plantas de milho em função das doses de fósforo com ajuste de um modelo linear.

A falta de ajuste também foi significativa a 5% de probabilidade, implicando que se poderia tentar ajustar um outro modelo, mais adequado à natureza dos dados, como por exemplo o quadrático:

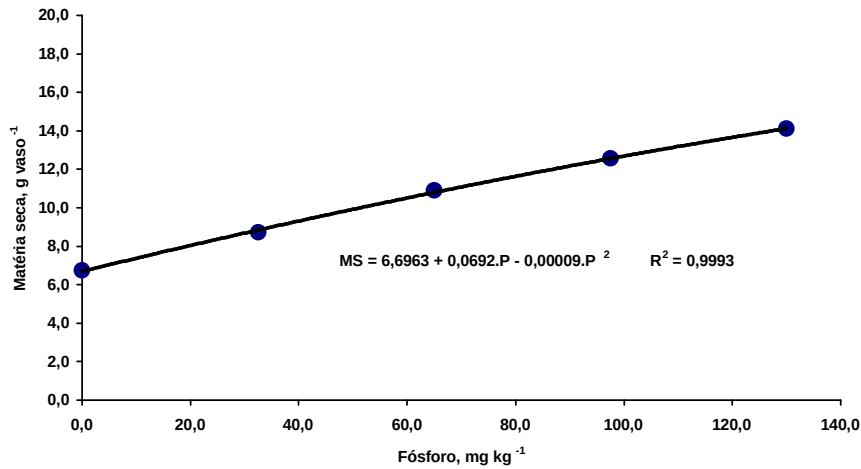


Figura 14.11 – Matéria seca da parte aérea das plantas de milho em função das doses de fósforo com ajuste de um modelo quadrático.

Apresentação final dos resultados para o modelo linear:

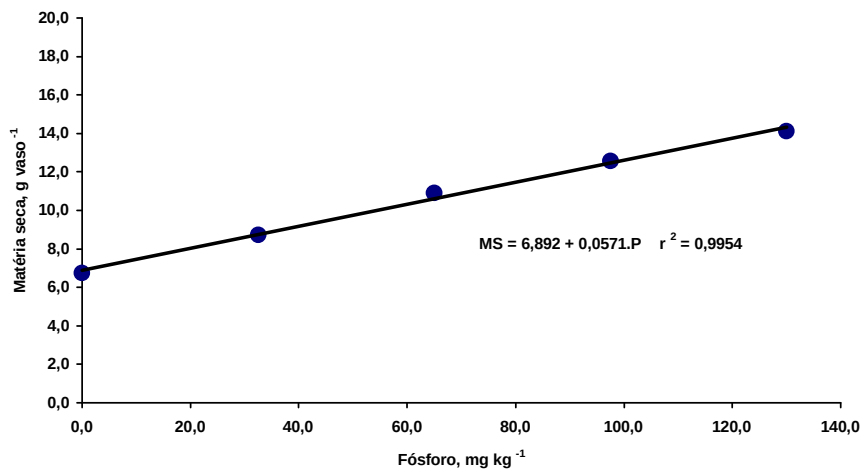


Figura 14.12 – Matéria seca da parte aérea das plantas de milho em função das doses de fósforo.

Quadro 14.6 – Análise de variância Matéria seca da parte aérea das plantas de milho em função das doses de fósforo

Causa da variação	GL	QMD	Pr
Bloco	4	0,004	0,9685
Tratamentos	(4)	43,303	< 0,0001
Dev. regressão	1	172,422	< 0,0001
Ind. regressão	3	0,263	0,0010
Resíduo	16	0,027	
Total	24		

15. Transformação de dados

15.1. Introdução

Em muitas situações, após o pesquisador ter coletado os dados, no início das análises estatísticas, verifica que os mesmos não atendem aos pressupostos requeridos pela análise a ser utilizada. Por exemplo, para realizar uma análise de variância (ANOVA) aos dados experimentais, são aplicados testes estatísticos preliminares para verificar a adequação, ou não, dos dados aos pressupostos desta análise. Quando esses pressupostos não são atendidos, uma das alternativas consiste na transformação dos dados originais em uma outra quantidade, de modo a que os pressupostos sejam, pelo menos em parte, ou no todo, atendidos. Este procedimento possibilita inferências mais adequadas e seguras que as que seriam obtidas a partir dos dados originais.

