

# Summary statistics for binary data

Tabriz University of Medical  
Sciences  
Standard Workshop on Systematic  
Reviews \_ October 2012

Dr. Shayesteh Jahanfar,  
University of British Columbia

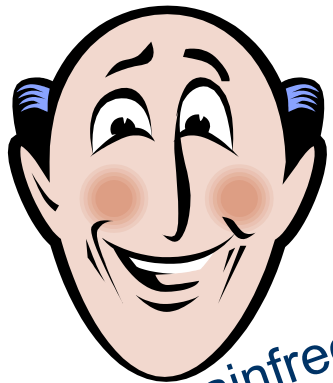
# Outline

- identify binary outcomes
- be familiar with ways of expressing chance of an event when using binary outcomes
- understand how to express and interpret the relative chance of an event when comparing groups
- select effect measure



# What is a binary outcome?

- e.g. dead or alive, pain free or in pain, smoking or not smoking
- each participant is in one of two possible, mutually exclusive, states



painfree



in pain

# What were the chances of that?

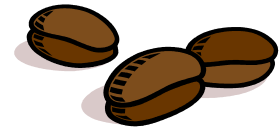
- risks and odds are just ways of expressing chance in numbers
- for binary events, they just express the chance of being in one of the two states



# Risk

- 24 people drank an espresso, and 6 fell asleep
- risk of falling asleep  
= 6 asleep/24 who could have fallen asleep  
=  $6/24 = 1/4 = 0.25 = 25\%$

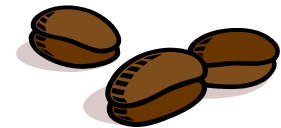
risk = number of events of interest  
total number of observations



# Odds

- 24 people drank an espresso, and 6 fell asleep
- odds of falling asleep
  - = 6 asleep/18 did not fall asleep
  - =  $6/18 = 1/3 = 0.33$  (not usually expressed as %)

odds = number of events of interest  
number without the event



# Expressing it in words

- risk
  - the chances of falling asleep were one in four, or 25%
- odds
  - the chances of falling asleep were one third of the chances of staying awake
  - one person fell asleep for every three that stayed awake
  - the chances of falling asleep were 3 to 1 against

# Do risks and odds differ much?

## 2 examples from caffeine trials

- 130 people 'still awake' out of 164
- chance of still being awake

➤ **risk =  $130/164 = 0.79$ ;                      odds =  $130/34 = 3.82$**

- 4 people with 'headaches' out of 63
- chance of having a headache

➤ **risk =  $4/63 = 0.063$ ;                      odds =  $4/59 = 0.068$**





# Comparing groups – 2x2 table

|                                                                                   |        |       |                     |
|-----------------------------------------------------------------------------------|--------|-------|---------------------|
|  | Asleep | Awake | Total<br>(by group) |
| Caffeine                                                                          | 12     | 48    | 60                  |
| Decaf                                                                             | 16     | 33    | 49                  |
| Total<br>(by event)                                                               | 28     | 81    | 109                 |

➤ to express the relative chance of an event

# Meta-analysis of binary data

- calculate a single summary statistic to represent the effect found in each study
- 3 options
  - risk ratio (relative risk)
  - odds ratio
  - risk difference



# Risk ratio

- risk of event on treatment  
= **12/60**
- risk of event on control  
= **16/49**

|          | Asleep    | Awake | Total     |
|----------|-----------|-------|-----------|
| Caffeine | <b>12</b> | 48    | <b>60</b> |
| Decaf    | <b>16</b> | 33    | <b>49</b> |
| Total    | 28        | 81    | 109       |

- risk ratio =  $\frac{\text{risk on treatment}}{\text{risk on control}}$   
=  $\frac{\mathbf{12/60}}{\mathbf{16/49}} = \frac{0.2}{0.327} = 0.61$

Where risk ratio = 1, this implies no difference in effect



# Expressing risk ratios in words

- risk ratio 0.61
  - the risk of falling asleep on treatment (caffeine) was about 61% of the risk on placebo (decaf)
  - caffeine reduced the risk to about 60% of what it was
  - the risk of falling asleep on caffeine is 39% lower compared to decaf
  - caffeine reduced the risk by 39%



# Odds ratio

- odds of event on treatment

$$= 12/48$$

- odds of event on control

$$= 16/33$$

|          | Asleep | Awake | Total |
|----------|--------|-------|-------|
| Caffeine | 12     | 48    | 60    |
| Decaf    | 16     | 33    | 49    |
| Total    | 28     | 81    | 109   |

- odds ratio =  $\frac{\text{odds on treatment}}{\text{odds on control}}$   
=  $\frac{12/48}{16/33} = \frac{0.25}{0.485} = 0.52$

Where odds ratio = 1, this implies no difference in effect



# Expressing odds ratios in words

- odds ratio 0.52
  - caffeine reduced the odds of falling asleep to 52% of what they were
  - the odds of falling asleep on caffeine is 48% lower compared to decaf
  - caffeine reduced the odds by 48%



# Risk difference

- risk of event on treatment

$$= 12/60$$

- risk of event on control

$$= 16/49$$

|          | Asleep | Awake | Total |
|----------|--------|-------|-------|
| Caffeine | 12     | 48    | 60    |
| Decaf    | 16     | 33    | 49    |
| Total    | 28     | 81    | 109   |

- risk difference = risk on control - risk on treatment  
=  $16/49 - 12/60 = 0.327 - 0.2 = 0.127$
- usually expressed as a %, 13%

# Expressing risk difference in words

- risk difference 13%
  - caffeine reduced the risk of falling asleep by about 13 percentage points





# Number needed to treat

- this is often expressed as how many we expect to treat, on average, before one extra person is helped
- $NNT = 1/RD$
- e.g.  $= 1/0.127 = 8$  (*round up* to whole people)
- we would need to give 8 people caffeine to keep one extra person from falling asleep
- not used directly for meta-analysis as there is no useful variance formula



# Choosing the effect measure

Criteria to consider when selecting a summary statistic

1. communication of effect
2. consistency of effect across studies
3. mathematical properties



# Summary

|               | OR | RR | RD |
|---------------|----|----|----|
| Communication | -  | +  | ++ |
| Consistency   | +  | +  | -  |
| Mathematics   | ++ | -  | -  |



# Take home message

- risks and odds are just ways of expressing chance
- risk ratios and odds ratios are ways of comparing chances in more than one setting/group
- RR and OR differ when the event is common



# Take home message

- risk difference shows the amount of change from baseline in absolute terms
- NNT communicates how many people would need to be treated for one extra to be helped
- ALL these estimates of treatment effect are uncertain, and should be presented with a confidence interval



# Summary statistics for continuous data

# Outline

- identify continuous outcomes
- understand how to summarise continuous data and pool studies with:
  - measures on the same scale
  - measures on different scales
- recognise some of the challenges of continuous data



# Types of data

- Binary data
- Counts of infrequent events (e.g. number of strokes)
- Short ordinal scales (e.g. pain grades: none/mild/moderate/severe)
- Long ordinal scales (e.g. disability scales)
- Continuous data (e.g. blood pressure)
- Censored data (e.g. survival times)





# What are continuous data?

- data with an infinite number of values that are equally spaced
- example: height - it can be measured along a numerical continuum of centimetres, metres or inches, feet
  - a person can be 175.24678cm tall, assuming the measurement instrument is accurate enough
  - the difference between 160 and 161cm, and 180 and 181cm, is the same



