

Reflection Questions for Math 58B

Johanna Hardin

Spring 2017

Chapter 1, Section 1 – binomial probabilities

1. What is a p-value?
2. What is the difference between a one- and two-sided hypothesis?
3. What is the difference between type I and type II errors?
4. Why is it never okay to accept H_0 ?
5. What is power? How is power calculated? What does power depend on?

Chapter 1, Section 2 – norm approx to binom

1. What does it mean for a variable to have a normal distribution?
2. What does the Central Limit Theorem say?
3. How is the normal distribution different from the binomial distribution? (one answer is that the normal describes a continuous variable and the binomial describes a discrete variable. what does that mean? what is another distinction?)
4. What are the technical conditions allowing the normal distribution to approximate the binomial distribution?
5. What does a z-score measure?
6. What is the difference between a z-critical value (z^*) and a z test statistics ($z_0 = \frac{\hat{p} - \pi_0}{\sqrt{\frac{\pi_0(1-\pi_0)}{n}}}$)?
7. What does it mean to say that to find normal probabilities compute the area under the normal curve?
8. When computing a confidence interval (i.e., when we don't have a preconceived idea for π), how is the standard deviation of $\hat{p}$ estimated?

9. When using the normal distribution to create a confidence interval for π , how is the critical value for, say, a 94.7% interval calculated?
10. What is one reason to choose to use the normal distribution?
11. What is one reason to choose to use the binomial distribution?

Chapter 1, Section 3 – sampling

1. Why is it better to take random samples?
2. What is a simple random sample?
3. Why don't researchers always take random samples?
4. What does a large sample provide?
5. What is the difference between practical significance and statistical significance? (See Inv 1.17)

Chapter 2, Section 1 – quantitative variables

1. What changed about the studies (data structure) from Chapter 1?
2. What is the statistic we use now?
3. What is the difference between the distribution of the data and the distribution of the statistic? There is a theoretical difference as well as a computational difference.

Chapter 2, Section 2 – inference on 1 mean

1. What is the sampling distribution of the statistic? (Note, the answer here is only for big samples, that is the Central *Limit* Theorem works only where there is a limit... i.e., the sample size is big.)
2. What is the difference between the distribution of the data and the distribution of the statistic? There is a theoretical difference as well as a computational difference.
3. What is the difference between a normal distribution and a t distribution?
4. When do we use a z and when do we use a t?
5. When would you use a confidence interval and when would you use a hypothesis test?
6. What different information does a boxplot give versus a histogram?

Chapter 3, Section 1 – 2 pop proportions

1. How does the inference *change* now that there is binary (response) data taken from two populations?
2. How does the inference *stay the same* now that there is binary (response) data taken from two populations?
3. What does the Central Limit Theorem say?
4. When is it appropriate to use a hypothesis test to evaluate the data? And when is it appropriate to use a confidence interval to evaluate the data?
5. Which formula is used for $SD(\hat{p}_1 - \hat{p}_2)$ under which conditions / scenarios?
6. What technical conditions must hold for the Central Limit Theorem to apply?

Chapter 3, Section 2 – types of studies

1. What is the difference between an observational study and an experiment?
2. Why aren't all studies done as experiments?
3. What is a confounding variable?
4. Have you looked at page 139? Do you understand that page? [Random sampling vs. Random allocation]
5. how is the statistical meaning of the word *cause* different from the usage in the sentence *The ball that hit me in the head caused me to get a headache.*
6. What are the meanings of the words: randomized, double-blind (single-blind), control, placebo, significant, and comparative. Why are these ideas important to interpreting study results?

Chapter 3, Section 3 – prob vs prop

1. How is the random process different from the examples at the beginning of chapter 2? And, therefore, how is the inference different?
2. What (exact) probability model is best used for data that has been randomly allocated? [note: there is no exact model at the beginning of chapter 2, but the simulation using binomial probabilities allowed inference without implementing the CLT.]
3. Fisher's exact test has an additional advantage that hasn't been discussed in class: it works for small sample sizes. Which method *doesn't* work for small sample sizes?

