

Probability

Contents

Configuring R	1
1 Core properties of probabilities	2
1.1 Defining probabilities	2
1.2 Conditional probability	3
2 Key probability distributions	5
2.1 The Bernoulli distribution	6
2.2 The Poisson distribution	7
2.3 The Negative-Binomial distribution	14
2.4 Weibull Distribution	15
3 Characteristics of probability distributions	15
3.1 Probability density function	15
3.2 Hazard function	16
3.3 Expectation	16
3.3.1 Conditional distributions and expectations	20
3.4 Fubini–Tonelli for expectations	28
3.5 Deviation, error, and noise	42
3.6 Variance and related characteristics	42
3.7 The Central Limit Theorem	48
4 Additional resources	50
References	50

Configuring R

Functions from these packages will be used throughout this document:

```
library(conflicted) # check for conflicting function definitions
# library(printr) # inserts help-file output into markdown output
library(rmarkdown) # Convert R Markdown documents into a variety of formats.
library(pander) # format tables for markdown
library(ggplot2) # graphics
library(ggfortify) # help with graphics
library(dplyr) # manipulate data
library(tibble) # `tibble`s extend `data.frame`s
library(magrittr) # `%>%` and other additional piping tools
library(haven) # import Stata files
library(knitr) # format R output for markdown
library(tidyr) # Tools to help to create tidy data
library(plotly) # interactive graphics
library(dobson) # datasets from Dobson and Barnett 2018
library(parameters) # format model output tables for markdown
```

```
library(haven) # import Stata files
library(latex2exp) # use LaTeX in R code (for figures and tables)
library(fs) # filesystem path manipulations
library(survival) # survival analysis
library(survminer) # survival analysis graphics
library(KMsurv) # datasets from Klein and Moeschberger
library(parameters) # format model output tables for
library(webshot2) # convert interactive content to static for pdf
library(forcats) # functions for categorical variables ("factors")
library(stringr) # functions for dealing with strings
library(lubridate) # functions for dealing with dates and times
library(broom) # Summarizes key information about statistical objects in tidy tibbles
library(broom.helpers) # Provides suite of functions to work with regression model 'broom::tidy()' t
```

Here are some R settings I use in this document:

```
rm(list = ls()) # delete any data that's already loaded into R

conflicts_prefer(dplyr::filter)
ggplot2::theme_set(
  ggplot2::theme_bw() +
    # ggplot2::labs(col = "") +
    ggplot2::theme(
      legend.position = "bottom",
      text = ggplot2::element_text(size = 12, family = "serif")))

knitr::opts_chunk$set(message = FALSE)
options('digits' = 6)

panderOptions("big.mark", ",")
pander::panderOptions("table.emphasize.rownames", FALSE)
pander::panderOptions("table.split.table", Inf)
conflicts_prefer(dplyr::filter) # use the `filter()` function from dplyr() by default
legend_text_size = 9
run_graphs = TRUE
```

Most of the content in this chapter should be review from UC Davis Epi 202.

1 Core properties of probabilities

1.1 Defining probabilities

Definition 1.1 (Probability measure). A **probability measure**, often denoted $\Pr()$ or $P()$, is a function whose domain is a σ -algebra^a of possible outcomes, $\mathcal{S}$, and which satisfies the following properties:

1. For any statistical event $A \in \mathcal{S}$, $\Pr(A) \geq 0$.
2. The probability of the union of all outcomes ($\Omega \stackrel{\text{def}}{=} \cup \mathcal{S}$) is 1:

$$\Pr(\Omega) = 1$$

3. The probability of the union of countably many mutually disjoint events $A_1, A_2, \dots$ (where $A_i \cap A_j = \emptyset$ for all $i \neq j$) is equal to the sum of their probabilities (*countable additivity* or *sigma-additivity*):

$$\Pr\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} \Pr(A_i)$$

^a<https://en.wikipedia.org/wiki/%CE%A3-algebra>

Property 3 (*countable additivity*) is stronger than *finite additivity*, which only requires

$$\Pr(A_1 \cup \dots \cup A_n) = \sum_{i=1}^n \Pr(A_i)$$

for every finite collection of mutually disjoint events. Countable additivity implies finite additivity (set $A_{n+1} = A_{n+2} = \dots = \emptyset$ in property 3, using $\Pr(\emptyset) = 0$), but not vice versa: there exist set functions that satisfy finite additivity but fail countable additivity (see Wikipedia: Sigma-additive set function — An additive function which is not σ -additive¹). Requiring countable additivity enables results such as the continuity of probability (if $A_1 \supseteq A_2 \supseteq \dots$ with $\bigcap_i A_i = \emptyset$, then $\Pr(A_i) \rightarrow 0$) and underpins the Theorem 1.4 for countable partitions.

Theorem 1.1 (Probability of a subset's intersection). *If A and B are statistical events and $A \subseteq B$, then $\Pr(A \cap B) = \Pr(A)$.*

i Proof

Proof. Left to the reader for now. □

Theorem 1.2 (An event and its complement sum to 1).

$$\Pr(A) + \Pr(\neg A) = 1$$

i Proof

Proof. By properties 2 and 3 of Definition 1.1. □

Corollary 1.1 (Complement rule).

$$\Pr(\neg A) = 1 - \Pr(A)$$

i Proof

Proof. By Theorem 1.2 and algebra. □

¹https://en.wikipedia.org/wiki/Sigma-additive_set_function#An_additive_function_which_is_not_%CF%83-additive

Corollary 1.2 (Complement rule in probability (π) notation). *If the probability of an outcome A is $\Pr(A) = \pi$, then the probability that A does not occur is:*

$$\Pr(\neg A) = 1 - \pi$$

i Proof

Proof. Using Corollary 1.1:

$$\begin{aligned}\Pr(\neg A) &= 1 - \Pr(A) \\ &= 1 - \pi\end{aligned}$$

□

1.2 Conditional probability

Definition 1.2 (Conditional probability). For two events A and B with $\Pr(B) > 0$, the **conditional probability** of A given B , denoted $\Pr(A \mid B)$, is:

$$\Pr(A \mid B) \stackrel{\text{def}}{=} \frac{\Pr(A \cap B)}{\Pr(B)}$$

Theorem 1.3 (Law of conditional probability). *For any two events A and B with $\Pr(B) > 0$:*

$$\Pr(A \cap B) = \Pr(A \mid B) \cdot \Pr(B)$$

i Proof

Proof. Rearranging Definition 1.2:

$$\begin{aligned}\Pr(A \mid B) &= \frac{\Pr(A \cap B)}{\Pr(B)} && \text{(definition of conditional probability)} \\ \Pr(A \cap B) &= \Pr(A \mid B) \cdot \Pr(B) && \text{(multiply both sides by } \Pr(B))\end{aligned}$$

□

Exm

Example 1.1 (Applying the law of conditional probability). Suppose 30% of adults exercise regularly ($\Pr(E) = 0.30$), and among adults who exercise regularly, 60% have low blood pressure ($\Pr(L \mid E) = 0.60$).

Then the probability that a randomly selected adult both exercises regularly and has low blood pressure is:

$$\begin{aligned}\Pr(L \cap E) &= \Pr(L \mid E) \cdot \Pr(E) \\ &= 0.60 \cdot 0.30 \\ &= 0.18\end{aligned}$$

Theorem 1.4 (Law of total probability). *If $B_1, B_2, \dots$ is a countable partition of the sample space (i.e., countably many mutually exclusive events whose union is the entire sample space),*

then for any event A :

$$\Pr(A) = \sum_{i=1}^{\infty} \Pr(A \mid B_i) \cdot \Pr(B_i)$$

i Proof

Proof. Since $B_1, B_2, \dots$ partition the sample space, the events $A \cap B_1, A \cap B_2, \dots$ are mutually exclusive and their union is A . By property 3 of Definition 1.1 (countable additivity), and then by Theorem 1.3:

$$\begin{aligned} \Pr(A) &= \sum_{i=1}^{\infty} \Pr(A \cap B_i) && \text{(countable additivity for partition of } A) \\ &= \sum_{i=1}^{\infty} \Pr(A \mid B_i) \cdot \Pr(B_i) && \text{(law of conditional probability)} \end{aligned}$$

□

Theorem 1.5 (Bayes' theorem). *For any two events A and B with $\Pr(A) > 0$ and $\Pr(B) > 0$:*

$$\Pr(A \mid B) = \frac{\Pr(B \mid A) \cdot \Pr(A)}{\Pr(B)}$$

i Proof

Proof. Apply Definition 1.2 to both $\Pr(A \mid B)$ and $\Pr(B \mid A)$:

$$\begin{aligned} \Pr(A \mid B) &= \frac{\Pr(A \cap B)}{\Pr(B)} \\ &= \frac{\Pr(B \mid A) \cdot \Pr(A)}{\Pr(B)} \end{aligned}$$

The second equality follows from Theorem 1.3 applied to $\Pr(B \cap A) = \Pr(B \mid A) \cdot \Pr(A)$. □

Exm

Example 1.2 (Positive predictive value of a medical test). Suppose a disease test has 99% sensitivity and 99% specificity, and the prevalence of the disease in the population is 7%. Let D be the event “person has the disease” and $+$ be the event “test is positive”. Then:

- $\Pr(+ \mid D) = 0.99$ (sensitivity)
- $\Pr(\neg+ \mid \neg D) = 0.99$ (specificity), so the false positive rate is $\Pr(+ \mid \neg D) = 1 - 0.99 = 0.01$
- $\Pr(D) = 0.07$ (prevalence)

By Bayes' theorem (Theorem 1.5) and the law of total probability (Theorem 1.4):

$$\begin{aligned}
\Pr(D \mid +) &= \frac{\Pr(+ \mid D) \cdot \Pr(D)}{\Pr(+)} \\
&= \frac{\Pr(+ \mid D) \cdot \Pr(D)}{\Pr(+ \mid D) \cdot \Pr(D) + \Pr(+ \mid \neg D) \cdot \Pr(\neg D)} \\
&= \frac{0.99 \cdot 0.07}{0.99 \cdot 0.07 + 0.01 \cdot 0.93} \\
&= \frac{0.0693}{0.0693 + 0.0093} \\
&= \frac{0.0693}{0.0786} \\
&\approx 0.88
\end{aligned}$$

Even with a highly accurate test (99% sensitive and 99% specific), only about 88% of people who test positive actually have the disease, because the disease prevalence is relatively low (7%).

2 Key probability distributions

Some distributions are typically used for outcome models (Table 1); other distributions are typically used for test statistics (Table 2).

Table 1: Distributions typically used for outcome models

Distribution	Uses
Bernoulli	Binary outcomes
Binomial	Sums of Bernoulli outcomes
Poisson	unbounded count outcomes
Geometric	Counts of non-events before an event occurs
Negative binomial	Mixtures of Poisson distributions, counts of non-events until a given number of events occurs
Normal (Gaussian)	Continuous outcomes without a more specific distribution
exponential	Time to event outcomes
Gamma	Time to event outcomes
Weibull	Time to event outcomes
Log-normal	Time to event outcomes

Table 2: Distributions typically used for test statistics

Distribution	Uses
χ^2	Regression comparisons (asymptotic), contingency table independence tests, goodness-of-fit tests
F	Gaussian model comparisons (exact)
Z (standard normal)	Proportions, means, regression coefficients (asymptotic)
T	Means, regression coefficients in Gaussian outcome models (exact)

2.1 The Bernoulli distribution

Definition 2.1 (Bernoulli distribution). The **Bernoulli distribution** family for a random variable X is defined as:

$$\begin{aligned}\Pr(X = x) &\stackrel{\text{def}}{=} 1_{x \in \{0,1\}} \pi^x (1 - \pi)^{1-x} \\ &= \begin{cases} \pi, & x = 1 \\ 1 - \pi, & x = 0 \end{cases}\end{aligned}$$

2.2 The Poisson distribution



(a) Siméon Denis Poisson



(b) Les Poissons^a

^a<https://youtu.be/UoJxBEQRLd0?t=12>

Figure 1: “Les Poissons”

Exercise 2.1. Define the Poisson distribution.

Solution

Solution 2.1.

Def

Definition 2.2 (Poisson distribution).

$$P(Y = y) \stackrel{\text{def}}{=} \frac{\mu^y e^{-\mu}}{y!}, \quad y \in \mathbb{N} \quad (1)$$

(see Figure 2)

Exercise 2.2. What is the range of possible values for a Poisson distribution?

Solution

Solution 2.2.

$$\mathcal{R}(Y) \stackrel{\text{def}}{=} \{0, 1, 2, \dots\} = \mathbb{N}$$

Theorem 2.1 (CDF of Poisson distribution).

$$P(Y \leq y) = e^{-\mu} \sum_{j=0}^{\lfloor y \rfloor} \frac{\mu^j}{j!} \quad (2)$$

i Proof

Proof. For any $y \geq 0$, the event $\{Y \leq y\}$ is the disjoint union of events $\{Y = j\}$ for all non-negative integers $j \leq \lfloor y \rfloor$. Applying the Poisson PMF (Equation 1) and factoring out the common term $e^{-\mu}$:

$$\begin{aligned} P(Y \leq y) &= \sum_{j=0}^{\lfloor y \rfloor} P(Y = j) && \text{(disjoint union of events } Y = j) \\ &= \sum_{j=0}^{\lfloor y \rfloor} \frac{\mu^j e^{-\mu}}{j!} && \text{(definition of Poisson PMF)} \\ &= e^{-\mu} \sum_{j=0}^{\lfloor y \rfloor} \frac{\mu^j}{j!} && \text{(factoring out } e^{-\mu} \text{ constant wrt } j) \end{aligned}$$

□

Exm

Example 2.1 (Example: Computing Poisson cumulative probabilities). For a Poisson random variable $X \sim \text{Pois}(\mu = 2)$, the probability of observing at most 2 events is computed as:

$$\begin{aligned} P(X \leq 2) &= e^{-2} \sum_{j=0}^2 \frac{2^j}{j!} && \text{(apply CDF formula with } \mu = 2, y = 2) \\ &= e^{-2} \left(\frac{2^0}{0!} + \frac{2^1}{1!} + \frac{2^2}{2!} \right) && \text{(expand terms for } j = 0, 1, 2) \\ &= e^{-2} (1 + 2 + 2) && \text{(simplify factorials and powers)} \\ &= 5e^{-2} \approx 0.677 && \text{(evaluate numerical value)} \end{aligned}$$

(see Figure 3)

```
library(dplyr)
pois_dists <- tibble(
  mu = c(0.5, 1, 2, 5, 10, 20)
) |>
  reframe(
    .by = mu,
    x = 0:30
  ) |>
  mutate(
    `P(X = x)` = dpois(x, lambda = mu),
    `P(X <= x)` = ppois(x, lambda = mu),
    mu = factor(mu)
  )

library(ggplot2)
library(latex2exp)

plot0 <- pois_dists |>
  ggplot(
    aes(
      x = x,
      y = `P(X = x)`,
      fill = mu,
      col = mu
    )
  ) +
```

```

theme(legend.position = "bottom") +
labs(
  fill = latex2exp::TeX("$\\mu$"),
  col = latex2exp::TeX("$\\mu$"),
  y = latex2exp::TeX("$\\Pr_{\\mu}(X = x)$")
)

plot1 <- plot0 +
  geom_segment(yend = 0) +
  facet_wrap(~mu)

print(plot1)