# Long ordinal scales

- sometimes treated as continuous data
- but not true continuous because
  - they have a finite number of distinct values
  - there are gaps in the continuum
- have multiple, ordered categories which imply magnitude
  - e.g. one category is greater or lesser than another
- spacing between categories is not numerically equivalent
- approach 'continuous' with increasing categories



# What continuous data can we combine?

- data represent continuous measures
- the mean value is in the middle (distribution is roughly symmetrical)
- measurements are made on all participants (not censored or survival type data)
- data are available for both groups in each trial



# What data is needed?

|           | Mean  | SD     | Sample size |
|-----------|-------|--------|-------------|
| Treatment | $m_t$ | $sd_t$ | $n_t$       |
| Control   | $m_c$ | $sd_c$ | $n_c$       |



# Meta-analysis of continuous data

- calculate a single summary statistic to represent the effect found in each study
- Summary statistics combined in meta-analysis
- 2 options
  - mean difference
  - standardised mean difference



# Mean difference

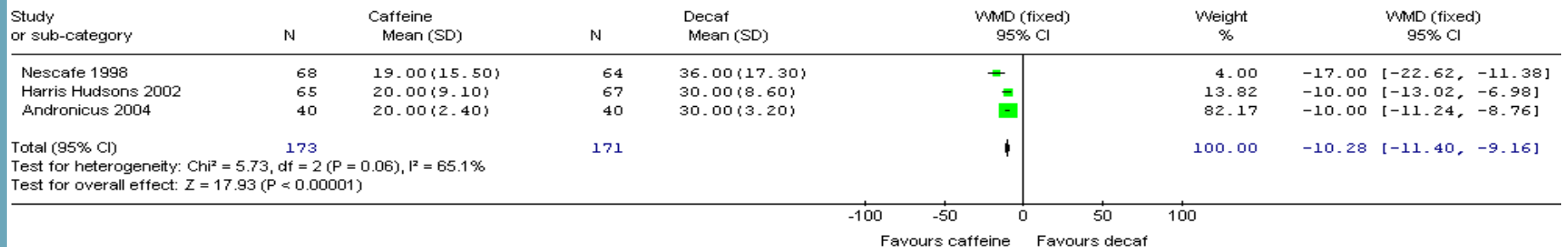
- outcomes measured in same unit using same scale (e.g. blood pressure as mmHg)
- pooled analysis in “natural units” and therefore easy to interpret
- studies weighted according to the inverse of the variance (a function of size and SD)

**MD = mean on treatment – mean on control**



# Mean difference: example

Review: Caffeine for daytime 'sluggishness'. (version with data)  
 Comparison: 01 Caffeinated Coffee versus Decaffeinated Coffee  
 Outcome: 03 Irritability at 30 minutes - INAS scale (1-50, high score worse)



# Standardised mean difference

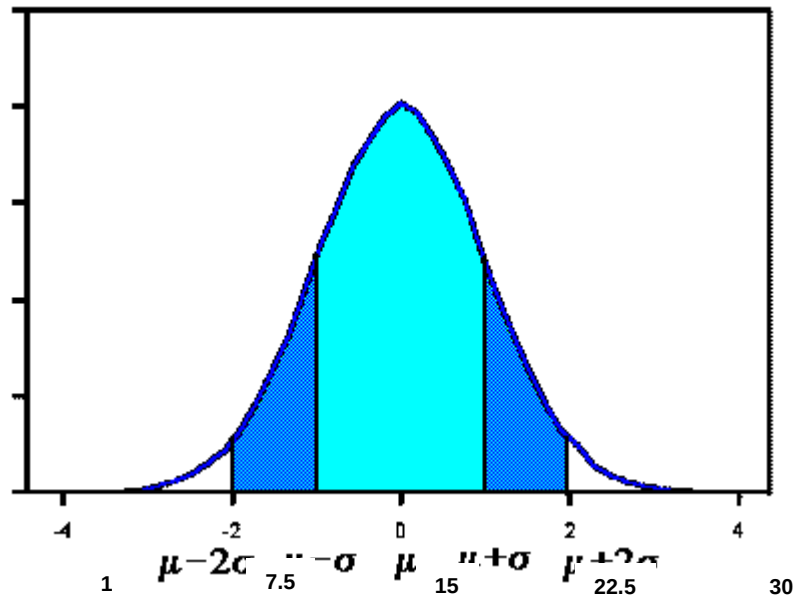
- Outcome is same concept measured on different scales, the values must be transformed to a common scale before pooling
- Sometimes scale factors are known and transformations are made directly (e.g weight)
- Standardised mean difference calculated as:

$$\frac{\text{Difference in means between groups}}{\text{Average standard deviation}}$$

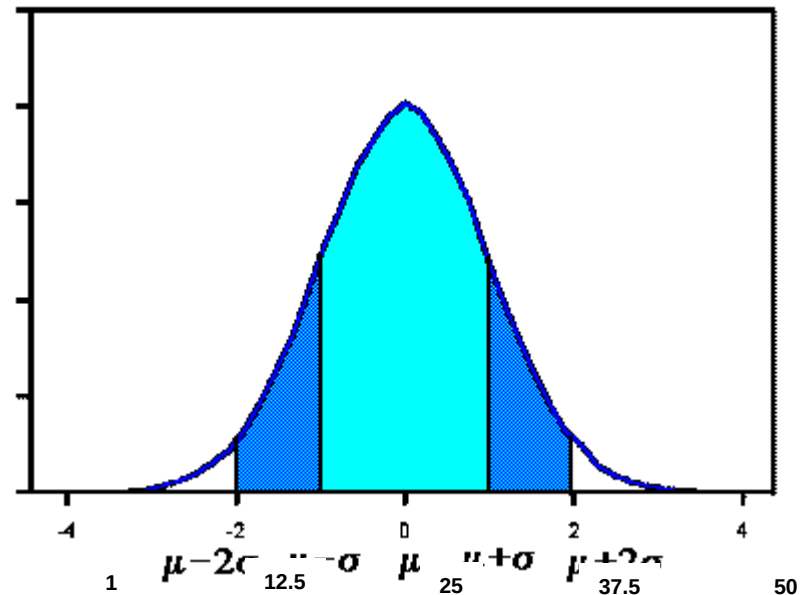


# Standardised mean difference

Beck Irritability Scale (1-30)



Irritability Negativity Affectivity Subscale (1-50)



Different scales but averages mean the same thing  
(i.e. average person is just as irritable!)

