



Clinical Epidemiology

Lecture notes 2015 - Part I

The background of the slide is a dense, overlapping pattern of stylized human figures. Each figure is composed of a rounded rectangular body and a circular head, both in various pastel colors including shades of blue, green, orange, pink, and purple. The figures are arranged in a way that creates a sense of a large, diverse crowd, with some figures appearing more prominent than others due to their position and size.

Lecture 1

Diagnostic Process

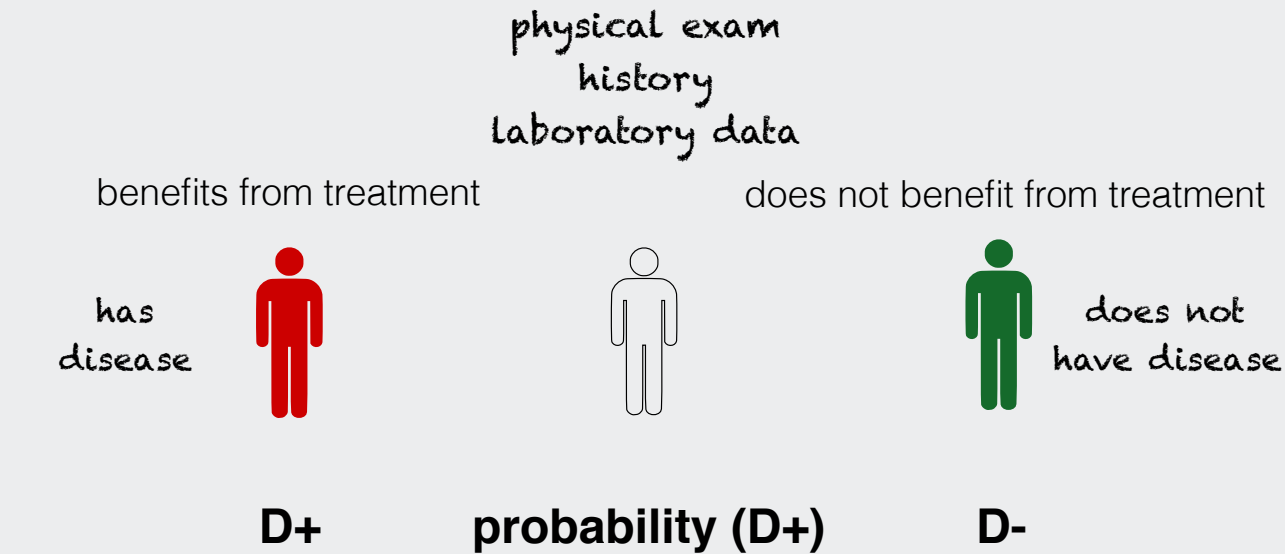
Interobserver agreement

Diagnostic Process

Simplified assumptions

Reality

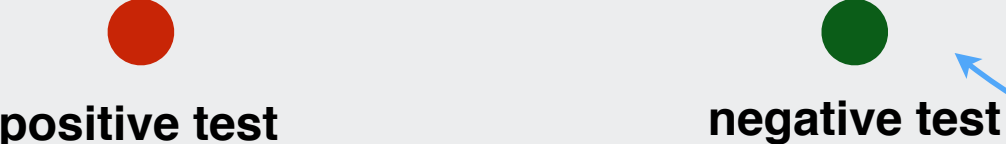
Patient



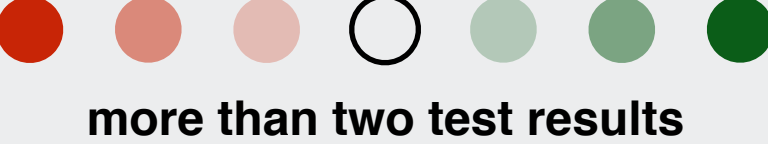
Patient



Dichotomous test



Multilevel and continuous test



continuous test can be dichotomized
with arbitrary cut-point
BUT information is lost

Evaluating diagnostic tests

Reliability

repetitions of the test
give the same result



reliable

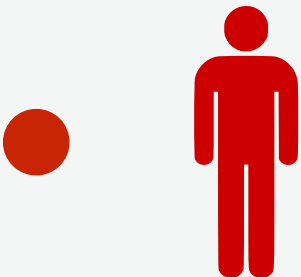


not reliable

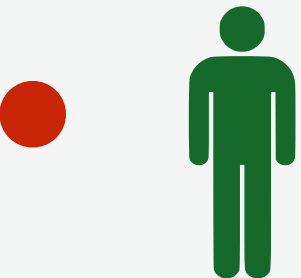
Inter-observer
reliability (KAPPA)