```

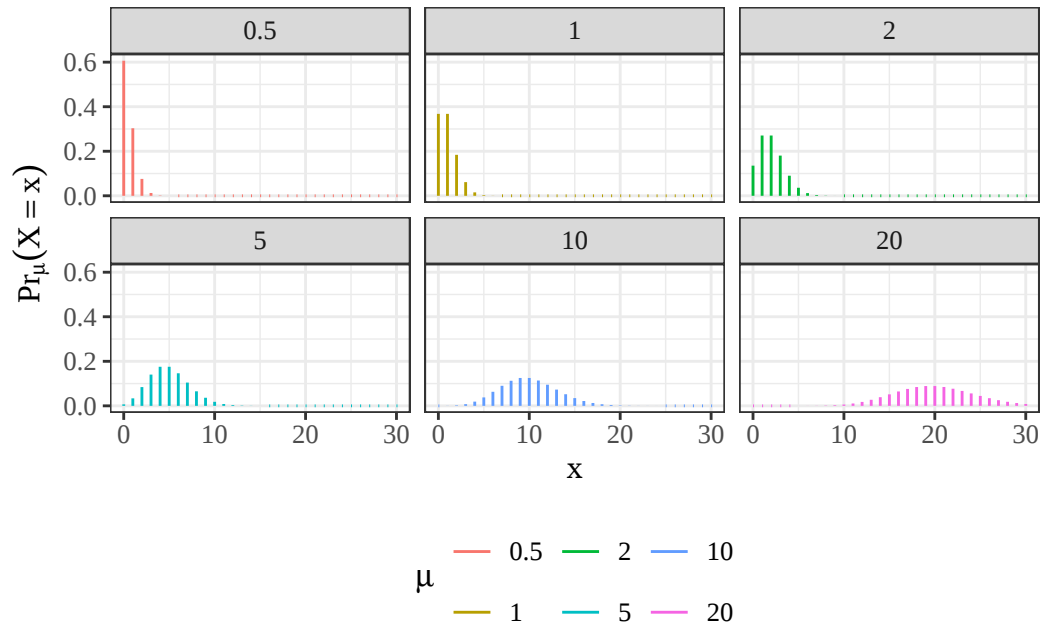


Figure 2: Poisson PMFs, by mean parameter μ

```

library(ggplot2)

plot2 <-
  plot0 +
  geom_step(alpha = 0.75) +
  aes(y = `P(X <= x)`) +
  labs(y = latex2exp::TeX("$\\Pr_{\\mu}(X \\leq x)$"))

print(plot2)

```

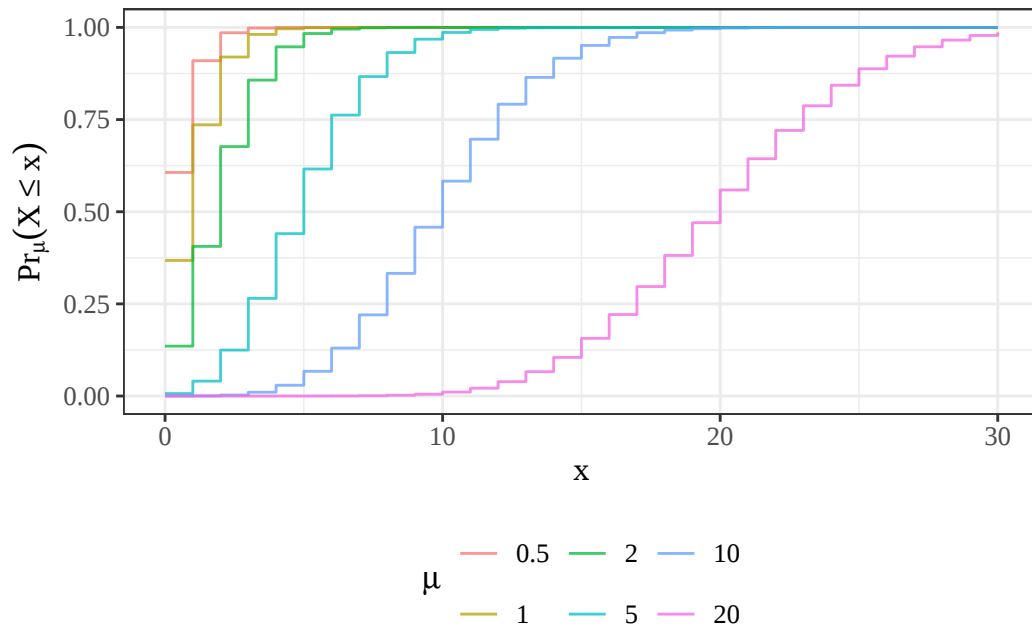


Figure 3: Poisson CDFs

Exercise 2.3 (Poisson distribution functions). Let $X \sim \text{Pois}(\mu = 3.75)$.

Compute:

- $P(X = 4 | \mu = 3.75)$
- $P(X \leq 7 | \mu = 3.75)$
- $P(X > 5 | \mu = 3.75)$

Solution

Solution.

- $P(X = 4) = 0.19378$
- $P(X \leq 7) = 0.962379$
- $P(X > 5) = 0.177117$

Theorem 2.2 (Properties of the Poisson distribution). If $X \sim \text{Pois}(\mu)$, then:

- $E[X] = \mu$
- $\text{Var}(X) = \mu$
- $P(X = x) = \frac{\mu}{x} P(X = x - 1)$
- For $x < \mu$, $P(X = x) > P(X = x - 1)$
- For $x = \mu$, $P(X = x) = P(X = x - 1)$
- For $x > \mu$, $P(X = x) < P(X = x - 1)$
- $\arg \max_x P(X = x) = \lfloor \mu \rfloor$

Exercise 2.4. Prove Theorem 2.2.

Solution

Solution.

$$\begin{aligned}
 E[X] &= \sum_{x=0}^{\infty} x \cdot P(X = x) && \text{(definition of expected value)} \\
 &= 0 \cdot P(X = 0) + \sum_{x=1}^{\infty} x \cdot P(X = x) && \text{(separate } x = 0 \text{ term)} \\
 &= \sum_{x=1}^{\infty} x \cdot \frac{\mu^x e^{-\mu}}{x!} && \text{(substitute Poisson PMF)} \\
 &= \sum_{x=1}^{\infty} x \cdot \frac{\mu^x e^{-\mu}}{x \cdot (x-1)!} && \text{(definition of factorial } x!) \\
 &= \sum_{x=1}^{\infty} \frac{\mu^x e^{-\mu}}{(x-1)!} && \text{(cancel factor of } x) \\
 &= \mu \cdot \sum_{x=1}^{\infty} \frac{\mu^{x-1} e^{-\mu}}{(x-1)!} && \text{(factor out one power of } \mu) \\
 &= \mu \cdot \sum_{y=0}^{\infty} \frac{\mu^y e^{-\mu}}{y!} && \text{(change index variable } y \stackrel{\text{def}}{=} x-1) \\
 &= \mu \cdot 1 && \text{(PMF sums to 1 over state space)} \\
 &= \mu && \text{(simplify)}
 \end{aligned}$$

See also <https://statproofbook.github.io/P/poiss-mean>.

For the variance, see <https://statproofbook.github.io/P/poiss-var>.

Accounting for exposure

Definition 2.3 (Exposure magnitude). For many count outcomes, there is some sense of an **exposure magnitude**, such as **population size**, or **duration of observation**, which multiplicatively rescales the expected (mean) count.

Exercise 2.5. What are some examples of exposure magnitudes?

Solution

Solution.

Table 3: Examples of exposure units

outcome	exposure units
disease incidence	number of individuals exposed; time at risk
car accidents	miles driven
worksite accidents	person-hours worked
population size	size of habitat

Exposure units are similar to the number of trials in a binomial distribution, but **in non-binomial count outcomes, there can be more than one event per unit of exposure**.

We can use t to represent continuous-valued exposures/observation durations, and n to represent discrete-valued exposures.

Definition 2.4 (Event rate).

For a count outcome Y with exposure magnitude t , the **event rate** (denoted λ) is defined as the mean of Y divided by the exposure magnitude. This mathematical relationship between mean and exposure magnitude is:

$$\begin{aligned}\mu &\stackrel{\text{def}}{=} \text{E}[Y|T = t] \\ \lambda &\stackrel{\text{def}}{=} \frac{\mu}{t}\end{aligned}\tag{3}$$

Event rate is somewhat analogous to odds in binary outcome models; it typically serves as an intermediate transformation between the mean of the outcome and the linear component of the model. However, in contrast with the odds function, the transformation $\lambda = \mu/t$ is *not* considered part of the Poisson model's link function, and it treats the exposure magnitude covariate differently from the other covariates.

Theorem 2.3 (Transformation function from event rate to mean). *For a count variable with mean μ , event rate λ , and exposure magnitude t :*

$$\mu = \lambda \cdot t\tag{4}$$

i Proof

Proof.

$$\begin{aligned}\lambda &\stackrel{\text{def}}{=} \frac{\mu}{t} && \text{(definition of event rate @eq-def-event-rate)} \\ \mu &= \lambda \cdot t && \text{(multiply both sides by } t > 0\text{)}\end{aligned}$$

□

Exm

Example 2.2 (Example: Calculating expected counts from event rates). Suppose a city records a disease event rate of $\lambda = 0.05$ cases per person-year. For a subpopulation with an exposure magnitude of $t = 100$ person-years, the expected count of cases is:

$$\begin{aligned}\mu &= \lambda \cdot t && \text{(apply transformation formula @eq-lambda-to-mu)} \\ &= 0.05 \times 100 && \text{(substitute } \lambda = 0.05 \text{ and } t = 100\text{)} \\ &= 5 \text{ cases} && \text{(evaluate expected count)}\end{aligned}$$

Solution

Solution. Start from definition of event rate and use algebra to solve for μ .

Equation 4 is analogous to the inverse-odds function for binary variables.

Theorem 2.4 (No exposure means no expected events). *When the exposure magnitude is 0, there is no opportunity for events to occur:*

$$\text{E}[Y|T = 0] = 0$$

i Proof

Proof.

$$\begin{aligned} E[Y \mid T = 0] &= \lambda \cdot 0 && \text{(apply transformation @eq-lambda-to-mu at } t = 0\text{)} \\ &= 0 && \text{(multiplication by zero)} \end{aligned}$$

□

Exm

Example 2.3 (Example: Zero exposure time). If a subject is observed for $t = 0$ person-years, no follow-up time has elapsed, so the expected number of incident events is $E[Y \mid T = 0] = 0$.

! Important

The exposure magnitude, T , is *similar* to a covariate in linear or logistic regression. However, there is an important difference: in count regression, **there is no intercept corresponding to** $E[Y \mid T = 0]$. In other words, this model assumes that if there is no exposure, there can't be any events.

Theorem 2.5 (Exposure is additive on the log scale). *If $\mu = \lambda \cdot t$, then:*

$$\log \mu = \log \lambda + \log t$$

i Proof

Proof.

$$\begin{aligned} \log \mu &= \log(\lambda \cdot t) && \text{(substitute } \mu = \lambda \cdot t \text{ from @eq-lambda-to-mu)} \\ &= \log \lambda + \log t && \text{(logarithm product rule)} \end{aligned}$$

□

Exm

Example 2.4 (Example: Log-linear representation of expected counts). If a clinic sees an event rate of $\lambda = 0.02$ events/day and $t = 30$ days of observation:

$$\begin{aligned} \log \mu &= \log(0.02) + \log(30) && \text{(apply log-scale formula)} \\ &= -3.912 + 3.401 && \text{(evaluate natural logarithms)} \\ &= -0.511 && \text{(sum terms)} \end{aligned}$$

Exponentiating yields $\mu = \exp\{-0.511\} \approx 0.60$ expected events.

Definition 2.5 (Offset). When the linear component of a model involves a term without an unknown coefficient, that term is called an **offset**.

Theorem 2.6 (Sum of independent Poisson random variables). *If X and Y are independent Poisson random variables with means μ_X and μ_Y , their sum, $Z = X + Y$, is also a Poisson random variable, with mean $\mu_Z = \mu_X + \mu_Y$.*

i Proof

Proof. Using the probability-generating function or PMF convolution for independent non-negative integer random variables:

$$\begin{aligned} P(Z = z) &= \sum_{k=0}^z P(X = k) P(Y = z - k) && \text{(independence and convolution formula)} \\ &= \sum_{k=0}^z \frac{\mu_X^k e^{-\mu_X}}{k!} \frac{\mu_Y^{z-k} e^{-\mu_Y}}{(z-k)!} && \text{(substitute Poisson PMFs)} \\ &= \frac{e^{-(\mu_X + \mu_Y)}}{z!} \sum_{k=0}^z \frac{z!}{k!(z-k)!} \mu_X^k \mu_Y^{z-k} && \text{(factor out } e^{-(\mu_X + \mu_Y)}/z! \text{ and multiply by } z!/z!) \\ &= \frac{e^{-(\mu_X + \mu_Y)}}{z!} (\mu_X + \mu_Y)^z && \text{(binomial theorem)} \end{aligned}$$

This probability expression matches the PMF of a $\text{Pois}(\mu_X + \mu_Y)$ random variable (see also https://web.stanford.edu/class/archive/cs/cs109/cs109.1206/lectureNotes/LN12_independent_rvs.pdf, Example 3). $\square$

Exm

Example 2.5 (Example: Aggregating independent region counts). Suppose Region A records $X \sim \text{Pois}(\mu_X = 12)$ cases and Region B records $Y \sim \text{Pois}(\mu_Y = 18)$ cases independently. The combined total count $Z = X + Y$ follows a Poisson distribution:

$$\begin{aligned} Z &\sim \text{Pois}(\mu_X + \mu_Y) && \text{(apply sum theorem @thm - sum - pois)} \\ &= \text{Pois}(12 + 18) && \text{(substitute region means)} \\ &= \text{Pois}(30) && \text{(evaluate sum)} \end{aligned}$$

2.3 The Negative-Binomial distribution

Definition 2.6 (Negative binomial distribution).

$$P(Y = y) \stackrel{\text{def}}{=} \frac{\mu^y}{y!} \cdot \frac{\Gamma(\rho + y)}{\Gamma(\rho) \cdot (\rho + \mu)^y} \cdot \left(1 + \frac{\mu}{\rho}\right)^{-\rho}$$

where ρ is an overdispersion parameter and $\Gamma(x) = (x-1)!$ for integers x .

You don't need to memorize or understand this expression.

As $\rho \rightarrow \infty$, the second term converges to 1 and the third term converges to $\exp\{-\mu\}$, which brings us back to the Poisson distribution.

