

Add-on Lecture 13

Results = is the last step in your research.

- In result you show us the product "نتاج" of all your efforts "جهود" , all the hassles "متاعب" , all the troubles "مشاكل" you have been through and present them to us.
- the result is consider to be the essence of your research , the main thing of your research.
- we do research cause we want to reach an outcomes , to reach conclusion , to reach some kind of result to help us for full our research objectives.

RESULTS IN GENERAL: (the structure of result)

THREE STATISTICAL THINGS TO DO WITH RESULTS

(a) Presentation :

- It's very important that you know how to present your result.
- In the result you try to make a thing very easy to reader ; specially when you use tables and graphs you need to explain them.
- if you have graphs and table you may Convert them into percentages "%" (ex ; in 6:15)

(b) Descriptive statistics : "التحليل الوصفي" (ex; in 7:22)

- in all result you always start with descriptive statistics (only describe the numbers of data)
- These are figures you (get the computer to) calculate from a lot of specific figures which arise from data. = use computer to calculate ;

فيه برامج مخصصه لحساب النسب والرسوم البيانيه نقدر نستعين فيها مثل الإكسل excel من مجموعة برامج المايكرو سوفت اوفس + برنامج احصائي اسمه spss

-- (b1) Measures of centrality :

- where you have cases that have been put in categories, is the category that the greatest proportion of people chose or fell in = the highest score that my participant selected.
- we want to know where our data are centered so we can describe them or we can depend on them to reach a result.

-- (b2) Measures of variation :

Standard Deviation = is how to separate the score from each other (score here is the participant answers ; are they close to each other or not)

- النقطة هذه راح تفهمونها اكثر بعد ما تسمعون المثال . Passed on the mean we can have a result.

-- (b3) Measures of difference:

- مهم جداً تسمعون المثال هنا 21:45 ex in

Level of significance : "مستوى الدلالة" in 24:40

It's should be something like ".05" It's called "p = Probability significance" → we get this number by computer (special programmes)

- If you get this number ".05" and less then it's called in research significance difference ← this is what we are looking for in research.
- If you get this number ".05" and above then NO!! it's not significance difference ; we can't say there is a different (ex in 27:37)

-- (b4) Measures of relationship : Ex in 29:37

**These quantify the amount of relationship between two (or more) variables as measured in the same group of people or whatever. بهذا الجزء هذه الجملة فقط هي المطلوبه معنا والباقي قراءه لزيادة المعلومات.

(c) Inferential statistics :

It's deal with p value >> هنا ما تكلم عنها كثير قال ان اللي يأخذون كلاسات الستاتيك راح يعرفون الاسماء والخ والخ

the level of certainty :

is about what inferential stats tells you that you will be satisfied with.

** No inferential stats give you 100% certainty of anything → after we do research we can't be 100% certain that are these result 100% perfect NO!! ; you should have some margin of errors and this calculated to be 5% at least in your research ... you can be 95% that your result are good.

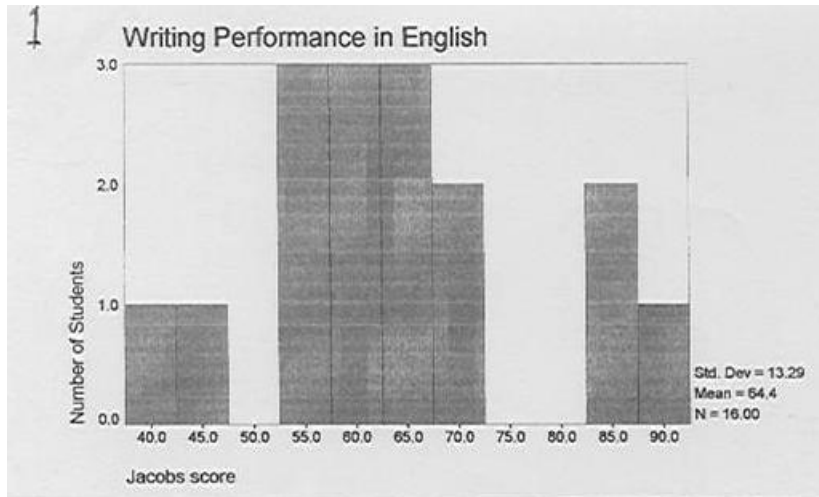
** 95% is commonly taken as adequate in language research: this is the same as choosing the .05 (or 5%) level of significance as the one you will be satisfied with. → 95% is the maximum that we can reach in certain of ore research.