Accuracy

test gives
the right answer



accurate

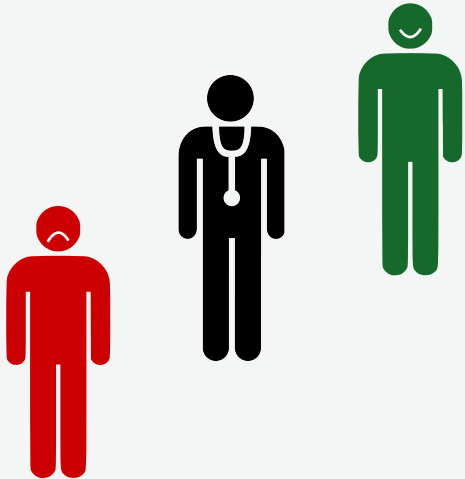


inaccurate

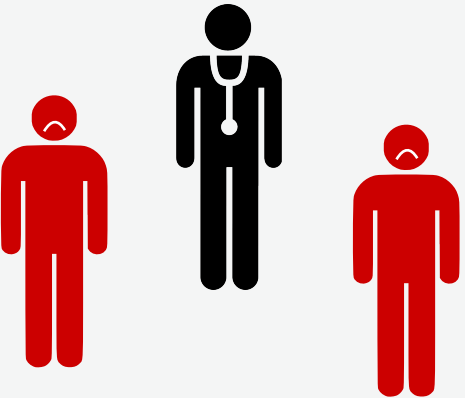
Sensitivity
Specificity
ROC curve

Usefulness

right answer improves
outcome by favorably
affecting decision



useful



useless

Treating threshold

Value

expected improvement
in health outcomes
justifies
the risk and costs



Evaluating diagnostic tests



inter-observer agreement

Observer 1

	result 1	result 2	result 3	Total
result 1	a	b	c	R1
result 2	d	e	f	R2
result 3	g	h	i	R3
Total	C1	C2	C3	N

assess for unbalanced disagreement by looking for asymmetry across the diagonal

i.e. observer 1 is the more concerned radiologist and is more likely to diagnose abnormalities on CXR's

Kappa is not particularly useful in assessing systematic disagreement



Observer 2

add along the diagonal

observed agreement (%) = concordance = $(a+e+i)/N$

observed agreement can be misleading: if two observers both know an abnormality is uncommon, they can have nearly perfect agreement just by never or rarely saying that it is present

this is based on the marginals

expected agreement (%) = $(R_1 \times C_1/N + R_2 \times C_2/N + R_3 \times R_3/N)/N$

the amount of agreement beyond what would be expected from the marginals (row and column totals)

observed agreement - expected agreement

Kappa = $\frac{\text{observed agreement} - \text{expected agreement}}{100\% - \text{expected agreement}}$