Theorem 2.7 (Mean and variance of the negative binomial distribution). *If $Y \sim \text{NegBin}(\mu, \rho)$, then:*

- $E[Y] = \mu$
- $\text{Var}(Y) = \mu + \frac{\mu^2}{\rho} > \mu$

2.4 Weibull Distribution

$$\begin{aligned}p(t) &= \alpha \lambda x^{\alpha-1} e^{-\lambda x^\alpha} \\ \lambda(t) &= \alpha \lambda x^{\alpha-1} \\ S(t) &= e^{-\lambda x^\alpha} \\ E(T) &= \Gamma(1 + 1/\alpha) \cdot \lambda^{-1/\alpha}\end{aligned}$$

When $\alpha = 1$, the Weibull distribution reduces to the exponential distribution. When $\alpha > 1$ the hazard is increasing and when $\alpha < 1$ the hazard is decreasing. This property provides more flexibility than the exponential.

We will see more of this distribution later.

3 Characteristics of probability distributions

3.1 Probability density function

Definition 3.1 (probability density). If X is a continuous random variable, then the **probability density** of X at value x , denoted $f(x)$, $f_X(x)$, $p(x)$, $p_X(x)$, or $p(X = x)$, is defined as the limit of the **probability** (mass) that X is in an interval around x , divided by the width of that interval, as that width reduces to 0.

$$f(x) \stackrel{\text{def}}{=} \lim_{\Delta \rightarrow 0} \frac{P(X \in [x, x + \Delta])}{\Delta}$$

See also Rothman et al. (2021) (Chapter 22, p. 535) and https://en.wikipedia.org/wiki/Probability_density_function#Formal_definition

Definition 3.2 (Cumulative distribution function (CDF)). For a random variable X , its population CDF is

$$F(t) \stackrel{\text{def}}{=} \Pr(X \leq t), \quad t \in \mathbb{R}.$$

Definition 3.3 (Quantile function (population inverse CDF)). For a random variable X with **cumulative distribution function (CDF)** F , its population quantile function (generalized inverse of F) is

$$Q(p) \stackrel{\text{def}}{=} \inf\{t : F(t) \geq p\}, \quad 0 < p \leq 1.$$

Theorem 3.1 (Density function is derivative of CDF). *The density function $f(t)$ or $p(T = t)$ for a random variable T at value t is equal to the derivative of the cumulative probability function $F(t) \stackrel{\text{def}}{=} P(T \leq t)$; that is:*

$$f(t) \stackrel{\text{def}}{=} \frac{\partial}{\partial t} F(t)$$

Theorem 3.2 (Density functions integrate to 1). *For any density function $f(x)$,*

$$\int_{x \in \mathcal{R}(X)} f(x) dx = 1$$

3.2 Hazard function

Definition 3.4 (Hazard function, hazard rate, hazard rate function).

The **hazard function**, **hazard rate**, **hazard rate function**, for a random variable T at value t , typically denoted as $h(t)$ ² or $\lambda(t)$, ³ is the conditional density^a of T at t , given $T \geq t$. Specifically, the hazard function is defined as:

$$\lambda(t) \stackrel{\text{def}}{=} p(T = t | T \geq t)$$

If T represents the time at which an event occurs, then $\lambda(t)$ is the probability that the event occurs at time t , given that it has not occurred prior to time t .

The name “hazard” carries a connotation that the event is undesirable — death, relapse, equipment failure, etc. When the event in question is neutral or desirable (recovery, conception, graduation, response to treatment), the same quantity $\lambda(t)$ is often called the **event incidence rate** instead. This terminology is parallel to the convention that conditional probabilities of undesirable events are called **risks**, while the same conditional probabilities for neutral/desirable events are simply called **probabilities**. The math is identical; only the name changes with the valence of the event.

^a[probability.qmd#def-pdf](#)

Table 4: Probability distribution functions

Name	Symbols	Definition
Probability density function (PDF)	$f(t), p(t)$	$p(T = t)$
Cumulative distribution function (CDF)	$F(t), P(t)$	$P(T \leq t)$
Survival function	$S(t), \bar{F}(t)$	$P(T > t)$
Hazard function	$\lambda(t), h(t)$	$p(T = t T \geq t)$
Cumulative hazard function	$\Lambda(t), H(t)$	$\int_{u=-\infty}^t \lambda(u) du$
Log-hazard function	$\eta(t)$	$\log\{\lambda(t)\}$

$$f(t) \xleftarrow{\frac{-S'(t)}{S(t)\lambda(t)}} S(t) \xleftarrow{\exp\{-\Lambda(t)\}} \Lambda(t) \xleftarrow{\int_{u=0}^t \lambda(u) du} \lambda(t) \xleftarrow{\exp\{\eta(t)\}} \eta(t)$$

$$f(t) \xrightarrow[\int_{u=t}^{\infty} f(u) du]{f(t)/\lambda(t)} S(t) \xrightarrow[-\log S(t)]{} \Lambda(t) \xrightarrow[\Lambda'(t)]{} \lambda(t) \xrightarrow[\log\{\lambda(t)\}]{} \eta(t)$$

3.3 Expectation

Definition 3.5 (Expectation, expected value, population mean). The **expectation**, **expected value**, or **population mean** of a *continuous* random variable X , denoted $E[X]$, $\mu(X)$, or μ_X , is the weighted mean of X ’s possible values, weighted by the [probability density function](#) of those values:

$$E[X] \stackrel{\text{def}}{=} \int_{x \in \mathcal{R}(X)} x \cdot p(X = x) dx$$

The **expectation**, **expected value**, or **population mean** of a *discrete* random variable X , denoted $E[X]$, $\mu(X)$, or μ_X , is the mean of X ’s possible values, weighted by the probability mass function of those values:

³for example in Dobson and Barnett (2018), Vittinghoff et al. (2012), Klein and Moeschberger (2003), and Kleinbaum and Klein (2012)

³for example, in Rothman et al. (2021) and Kalbfleisch and Prentice (2011)

$$E[X] \stackrel{\text{def}}{=} \sum_{x \in \mathcal{R}(X)} x \cdot P(X = x)$$

(c.f. https://en.wikipedia.org/wiki/Expected_value)

Theorem 3.3 (Expectation of the Bernoulli distribution). *The expectation of a Bernoulli random variable with parameter π is:*

$$E[X] = \pi$$

i Proof

Proof.

$$\begin{aligned} E[X] &= \sum_{x \in \mathcal{R}(X)} x \cdot P(X = x) && \text{(definition of expectation for discrete r.v.)} \\ &= \sum_{x \in \{0,1\}} x \cdot P(X = x) && \text{(range of Bernoulli r.v. is } \{0,1\}) \\ &= (0 \cdot P(X = 0)) + (1 \cdot P(X = 1)) && \text{(expand sum over } x = 0 \text{ and } x = 1) \\ &= (0 \cdot (1 - \pi)) + (1 \cdot \pi) && \text{(substitute Bernoulli PMF values)} \\ &= 0 + \pi && \text{(multiplication by 0 and 1)} \\ &= \pi && \text{(addition of 0)} \end{aligned}$$

□

Theorem 3.4 (Expectation of time-to-event variables). *If T is a non-negative random variable, then:*

$$E[T] = \int_{t=0}^{\infty} S(t) dt$$

i Proof

Proof. We prove the continuous case, in which T has a density f . The result follows from applying Tonelli's theorem (condition (a) of Fubini–Tonelli^a) to the function $g(t, u) = f(u) \cdot \mathbb{1}(0 \leq t \leq u)$ on the product space $[0, \infty) \times [0, \infty)$: g is nonnegative everywhere and vanishes outside the (unbounded) triangular region $D = \{(t, u) : 0 \leq t \leq u < \infty\}$, so the iterated integrals over D are exchangeable.

Since $f(u) \geq 0$, condition (a) of Fubini–Tonelli^b (the nonnegative case, **Tonelli's theorem**) applies, and we may exchange the order of integration over D :

$$\begin{aligned}
E[T] &= \int_{u=0}^{\infty} u f(u) du \\
&= \int_{u=0}^{\infty} \left(\int_{t=0}^u 1 dt \right) f(u) du \\
&= \int_{u=0}^{\infty} \int_{t=0}^u f(u) dt du \\
&= \int_{t=0}^{\infty} \int_{u=t}^{\infty} f(u) du dt \\
&= \int_{t=0}^{\infty} P(T > t) dt \\
&= \int_{t=0}^{\infty} S(t) dt.
\end{aligned}$$

See (Soch 2020) for an alternative presentation of this proof. □

^a[math-prereqs.qmd#thm-fubini-tonelli](#)
^b[math-prereqs.qmd#thm-fubini-tonelli](#)

Exm

Example 3.1 (Mean of an exponential random variable via survival function). Let $T \sim \text{Exponential}(\lambda)$, so $S(t) = e^{-\lambda t}$ for $t \geq 0$. By Theorem 3.4:

$$\begin{aligned}
E[T] &= \int_0^{\infty} S(t) dt \\
&= \int_0^{\infty} e^{-\lambda t} dt \\
&= \left[-\frac{1}{\lambda} e^{-\lambda t} \right]_0^{\infty} \\
&= \frac{1}{\lambda},
\end{aligned}$$

confirming the standard result $E[T] = 1/\lambda$.

Theorem 3.5 (Law of the Unconscious Statistician (LOTUS)). ***Discrete case.** For any function g of a discrete random variable X :*

$$E[g(X)] = \sum_{x \in \mathcal{R}(X)} g(x) \cdot P(X = x)$$

***Continuous case.** For any function g of a continuous random variable X with density $p(X = x)$:*

$$E[g(X)] = \int_{x \in \mathcal{R}(X)} g(x) \cdot p(X = x) dx$$

i Proof

Proof. We prove the discrete case.

Let $Y = g(X)$. By Definition 3.5 applied to Y :

$$\begin{aligned}
\mathbb{E}[g(X)] &= \mathbb{E}[Y] && \text{(substitution } Y = g(X)) \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot \mathbb{P}(Y = y) && \text{(definition of expectation for discrete r.v.)} \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot \mathbb{P}(g(X) = y) && \text{(substitute } Y = g(X) \text{ in probability expression)} \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot \sum_{\substack{x \in \mathcal{R}(X) \\ g(x) = y}} \mathbb{P}(X = x) && \text{(law of total probability over } \{x : g(x) = y\}) \\
&= \sum_{x \in \mathcal{R}(X)} g(x) \cdot \mathbb{P}(X = x) && \text{(rearrange double sum grouping by } g(x))
\end{aligned}$$

where the last equality follows by rearranging the double sum, grouping each term x by its image $y = g(x)$.

The continuous case is the density-weighted analogue of this argument, but a fully rigorous proof needs the general change-of-variables theorem for integrals against a pushforward measure — approximating g by simple functions and passing to the limit — which is beyond Epi 204's scope (Billingsley 1995; Gut 2013; Casella and Berger 2002; Wikipedia contributors 2026). $\square$

LOTUS says that to compute $\mathbb{E}[g(X)]$, we do not need to first find the distribution of $g(X)$; we can compute the expectation directly using the distribution of X .

Exm

Example 3.2 (Expected value of X^2 for a Bernoulli variable). Let $X \sim \text{Ber}(\pi)$. By LOTUS (Theorem 3.5, discrete case):

$$\begin{aligned}
\mathbb{E}[X^2] &= \sum_{x \in \{0,1\}} x^2 \cdot \mathbb{P}(X = x) \\
&= 0^2 \cdot \mathbb{P}(X = 0) + 1^2 \cdot \mathbb{P}(X = 1) \\
&= 0^2 \cdot (1 - \pi) + 1^2 \cdot \pi \\
&= 0 + \pi \\
&= \pi
\end{aligned}$$

Exm

Example 3.3 (Expected value of X^2 for a Uniform(0,1) variable). Let $X \sim \text{Uniform}(0,1)$, so $p(X = x) = 1$ for $x \in [0,1]$. By LOTUS (Theorem 3.5, continuous case):

$$\begin{aligned}
\mathbb{E}[X^2] &= \int_0^1 x^2 \cdot p(X = x) \, dx && \text{(LOTUS, continuous case)} \\
&= \int_0^1 x^2 \cdot 1 \, dx && \text{(p(X = x) = 1 on [0, 1])} \\
&= \left[\frac{x^3}{3} \right]_0^1 && \text{(antiderivative of } x^2) \\
&= \frac{1}{3}. && \text{(evaluate at the bounds)}
\end{aligned}$$

3.3.1 Conditional distributions and expectations

Definition 3.6 (Conditional probability mass function). Let X and Y be jointly distributed discrete random variables. The **conditional probability mass function** of Y given $X = x$ (for values of x with $P(X = x) > 0$) is:

$$P(Y = y \mid X = x) \stackrel{\text{def}}{=} \frac{P(X = x, Y = y)}{P(X = x)}$$

Exm

Example 3.4 (Conditional PMF from a vaccine trial). The `vaccine` dataset in the `dobson` package (Dobson and Barnett (2018), Table 9.6) records the response to treatment in a vaccine trial: X is treatment group (placebo or vaccine) and Y is response severity (small, moderate, or large).

```
library(dobson)
library(dplyr)
data(vaccine)
response_levels <- c("small", "moderate", "large")
vaccine_tab <- vaccine |>
  mutate(response = factor(response, levels = response_levels)) |>
  xtabs(frequency ~ treatment + response, data = _)
pander::pander(as.data.frame.matrix(vaccine_tab))
```

	small	moderate	large
placebo	25	8	5
vaccine	6	18	11

The joint PMF of (X, Y) is the table of frequencies divided by the total $n = 73$ trial participants. The marginal probability $P(X = \text{placebo})$ is:

$$\begin{aligned} P(X = \text{placebo}) &= P(X = \text{placebo}, Y = \text{small}) + P(X = \text{placebo}, Y = \text{moderate}) + P(X = \text{placebo}, Y = \text{large}) \\ &= \frac{25}{73} + \frac{8}{73} + \frac{5}{73} \\ &= \frac{38}{73} \end{aligned}$$

By Definition 3.6, the conditional PMF of Y given $X = \text{placebo}$ is:

$$\begin{aligned} P(Y = \text{small} \mid X = \text{placebo}) &= \frac{P(X = \text{placebo}, Y = \text{small})}{P(X = \text{placebo})} \\ &= \frac{25/73}{38/73} \\ &= \frac{25}{38} \end{aligned}$$

Definition 3.7 (Conditional probability density function). Let X and Y be jointly distributed continuous random variables with joint density $p(X = x, Y = y)$ and marginal density $p(X = x)$. The **conditional probability density function** of Y given $X = x$ (for values of x with $p(X = x) > 0$) is:

$$p(Y = y \mid X = x) \stackrel{\text{def}}{=} \frac{p(X = x, Y = y)}{p(X = x)}$$