Significance tests :

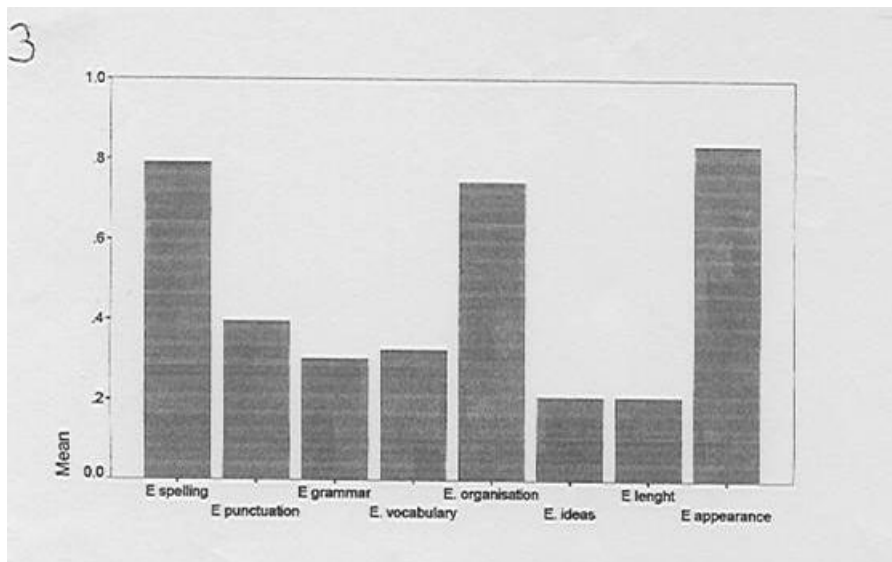
We use it for just to give us a poised or confidence that's we can say (yes we are have a real difference or a real relationship) << "How??" through the Level of significance

**** <http://privatewww.essex.ac.uk/~scholp/onevardesc.htm>**

من الرابط هذا نزل على طول للرسم البياني وجلس يوضح عليه كذا نقطه في الدقيقه ٤٠ تقريباً



وتكلم عن هذا في الدقيقه ٤٥



تم بفضل الله .. اللهم اكتب لنا توفيقاً من عندك ♥

The End ..

اضافته اخيره .. هنا نسخه لمحتوى الرابط المضاف في الشريحة رقم ٧ للي حاب يطبعه ^_^

One variable: graphs and descriptive statistics

WHEN DO YOU NEED THEM?

When does one ever want to look at results for just one variable? True, most classical research involves questions/hypotheses that entail looking at relationships between at least two variables. But here are some common situations where you want to look at graphs and descriptive statistics for cases on one variable:

- Some studies look directly at more or less all of an entire population of interest, and have questions about single variables, so the results just need to be given as graphs and descriptive statistics. E.g.
 - A census survey to find out how many people in Wales speak Welsh shows that 20.8% claim to.
 - A teacher has the feeling (hypothesis) that her vocab teaching (or the students' learning of vocab) is not very successful. She tests her class to see how much of the vocabulary of the last ten lessons they have learnt, as part of an action research project to improve her vocabulary teaching.
- In any study, however many variables are involved, it is often valuable to describe our cases in terms of each of a set of variables that are not central to the investigation, as part of your control of unwanted factors (cf CVs and SAMPLING). E.g.
 - You are interested in the opinions of Greek learners in private schools in Greece about their English course materials, so you send 50 questionnaires to your friend who works in one to distribute for you. You get thirty-one back. On the questionnaire, apart from questions about the central variables of your study, you might ask questions to elicit their first language, age, gender, experience of English outside school etc. You then check each of these separately to see if it suggests your sample is unusual in any way (?unexpectedly many girls), has odd cases in it (?two who say their first language is Bulgarian) etc. You might well display proportions and graphs for genders, age groups, respondents vs non-respondents etc. as part of your report on your subjects.
- In any study, however many variables are involved, you may be in the position of deciding groups of cases on the basis of information gathered about them, rather than in advance. E.g.
 - In the above example you might actually want to make up groups out of your subjects, using their questionnaire responses, to use as EVs. Examining the 'experience of English outside school' variable you find there are ten who have been abroad to English speaking countries, so you might set up a two category EV on this basis to see if opinions about course materials relate to having/not having this experience at all. If so, would you have any hypothesis about the answer? Anyway, you might display a graph or table showing the proportions.
- In any study, however many variables are involved, it is often valuable to look at results for all cases on each dependent variable or condition separately and/or in each group separately as well as doing what is necessary to establish relationships between variables. E.g.
 - In our example of gender and attitude to RP, with the hypothesis that there is an attitude difference between genders. Apart from doing the relevant two-variable graphs and statistics one would do well to explore the data with histograms for each group separately and the whole set of

subjects as if one group. One does not necessarily report in the write-up every statistic or graph one calculates if it does not prompt anything interesting to say about it, but even statisticians often comment that researchers often 'don't look at their data enough, but just want to do a significance test and get on to the next thing'.

DECISIONS ABOUT THE RIGHT GRAPHS AND DESCRIPTIVE STATISTICS

If you are not into one variable inferential stats (which we are omitting here), the choices to be made are simple:

Scale type of Variable:

	Interval	Rank order	Categories of any sort	Counts in a continuum
Graphic Presentation:	<i>histogram of scores, frequency polygon</i>	ordered list of cases	<i>bar chart, pie chart (of frequencies or percent)</i>	single bar
Centrality statistic:	<i>mean, median score, modal score</i>	median rank	<i>modal category</i>	frequency
Variation statistic:	<i>standard deviation, range</i>	quartile deviation	index of commonality	

- Only the italicised ones are commonly met and will be dealt with here.
- The modal category is simply the one with the most cases in it - the most popular one.
- The mean (denoted by M or X-bar) is what we usually call the average in everyday English.
- The standard deviation (SD) is a measure of the spread of scores. Roughly it is the average of the differences between each score and the mean score (see any stats book for the formula). So if everyone in a group scores the same, which will be the mean for the group, then the average of the differences of each score from the mean is 0 (SD=0; no variation). The more each score differs from the mean, the higher the SD gets, indicating more variation or 'disagreement' in the group. Usually one 'wants' small SDs.
- Similar concepts to SD, calculated in various ways, are called 'error' and 'variance' in statistics.