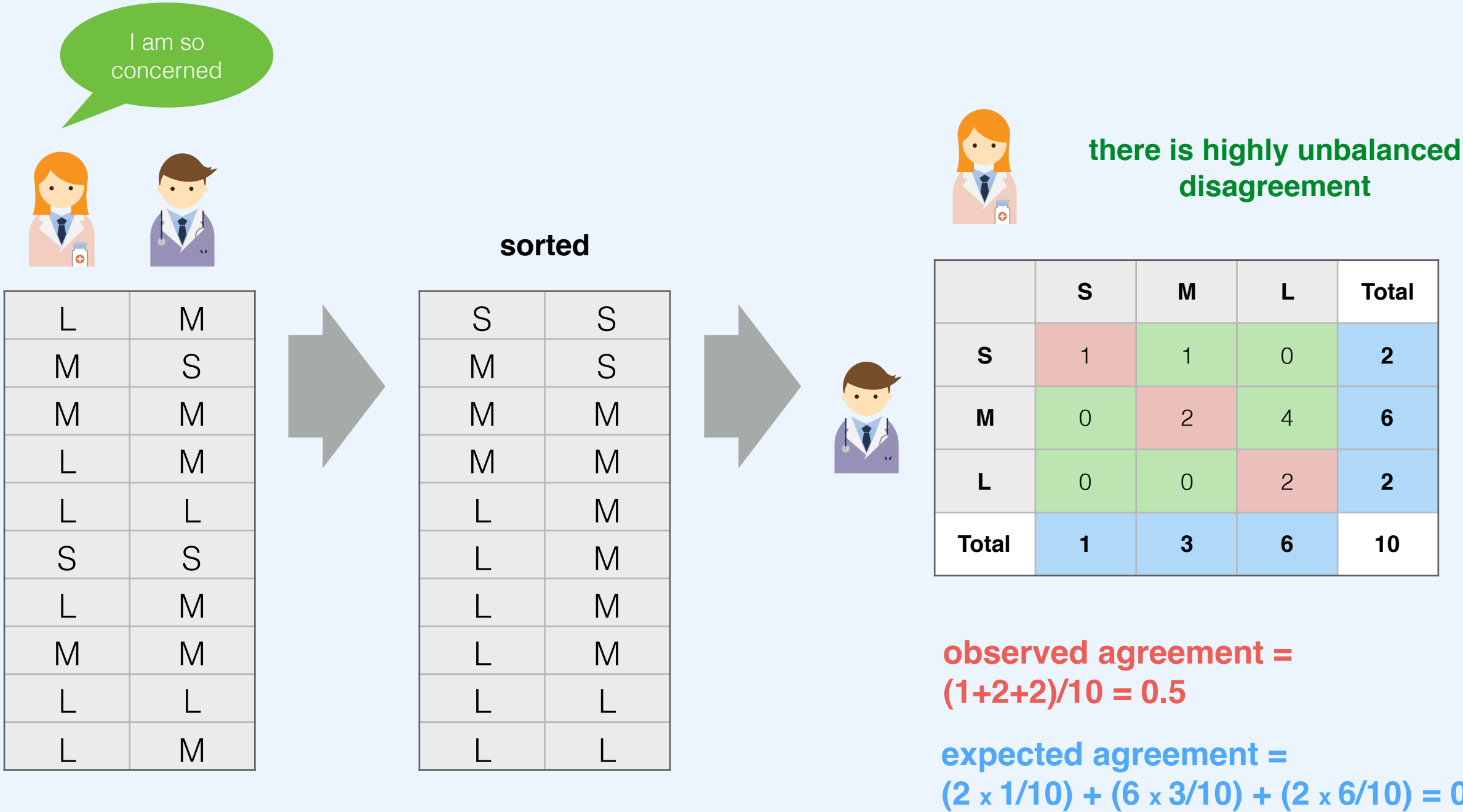
weighted Kappa

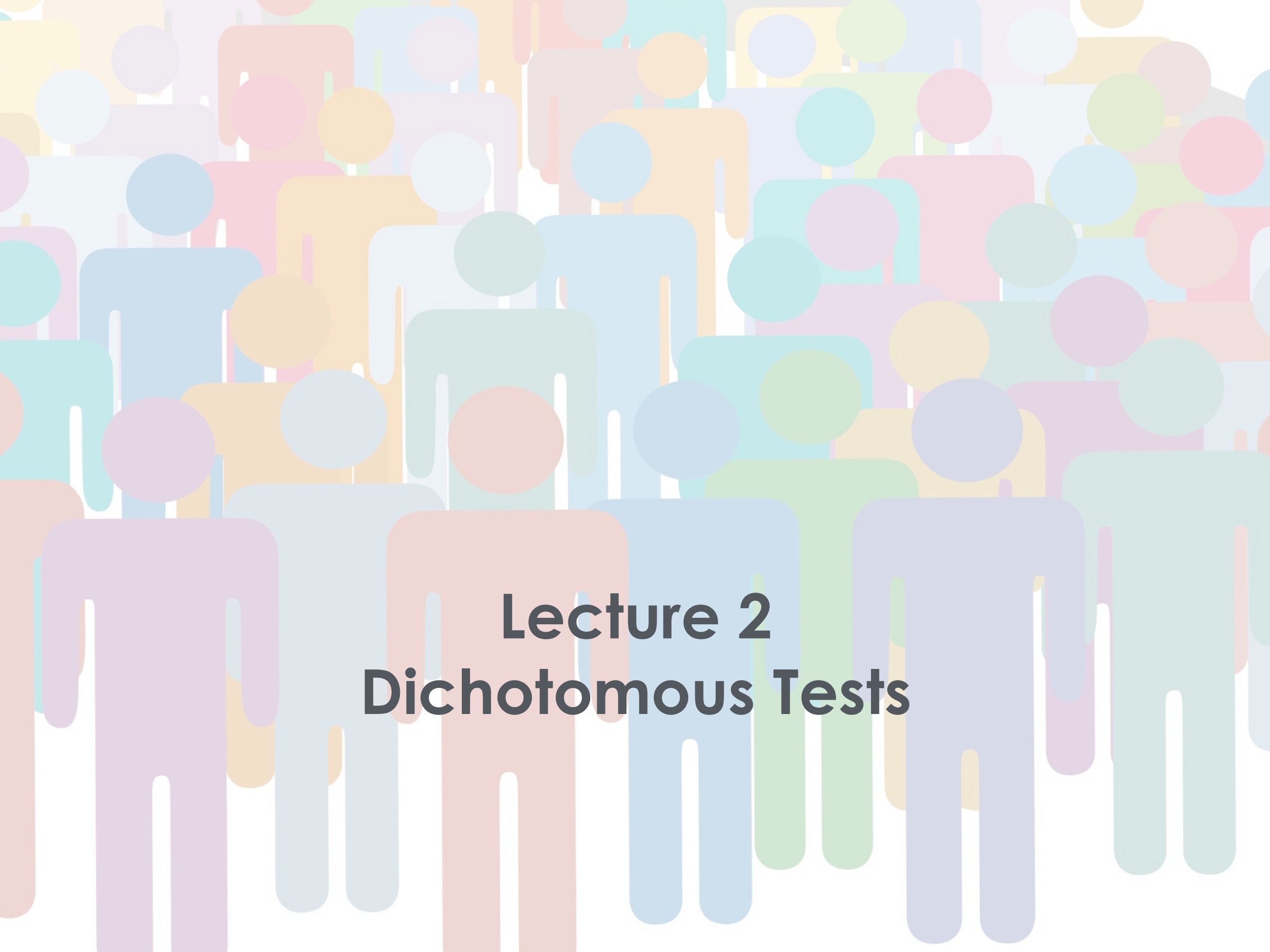
for **ordinal** variable with more than 2 values, gives partial credit for 2 observers being close
i.e. partial credit for b, d, f and h