Exm

Example 3.5 (Conditional PDF from a bivariate normal model of birthweight data). The `birthweight` dataset in the `dobson` package (Dobson and Barnett (2018), Table 2.3) records gestational age (weeks) and birthweight (grams) for 12 boys and 12 girls. Let X be gestational age and Y be birthweight, pooling both sexes into $n = 24$ observations.

```
library(dobson)
data(birthweight)
ga <- c(birthweight[["boys gestational age"]],
        birthweight[["girls gestational age"]])
wt <- c(birthweight[["boys weight"]], birthweight[["girls weight"]])
mu_x <- mean(ga)
sigma_x <- sd(ga)
mu_y <- mean(wt)
sigma_y <- sd(wt)
rho <- cor(ga, wt)
beta <- cov(ga, wt) / var(ga)
intercept <- mu_y - beta * mu_x
x0 <- 40
marg_dens_x0 <- dnorm(x0, mean = mu_x, sd = sigma_x)
cond_mean <- intercept + beta * x0
cond_sd <- sigma_y * sqrt(1 - rho^2)
pander::pander(data.frame(
  parameter = c("mu_x", "sigma_x", "mu_y", "sigma_y", "rho"),
  estimate = round(c(mu_x, sigma_x, mu_y, sigma_y, rho), 4)
))
```

parameter	estimate
mu_x	38.54
sigma_x	1.817
mu_y	2,968
sigma_y	282.1
rho	0.7443

Modeling (X, Y) as bivariate normal with parameters set equal to these sample moments, the joint density is (a standard result; e.g. Casella and Berger (2002)):

$$p(X = x, Y = y) = \frac{1}{2\pi\sigma_X\sigma_Y\sqrt{1-\rho^2}} e^{-\frac{1}{2(1-\rho^2)} \left[\frac{(x-\mu_X)^2}{\sigma_X^2} - \frac{2\rho(x-\mu_X)(y-\mu_Y)}{\sigma_X\sigma_Y} + \frac{(y-\mu_Y)^2}{\sigma_Y^2} \right]}$$

A further standard fact about the bivariate normal (Casella and Berger (2002)) is that the marginal distribution of X is $X \sim N(\mu_X, \sigma_X^2)$. At $x = 40$ weeks (a full-term pregnancy), $\mu_X = 38.5417$ and $\sigma_X = 1.8173$, so:

$$\begin{aligned} p(X = 40) &= \frac{1}{\sigma_X\sqrt{2\pi}} e^{-\frac{(40-\mu_X)^2}{2\sigma_X^2}} \\ &= \frac{1}{1.8173\sqrt{2\pi}} e^{-\frac{(40-38.5417)^2}{2(1.8173)^2}} \\ &\approx 0.1591 \end{aligned}$$

By Definition 3.7, dividing the joint density by this marginal density and simplifying the exponent (completing the square in y ; Casella and Berger (2002)) gives the conditional PDF of Y given $X = 40$, which is itself normal with mean shifted along the regression line and variance reduced by a factor of $1 - \rho^2$:

$$\begin{aligned} p(Y = y \mid X = 40) &= \frac{p(X = 40, Y = y)}{p(X = 40)} \\ &= \frac{1}{\sigma_Y\sqrt{2\pi(1-\rho^2)}} e^{-\frac{1}{2(1-\rho^2)} \left[\frac{(40-\mu_X)^2}{\sigma_X^2} - \frac{2\rho(40-\mu_X)(y-\mu_Y)}{\sigma_X\sigma_Y} + \frac{(y-\mu_Y)^2}{\sigma_Y^2} \right] + \frac{(40-\mu_X)^2}{2\sigma_X^2}} \\ &= \frac{1}{\sigma_Y\sqrt{2\pi(1-\rho^2)}} e^{-\frac{\rho^2(40-\mu_X)^2}{2\sigma_X^2(1-\rho^2)} + \frac{\rho(40-\mu_X)(y-\mu_Y)}{\sigma_X\sigma_Y(1-\rho^2)} - \frac{(y-\mu_Y)^2}{2\sigma_Y^2(1-\rho^2)}} \\ &= \frac{1}{\sigma_Y\sqrt{2\pi(1-\rho^2)}} e^{-\frac{1}{2\sigma_Y^2(1-\rho^2)} \left[\rho^2 \frac{\sigma_Y^2}{\sigma_X^2} (40-\mu_X)^2 - 2\rho \frac{\sigma_Y}{\sigma_X} (40-\mu_X)(y-\mu_Y) + (y-\mu_Y)^2 \right]} \\ &= \frac{1}{\sigma_Y\sqrt{2\pi(1-\rho^2)}} e^{-\frac{(y - [\mu_Y + \rho \frac{\sigma_Y}{\sigma_X} (40-\mu_X)])^2}{2\sigma_Y^2(1-\rho^2)}} \end{aligned}$$

The first equality substitutes the joint and marginal densities and combines their prefactors and exponents; the second equality combines the two $(40 - \mu_X)^2$ terms using $1 - \frac{1}{1-\rho^2} = \frac{-\rho^2}{1-\rho^2}$; the third equality factors $-\frac{1}{2\sigma_Y^2(1-\rho^2)}$ out of the bracket (multiplying through confirms it reproduces the previous line); and the fourth equality completes the square in y , since expanding $(y - [\mu_Y + \rho \frac{\sigma_Y}{\sigma_X} (40 - \mu_X)])^2$ reproduces exactly the bracketed quadratic in the third equality.

So $Y \mid X = 40 \sim N(3136.15, 188.37^2)$: the conditional mean, 3136.15 g, is exactly the fitted regression line's prediction at $x = 40$ ($-1484.9846 + 115.5283 \times 40 = 3136.15$), matching R's `lm(wt ~ ga)` fit directly:

```
pander::pander(coef(lm(wt ~ ga)))
```

(Intercept)	ga
-1,485	115.5

Definition 3.8 (Conditional expectation). **Discrete case.** Let X and Y be jointly distributed discrete random variables. The **conditional expectation** of Y given $X = x$, using Definition 3.6, is:

$$E[Y | X = x] \stackrel{\text{def}}{=} \sum_{y \in \mathcal{R}(Y)} y \cdot P(Y = y | X = x)$$

Continuous case. Let X and Y be jointly distributed continuous random variables. The **conditional expectation** of Y given $X = x$, using Definition 3.7, is:

$$E[Y | X = x] \stackrel{\text{def}}{=} \int_{y \in \mathcal{R}(Y)} y \cdot p(Y = y | X = x) dy$$

Exm

Example 3.6 (Conditional expectation from real trial and birthweight data). **Discrete case.** Continuing Example 3.4, score the vaccine trial's response levels small = 1, moderate = 2, large = 3. The conditional PMF of Y given $X = \text{placebo}$ (from Example 3.4) is:

$$\begin{aligned} P(Y = \text{small} | X = \text{placebo}) &= \frac{25}{38} \\ P(Y = \text{moderate} | X = \text{placebo}) &= \frac{8}{38} \\ P(Y = \text{large} | X = \text{placebo}) &= \frac{5}{38} \end{aligned}$$

so:

$$\begin{aligned} E[Y | X = \text{placebo}] &= 1 \cdot \frac{25}{38} + 2 \cdot \frac{8}{38} + 3 \cdot \frac{5}{38} \\ &= \frac{25 + 16 + 15}{38} \\ &= \frac{56}{38} \\ &= \frac{28}{19} \\ &\approx 1.47 \end{aligned}$$

Continuous case. Continuing Example 3.5, $Y | X = 40 \sim N(3136.15, 188.37^2)$. Since the mean of a normal random variable is its own location parameter, integrating y against this conditional density directly gives:

$$\begin{aligned} E[Y | X = 40] &= \int_{-\infty}^{\infty} y \cdot p(Y = y | X = 40) dy \\ &= 3136.15 \text{ g} \end{aligned}$$

matching the fitted regression line's prediction at $x = 40$ weeks exactly, as expected since the conditional mean of a bivariate normal *is* the linear regression of Y on X .

Definition 3.9 (Conditional expectation: mixed case). Suppose exactly one of X, Y is discrete and the other is continuous. Write $p(X = x, Y = y)$ for their **joint density-mass function**: a probability density in the continuous variable and a probability mass in the discrete variable. **X discrete, Y continuous.** Here $p(X = x, Y = y)$ is, for each fixed x , a probability density in y , scaled so that $\int_y p(X = x, Y = y) dy = P(X = x)$. The conditional PDF of Y given $X = x$ (for values of x with $P(X = x) > 0$) is:

$$p(Y = y | X = x) \stackrel{\text{def}}{=} \frac{p(X = x, Y = y)}{P(X = x)}$$

and the conditional expectation of Y given $X = x$ is:

$$E[Y \mid X = x] \stackrel{\text{def}}{=} \int_{y \in \mathcal{R}(Y)} y \cdot p(Y = y \mid X = x) dy$$

X continuous, Y discrete. Here $p(X = x, Y = y)$ is, for each fixed y , a probability density in x , scaled so that $\sum_y p(X = x, Y = y) = p(X = x)$. The conditional PMF of Y given $X = x$ (for values of x with $p(X = x) > 0$) is:

$$P(Y = y \mid X = x) \stackrel{\text{def}}{=} \frac{p(X = x, Y = y)}{p(X = x)}$$

and the conditional expectation of Y given $X = x$ is:

$$E[Y \mid X = x] \stackrel{\text{def}}{=} \sum_{y \in \mathcal{R}(Y)} y \cdot P(Y = y \mid X = x)$$

Exm

Example 3.7 (Conditional expectation, one discrete variable and one continuous variable). **X discrete, Y continuous.** The `plasma` dataset in the `dobson` package (Dobson and Barnett (2018), Table 6.25) records plasma inorganic phosphate levels (mg/dL) one hour after a glucose tolerance test, for hyperinsulinemic obese (H-O) and control (C) participants. Let X be group and Y be phosphate level.

```
library(dobson)
library(dplyr)
data(plasma)
plasma_subset <- plasma |> filter(Group %in% c("H-O", "C"))
plasma_summary <- plasma_subset |>
  summarise(n = n(), mean = mean(phosphate), sd = sd(phosphate), .by = Group)
n_ho <- plasma_summary$n[plasma_summary$Group == "H-O"]
n_c <- plasma_summary$n[plasma_summary$Group == "C"]
mean_ho <- plasma_summary$mean[plasma_summary$Group == "H-O"]
sd_ho <- plasma_summary$sd[plasma_summary$Group == "H-O"]
mean_c <- plasma_summary$mean[plasma_summary$Group == "C"]
sd_c <- plasma_summary$sd[plasma_summary$Group == "C"]
pander::pander(plasma_summary)
```

Group	n	mean	sd
H-O	11	3.945	0.7776
C	12	2.783	0.4086

Modeling phosphate as approximately normal within each group, with parameters set equal to each group's sample mean and SD, and $P(X = \text{H-O}) = 11/23$:

$$p(Y = y \mid X = \text{H-O}) \approx N(3.95, 0.78^2)$$

$$p(Y = y \mid X = \text{C}) \approx N(2.78, 0.41^2)$$

By Definition 3.9, since the mean of a normal random variable is its own location parameter:

$$E[Y \mid X = \text{H-O}] = 3.95 \text{ mg/dL}, \quad E[Y \mid X = \text{C}] = 2.78 \text{ mg/dL}$$

X continuous, Y discrete. The **senility** dataset in the **dobson** package (Dobson and Barnett (2018), Table 7.8) records, for 54 elderly people, a WAIS (Wechsler Adult Intelligence Scale) score and whether symptoms of senility were present. Let X be WAIS score and Y indicate senility symptoms.

```
data(senility)
senility_fit <- glm(s ~ x, data = senility, family = binomial)
b0 <- coef(senility_fit)[["(Intercept)"]]
b1 <- coef(senility_fit)[["x"]]
pander::pander(coef(senility_fit))
```

(Intercept)	x
2.404	-0.3235

Modeling $P(Y = 1 \mid X = x)$ with the fitted logistic regression:

$$\text{logit } P(Y = 1 \mid X = x) = 2.4040 - 0.3235 x$$

By Definition 3.9, since $Y \mid X = x$ is Bernoulli:

$$\begin{aligned} E[Y \mid X = x] &= 0 \cdot P(Y = 0 \mid X = x) + 1 \cdot P(Y = 1 \mid X = x) \\ &= P(Y = 1 \mid X = x) \\ &= \frac{1}{1 + e^{-(2.4040 - 0.3235 x)}} \end{aligned}$$

```
senility_p10 <- predict(senility_fit,
                        newdata = data.frame(x = 10), type = "response")
senility_p15 <- predict(senility_fit,
                        newdata = data.frame(x = 15), type = "response")
```

At $x = 10$: $E[Y \mid X = 10] = 0.3$. At $x = 15$: $E[Y \mid X = 15] = 0.08$ — a lower WAIS score is associated with higher predicted probability of senility symptoms.

Definition 3.10 (Conditional expectation function). The **conditional expectation function** $E[Y \mid X]$ is the function (and hence random variable) of X obtained by evaluating $E[Y \mid X = x]$ at X ; specifically, $E[Y \mid X] = g(X)$ where $g(x) \stackrel{\text{def}}{=} E[Y \mid X = x]$.

Exm

Example 3.8 (The conditional expectation function of the birthweight model). Continuing Example 3.5 and Example 3.6, (X, Y) (gestational age, birthweight) is modeled as bivariate normal. The general form of the conditional mean derived in Example 3.5, evaluated at an arbitrary x instead of just $x = 40$, gives the conditional expectation function directly:

$$\begin{aligned}
g(x) &= \mu_Y + \rho \frac{\sigma_Y}{\sigma_X} (x - \mu_X) \\
&= -1484.9846 + 115.5283 x
\end{aligned}$$

which is exactly the fitted regression line (the general algebraic identity $\mu_Y + \rho \frac{\sigma_Y}{\sigma_X} (x - \mu_X) = (\mu_Y - \rho \frac{\sigma_Y}{\sigma_X} \mu_X) + \rho \frac{\sigma_Y}{\sigma_X} x$ is what makes the bivariate-normal conditional mean linear in x). As a check, evaluating at $x = 40$ recovers Example 3.6's result:

$$g(40) = -1484.9846 + 115.5283 \times 40 = 3136.15 \text{ g}$$

Exercise 3.1 (Expectation of a sum, given a joint PMF). Let (X, Y) be discrete with joint probability mass function:

	$Y = 0$	$Y = 1$
$X = 0$	0.2	0.3
$X = 1$	0.1	0.4

Compute $E[X + Y]$.