# Measurements on different scales

Comparing irritability at 30 minutes between caffeinated coffee and decaffe coffee

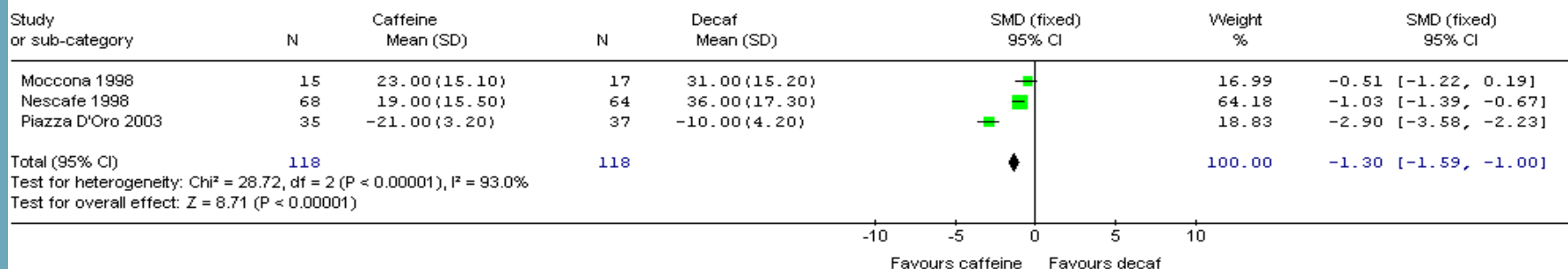
| Trial             | Caffeinated<br>N. mean (SD) | Decafe<br>N. mean (SD) | Irritability<br>scale |
|-------------------|-----------------------------|------------------------|-----------------------|
| Moccona 1998      | 15 23.0 (15.1)              | 17 31.0 (15.2)         | INAS                  |
| Nescafe 1998      | 68 19.0 (15.5)              | 64 36.0 (17.3)         | INAS                  |
| Piazza D'oro 2003 | 35 21.0 (3.2)               | 37 10.0 (4.20)         | BII                   |

*High scores on the Beck Irritability Scale (BII) (1-30) good outcomes, while high scores on the Irritability Negative Affectivity Subscale (INAS) (1-50) are poor outcomes*



# SMD: example

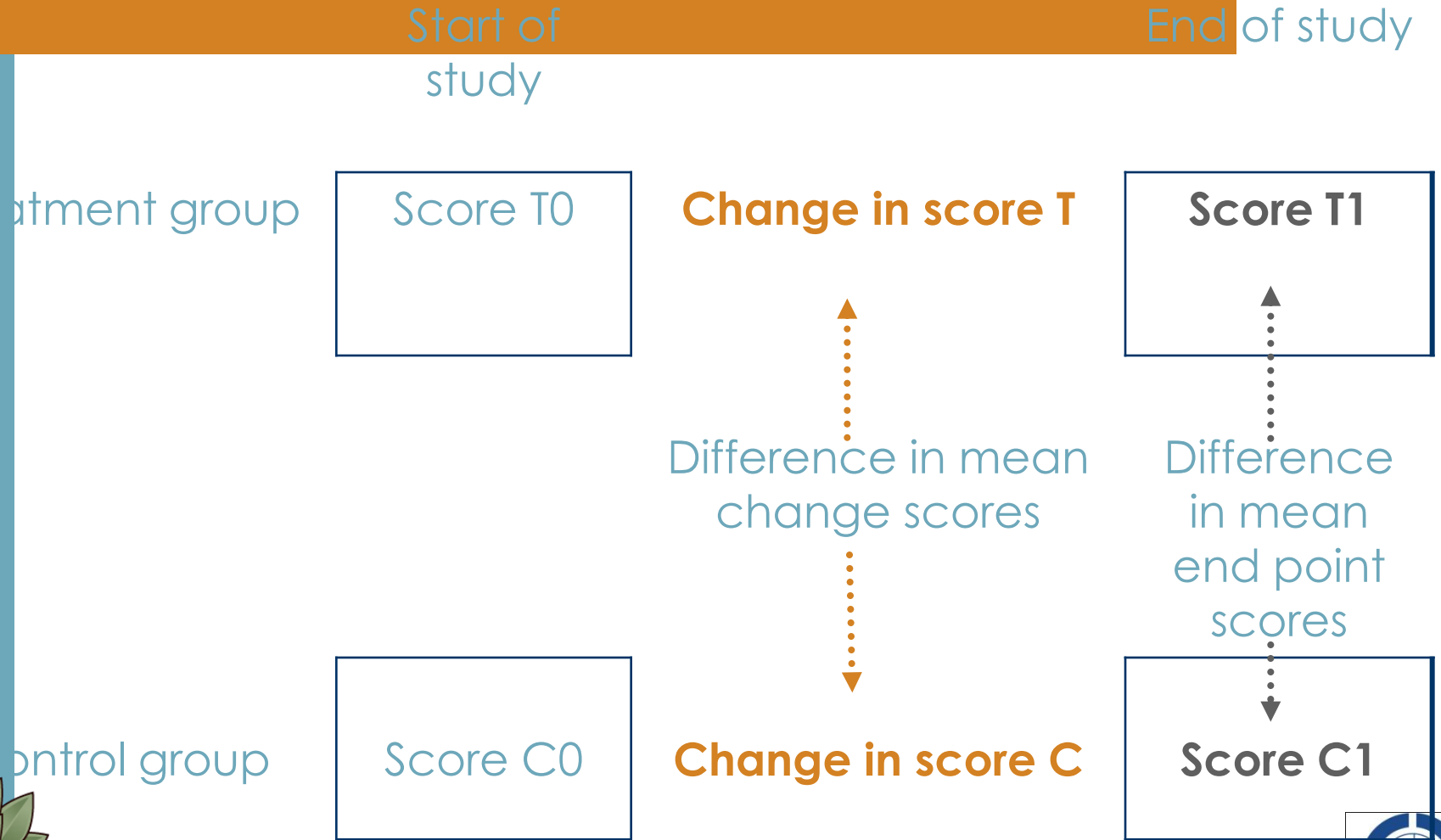
Review: Caffeine for daytime 'sluggishness'. (version with data)  
 Comparison: 01 Caffeinated Coffee versus Decaffeinated Coffee  
 Outcome: 06 Irritability at 30 minutes



# RevMan exercise



# Change vs endpoint scores



# Problems with MD and SMD

- what constitutes a clinically important change?
- restrictive eligibility criteria results in smaller standard deviations; therefore these trials given more weight
- mean difference
  - measurements on the same scale are not always comparable (e.g. health care costs in different places, process of care measures)
- standardised mean difference
  - difficult to interpret outcomes in units of SD, but can transform back to units of the scale
  - estimates of variation may not always be comparable making the SD a poor scaling factor



# Take home message

- pooling continuous data – use mean difference or standardised mean difference
- check data for skewness
- can calculate SDs from other statistics
- can use either endpoint or change scores



# Heterogeneity



# Outline

- what is heterogeneity?
- causes of heterogeneity
- identifying heterogeneity
- dealing with heterogeneity
- fixed and random effects meta-analysis



# What is heterogeneity?

- Heterogeneity is variation between the results of a set of studies



# Causes of heterogeneity:

## clinical

*Differences between studies with respect to:*

- participants
  - conditions under investigation, eligibility criteria for trials, geographical variation
- interventions
  - e.g. type of drug, intensity, dose, duration, mode of administration, experience of practitioners, nature of control (placebo, none, standard care)
- outcomes
  - e.g. type, follow-up duration, ways of measuring outcomes, definition of an event



# Causes of heterogeneity: methodological

*Differences between studies with respect to:*

- design
  - e.g. randomised vs non-randomised, parallel group vs crossover vs cluster randomised, length
- conduct
  - e.g. allocation concealment, blinding, approach to analysis, imputation methods for missing data