Evaluating diagnostic tests

inter-observer agreement: example

two doctors are estimating the size of a lesion





Lecture 2

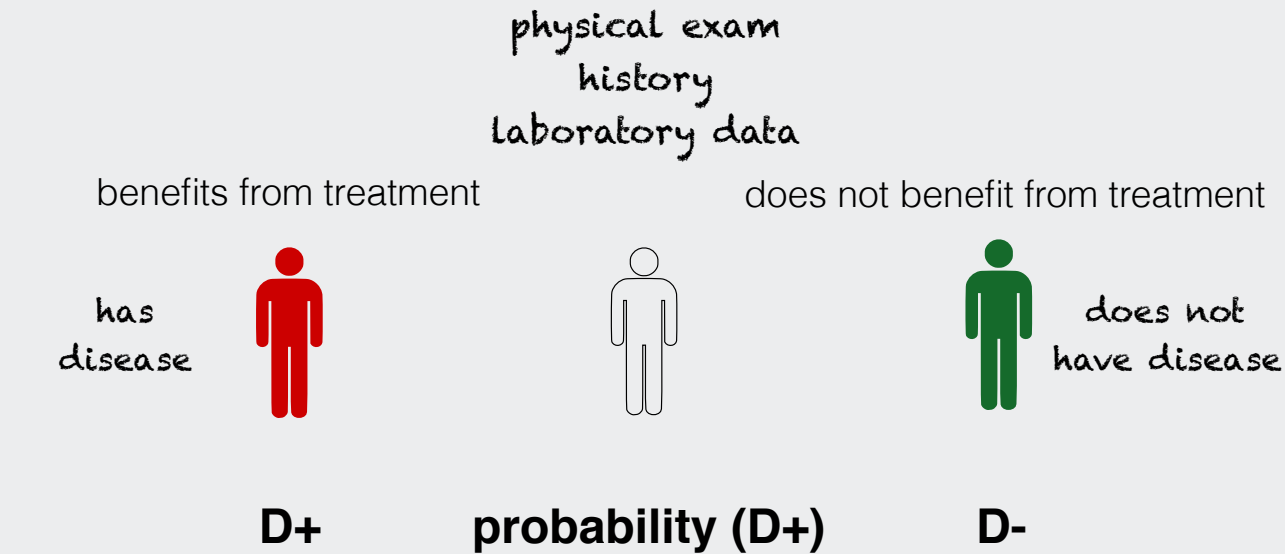
Dichotomous Tests

Diagnostic Process

Simplified assumptions

Reality

Patient



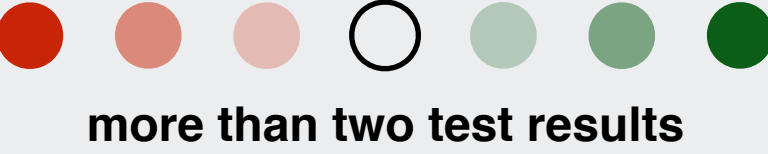
Patient



Dichotomous test

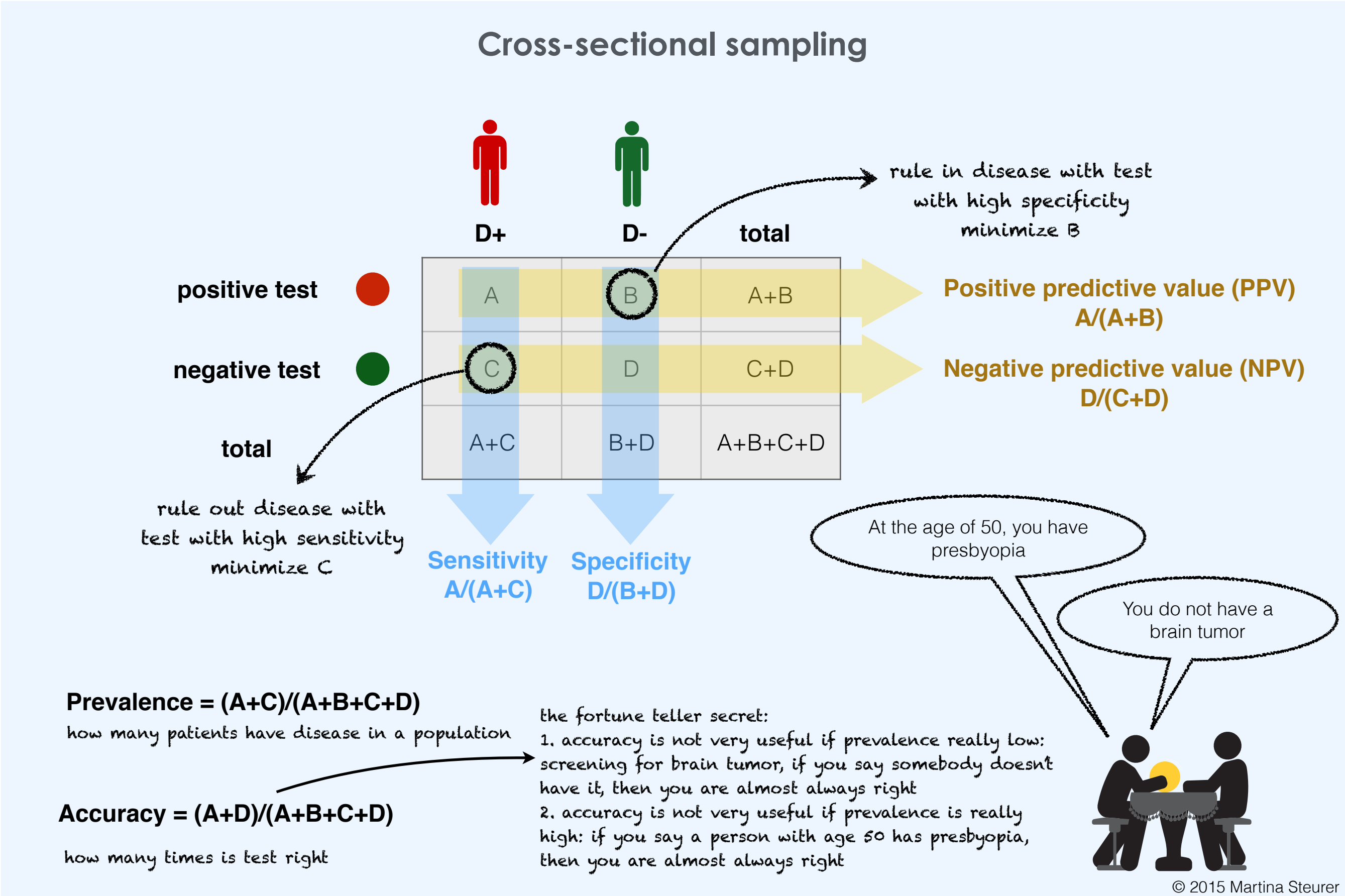


Multilevel and continuous test



continuous test can be dichotomized
with arbitrary cut-point
BUT information is lost