Solution

Solution. Treating (X, Y) as a single discrete random object taking one of the four values $(0, 0)$, $(0, 1)$, $(1, 0)$, $(1, 1)$, LOTUS (Theorem 3.5) applied to $g(x, y) = x + y$ gives directly:

$$\begin{aligned}
E[X + Y] &= \sum_{x \in \{0, 1\}} \sum_{y \in \{0, 1\}} (x + y) P(X = x, Y = y) \\
&= (0+0)(0.2) + (0+1)(0.3) + (1+0)(0.1) + (1+1)(0.4) \\
&= 0 + 0.3 + 0.1 + 0.8 \\
&= 1.2
\end{aligned}$$

As a check: $E[X] = 0(0.5) + 1(0.5) = 0.5$ and $E[Y] = 0(0.3) + 1(0.7) = 0.7$, so $E[X + Y] = E[X] + E[Y] = 1.2$.

```

x_labs <- c("X=0", "X=0", "X=1", "X=1")
y_labs <- c("Y=0", "Y=1", "Y=0", "Y=1")
probs <- c(0.2, 0.3, 0.1, 0.4)

plotly::plot_ly(
  x = ~y_labs, y = ~probs, color = ~x_labs,
  colors = c("steelblue", "tomato"),
  type = "bar"
) |>
plotly::layout(
  barmode = "group",
  xaxis = list(title = "Y"),
  yaxis = list(title = "P(X = x, Y = y)", range = c(0, 0.5)),
  legend = list(title = list(text = "X value"))
)

```

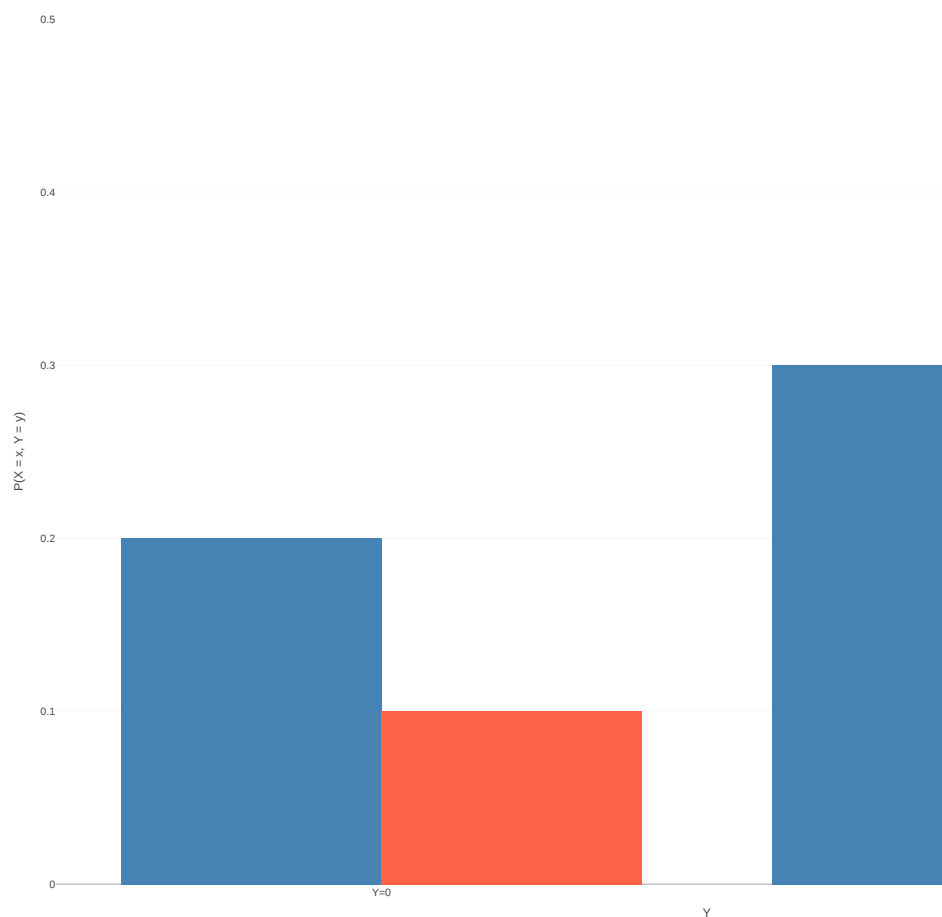


Figure 4: Joint probability mass function $P(X = x, Y = y)$. Marginal totals: $P(X = 0) = 0.5$, $P(X = 1) = 0.5$, $P(Y = 0) = 0.3$, $P(Y = 1) = 0.7$.

Note that X and Y are **not** independent here: $P(X = 0, Y = 0) = 0.2 \neq 0.15 = P(X = 0) P(Y = 0)$. LOTUS applies regardless, since it requires only the *actual* joint mass function, not independence.

There are only four (x, y) pairs here, so summing them in any order — row by row, column by column, or any other listing — gives the same total by ordinary commutativity and associativity of addition; no result about interchanging summation order is needed. Fubini–Tonelli, covered next, is what’s needed once the support is countably infinite instead — see Exercise 3.2.

3.4 Fubini–Tonelli for expectations

The Riemann version of Fubini’s theorem⁴, stated in the math-prereqs chapter⁵, lets us switch the order of integration for continuous integrands on simple regions. For expectations against probability measures we use its measure-theoretic form⁶, which holds on any σ -finite measure space. The σ -finiteness condition is automatic for probability measures (every probability measure is finite, hence σ -finite), so Fubini–Tonelli⁷ yields Corollary 3.1 directly.

Corollary 3.1 (Joint-distribution form (without independence; corollary of Fubini–Tonelli)). *Let (X, Y) be jointly distributed random variables whose joint distribution has a density $f_{X,Y}$ with respect to a product of σ -finite reference measures $\mu_X \otimes \mu_Y$ on $\mathcal{R}(X) \times \mathcal{R}(Y)$, and let $h : \mathcal{R}(X) \times \mathcal{R}(Y) \rightarrow \mathbb{R}$ be measurable. If either*

- (a) $h(X, Y) \geq 0$ almost surely, or
- (b) $E[|h(X, Y)|] < \infty$,

then the expectation of $h(X, Y)$ can be written as an iterated integral against $f_{X,Y}$, with the order of integration exchangeable:

$$\begin{aligned} E[h(X, Y)] &= \int_{\mathcal{R}(X)} \left(\int_{\mathcal{R}(Y)} h(x, y) f_{X,Y}(x, y) d\mu_Y(y) \right) d\mu_X(x) \\ &= \int_{\mathcal{R}(Y)} \left(\int_{\mathcal{R}(X)} h(x, y) f_{X,Y}(x, y) d\mu_X(x) \right) d\mu_Y(y). \end{aligned}$$

The choice of reference measures covers three cases:

- **Both continuous:** $\mu_X = \mu_Y = \text{Lebesgue measure}$; $f_{X,Y}$ is the joint probability density function (PDF), and $\int g(x) d\mu_X(x) = \int g(x) dx$.
- **Both discrete:** $\mu_X = \mu_Y = \text{counting measure}$; $f_{X,Y}(x, y) = P(X = x, Y = y)$ is the joint probability mass function (PMF), and $\int g(x) d\mu_X(x) = \sum_{x \in \mathcal{R}(X)} g(x)$.
- **Mixed** (one continuous, one discrete): one reference measure is Lebesgue and the other is counting; $f_{X,Y}(x, y) = f_{X|Y}(x | y) P(Y = y)$ (or $P(X = x | Y = y) f_Y(y)$ if X is discrete and Y continuous), and the iterated integrals combine an ordinary integral with a sum. The conditional densities/PMFs here are defined the same way as in Definition 3.9, just conditioning on Y instead of X .

Proof

Proof. Apply Fubini–Tonelli^a with $\mu_1 = \mu_X$ and $\mu_2 = \mu_Y$ to the integrand $h(x, y) f_{X,Y}(x, y)$ on $\mathcal{R}(X) \times \mathcal{R}(Y)$. Lebesgue measure and counting measure on a countable set are each σ -finite, so $\mu_X \otimes \mu_Y$ is σ -finite in all three cases. The relevant condition is (a) when $h \geq 0$ and (b) when $E[|h(X, Y)|] < \infty$. Independence is not required. When X and Y are independent, $f_{X,Y}(x, y) = f_X(x) f_Y(y)$ (or $P(X = x, Y = y) = P(X = x) P(Y = y)$ in the discrete case), and the iterated integrals factor into separate integrals over the marginals. $\square$

^a[math-prereqs.qmd#thm-fubini-tonelli](#)

⁴[math-prereqs.qmd#thm-fubini](#)

⁵[math-prereqs.qmd](#)

⁶[math-prereqs.qmd#thm-fubini-tonelli](#)

⁷[math-prereqs.qmd#thm-fubini-tonelli](#)

Exm

Example 3.9 (Expectation of a product of independent variables). Let $X \sim \text{Uniform}(0, 1)$ and $Y \sim \text{Uniform}(0, 2)$, independently distributed. Compute $E[XY]$.

We apply Corollary 3.1 (both-continuous case) with $h(x, y) = xy$. Since X and Y are independent with densities $f_X(x) = 1$ on $[0, 1]$ and $f_Y(y) = \frac{1}{2}$ on $[0, 2]$, the joint density factors as $f_{X,Y}(x, y) = f_X(x) f_Y(y) = \frac{1}{2}$, and $\mu_X = \mu_Y = \text{Lebesgue measure}$:

$$\begin{aligned} E[XY] &= \int_0^1 \left(\int_0^2 xy \cdot \frac{1}{2} dy \right) dx \\ &= \int_0^1 x \left(\frac{1}{2} \int_0^2 y dy \right) dx \\ &= \int_0^1 x \cdot \frac{1}{2} \cdot \left[\frac{y^2}{2} \right]_0^2 dx \\ &= \int_0^1 x \cdot \frac{1}{2} \cdot 2 dx \\ &= \int_0^1 x dx \\ &= \frac{1}{2} \end{aligned}$$

As a check: $E[X] = \frac{1}{2}$, $E[Y] = 1$, and $E[X] E[Y] = \frac{1}{2}$.

Exm

Example 3.10 (When independence fails: a counterexample). Correctly applying Corollary 3.1 requires the *actual* joint density $f_{X,Y}$ — not the product of marginals $f_X(x) f_Y(y)$, which is valid only when X and Y are independent. Using the wrong joint density gives the wrong answer.

Let $X \sim \text{Uniform}(0, 1)$ and set $Y = X$ (so X and Y are perfectly correlated and **not** independent).

True expectation:

$$E[XY] = E[X \cdot X] = E[X^2] = \int_0^1 x^2 dx = \frac{1}{3}$$

Erroneously applying the product-measure formula:

Note that Fubini–Tonelli’s own conditions still hold here ($h(x, y) = xy$ is nonnegative and integrable), so the error is not a failure of Fubini–Tonelli^a. Rather, the error is using the *wrong measure*: the joint distribution of (X, X) is concentrated on the diagonal $\{(x, x) : x \in [0, 1]\} \subset [0, 1]^2$, which has Lebesgue measure zero in $\mathbb{R}^2$. The joint distribution is therefore **not** absolutely continuous with respect to two-dimensional Lebesgue measure, so **no joint density $f_{X,Y}$ on $[0, 1]^2$ exists**, which is the reference density Corollary 3.1 requires.

The following calculation is what someone would *erroneously* write if they assumed independence and used $f_X(x) f_Y(y)$ as a “joint density” — a function that does not in fact correspond to the joint distribution of (X, X) . The marginals $X \sim \text{Uniform}(0, 1)$ and $Y \sim \text{Uniform}(0, 1)$ do have densities $f_X = f_Y = 1$, but the *product* $f_X(x) f_Y(y) = 1$ on $[0, 1]^2$ is the density of an *independent* pair, not of (X, X) :

$$\begin{aligned}
\int_0^1 \int_0^1 xy \cdot f_X(x) \cdot f_Y(y) \, dy \, dx &= \int_0^1 \int_0^1 xy \, dy \, dx \\
&= \int_0^1 x \left(\int_0^1 y \, dy \right) dx \\
&= \int_0^1 x \cdot \frac{1}{2} dx \\
&= \frac{1}{4}
\end{aligned}$$

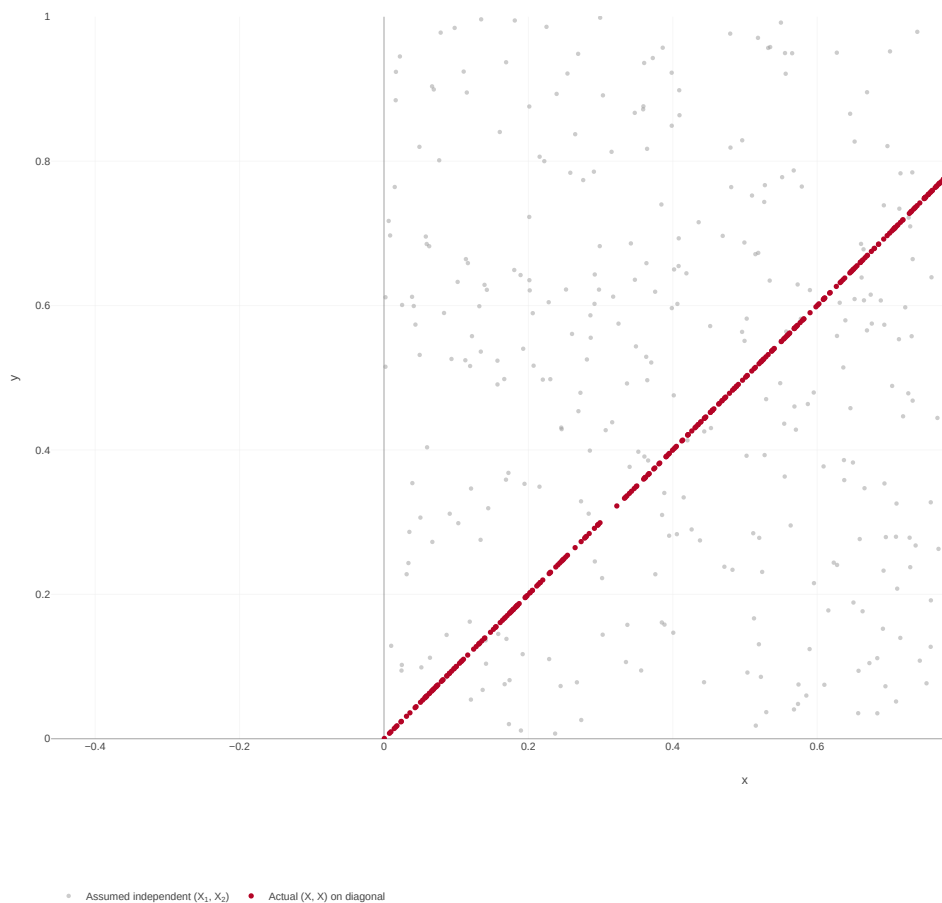
This calculation recovers $E[XY]$ for *independent* uniforms ($\frac{1}{4}$), not $E[XX]$ for the perfectly correlated pair ($\frac{1}{3}$). The lesson is that Corollary 3.1 requires the *actual* joint density $f_{X,Y}$. For independent (X, Y) , this factors as $f_X(x) f_Y(y)$; for dependent (X, Y) , $f_{X,Y}$ need not factor — and for (X, X) , no joint density on $\mathbb{R}^2$ exists at all, so Corollary 3.1 simply does not apply.

```

set.seed(204)
n <- 400
x_dep <- runif(n)
y_dep <- x_dep
x_ind <- runif(n)
y_ind <- runif(n)

plotly::plot_ly() |>
  plotly::add_trace(
    type = "scatter", mode = "markers",
    x = x_ind, y = y_ind,
    name = "Assumed independent (X<sub>1</sub>, X<sub>2</sub>)",
    marker = list(size = 5, color = "#999999", opacity = 0.5)
  ) |>
  plotly::add_trace(
    type = "scatter", mode = "markers",
    x = x_dep, y = y_dep,
    name = "Actual (X, X) on diagonal",
    marker = list(size = 6, color = "#b40426")
  ) |>
  plotly::layout(
    xaxis = list(title = "x", range = c(0, 1), scaleanchor = "y"),
    yaxis = list(title = "y", range = c(0, 1)),
    legend = list(orientation = "h", y = -0.2)
  )