Uma vez transformados os dados a análise prossegue normalmente, ou seja, são realizados todos os cálculos sobre os valores transformados e feitas todas as inferências. Para a apresentação final dos resultados, entretanto, as médias de tratamentos devem ser apresentadas com seus valores originais, não transformados, pois os valores transformados representam quantidades abstratas.

15.2. Transformação angular

$$\text{arc sen} \sqrt{\frac{p\%}{100}}$$

15.2.1. Pressuposições

Dados provenientes de populações com distribuição Binomial (experimentos que apresentam apenas dois resultados: sucesso e fracasso) onde a variância está intimamente relacionada à média. Se forem retiradas amostras de várias distribuições binomiais, as médias dos tratamentos e as variâncias, não são independentes.

15.2.2. Uso

Homogeneizar a variância residual de dados de proporção

$$\frac{y}{n}$$

ou percentagens

$$100 \cdot \frac{y}{n}$$

15.2.3. Recomendações

Especialmente recomendada quando as porcentagens cobrem grandes amplitudes de valores. Se as porcentagens estiverem todas entre 30% e 70%, a transformação será desnecessária, pois ela produzirá sensíveis alterações nos valores que estiverem entre 0 e 30% e 70 e 100%:

- porcentagem de plantas doentes
- número de estacas enraizadas
- número de plantas não atacadas por determinada doença, etc.

15.1. Transformação raiz quadrada

15.1.1. Pressuposições

Dados provenientes de populações com distribuição Poisson, ou seja, experimentos em que se conhece apenas o número de sucessos

$$\sigma_y = \sigma_y^2$$

15.1.2. Uso

Homogeneizar a variância residual de dados e torná-la independente da média.

15.1.3. Recomendações

Especialmente recomendada quando os dados são provenientes de contagens:

- número de galhos secos em função de diversos adubos utilizados
- contagem de árvores doentes, acidentes ou defeitos, ervas daninhas
- número de bactérias por placa, plantas ou insetos em determinada área, etc

18.1.1. Dicas úteis

Quando nos dados ocorrem valores pequenos, inferiores a 10 e, principalmente, zeros (0) as transformações abaixo:

$$\sqrt{y + 0,5}$$

$$\sqrt{y + 1}$$

$$\sqrt{y} + \sqrt{y + 1}$$

estabilizam a variância mais efetivamente que $\sqrt{y}$.

15.2. Transformação Logarítmica

15.2.1. Pressuposições

Quando o desvio padrão na escala original varia diretamente com a média, ou seja, o coeficiente de variação é constante de tratamento para tratamento ou dados provenientes de populações com distribuição exponencial

$$\sigma_y = \sigma_y$$

15.2.2. Uso

Este tipo de relação entre média e desvio padrão é encontrado, geralmente, quando os efeitos são multiplicativos em lugar de aditivos. Nesta situação, tal transformação, além de estabilizar a variância residual, produz aditividade nos efeitos e tende a normalizar a distribuição dos erros.

15.2.3. Recomendações

Esta relação entre média e desvio padrão são freqüentes nos casos de:

- contagem do número de raízes por plântula, árvores por hectare e observações biológicas
- medição dos comprimentos totais de raízes por plântulas, etc.

15.2.1. Dicas úteis

Para números inteiros positivos que cobrem uma grande amplitude. Seria necessário uma transformação equivalente a

$$\sqrt{y}$$

para valores pequenos e a

$$\text{Log}(y)$$

para valores grandes de y. A transformação que mais se aproxima da desejada é

$$\text{Log}(x+1)$$

quando ocorrem zeros (0) ou valores negativos (< 1), pode-se adicionar um valor constante a cada observação da variável antes da transformação, de modo a tornar positivos todos os valores.

A base 10 para logaritmo é a mais utilizada, por conveniência, contudo, qualquer base é satisfatória.

16. Tabelas estatísticas