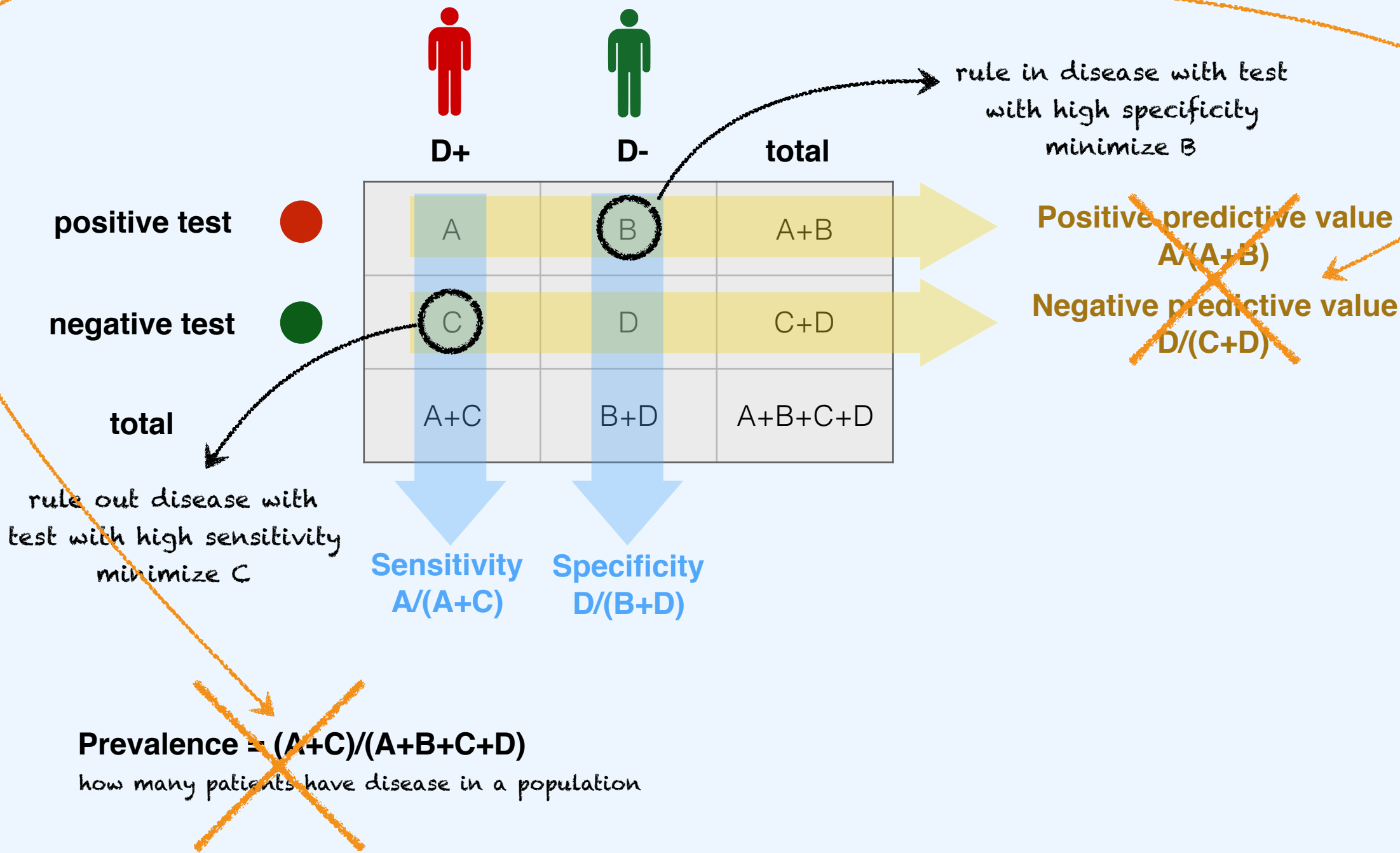
Dichotomous test



Dichotomous test

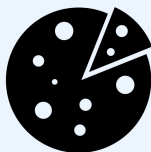
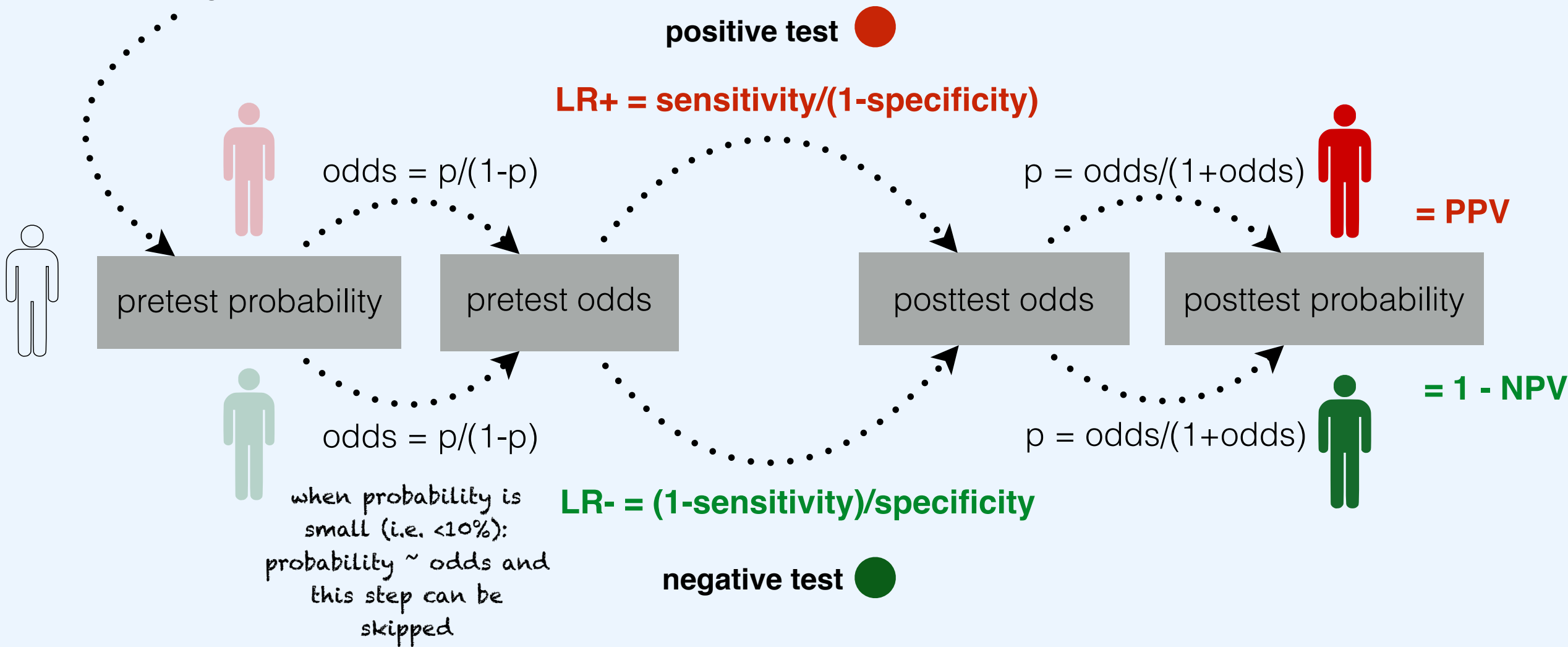
Case-control sampling