```



Exm

Example 3.11 (Both-continuous case: joint PDF on a non-rectangular support). Let (X, Y) have joint density $f_{X,Y}(x, y) = 2$ for $0 \leq x \leq y \leq 1$ (and 0 otherwise). Compute $E[X + Y]$. By Corollary 3.1:

$$\begin{aligned} E[X + Y] &= \int_0^1 \int_0^y (x + y) \cdot 2 \, dx \, dy \\ &= 2 \int_0^1 \left[\frac{x^2}{2} + xy \right]_{x=0}^{x=y} dy \\ &= 2 \int_0^1 \left(\frac{y^2}{2} + y^2 \right) dy \\ &= 2 \int_0^1 \frac{3y^2}{2} dy \\ &= 3 \int_0^1 y^2 dy \\ &= 3 \cdot \frac{1}{3} \\ &= 1 \end{aligned}$$

```

n_grid <- 51
x_seq <- seq(0, 1, length.out = n_grid)
y_seq <- seq(0, 1, length.out = n_grid)

z_mat <- outer(x_seq, y_seq, function(x, y) {
  z <- rep(2, length(x))
  z[x > y] <- NA
  z
})

plotly::plot_ly(x = ~x_seq, y = ~y_seq, z = ~t(z_mat)) |>
  plotly::add_surface(showscale = FALSE) |>
  plotly::layout(scene = list(
    xaxis = list(title = "x"),
    yaxis = list(title = "y"),
    zaxis = list(title = "f(x, y)", range = c(0, 2.5)),
    camera = list(eye = list(x = 1.6, y = -1.6, z = 0.8))
  ))

```

Exercise 3.2 (Both-discrete case, infinite support: joint PMF). Let X and Y be independent, each Geometric on $\mathcal{R}(X) = \mathcal{R}(Y) = \{0, 1, 2, \dots\}$ with $P(X = x) = (1 - p)p^x$ for a fixed $p \in (0, 1)$ (X counts the number of failures before the first success in a sequence of independent trials with success probability $1 - p$; likewise for Y). Unlike Exercise 3.1, the support here is countably infinite. The joint PMF is $P(X = x, Y = y) = (1 - p)^2 p^{x+y}$. Compute $E[X + Y]$.

Solution

Solution. Compute $E[X + Y]$ using Corollary 3.1 with $\mu_X = \mu_Y =$ counting measure and $h(x, y) = x + y$. Since $h(x, y) = x + y \geq 0$ for every (x, y) in this support, condition (a) holds, so Corollary 3.1 (via Tonelli's theorem) guarantees the order of this now-infinite double sum is exchangeable — unlike the finite case, elementary algebra alone could not establish this. By Corollary 3.1 (both-discrete case), summing over y first for each fixed x :

$$\begin{aligned}
 E[X + Y] &= \sum_{x=0}^{\infty} \sum_{y=0}^{\infty} (x + y) P(X = x, Y = y) \\
 &= \sum_{x=0}^{\infty} \sum_{y=0}^{\infty} (x + y) (1 - p)^2 p^{x+y} \\
 &= \sum_{x=0}^{\infty} (1 - p)^2 p^x \sum_{y=0}^{\infty} (x + y) p^y \\
 &= \sum_{x=0}^{\infty} (1 - p)^2 p^x \left(x \sum_{y=0}^{\infty} p^y + \sum_{y=0}^{\infty} y p^y \right) \\
 &= \sum_{x=0}^{\infty} (1 - p)^2 p^x \left(\frac{x}{1 - p} + \frac{p}{(1 - p)^2} \right) \\
 &= \sum_{x=0}^{\infty} p^x [x(1 - p) + p] \\
 &= (1 - p) \sum_{x=0}^{\infty} x p^x + p \sum_{x=0}^{\infty} p^x \\
 &= (1 - p) \cdot \frac{p}{(1 - p)^2} + p \cdot \frac{1}{1 - p} \\
 &= \frac{p}{1 - p} + \frac{p}{1 - p} \\
 &= \frac{2p}{1 - p}
 \end{aligned}$$

using the standard geometric-series facts $\sum_{y=0}^{\infty} p^y = \frac{1}{1 - p}$ and $\sum_{y=0}^{\infty} y p^y = \frac{p}{(1 - p)^2}$ (e.g. Casella and Berger (2002)).

As a check, $E[X] = E[Y] = \frac{p}{1 - p}$ (the mean of this Geometric distribution; Casella and Berger (2002)), so $E[X + Y] = E[X] + E[Y] = \frac{2p}{1 - p}$ by linearity, matching.

```

p <- 0.4
exact_sum <- 2 * p / (1 - p)

n_trunc <- 500
support <- 0:n_trunc
joint_pmf_fn <- function(x, y) (1 - p)^2 * p^(x + y)
joint_probs <- outer(support, support, joint_pmf_fn)
h_vals <- outer(support, support, function(x, y) x + y)

sum_y_first <- sum(rowSums(h_vals * joint_probs))
sum_x_first <- sum(colSums(h_vals * joint_probs))

pander::pander(data.frame(
  quantity = c("exact (closed form)", "sum over y first, then x",
               "sum over x first, then y"),
  value = round(c(exact_sum, sum_y_first, sum_x_first), 6)
))

```

quantity	value
exact (closed form)	1.333
sum over y first, then x	1.333
sum over x first, then y	1.333

With $p = 0.4$, summing y first then x , and summing x first then y , agree with each other and with the closed form $\frac{2p}{1-p} = 1.3333$ up to truncation error — exactly Tonelli's guarantee.

i Tonelli's nonnegativity condition, and what happens without it

The calculation in Exercise 3.2 only needed condition (a), $h(X, Y) \geq 0$, because $h(x, y) = x + y$ is nonnegative on this support. For a **signed** h , interchanging an infinite double sum is not automatically valid — Corollary 3.1's condition (b), $E[|h(X, Y)|] < \infty$, is what licenses it in that case. Without either condition, the two orders can genuinely disagree. A standard example (a signed array, not a probability distribution; see e.g. Rudin (1976) for the general theory of rearranging series): let $a_{m,n} = 1$ if $m = n$, $a_{m,n} = -1$ if $m = n + 1$, and $a_{m,n} = 0$ otherwise, for $m, n = 0, 1, 2, \dots$. Summing each row m first: row 0 has only the term $a_{0,0} = 1$ (there is no valid $n = -1$), so its row sum is 1; every row $m \geq 1$ has $a_{m,m} = 1$ and $a_{m,m-1} = -1$, so its row sum is 0. Summing the rows then gives $1 + 0 + 0 + \dots = 1$. Summing each column n first: every column $n \geq 0$ has $a_{n,n} = 1$ and $a_{n+1,n} = -1$, so its column sum is always 0, and summing the columns then gives $0 + 0 + \dots = 0$. The two orders give 1 and 0: genuinely different answers, confirming that a condition like (a) or (b) really is needed once the terms are no longer all nonnegative.

Exercise 3.3 (Mixed case: one continuous variable, one discrete variable). Let $Y \sim \text{Bernoulli}(0.6)$ and, given $Y = y$, let $X | Y = y \sim \text{Uniform}(0, y + 1)$. Compute $E[X]$.

Solution

Solution. Compute $E[X]$ using Corollary 3.1 with $\mu_X = \text{Lebesgue measure}$, $\mu_Y = \text{counting measure}$, and $h(x, y) = x$.

The joint density w.r.t. Lebesgue $\times$ counting measure is $f_{X,Y}(x, y) = f_{X|Y}(x | y) P(Y = y)$:

$$\begin{aligned} f_{X,Y}(x, 0) &= 1 \cdot 0.4 = 0.4 && \text{for } x \in [0, 1]; \\ f_{X,Y}(x, 1) &= \tfrac{1}{2} \cdot 0.6 = 0.3 && \text{for } x \in [0, 2]. \end{aligned}$$

By Corollary 3.1 (mixed case):

$$\begin{aligned} \mathbb{E}[X] &= \sum_{y \in \{0,1\}} \int_0^{y+1} x f_{X,Y}(x, y) dx \\ &= \int_0^1 x \cdot 0.4 dx + \int_0^2 x \cdot 0.3 dx \\ &= 0.4 \cdot \frac{1}{2} + 0.3 \cdot 2 \\ &= 0.2 + 0.6 = 0.8 \end{aligned}$$

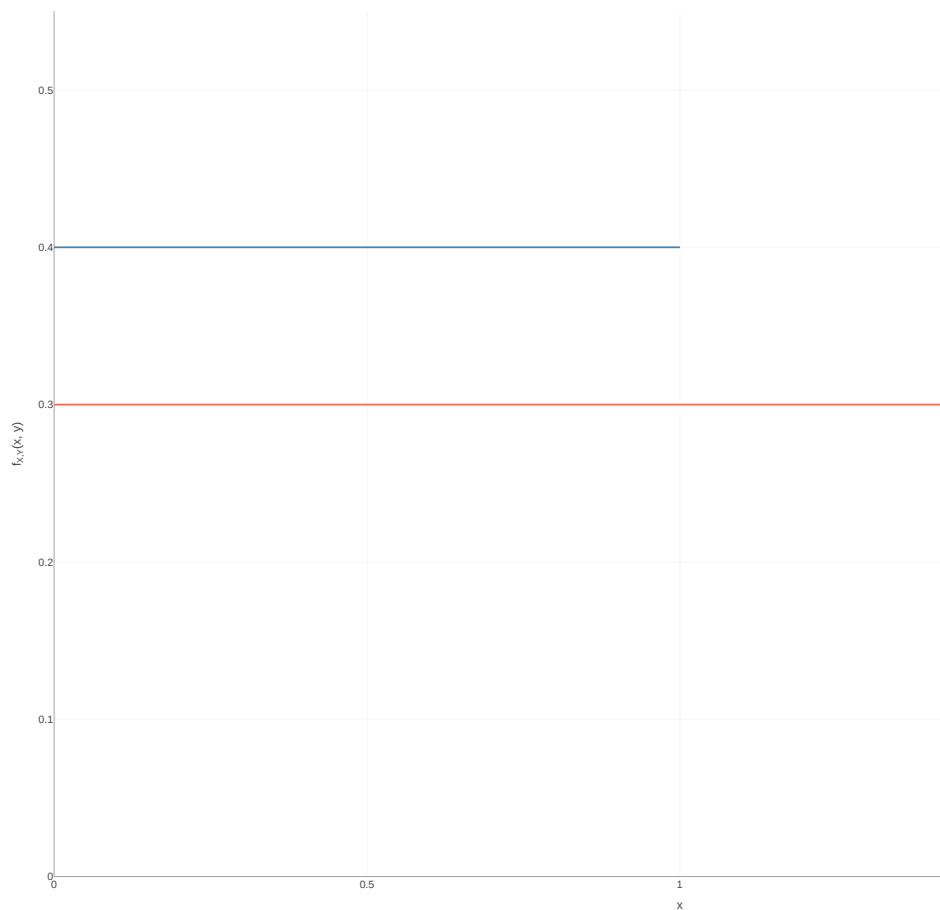
As a check using the law of total expectation: $\mathbb{E}[X \mid Y = 0] = \frac{1}{2}$ and $\mathbb{E}[X \mid Y = 1] = 1$, so $\mathbb{E}[X] = \frac{1}{2}(0.4) + 1(0.6) = 0.2 + 0.6 = 0.8$.

```

x_fine <- seq(0, 2, by = 0.005)
df <- data.frame(
  x = c(x_fine[x_fine <= 1], x_fine),
  density = c(rep(0.4, sum(x_fine <= 1)), rep(0.3, length(x_fine))),
  label = c(
    rep("Y = 0 (P = 0.4)", sum(x_fine <= 1)),
    rep("Y = 1 (P = 0.6)", length(x_fine))
  )
)

plotly::plot_ly(
  df, x = ~x, y = ~density, color = ~label,
  colors = c("steelblue", "tomato")
) |>
plotly::add_lines() |>
plotly::layout(
  xaxis = list(title = "x"),
  yaxis = list(title = "fX,Y(x, y)", range = c(0, 0.55)),
  legend = list(title = list(text = "Y value"))
)

```



Y takes only finitely many values here (two), so $\sum_{y \in \{0,1\}} \int_0^{y+1} x f_{X,Y}(x, y) dx$ is just linearity of the integral applied to a two-term sum — $\int g + \int k = \int(g + k)$ — not a genuine interchange of summation and integration order. Exercise 3.4 gives Y a countably infinite range instead, so the sum-of-integrals is genuinely infinite and Corollary 3.1's guarantee is load-bearing.

Exercise 3.4 (Mixed case, infinite discrete support). Let Y be Geometric on $\{0, 1, 2, \dots\}$ with $P(Y = y) = (1-q)q^y$ for a fixed $q \in (0, 1)$ and, given $Y = y$, let $X | Y = y \sim \text{Uniform}(0, y+1)$. Unlike Exercise 3.3, Y 's range here is countably infinite. Compute $E[X]$.

Solution

Solution. Compute $E[X]$ using Corollary 3.1 with $\mu_X =$ Lebesgue measure, $\mu_Y =$ counting measure, and $h(x, y) = x$.

The joint density w.r.t. Lebesgue $\times$ counting measure is $f_{X,Y}(x, y) = f_{X|Y}(x | y) P(Y = y) = \frac{(1-q)q^y}{y+1}$ for $x \in [0, y+1]$.

Since $h(x, y) = x \geq 0$ on this support, condition (a) holds, so Corollary 3.1 (via Tonelli's theorem) guarantees the now-infinite sum-of-integrals expression is valid — unlike in Exercise 3.3, this Fubini–Tonelli justification is required because the sum is infinite rather than finite.

By Corollary 3.1 (mixed case):

$$\begin{aligned}
 E[X] &= \sum_{y=0}^{\infty} \int_0^{y+1} x f_{X,Y}(x, y) dx \\
 &= \sum_{y=0}^{\infty} \frac{(1-q)q^y}{y+1} \int_0^{y+1} x dx \\
 &= \sum_{y=0}^{\infty} \frac{(1-q)q^y}{y+1} \cdot \frac{(y+1)^2}{2} \\
 &= \frac{1-q}{2} \sum_{y=0}^{\infty} (y+1) q^y \\
 &= \frac{1-q}{2} \left(\sum_{y=0}^{\infty} y q^y + \sum_{y=0}^{\infty} q^y \right) \\
 &= \frac{1-q}{2} \left(\frac{q}{(1-q)^2} + \frac{1}{1-q} \right) \\
 &= \frac{1-q}{2} \cdot \frac{q + (1-q)}{(1-q)^2} \\
 &= \frac{1-q}{2} \cdot \frac{1}{(1-q)^2} \\
 &= \frac{1}{2(1-q)}
 \end{aligned}$$

using the same geometric-series facts as Exercise 3.2 (e.g. Casella and Berger (2002)).

