

24/12/21

- Agenda:
- Assignment 8, Q 1, 4, 7.
 - Ross Ch 5, Ex 2(c), 4(c), 4(i), 5(f), 7(d).
 - A brief note on "rigorous" prob. theory.

Assigned Ques.

Q1) This follows from the definition of the Poisson (i) random variable. That is,

$$P(X=k) = \frac{e^{-\lambda} \lambda^k / k!}{e^{-\lambda} \lambda^{k-1} / (k-1)!} = \lambda / k$$

Now, if λ is an integer, then $P(X=\lambda) = P(X=\lambda-1)$ is the maximum. (i.e., $k > \lambda > k \Rightarrow P(X=k)$ is the maximum.

For the expectation,

$$E[X] = \sum_{k=1}^{\infty} k P(X=k) = \sum_{k=1}^{\infty} k \frac{e^{-\lambda} \lambda^k}{k!}$$

$$= \sum_{k=1}^{\infty} \lambda \frac{e^{-\lambda} \lambda^{k-1}}{(k-1)!} = \lambda \sum_{k=0}^{\infty} \frac{e^{-\lambda} \lambda^k}{k!} = \lambda$$

Key point: The "mode" of the $Poi(\lambda)$ r.v. is very close to its mean i.e. λ . This means that $Poi(\lambda)$ concentrates heavily around its mean. So: You expect mean & var. to be close to each other. Indeed, $E = Var = \lambda$, so it's binomial-like concentration!

Q4) Consider the situation $f_n(n) = c n^{-\alpha}$. Then, for this to be the pmf of a $\mathbb{N}$ -valued random variable, we require that $\sum_{n=1}^{\infty} f_n(n) = 1$, which becomes

$$c \sum_{n=1}^{\infty} n^{-\alpha} < \infty$$

By the integral test, $\sum_{n=1}^{\infty} n^{-\alpha}$ converges to some $K \in \mathbb{R}$. Then $c = 1/K$.

Now, note that if $P(X=n) = f_n(n)$, then

$$E[X] = \sum_{n=1}^{\infty} n P(X=n) = c \sum_{n=1}^{\infty} n^{1-\alpha}$$

Noting that $-\alpha+1 < 1$, we get that $\sum_{n=1}^{\infty} n^{1-\alpha}$ converges to a finite number. So: $E[X] < \infty$ is well-defined.

The point: Markov's inequality tells you that if $E[X] < \infty$ then $P(X > t) \leq E[X] / t$. In other words, if $E[X]$ is finite, then $P(X > t)$ decays like $1/t$. Here, we have the converse: if $P(X > t)$ decays like $1/t^{\alpha}$, $\alpha > 1$, then $E[X]$ exists. Indeed, for $f_n(n)$, $P(X > n) \sim 1/n^{\alpha+1}$, so this illustrates the reverse of the Markov inequality, in some sense.

Q7) First caveat: Using Fubini-Tonelli theorem here is a **FORMALITY**: it's a one-line obvious justification: you can always exchange summations when all terms are non-negative. So there is no need to worry here.

The idea is: indicators.

$$\begin{aligned} a) \sum_{j=1}^{\infty} (a_1 + \dots + a_j) P(X=j) &= \sum_{j=1}^{\infty} \sum_{k=1}^{\infty} \mathbb{1}_{\{k \leq j\}} a_k P(X=j) \\ &= \sum_{k=1}^{\infty} \sum_{j=1}^{\infty} \mathbb{1}_{\{k \leq j\}} a_k P(X=j) = \sum_{k=1}^{\infty} a_k \sum_{j=k}^{\infty} P(X=j) \\ &= \sum_{k=1}^{\infty} a_k P(X \geq k) \end{aligned}$$

For (b), use $a_i = 1$ in above. For (c), use $a_i = 2^i$ above for all i , along with the fact that

$$\sum_{j=1}^{\infty} (a_1 + \dots + a_j) P(X=j) = E[f(X)]$$

where $f(x) = a_1 + \dots + a_x$

Ross

2(c) Points to take away:

- Usage of the change of variable formula.
- The logic that is used to find the random variable of interest: identifying it as a function of the underlying randomness.
- Interesting question: is it possible to prove that $E[L_p(W)] \geq E[L_q(W)] \forall |q-1/2| \geq |p-1/2|$? WITHOUT looking at the distributions of L_p and L_q ?