Ratio between people with and without disease is artificially set by investigator



Prior and posterior probabilities

prevalence = pretest probability for screening test
prevalence and patient presentation = pretest probability in diagnostic test



odds = ratio of two portions
probability = proportion of the whole

likelihood ratio =
$$\frac{P(\text{Result} \mid \text{Disease})}{P(\text{Result} \mid \text{No Disease})}$$
 = WOWO (with over without disease)

Treatment and Testing thresholds

in the green area test is not useful because prior probability is so low that even after a positive test result your posterior probability of disease would still be below treatment threshold

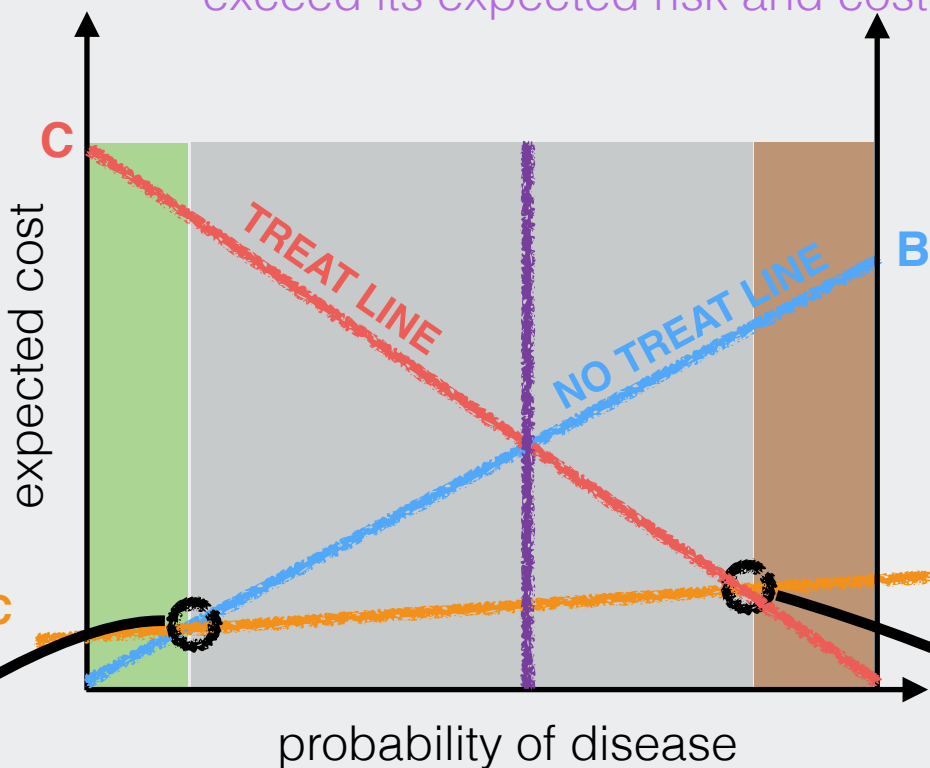
treatment threshold (PTT) = lowest probability of disease at which the expected benefits of treatment exceed its expected risk and costs

in the brown area test is not useful because prior probability is already so high that you can just treat (i.g. setting of an epidemic)

cost of treatment x probability of unnecessary treatment (no disease)
 $C \times (1-p)$

how bad is it to treat someone who does not have the disease ($=C$)

TEST



net benefit foregone x probability of failure to offer treatment (disease)
 $B \times p$

how bad is it not to treat someone who does have the disease ($=B$)

No treat test threshold:
below this threshold the prior probability is so low that it is impossible to get above the PTT after the test, thus it is not worth testing, just don't treat

in gray area test is useful because prior probability close enough to treatment threshold that test result will impact decision

Treat-test threshold:
above this it is not worth testing because the prior probability is already so high that it is ok just to treat