As a check using the law of total expectation: $E[X | Y = y] = \frac{y+1}{2}$ (the mean of $\text{Uniform}(0, y+1)$) and $E[Y] = \frac{q}{1-q}$ (the mean of this Geometric distribution; Casella and Berger (2002)), so:

$$\begin{aligned}
E[X] &= E[E[X \mid Y]] \\
&= E\left[\frac{Y+1}{2}\right] \\
&= \frac{E[Y] + 1}{2} \\
&= \frac{1}{2} \left(\frac{q}{1-q} + 1 \right) \\
&= \frac{1}{2} \cdot \frac{q + (1-q)}{1-q} \\
&= \frac{1}{2(1-q)}
\end{aligned}$$

matching.

```

q <- 0.5
exact_ex <- 1 / (2 * (1 - q))

n_trunc <- 2000
y_vals <- 0:n_trunc
py_vals <- (1 - q) * q^y_vals
cond_mean <- (y_vals + 1) / 2
trunc_ex <- sum(cond_mean * py_vals)

pander::pander(data.frame(
  quantity = c("exact (closed form)", "truncated sum-of-integrals"),
  value = round(c(exact_ex, trunc_ex), 6)
))

```

quantity	value
exact (closed form)	1
truncated sum-of-integrals	1

With $q = 0.5$, the truncated sum-of-integrals matches the closed form $\frac{1}{2(1-q)} = 1$.

Theorem 3.6 (Law of iterated expectations). *For any two random variables X and Y :*

$$E[Y] = E[E[Y \mid X]]$$

*Alternate names for this identity include: the **tower rule**, the **tower property**, the **law of total expectation**, and the **smoothing theorem**.*

i Proof

Proof. Discrete case. When X and Y are discrete, applying Definition 3.5 to $E[E[Y \mid X]]$ and then the law of total probability (Theorem 1.4) applied to the countable partition $\{X = x : x \in \mathcal{R}(X)\}$:

$$\begin{aligned}
E[E[Y | X]] &= \sum_{x \in \mathcal{R}(X)} E[Y | X = x] \cdot P(X = x) && \text{(definition of expectation of } E[Y | X]) \\
&= \sum_{x \in \mathcal{R}(X)} \left(\sum_{y \in \mathcal{R}(Y)} y \cdot P(Y = y | X = x) \right) \cdot P(X = x) && \text{(definition of conditional expectation } E[Y | X = x]) \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot \sum_{x \in \mathcal{R}(X)} P(Y = y | X = x) \cdot P(X = x) && \text{(exchange order of summation)} \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot P(Y = y) && \text{(law of total probability over } X) \\
&= E[Y] && \text{(definition of expectation of } Y)
\end{aligned}$$

Continuous case. When X and Y are continuous, applying Definition 3.5 to $E[E[Y | X]]$ and then using Definition 3.8 for $E[Y | X = x]$:

$$\begin{aligned}
E[E[Y | X]] &= \int_{x \in \mathcal{R}(X)} E[Y | X = x] \cdot p(X = x) dx \\
&= \int_{x \in \mathcal{R}(X)} \left(\int_{y \in \mathcal{R}(Y)} y \cdot p(Y = y | X = x) dy \right) \cdot p(X = x) dx \\
&= \int_{y \in \mathcal{R}(Y)} y \cdot \left(\int_{x \in \mathcal{R}(X)} p(Y = y | X = x) \cdot p(X = x) dx \right) dy \\
&= \int_{y \in \mathcal{R}(Y)} y \cdot p(Y = y) dy \\
&= E[Y]
\end{aligned}$$

where the third equality exchanges the order of integration by condition (b) of Fubini–Tonelli^a (the absolute-integrability case, **Fubini’s theorem**); this theorem requires $E[|Y|] < \infty$, which is implicit in $E[Y]$ being defined, and the fourth equality uses $\int_x p(Y = y | X = x) \cdot p(X = x) dx = \int_x p(X = x, Y = y) dx = p(Y = y)$ (marginalization of the joint density). $\square$

^a[math-prereqs.qmd#thm-fubini-tonelli](#)

Theorem 3.7 (Conditional law of iterated expectations). *For random variables X , Y , and Z :*

$$E[Y | Z] = E[E[Y | X, Z] | Z]$$

This identity is the tower rule applied conditionally on Z .

i Proof

Proof. For each fixed value z with positive probability or density:

Discrete case. Conditioning on $Z = z$, and applying the law of total probability to the partition $\{X = x : x \in \mathcal{R}(X)\}$ under the conditional distribution given $Z = z$:

$$\begin{aligned}
E[E[Y | X, Z] | Z = z] &= \sum_{x \in \mathcal{R}(X)} E[Y | X = x, Z = z] \cdot P(X = x | Z = z) && \text{(definition of expectation)} \\
&= \sum_{x \in \mathcal{R}(X)} \left(\sum_{y \in \mathcal{R}(Y)} y \cdot P(Y = y | X = x, Z = z) \right) \cdot P(X = x | Z = z) && \text{(definition of conditional expectation)} \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot \sum_{x \in \mathcal{R}(X)} P(Y = y | X = x, Z = z) \cdot P(X = x | Z = z) && \text{(exchange order of summation)} \\
&= \sum_{y \in \mathcal{R}(Y)} y \cdot P(Y = y | Z = z) && \text{(law of total probability given } Z = z) \\
&= E[Y | Z = z] && \text{(definition of conditional expectation)}
\end{aligned}$$

Continuous case. Conditioning on $Z = z$, and integrating over X under the conditional density $p(X = x | Z = z)$:

$$\begin{aligned}
E[E[Y | X, Z] | Z = z] &= \int_{x \in \mathcal{R}(X)} E[Y | X = x, Z = z] \cdot p(X = x | Z = z) dx \\
&= E[Y | Z = z]
\end{aligned}$$

Therefore, as random variables of Z , $E[Y | Z] = E[E[Y | X, Z] | Z]$. □

Exm

Example 3.12 (Marginal expectation from conditional expectations). Suppose X is a binary random variable indicating treatment assignment ($X = 1$ treated, $X = 0$ control), with $P(X = 1) = 0.5$, and suppose the outcome Y has conditional expectations:

$$E[Y | X = 1] = 10, \quad E[Y | X = 0] = 6$$

By the law of iterated expectations (Theorem 3.6):

$$\begin{aligned}
E[Y] &= E[E[Y | X]] \\
&= E[Y | X = 1] \cdot P(X = 1) + E[Y | X = 0] \cdot P(X = 0) \\
&= 10 \cdot 0.5 + 6 \cdot 0.5 \\
&= 5 + 3 \\
&= 8
\end{aligned}$$

Definition 3.11 (Expectation of a random matrix). For a random matrix $\mathbf{A}$ of size $m \times n$ with (i, j) -th element A_{ij} , the **expectation** $E \mathbf{A}$ is the $m \times n$ matrix whose (i, j) -th element is $E[A_{ij}]$:

$$E \mathbf{A} \stackrel{\text{def}}{=} \begin{pmatrix} E[A_{11}] & E[A_{12}] & \cdots & E[A_{1n}] \\ E[A_{21}] & E[A_{22}] & \cdots & E[A_{2n}] \\ \vdots & \vdots & \ddots & \vdots \\ E[A_{m1}] & E[A_{m2}] & \cdots & E[A_{mn}] \end{pmatrix}$$

In other words, expectation is applied **element-wise** to a random matrix.

3.5 Deviation, error, and noise

Definition 3.12 (Deviation). A **deviation** is the difference between a value and a reference value. For any quantity z and reference value r :

$$z - r$$

In probability and statistics, “deviation” often means deviation from a **population mean**. For a random variable Y :

$$Y - E[Y]$$

See: Wikipedia: Deviation (statistics)^a

^a[https://en.wikipedia.org/wiki/Deviation_\(statistics\)](https://en.wikipedia.org/wiki/Deviation_(statistics))

Definition 3.13 (Deviation from a population or subpopulation mean). In probabilistic models, we call this quantity a **deviation from a mean**. It is often also called an **error** or **noise term** in other sources. For the random variable Y , define the **deviation** from its mean as:

$$e(Y) \stackrel{\text{def}}{=} Y - E[Y]$$

For a realized observation y :

$$e(y) \stackrel{\text{def}}{=} y - E[Y]$$

In regression settings, the reference mean is often conditional on covariates: $e(y_i) \stackrel{\text{def}}{=} y_i - E[Y_i | X_i]$.

In this course, we prefer “deviation” for this mean-deviation quantity. The terms “error” and “noise” are common aliases. We use “residual” (defined in the Linear regression chapter^a) for deviations from fitted values. For notation in this course, we use $e(\cdot)$ for these model/data deviations, and reserve $\varepsilon(\cdot)$ for estimator-to-estimand deviations (see Estimation^b).

See:

- Wikipedia: Errors and residuals^c
- Wikipedia: Deviation (statistics)^d
- Wikipedia: Linear regression — Notation and terminology^e

^a[Linear-models-overview.qmd#def-resid-fitted](#)

^b[estimation.qmd#def-estimation-error](#)

^chttps://en.wikipedia.org/wiki/Errors_and_residuals

^d[https://en.wikipedia.org/wiki/Deviation_\(statistics\)](https://en.wikipedia.org/wiki/Deviation_(statistics))

^ehttps://en.wikipedia.org/wiki/Linear_regression#Notation_and_terminology

3.6 Variance and related characteristics

Definition 3.14 (Variance). The variance of a random variable X is the **expectation** of the squared **deviation from the mean**; that is:

$$\text{Var}(X) \stackrel{\text{def}}{=} E[[e(X)]^2]$$

Theorem 3.8 (Variance as expected squared deviation from the mean).

$$\text{Var}(X) = E[(X - E[X])^2]$$

Proof

Proof. Substituting the definition of $e(X)$ from Definition 3.13 into Definition 3.14:

$$\text{Var}(X) \stackrel{\text{def}}{=} E[[e(X)]^2] = E[(X - E[X])^2].$$

□

Theorem 3.9 (Simplified expression for variance).

$$\text{Var}(X) = \mathbb{E}[X^2] - (\mathbb{E}[X])^2$$

Proof

Proof. By linearity of expectation, we have:

$$\begin{aligned} \text{Var}(X) &\stackrel{\text{def}}{=} \mathbb{E}[(e(X))^2] && \text{(definition of variance)} \\ &= \mathbb{E}[(X - \mathbb{E}[X])^2] && \text{(definition of deviation from mean)} \\ &= \mathbb{E}[X^2 - 2X\mathbb{E}[X] + (\mathbb{E}[X])^2] && \text{(expand binomial square)} \\ &= \mathbb{E}[X^2] - \mathbb{E}[2X\mathbb{E}[X]] + \mathbb{E}[(\mathbb{E}[X])^2] && \text{(linearity of expectation)} \\ &= \mathbb{E}[X^2] - 2\mathbb{E}[X]\mathbb{E}[X] + (\mathbb{E}[X])^2 && \text{(constants factor out of expectation)} \\ &= \mathbb{E}[X^2] - (\mathbb{E}[X])^2 && \text{(algebraic simplification)} \end{aligned}$$

□

Theorem 3.10 (Law of total variance). *For random variables X and Y :*

$$\text{Var}(Y) = \mathbb{E}[\text{Var}(Y | X)] + \text{Var}(\mathbb{E}[Y | X])$$

where $\text{Var}(Y | X) \stackrel{\text{def}}{=} \mathbb{E}[(Y - \mathbb{E}[Y | X])^2 | X]$.

*Alternate names include: the **conditional variance formula**, **Eve's law**, and the **variance decomposition formula**.*

i Proof

Proof. Write $Y - \mathbb{E}[Y] = (Y - \mathbb{E}[Y | X]) + (\mathbb{E}[Y | X] - \mathbb{E}[Y])$. Then:

$$(Y - \mathbb{E}[Y])^2 = (Y - \mathbb{E}[Y | X])^2 + (\mathbb{E}[Y | X] - \mathbb{E}[Y])^2 + 2(Y - \mathbb{E}[Y | X])(\mathbb{E}[Y | X] - \mathbb{E}[Y])$$

Taking expectation:

$$\begin{aligned} \text{Var}(Y) &= \mathbb{E}[(Y - \mathbb{E}[Y | X])^2] + \mathbb{E}[(\mathbb{E}[Y | X] - \mathbb{E}[Y])^2] \\ &\quad + 2\mathbb{E}[(Y - \mathbb{E}[Y | X])(\mathbb{E}[Y | X] - \mathbb{E}[Y])] \end{aligned}$$

For the cross-term:

Discrete case.

$$\begin{aligned} \mathbb{E}[(Y - \mathbb{E}[Y | X])(\mathbb{E}[Y | X] - \mathbb{E}[Y])] &= \sum_{x \in \mathcal{R}(X)} \mathbb{E}[(Y - \mathbb{E}[Y | X])(\mathbb{E}[Y | X] - \mathbb{E}[Y]) | X = x] \cdot \mathbb{P}(X = x) \\ &= \sum_{x \in \mathcal{R}(X)} (\mathbb{E}[Y | X = x] - \mathbb{E}[Y]) \cdot \mathbb{E}[Y - \mathbb{E}[Y | X = x] | X = x] \cdot \mathbb{P}(X = x) \\ &= 0 \end{aligned}$$

Continuous case.

$$\begin{aligned}
E[(Y - E[Y | X])(E[Y | X] - E[Y])] &= \int_{x \in \mathcal{R}(X)} E[(Y - E[Y | X])(E[Y | X] - E[Y]) | X = x] \cdot p(X = x) dx \\
&= \int_{x \in \mathcal{R}(X)} (E[Y | X = x] - E[Y]) \cdot E[Y - E[Y | X = x] | X = x] \cdot p(X = x) dx \\
&= 0
\end{aligned}$$

Therefore:

$$\begin{aligned}
\text{Var}(Y) &= E[(Y - E[Y | X])^2] + E[(E[Y | X] - E[Y])^2] \\
&= E[\text{Var}(Y | X)] + \text{Var}(E[Y | X])
\end{aligned}$$

□

Definition 3.15 (Precision). The **precision** of a random variable X , often denoted $\tau(X)$, τ_X , or shorthand as τ , is the inverse of that random variable's **variance**; that is:

$$\tau(X) \stackrel{\text{def}}{=} (\text{Var}(X))^{-1}$$

Definition 3.16 (Standard deviation). The standard deviation of a random variable X is the square-root of the **variance** of X :

$$\text{SD}(X) \stackrel{\text{def}}{=} \sqrt{\text{Var}(X)}$$

Definition 3.17 (Covariance). For any two one-dimensional random variables, X, Y :

$$\text{Cov}(X, Y) \stackrel{\text{def}}{=} E[(X - E[X])(Y - E[Y])]$$

Theorem 3.11 (Alternative formula for covariance).

$$\text{Cov}(X, Y) = E[XY] - E[X] E[Y]$$

Theorem 3.12 (Law of total covariance). For random variables X, Y , and Z :

$$\text{Cov}(Y, Z) = E[\text{Cov}(Y, Z | X)] + \text{Cov}(E[Y | X], E[Z | X])$$

where $\text{Cov}(Y, Z | X) \stackrel{\text{def}}{=} E[(Y - E[Y | X])(Z - E[Z | X]) | X]$.