4. How is the normal approximation for two sample proportions different when the data are a random sample as compared to when they come from a randomized experiment? How is it the same?

Chapter 3, Section 4 – odds ratios and RR

1. What is the differences between cross-classification, cohort, and case-control studies?
2. When is it not appropriate to calculate differences or ratios of proportions? Why isn't it appropriate?
3. How are odds calculated? How is OR calculated?
4. What do we do when we we can't calculate statistics based on proportions? Why does this "fix" work?
5. Why do we look at the natural log of the RR and the natural log of the OR when finding confidence intervals for the respective parameters?
6. How do you calculate the SE for the $\ln(\hat{RR})$ and $\ln(\hat{OR})$?
7. Once you have the CI for $\ln(RR)$ or for $\ln(OR)$, what do you do? Why does that process work?

Chapter 4, Section 1&2&3 – 2 means

1. What changed about the studies (data structure) from chapter 2?
2. What is the statistic of interest now? What is the parameter of interest?
3. What is the sampling distribution for the statistic of interest?
4. Why does the t-distribution become relevant?
5. How does the t-distribution become relevant?
6. What are degrees of freedom in general? What are the actual degrees of freedom for the test in section 3.2?
7. How is the null mechanism different across the three analysis methods in section 3.2: randomization test, two-sample t-test, random sampling test (n.b. this is also called the parametric bootstrap)?
8. How do you create a CI? How do you interpret the CI?
9. What if your data are NOT normal? What strategies can you try out?

Chapter 4, Section 4 – paired samples

1. What changed about the studies (data structure) in section 3.3?
2. What is the statistic of interest now? What is the parameter of interest?
3. What is the sampling distribution for the statistic of interest?
4. What does pairing do for us?
5. What happens if we analyze a paired study as if we had an independent two sample study? (What happens to the p-value? What happens to the CI?)
6. What is the easiest way to think of / analyze paired data?

Chapter 5, Section 1 – 2 categorical variables

1. How would you describe the data we see in $r \times c$ tables?
2. Describe the applet mechanism that creates a sampling distribution under the assumption that the null hypothesis is true.
3. What is the test statistic (for both the randomization applet and the chi-square test!!)? Why do we need a complicated test statistic here and we didn't need one with 2×2 tables?
4. How do you compute the expected count? Why does that computation make sense?
5. What is one benefit that the two sample z-test of proportions has? That is, what is one thing we can do if we have a 2×2 table instead of an $r \times c$ table?
6. Describe the directionality of the test statistic. That is, what values of X^2 make you reject H_0 ?
7. What are the technical assumptions for the chi-square test? Why do you need the technical assumptions?
8. What are the null and alternative hypotheses?

Chapter 5, Sections 3 & 4 – correlation & regression

1. How do we find the values of b_0 and b_1 for estimating the least squares line?
2. Why is it dangerous to extrapolate?
3. How do we interpret R^2 ? Why is that? (Look at SSE values on page 325, Inv 4.8 (aa) & (bb).)

4. What does it mean to say that b_1 has a sampling distribution? Why is it that we would never talk about the sampling distribution of β_1 ?
5. Why do we need the LINE technical conditions for the inference parts of the analysis but not for the estimation parts of the analysis?
6. What are the LINE technical conditions?
7. What are the three factors that influence the $SE(b_1)$? (Note: when something influences $SE(b_1)$, that means the inference is also effected. If you have a huge $SE(b_1)$, it will be hard to tell if the slope is significant because the t value will be small.)
8. What does it mean to do a randomization test for the slope? That is, explain the process of doing a randomization test here. (See shuffle options in the Analyzing Two Quantitative Variables applet.)
9. Why would someone transform either of the variables?
10. What is the difference between a confidence interval and a prediction interval? Which is bigger? Why does that make sense? How do the centers of the intervals differ? (They don't. Why not?)