Lecture 3

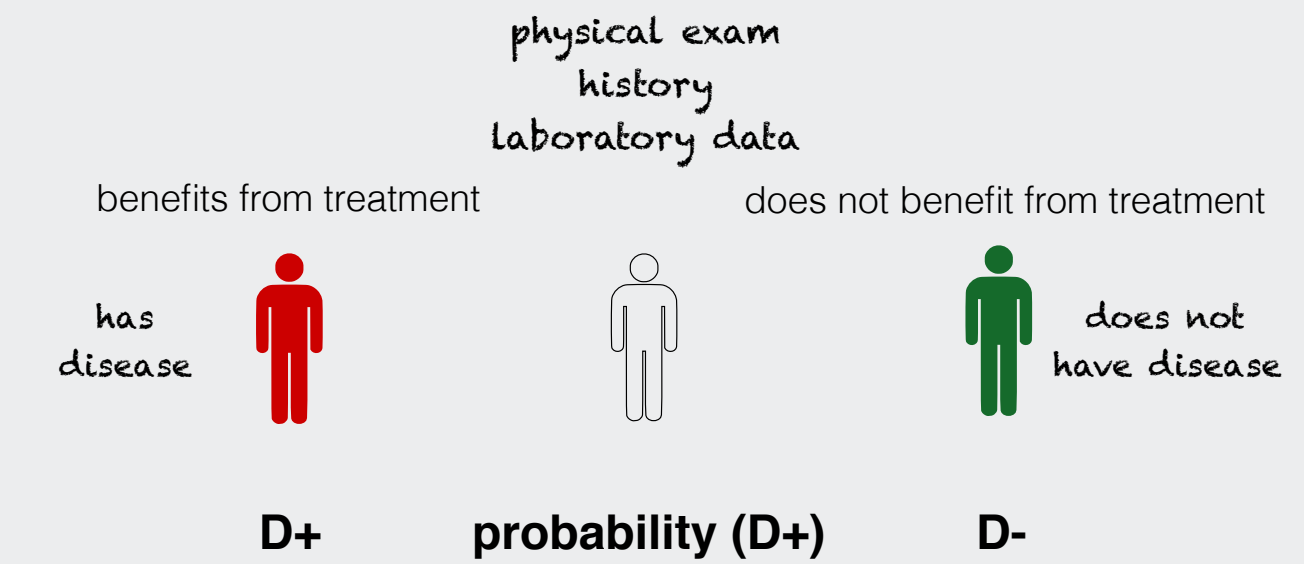
Multilevel and continuous test

Diagnostic Process

Simplified assumptions

Reality

Patient



Patient



Dichotomous test

$LR+ = \frac{\text{sensitivity}}{1 - \text{specificity}}$

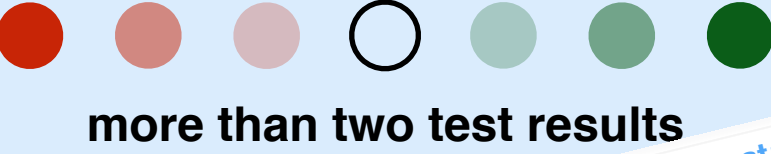
$LR- = \frac{1 - \text{sensitivity}}{\text{specificity}}$



continuous test can be dichotomized with arbitrary cut-point BUT information is lost

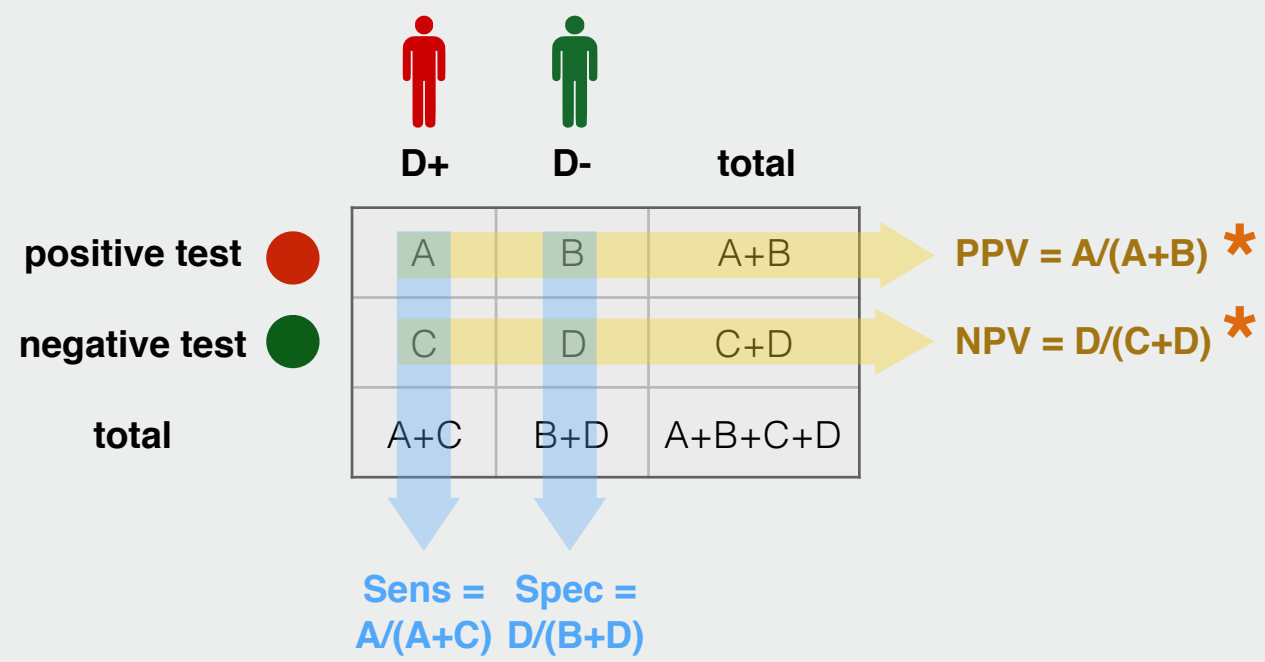
Multilevel and continuous test

prior probability might affect test characteristics



caveat:
never calculate LR+ and LR- for test with more than 2 results, in these cases calculate interval LR, LR equals slope of ROC curve in that interval