# Statistical heterogeneity

- excessive variation in the results of studies above that expected by chance



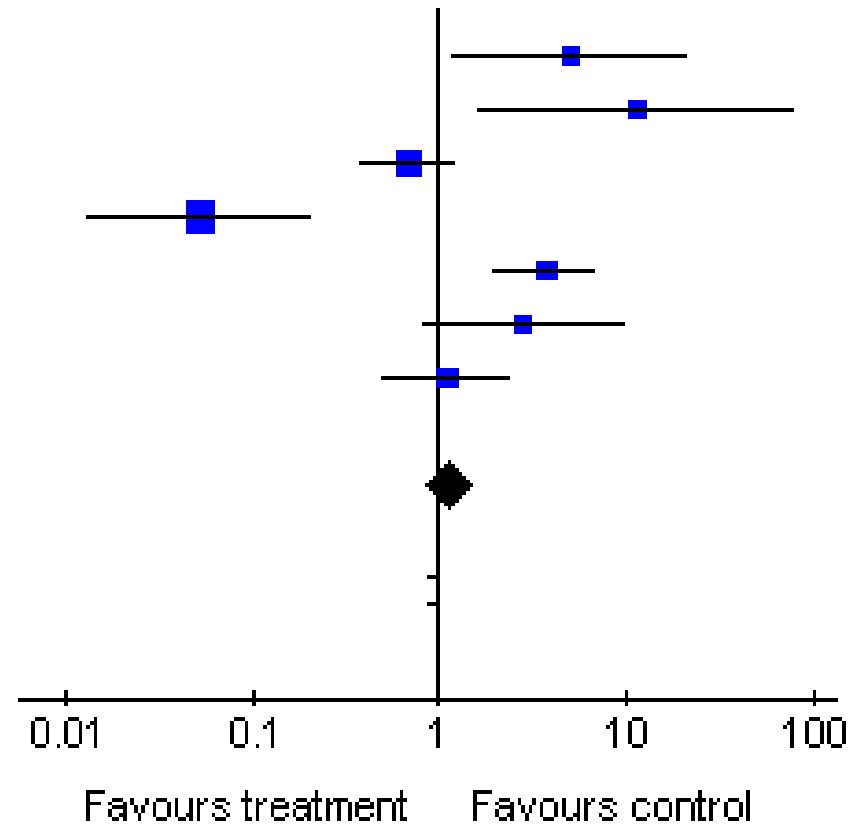
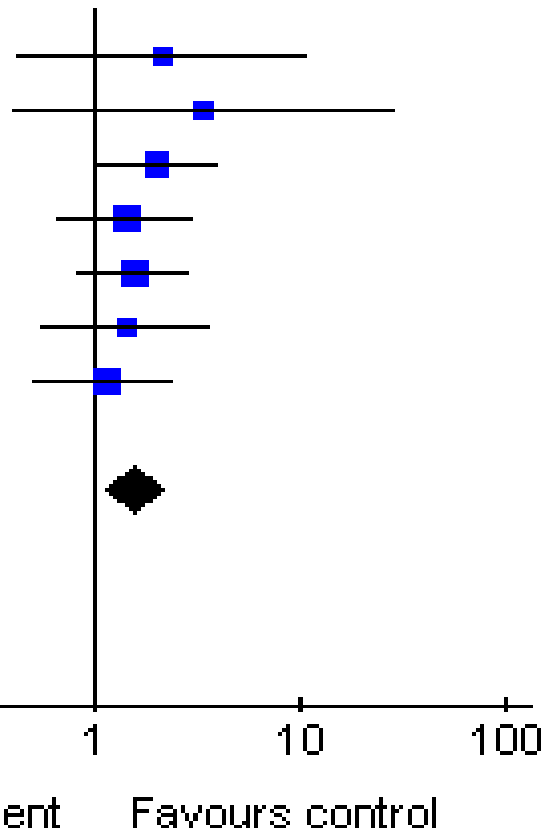
# Identifying heterogeneity

1. graphically – the eyeball test
2. numerically – the  $I^2$  test



# Forest plot A

# Forest plot B



# Quantifying heterogeneity

- $I^2$  describes the proportion of total variation across studies that is due to heterogeneity rather than chance
- based on Cochran Q test and its degrees of freedom
- $I^2 = \frac{(Q - df)}{Q} \times 100\%$  (df = the number of studies minus 1)



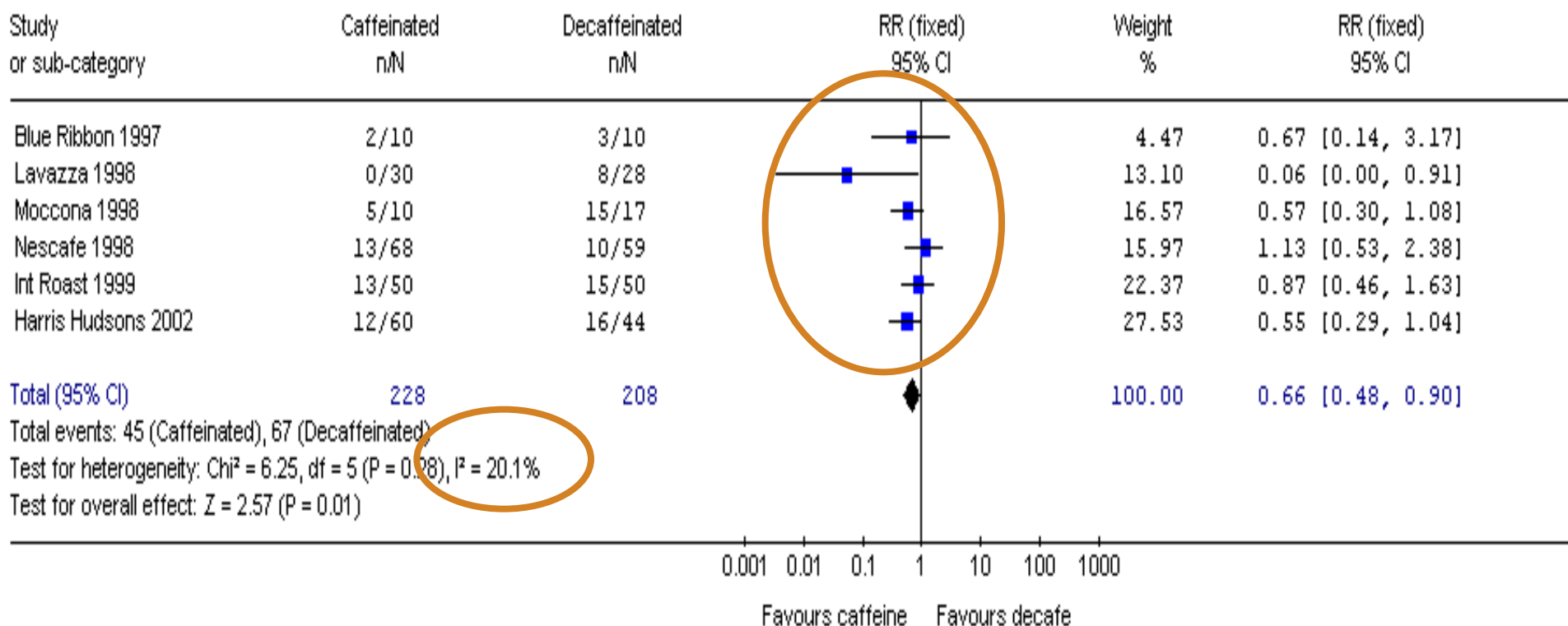


# Quantifying heterogeneity

- low (and negative) values of  $I^2$  indicate no, or little, heterogeneity
- larger values of  $I^2$  show increasing heterogeneity
- roughly, values of 25%, 50% and 75% correspond to low, moderate and high levels of heterogeneity (Higgins et al 2003, BMJ)