Joke from WWW: Most of us have A Greater Than Average Number of Legs

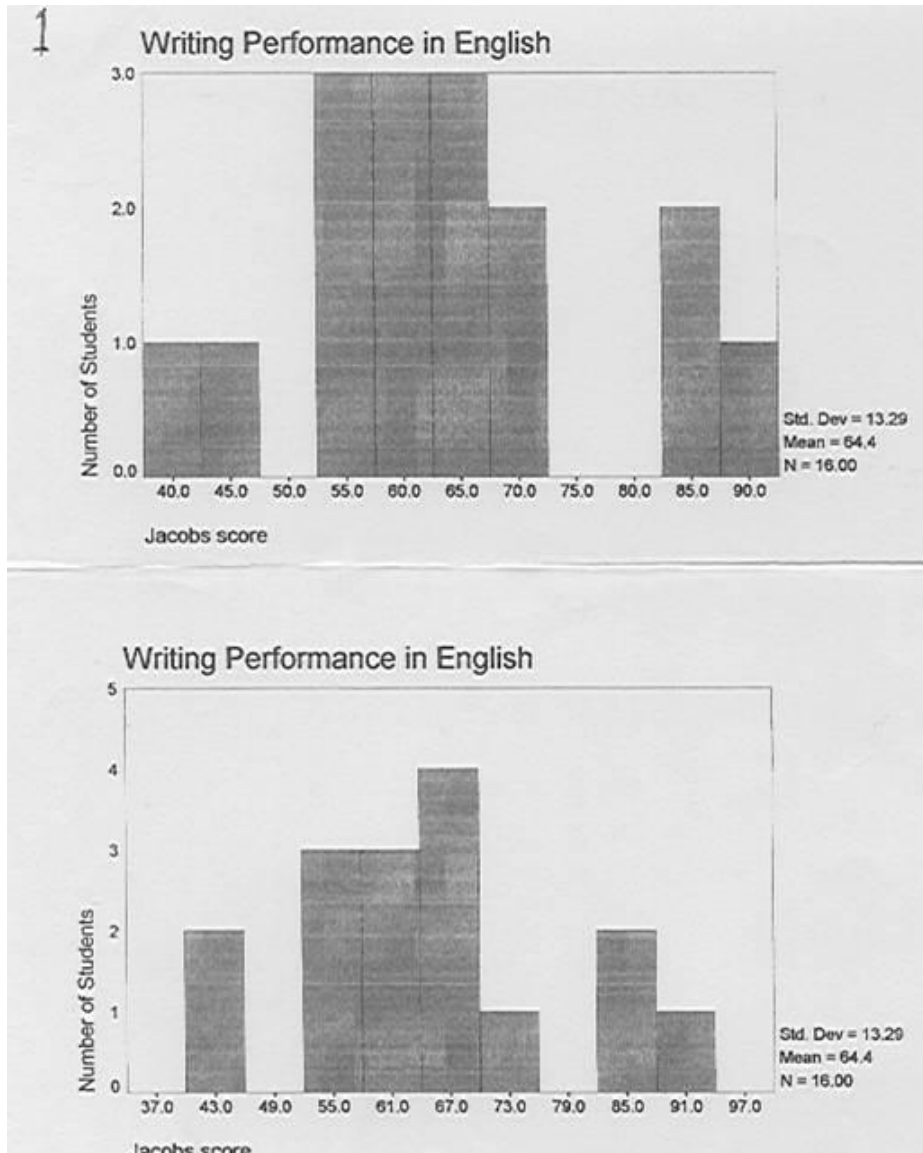
The great majority of people have more than the average number of legs. Amongst the 57 million people in Britain there are probably 5,000 people who have only one leg. Therefore the average number of legs is