4(c) Key takeaways:

- Exam marks can be modeled using a normal r.v.: to be more precise, the reason is that if independent questions are taken as Bernoullis of various unknown p s, then their sum can be approximated by a normal random variable, suitably normalized.
- Once you assume your data is normal, then there is a reasonable scheme of grading: this can be thought of as "discretizing" the normal, since the grades form a discrete set, but are a function of the normal.

7(i) Key takeaways:

- Comprehension: This question aims to test your understanding of the scenario and see "which" random variable is applicable; and how. Once you see it is the binomial, you are home. The fact that the diet has no effect on the diet level tells you that each person's diet level is likely to go any way (up or down) during the experiment i.e. $Ber(1/2)$ for each of 100 independent people.
- Using the normal approximation to the Binomial.

$$\left[\frac{\text{Bin}(n, p) - np}{\sqrt{np(1-p)}} \approx N(0, 1) \text{ [in distribution]} \right]$$

Ex 5(f): Key takeaways:

- Misunderstanding the hazard rate concept. The best explanation is that:

$$P[X < t + \Delta t | X > t] \approx \lambda_X(t) \Delta t \quad [\forall \Delta t \approx 0]$$

→ The Hazard Rate integration formula:

$$F(t) = 1 - \exp\left\{-\int_0^t \lambda_X(s) ds\right\}$$

[Remember that this works only for non-neg. random variables]

This extends to the more general:

$$P[X > B | X > A] = \exp\left\{-\int_A^B \lambda_X(t) dt\right\}$$

→ Interpreting the hazard rate when it appears in a question. The key word is rate here.

→ When you consider random variables that model "lifetimes" in general, it may be required to compare random variables that represent such concepts. One of these is the notion of stochastic dominance. Another, is the hazard rate order: so we can compare random variables reasonably using this quantity.

7(a) Key takeaways:

- If f is an increasing $\uparrow$ -function, then applying $P(f(X) \geq y) = P(X \geq f^{-1}(y))$ yields, instantly, the PDF / CDF of $F(X)$ in terms of that for X .
- If f is strictly decreasing, the above can be repeated in the opposite direction. The short of it is the importance of the above technique.

On rigorous probability theory

It is fair to wonder how the theory of sample spaces and random variables are expanded so as to fit infinite spaces and continuous random variables.

The answer, is to redefine a prob. space as follows:

- A sample space Ω of outcomes.
- A suitable set of subsets of Ω (or events), that we wish to study across our experiment. Call this σ .
- A way to study these, i.e. a "measure" $\mu: \sigma \rightarrow [0, 1]$ that reasonably assigns probabilities to each subset.

Given these three, you put them together to get a probability space.

The only condition that σ must satisfy is some consistency conditions. Stochastically, σ should be closed under complement & unions, since the ability to study an event implies we can study its complement. Such sets are called σ -algebras.

What happens in the discrete case is that every subset can be studied. Typically on $\mathbb{R}$, and $[0, 1]$ and often non-discrete sample spaces, you usually cannot study every subset if you wish to measure them in a certain manner. Thus, for some measures, you need to choose suitable σ -algebras for these measures.

For $\mathbb{R}$, or $[0, 1]$, the respective σ -algebra used is called the Borel σ -algebra. This is the σ -algebra generated by all intervals of the form $\{[a, b]: a \leq b\}$. (Generated means to take every set that can be generated from intervals using only unions and complements)

Remarkably enough, it can be shown that any measure, if defined on a generating set, extends to the entire σ -algebra uniquely. Using this, if we say:

$$\Omega = \mathbb{R}, \sigma = \text{Borel}(\mathbb{R}), \mu([a, b]) = \int_a^b e^{-x/2} \frac{1}{\sqrt{2\pi}} dx$$

Then μ is the prob. measure corresponding to a $N(0, 1)$ r.v. since if $X \sim N(0, 1)$ then $\mu([a, b]) = P[X \in [a, b]]$.

$\Omega = [0, 1], \sigma = \text{Borel}(\mathbb{R}), \mu([a, b]) = b - a \Rightarrow \mu \sim U[0, 1]$

$\Omega = [n], \sigma = \mathcal{P}([n])$ (generated by $\{ \{x\} : x \in [n] \}$), $\mu(\{x\}) = p^x (1-p)^{n-x} \binom{n}{x}$

Then $\mu \sim \text{Bin}(n, p)$

Thus, it is possible now to define r.v.s continuously. More on this in Klenke, Chapter 1.