Review: Caffeine for daytime 'sluggishness'. (version with data)  
 Comparison: 01 Caffeinated Coffee versus Decaffeinated Coffee  
 Outcome: 09 Asleep at the end of the lecture



# Dealing with heterogeneity

Options available to you:

1. check the data
2. don't pool studies
3. ignore heterogeneity: use fixed effect model
4. investigate reasons for heterogeneity
5. incorporate heterogeneity: use random effects model



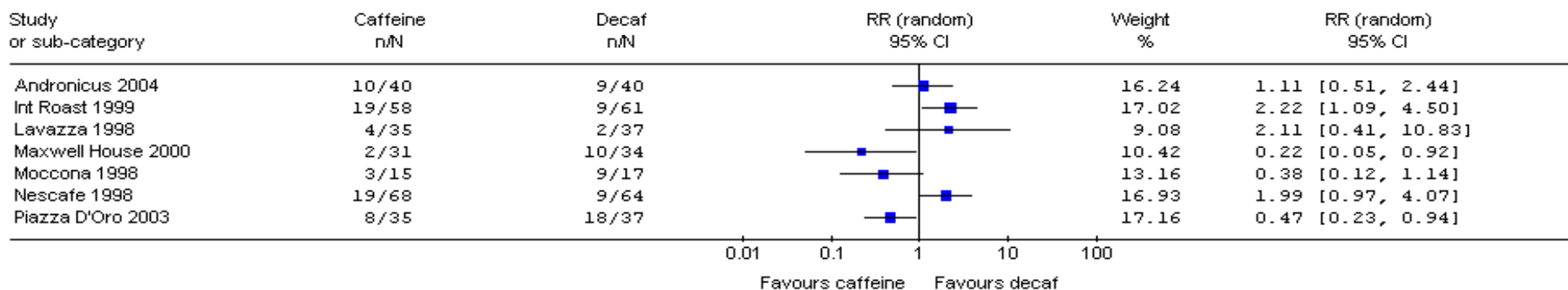
## Option 1: Check the data

- Check extracted data
- Check analyses of individual studies



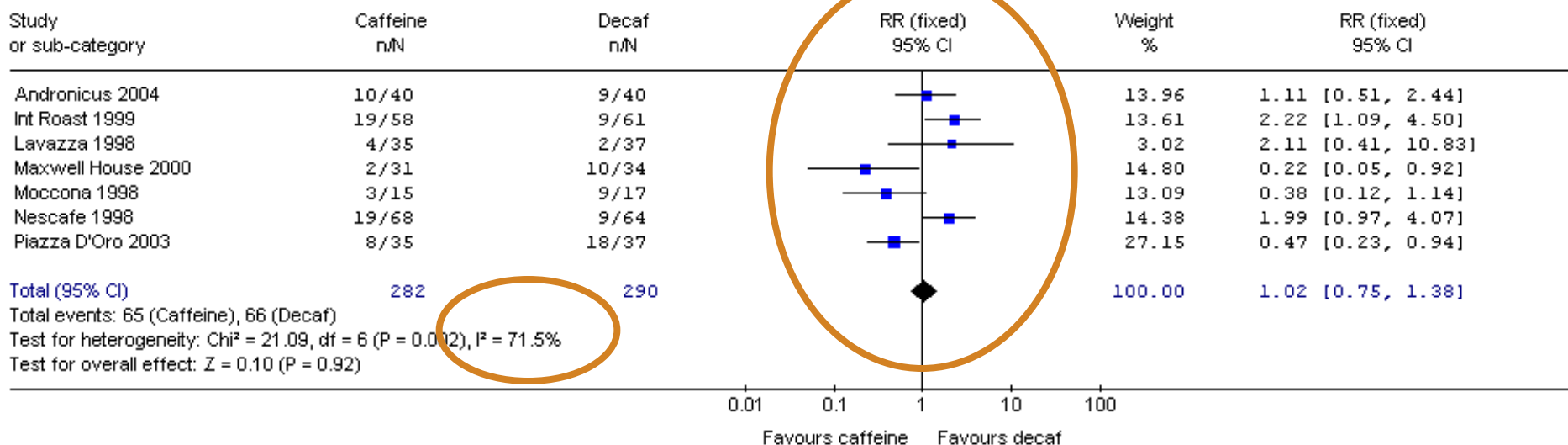
# Option 2: Don't pool studies

Review: Caffeine for daytime 'sluggishness'. (Version 251105)  
 Comparison: 01 Caffeinated Coffee versus Decaffeinated Coffee  
 Outcome: 02 Headache



# Option 3: Ignore heterogeneity

Review: Caffeine for daytime 'sluggishness'. (Version 251105)  
 Comparison: 01 Caffeinated Coffee versus Decaffeinated Coffee  
 Outcome: 02 Headache



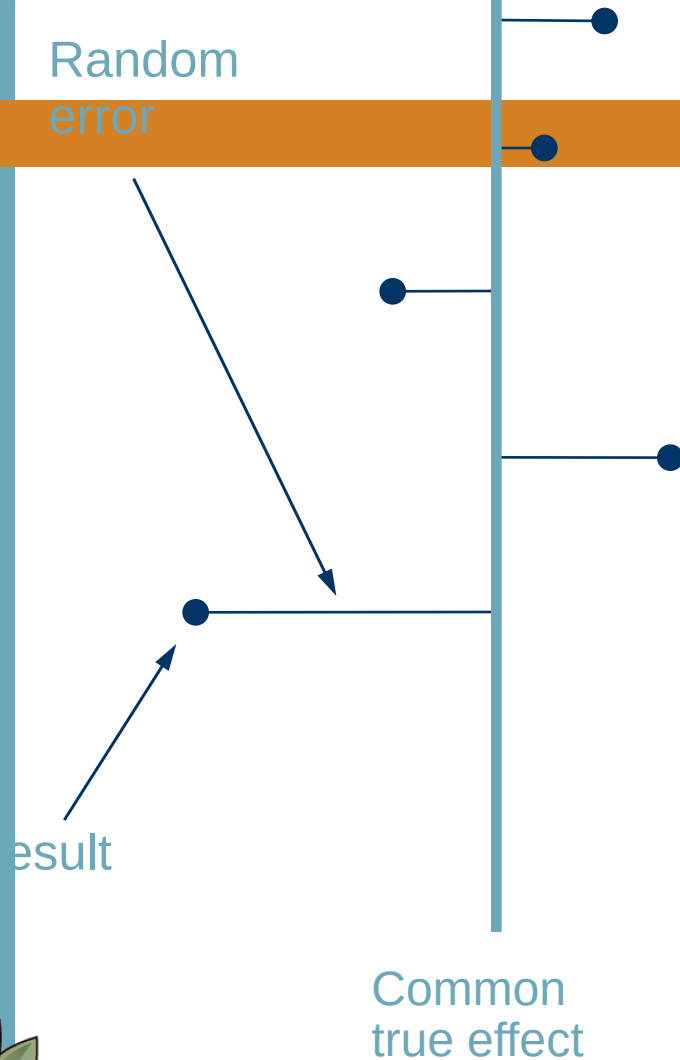
# Fixed effect model

Philosophy behind model:

- there is one real value for the treatment effect
- all trials are estimating this common treatment effect



# Fixed effect model



- assumes that all studies are evaluating the same treatment effect
- *i.e.* if they were all infinitely large they'd produce an identical result



## Option 4: Investigating heterogeneity

- as an objective of your review  
(should be pre-specified in your protocol)
- to determine causes of unexpected statistical heterogeneity
  - note. post hoc investigations should be reported as such and are hypothesis-generating at best



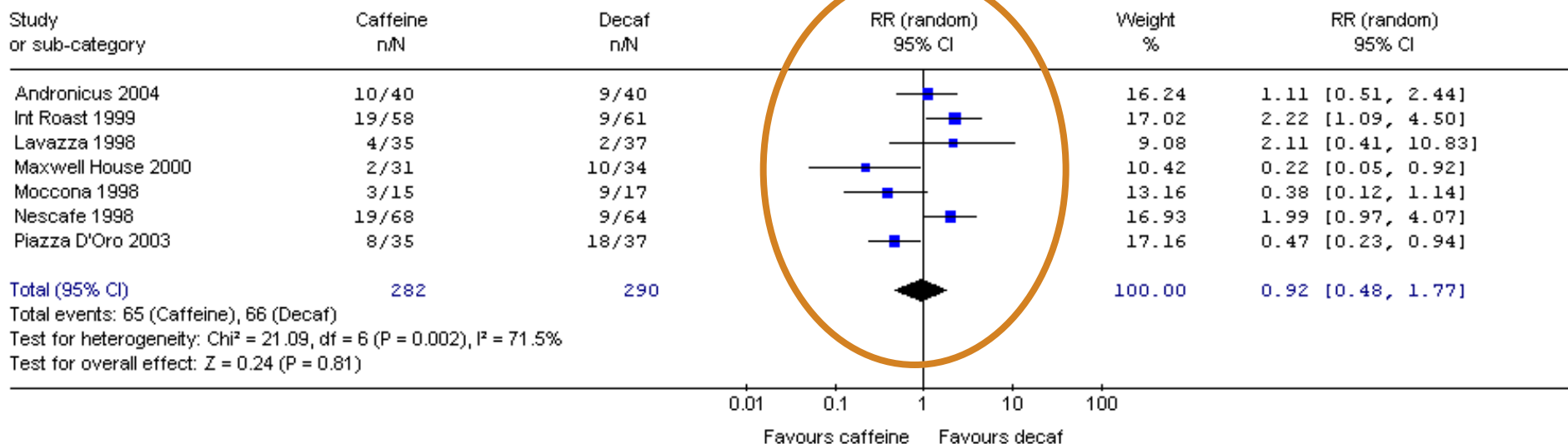
# Investigating heterogeneity: tools

- subgroup analysis
  - get answers to secondary questions concerning subsets of participants or interventions
  - can yield spurious findings if not used carefully
- meta-regression
  - examine relationship between treatment effect and a particular characteristic of the study (not patients)
    - not available in RevMan
- individual patient data (IPD) meta-analysis
  - investigate patient-level characteristics
  - time consuming and expensive



# Option 5: Incorporate heterogeneity

Review: Caffeine for daytime 'sluggishness'. (Version 251105)  
 Comparison: 01 Caffeinated Coffee versus Decaffeinated Coffee  
 Outcome: 02 Headache

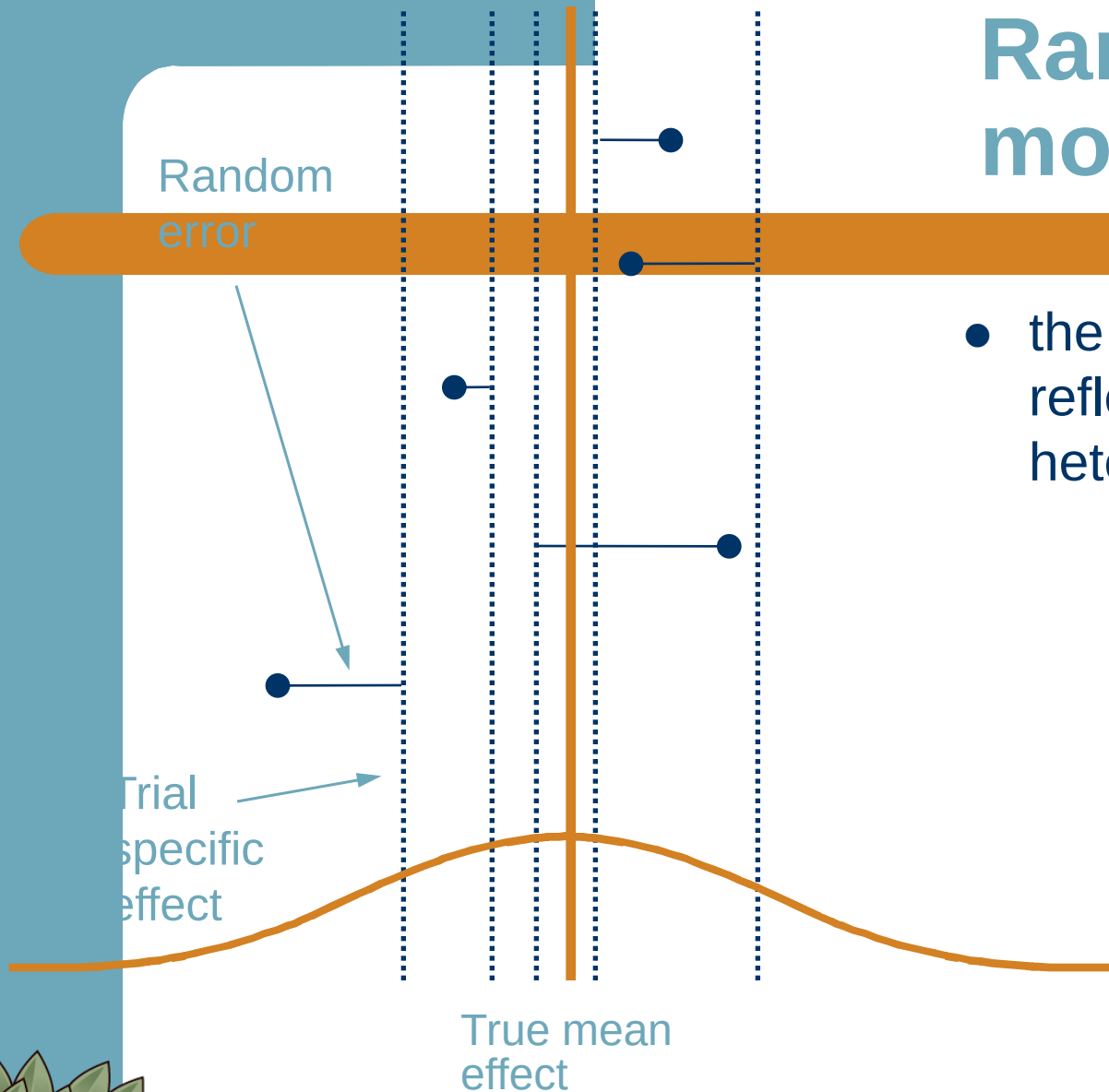


# Random effects model

- if heterogeneity cannot be explained by characteristics of the studies, it may be incorporated into the meta-analysis using the random-effects model
- the true treatment effects underlying the studies are allowed to differ and are assumed to be distributed around a central (mean) value
- weights are adjusted to account for both within-study and between-study variation