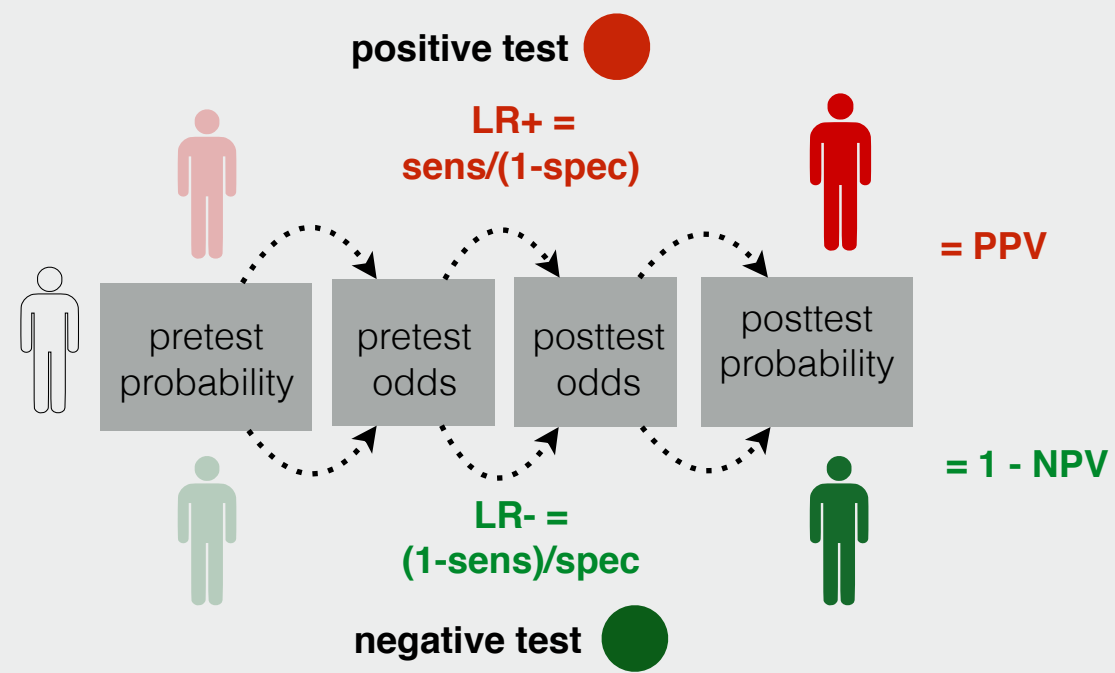
Dichotomous test



Prevalence = $(A+C)/(A+B+C+D) *$

Accuracy = $(A+D)/(A+B+C+D) *$

*not useful in case-control studies because prevalence set by investigator with case-control ratio



Multilevel and continuous test

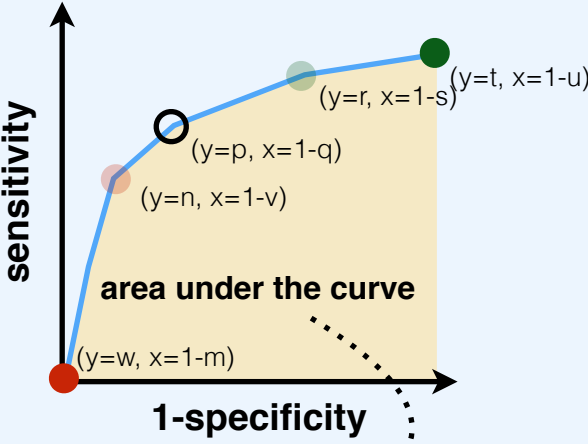
ROC table

test result sensitivity 1-specificity

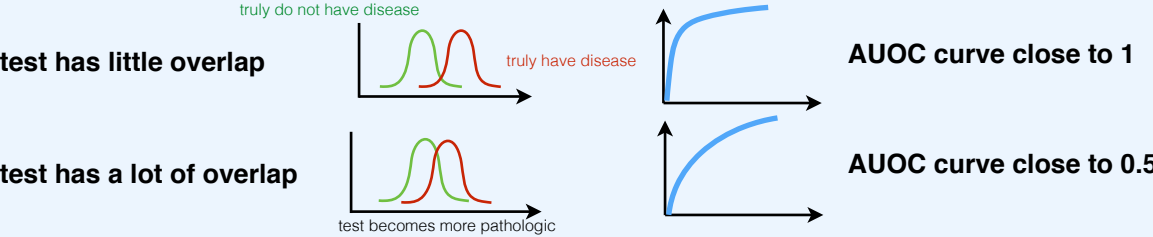
●	w	m
●	n	v
○	p	q
●	r	s
●	t	u

don't report LR in an ROC table, report LR in LR table

ROC curve

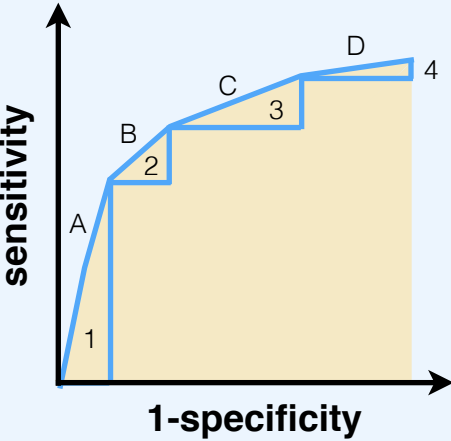


summary measure of discriminatory ability



LR table

test interval	% of D+	% of D-	interval LR
1	e	f	A
2	g	h	B
3	i	j	C
4	k	l	D



Interval LR represent slope of curve in this interval