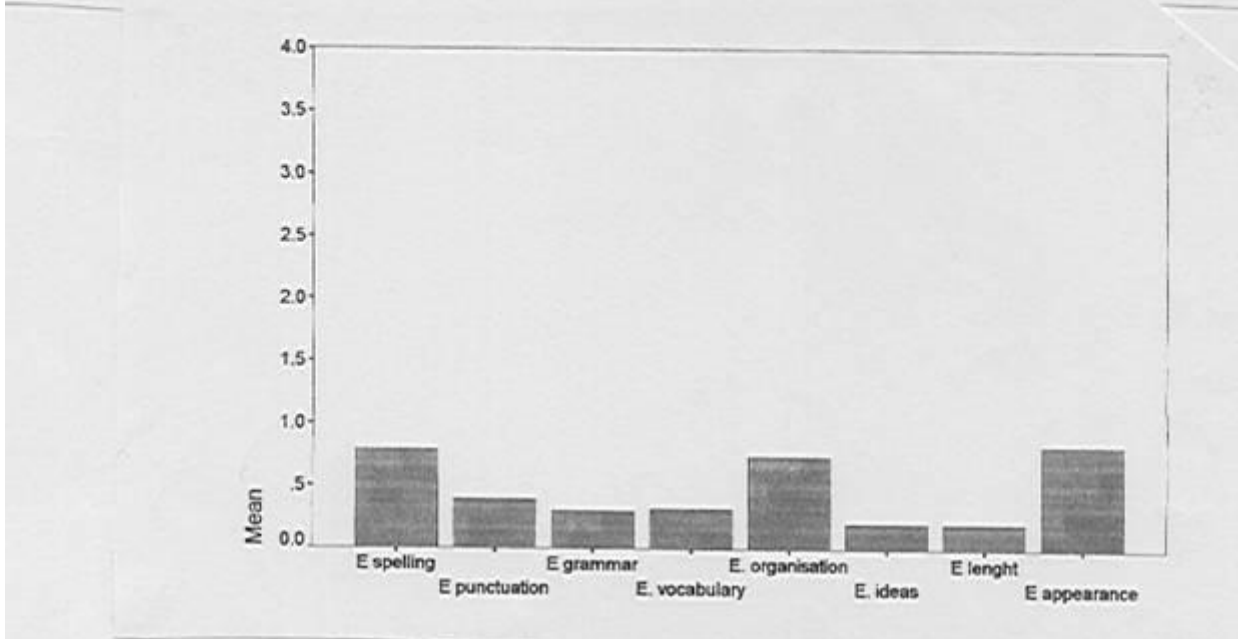
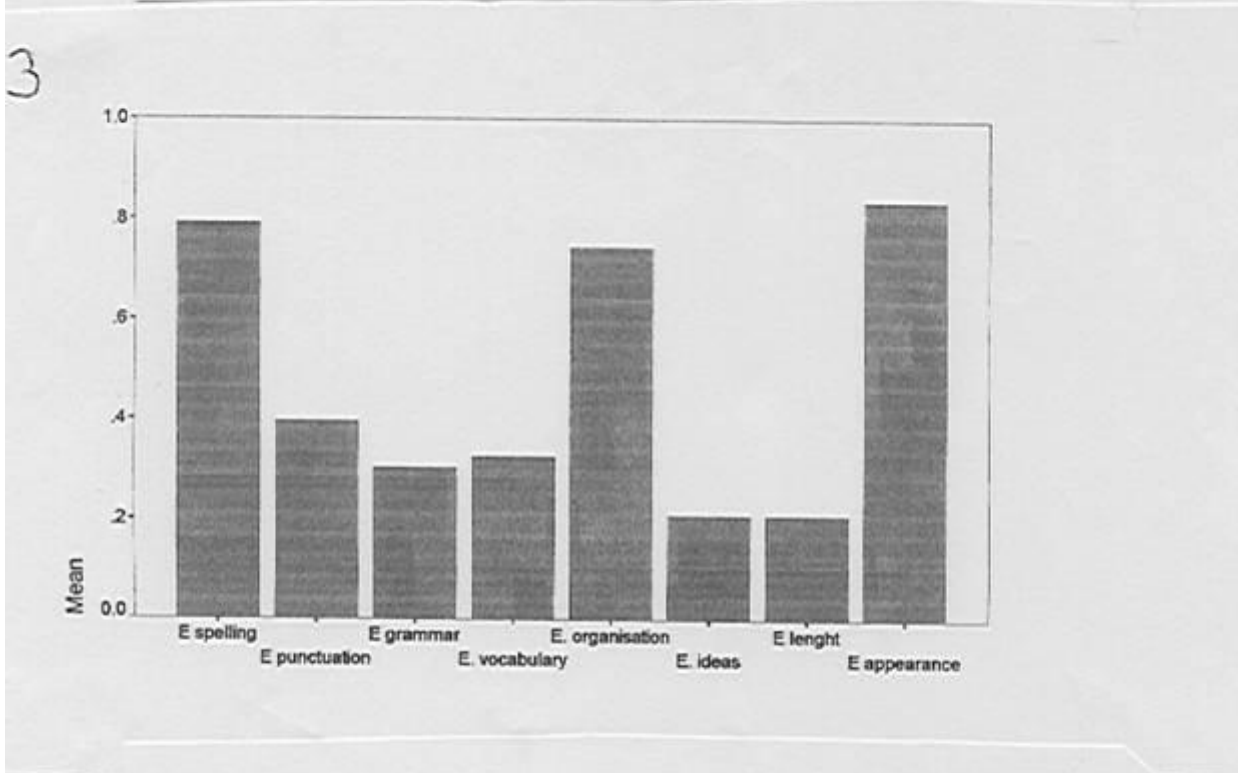
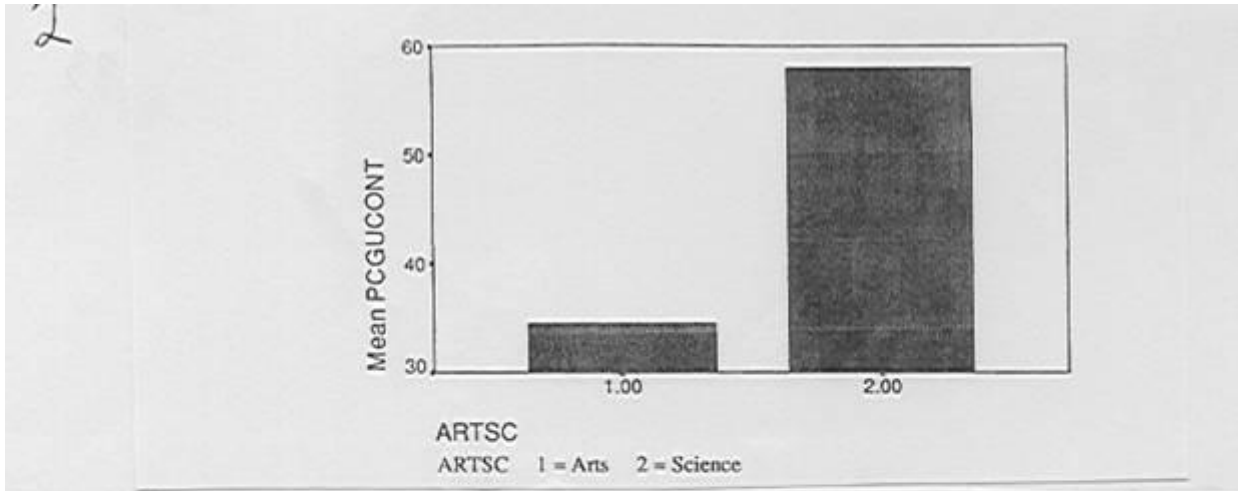
$$((5000 \times 1) + (56,995,000 \times 2)) / 57,000,000 = 1.9999123.$$