# Random effects model



- the width of the bell shape reflects the amount of heterogeneity

# Interpreting random effects meta-analyses

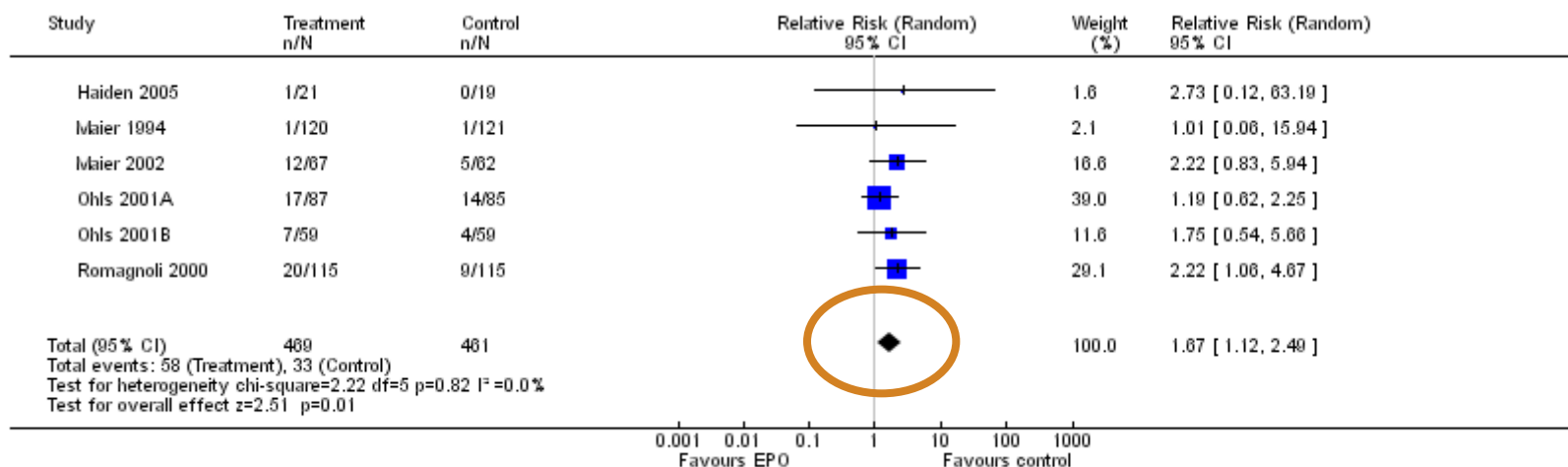
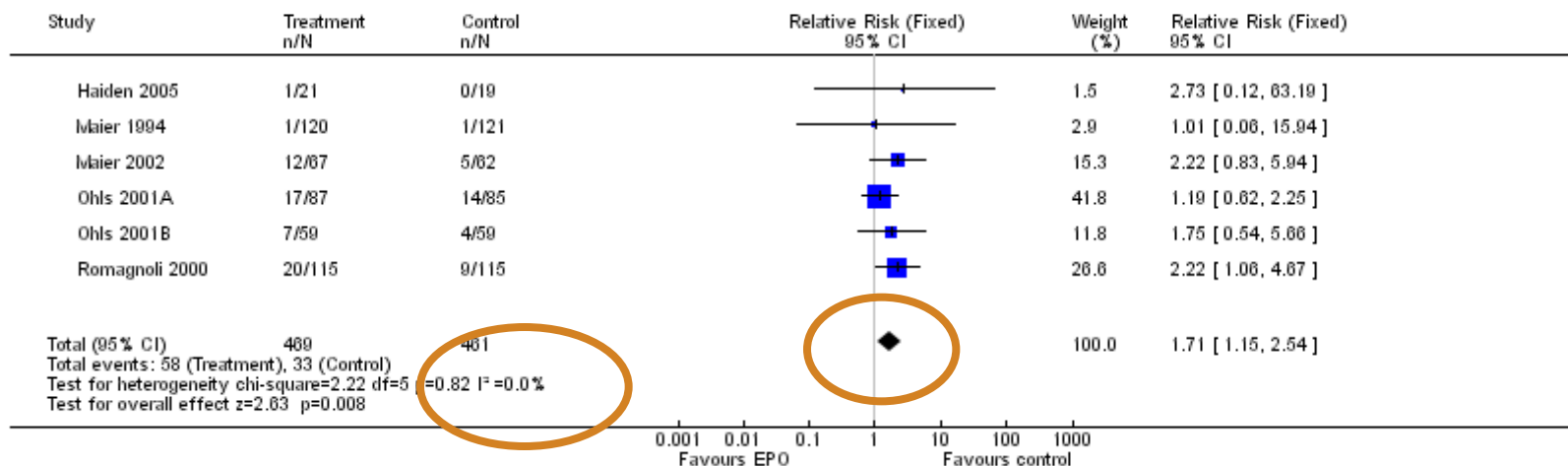
Random effects meta-analyses are...

- **identical** to fixed effect analyses when there is no clear heterogeneity
- **similar** to fixed effect meta-analyses but *with wider confidence intervals* when there is heterogeneity
- **different** from fixed effect meta-analyses when there is publication bias (or funnel plot asymmetry)
  - random effects analyses give relatively more weight to smaller studies



# Fixed versus random effects

Review: Early erythropoietin for preventing red blood cell transfusion in preterm and/or low birth weight infants  
 Comparison: 01 Erythropoietin vs. placebo or no treatment  
 Outcome: 09 Retinopathy of prematurity (stage  $\geq 3$ )

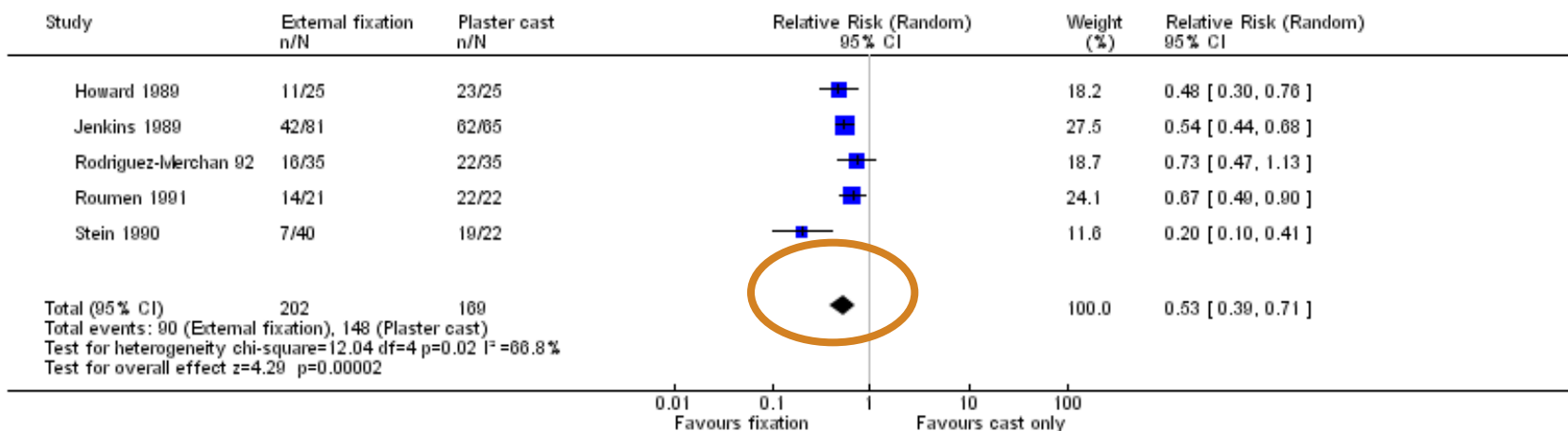
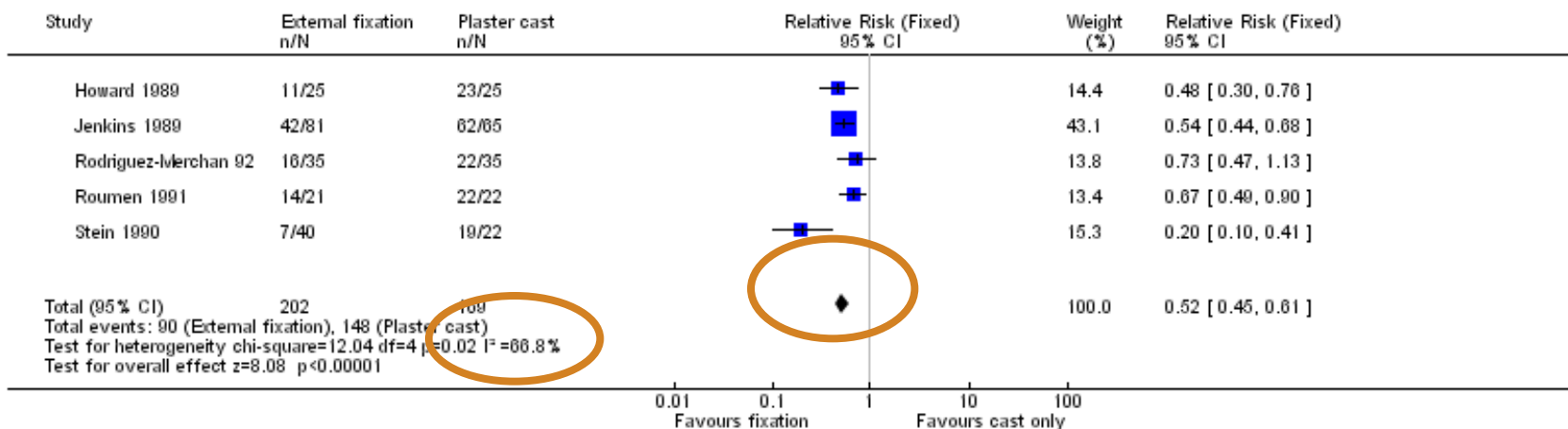


almost identical



# Fixed versus random effects

Review: Surgical interventions for treating distal radial fractures in adults  
 Comparison: 01 External fixation versus plaster cast  
 Outcome: 03 Anatomical grading: not excellent

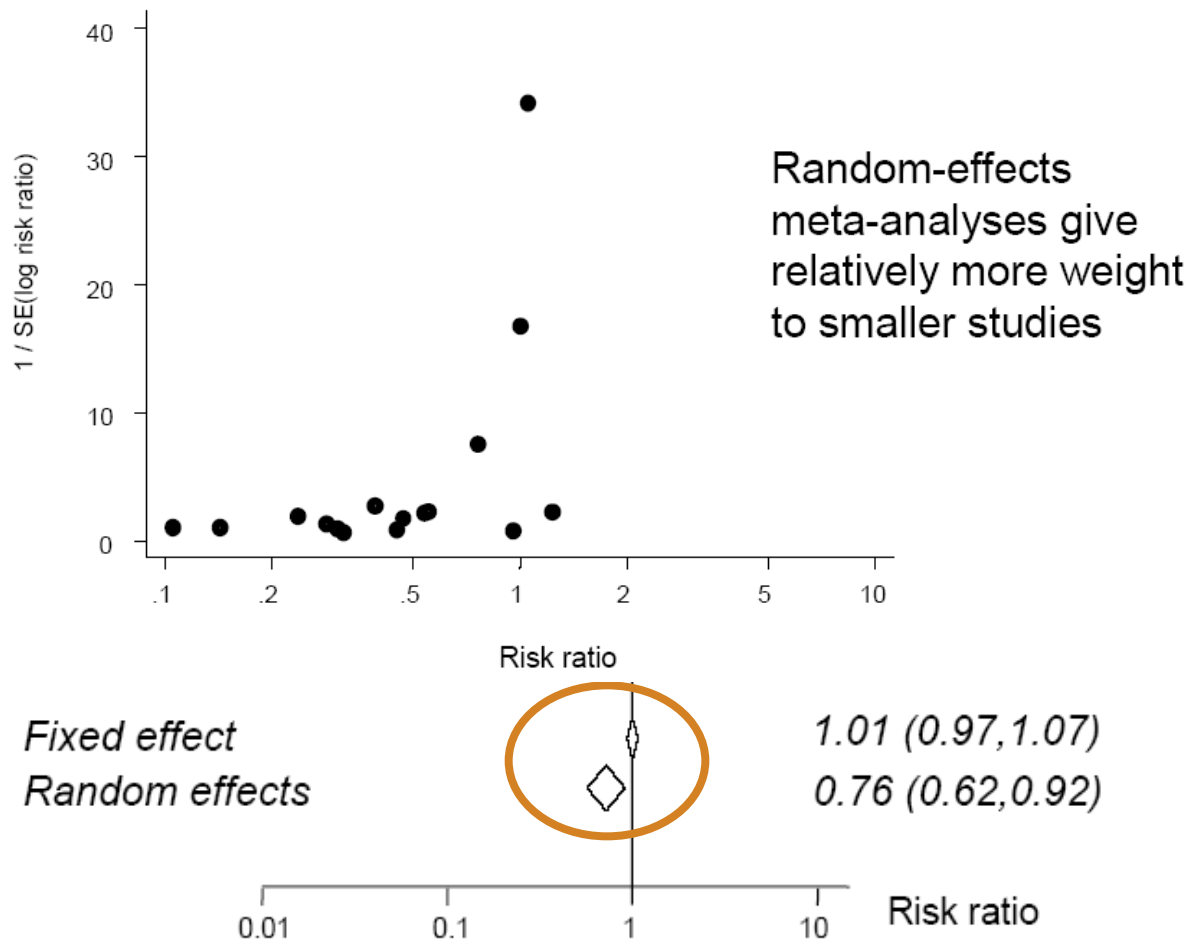


similar, but wider CIs





## Fixed versus random effects



very different results

source: with thanks to Julian Higgins



# Take home messages

- heterogeneity should be assessed and addressed
- statistical heterogeneity occurs when studies are not all evaluating the same treatment effect
- looking at overlap of confidence intervals on forest plot is a good way to identify statistical heterogeneity
- $I^2$  can quantify the degree of inconsistency across studies
- there are several options for dealing with heterogeneity
- methods to investigate heterogeneity should be pre-specified in the protocol
- random effects meta-analyses are useful for incorporating unexplained variability into a summary
- but random effects meta-analyses are not a panacea