Since most people have two legs... need I say more?

A FEW EXAMPLES OF SIMPLE GRAPHS: HOW TO MISLEAD

- Shows two versions of a **histogram** of results for one group of 16 learners on one variable ('interval' scores for quality of each subject's written composition). Which version is better and why, or is neither optimal? What distinguishes a histogram from a bar graph/chart (seen in 2)? When to use each?
- Is a **bar graph** (=bar chart) showing the broad subject specialism of participants in a study. I.e. it displays how all cases are categorised on a two category variable. How would you improve it for inclusion in a write-up?
- Shows two bar graphs for the same data – a set of several mean scores. One group of learners has given their ratings (on a five point scale) of how much they think eight different aspects of their compositions improved when done by word processing. The average ratings for each of these 8 variables are displayed together. Which is the better version and why?





SIMPLE PERCENTAGES: HOW TO MISLEAD

1) Which sounds more impressive, A or B?

A) 2 out of 4 subjects agreed B) 50% of subjects agreed

A) 40 out of 80 subjects agreed B) 50% of subjects agreed

OK, but which result would you actually trust more? How should one report such results?

2) What is unclear? How to restate this better?

In our survey we polled 50 people, though 10 declined to participate. 60% said yes to the question 'Do you like the English class?'...

3) Percentage scores versus group/aggregate percent.

Two ways of handling data arising from different numbers of potential occurrences for different people. Imaginary example of data where three subjects have been recorded in quasi-natural conversation, and counts have been made of their NS-like/correct use of third person –s. Why do the percent differ in A and B? Which would the statistician prefer and why?

A) Analysis with subjects as cases: percentage scores and their mean

Case	Correct	Incorrect	Total occurrences	Percent correct score	Mean percent correct
Learner 1	12	12	24	50	
Learner 2	8	12	20	40	
Learner 3	3	9	12	25	
Total	23	33	56		

B) Analysis with occurrences as cases: group percent

		Total frequency	Percent
	Correct	23	41.1%
	Incorrect	33	58.9%
	Total occurrences	56	