Alternate names include: the **covariance decomposition formula** and the **conditional covariance formula**.

i Proof

Proof. Write:

$$\begin{aligned}
Y - E[Y] &= (Y - E[Y | X]) + (E[Y | X] - E[Y]) \\
Z - E[Z] &= (Z - E[Z | X]) + (E[Z | X] - E[Z])
\end{aligned}$$

Then:

$$\begin{aligned}
\text{Cov}(Y, Z) &= E[(Y - E[Y])(Z - E[Z])] \\
&= E[(Y - E[Y | X])(Z - E[Z | X])] \\
&\quad + E[(Y - E[Y | X])(E[Z | X] - E[Z])] \\
&\quad + E[(E[Y | X] - E[Y])(Z - E[Z | X])] \\
&\quad + E[(E[Y | X] - E[Y])(E[Z | X] - E[Z])]
\end{aligned}$$

For the two mixed terms:

Discrete case.

$$\begin{aligned}
E[(Y - E[Y | X])(E[Z | X] - E[Z])] &= \sum_{x \in \mathcal{R}(X)} E[(Y - E[Y | X])(E[Z | X] - E[Z]) | X = x] \cdot P(X = x) \\
&= \sum_{x \in \mathcal{R}(X)} (E[Z | X = x] - E[Z]) \cdot E[Y - E[Y | X = x] | X = x] \cdot P(X = x) \\
&= 0
\end{aligned}$$

and similarly:

$$E[(E[Y | X] - E[Y])(Z - E[Z | X])] = 0.$$

Continuous case.

$$\begin{aligned}
E[(Y - E[Y | X])(E[Z | X] - E[Z])] &= \int_{x \in \mathcal{R}(X)} E[(Y - E[Y | X])(E[Z | X] - E[Z]) | X = x] \cdot p(X = x) dx \\
&= \int_{x \in \mathcal{R}(X)} (E[Z | X = x] - E[Z]) \cdot E[Y - E[Y | X = x] | X = x] \cdot p(X = x) dx \\
&= 0
\end{aligned}$$

and similarly:

$$E[(E[Y | X] - E[Y])(Z - E[Z | X])] = 0.$$

Hence:

$$\begin{aligned}
\text{Cov}(Y, Z) &= E[(Y - E[Y | X])(Z - E[Z | X])] + E[(E[Y | X] - E[Y])(E[Z | X] - E[Z])] \\
&= E[\text{Cov}(Y, Z | X)] + \text{Cov}(E[Y | X], E[Z | X])
\end{aligned}$$

□

Lemma 3.1 (The covariance of a variable with itself is its variance). *For any random variable X :*

$$\text{Cov}(X, X) = \text{Var}(X)$$

i Proof

Proof.

$$\begin{aligned}
\text{Cov}(X, X) &= E[XX] - E[X]E[X] \\
&= E[X^2] - (E[X])^2 \\
&= \text{Var}(X)
\end{aligned}$$

□

Definition 3.18 (Variance/covariance of a $p \times 1$ random vector). For a $p \times 1$ dimensional random vector $\tilde{X}$,

$$\begin{aligned}\text{Var}(\tilde{X}) &\stackrel{\text{def}}{=} \text{Cov}(\tilde{X}) \\ &\stackrel{\text{def}}{=} \mathbb{E}[(\tilde{X} - \mathbb{E} \tilde{X})(\tilde{X} - \mathbb{E} \tilde{X})^\top]\end{aligned}$$

Theorem 3.13 (Elements of the variance-covariance matrix are pairwise covariances). For a $p \times 1$ random vector $\tilde{X} = (X_1, \dots, X_p)^\top$, the (i, j) -th element of $\text{Var}(\tilde{X})$ is $\text{Cov}(X_i, X_j)$:

$$\text{Var}(\tilde{X}) = \begin{pmatrix} \text{Var}(X_1) & \text{Cov}(X_1, X_2) & \cdots & \text{Cov}(X_1, X_p) \\ \text{Cov}(X_2, X_1) & \text{Var}(X_2) & \cdots & \text{Cov}(X_2, X_p) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(X_p, X_1) & \text{Cov}(X_p, X_2) & \cdots & \text{Var}(X_p) \end{pmatrix}$$

Proof

Proof. Let $\mu_i = \mathbb{E}[X_i]$ for $i = 1, \dots, p$, so $\mathbb{E} \tilde{X} = (\mu_1, \dots, \mu_p)^\top$. By Definition 3.18:

$$\begin{aligned}\text{Var}(\tilde{X}) &= \mathbb{E}[(\tilde{X} - \mathbb{E} \tilde{X})(\tilde{X} - \mathbb{E} \tilde{X})^\top] \\ &= \mathbb{E}\left[\begin{pmatrix} X_1 - \mu_1 \\ \vdots \\ X_p - \mu_p \end{pmatrix} (X_1 - \mu_1 \quad \cdots \quad X_p - \mu_p)\right] \\ &= \mathbb{E}\left[\begin{pmatrix} (X_1 - \mu_1)(X_1 - \mu_1) & \cdots & (X_1 - \mu_1)(X_p - \mu_p) \\ \vdots & \ddots & \vdots \\ (X_p - \mu_p)(X_1 - \mu_1) & \cdots & (X_p - \mu_p)(X_p - \mu_p) \end{pmatrix}\right] \\ &= \begin{pmatrix} \mathbb{E}[(X_1 - \mu_1)(X_1 - \mu_1)] & \cdots & \mathbb{E}[(X_1 - \mu_1)(X_p - \mu_p)] \\ \vdots & \ddots & \vdots \\ \mathbb{E}[(X_p - \mu_p)(X_1 - \mu_1)] & \cdots & \mathbb{E}[(X_p - \mu_p)(X_p - \mu_p)] \end{pmatrix} \\ &= \begin{pmatrix} \text{Cov}(X_1, X_1) & \cdots & \text{Cov}(X_1, X_p) \\ \vdots & \ddots & \vdots \\ \text{Cov}(X_p, X_1) & \cdots & \text{Cov}(X_p, X_p) \end{pmatrix} \\ &= \begin{pmatrix} \text{Var}(X_1) & \cdots & \text{Cov}(X_1, X_p) \\ \vdots & \ddots & \vdots \\ \text{Cov}(X_p, X_1) & \cdots & \text{Var}(X_p) \end{pmatrix}\end{aligned}$$

where:

- the step from the third to fourth line uses Definition 3.11,
- the step from the fourth to fifth line uses Definition 3.17, and
- the last step uses Lemma 3.1.

□

Theorem 3.14 (Alternate expression for variance of a random vector).

$$\text{Var}(\tilde{X}) = \mathbb{E}[\tilde{X} \tilde{X}^\top] - (\mathbb{E} \tilde{X})(\mathbb{E} \tilde{X})^\top$$

i Proof

Proof.

$$\begin{aligned}\text{Var}(\tilde{X}) &= \mathbb{E}[(\tilde{X} - \mathbb{E}\tilde{X})(\tilde{X} - \mathbb{E}\tilde{X})^\top] \\ &= \mathbb{E}[\tilde{X}\tilde{X}^\top - \tilde{X}(\mathbb{E}\tilde{X})^\top - (\mathbb{E}\tilde{X})\tilde{X}^\top + (\mathbb{E}\tilde{X})(\mathbb{E}\tilde{X})^\top] \\ &= \mathbb{E}[\tilde{X}\tilde{X}^\top] - (\mathbb{E}\tilde{X})(\mathbb{E}\tilde{X})^\top - (\mathbb{E}\tilde{X})(\mathbb{E}\tilde{X})^\top + (\mathbb{E}\tilde{X})(\mathbb{E}\tilde{X})^\top \\ &= \mathbb{E}[\tilde{X}\tilde{X}^\top] - (\mathbb{E}\tilde{X})(\mathbb{E}\tilde{X})^\top\end{aligned}$$

□

Theorem 3.15 (Variance of a linear combination). *For any vector of random variables $\tilde{X} = (X_1, \dots, X_n)$ and corresponding vector of constants $\tilde{a} = (a_1, \dots, a_n)$, the variance of their linear combination is:*

$$\begin{aligned}\text{Var}(\tilde{a} \cdot \tilde{X}) &= \text{Var}\left(\sum_{i=1}^n a_i X_i\right) \\ &= \tilde{a}^\top \text{Var}(\tilde{X}) \tilde{a} \\ &= \sum_{i=1}^n \sum_{j=1}^n a_i a_j \text{Cov}(X_i, X_j)\end{aligned}$$

i Proof

Proof. Left to the reader...

□

Corollary 3.2 (Variance of a sum of two random variables). *For any two random variables X and Y and scalars a and b :*

$$\text{Var}(aX + bY) = a^2 \text{Var}(X) + b^2 \text{Var}(Y) + 2(a \cdot b) \text{Cov}(X, Y)$$

i Proof

Proof. Apply Theorem 3.15 with $n = 2$, $X_1 = X$, and $X_2 = Y$.
Or, see <https://statproofbook.github.io/P/var-lincomb.html>

□

Definition 3.19 (Homoskedastic and heteroskedastic). A random variable Y is **homoskedastic** (with respect to covariates X) if the **variance** of Y does not vary with X :

$$\text{Var}(Y|X = x) = \sigma^2, \forall x$$

Otherwise it is **heteroskedastic**.

Definition 3.20 (Statistical independence). A set of random variables $X_1, \dots, X_n$ are **statistically independent** if their joint **probability** is equal to the product of their marginal **probabilities**:

$$\Pr(X_1 = x_1, \dots, X_n = x_n) = \prod_{i=1}^n \Pr(X_i = x_i)$$

 **Tip**

The symbol for independence, $\perp$, is essentially just $\prod$ upside-down. So the symbol can remind you of its definition (Definition 3.20).

Definition 3.21 (Conditional independence). A set of random variables $Y_1, \dots, Y_n$ are **conditionally statistically independent** given a set of covariates $X_1, \dots, X_n$ if the joint **probability** of the Y_i s given the X_i s is equal to the product of their marginal **probabilities**:

$$\Pr(Y_1 = y_1, \dots, Y_n = y_n | X_1 = x_1, \dots, X_n = x_n) = \prod_{i=1}^n \Pr(Y_i = y_i | X_i = x_i)$$

Definition 3.22 (Identically distributed). A set of random variables $X_1, \dots, X_n$ are **identically distributed** if they have the same range $\mathcal{R}(X)$ and if their marginal distributions $P(X_1 = x_1), \dots, P(X_n = x_n)$ are all equal to some shared distribution $P(X = x)$:

$$\forall i \in \{1 : n\}, \forall x \in \mathcal{R}(X) : P(X_i = x) = P(X = x)$$

Definition 3.23 (Conditionally identically distributed). A set of random variables $Y_1, \dots, Y_n$ are **conditionally identically distributed** given a set of covariates $X_1, \dots, X_n$ if $Y_1, \dots, Y_n$ have the same range $\mathcal{R}(X)$ and if the distributions $P(Y_i = y_i | X_i = x_i)$ are all equal to the same distribution $P(Y = y | X = x)$:

$$P(Y_i = y | X_i = x) = P(Y = y | X = x)$$

Definition 3.24 (Independent and identically distributed). A set of random variables $X_1, \dots, X_n$ are **independent and identically distributed** (shorthand: “ X_i iid”) if they are **statistically independent** and **identically distributed**.

Definition 3.25 (Conditionally independent and identically distributed). A set of random variables $Y_1, \dots, Y_n$ are **conditionally independent and identically distributed** (shorthand: “ $Y_i | X_i$ ciid” or just “ $Y_i | X_i$ iid”) given a set of covariates $X_1, \dots, X_n$ if $Y_1, \dots, Y_n$ are **conditionally independent** given $X_1, \dots, X_n$ and $Y_1, \dots, Y_n$ are **conditionally identically distributed** given $X_1, \dots, X_n$.

3.7 The Central Limit Theorem

The sum of many independent or nearly-independent random variables with small variances (relative to the number of RVs being summed) produces bell-shaped distributions.

For example, consider the sum of five dice (Figure 8).

```
library(dplyr)
dist =
  expand.grid(1:6, 1:6, 1:6, 1:6, 1:6) |>
  rowwise() |>
```

```

mutate(total = sum(c_across(everything())) |>
ungroup() |>
count(total) |>
mutate(`p(X=x)` = n/sum(n))

library(ggplot2)

dist |>
  ggplot() +
  aes(x = total, y = `p(X=x)`) +
  geom_col() +
  xlab("sum of dice (x)") +
  ylab("Probability of outcome, Pr(X=x)") +
  expand_limits(y = 0)

```

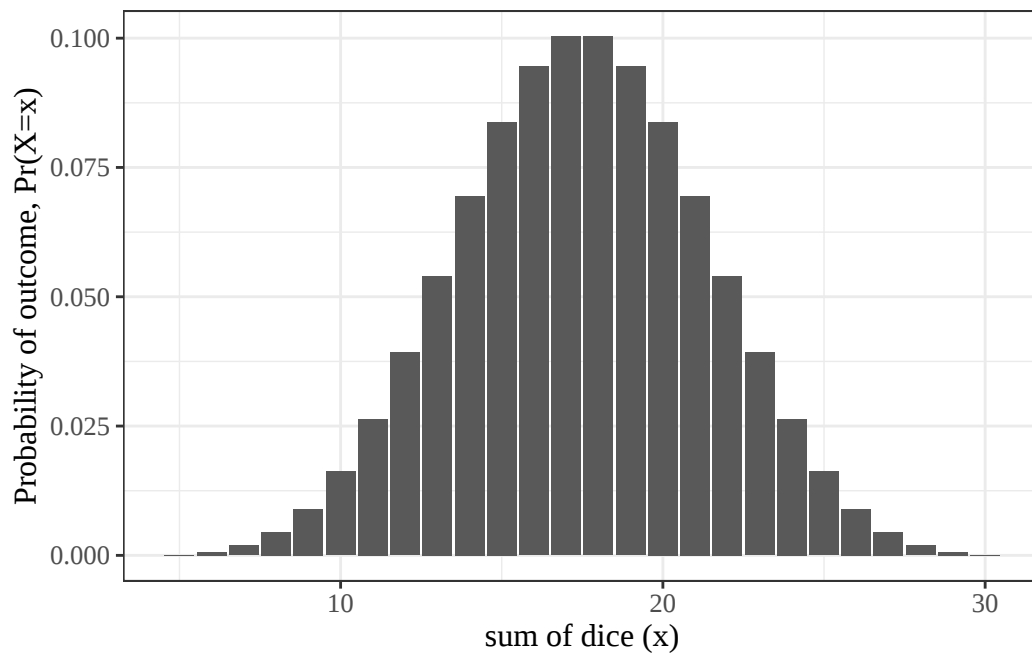


Figure 8: Distribution of the sum of five dice

In comparison, the outcome of just one die is not bell-shaped (Figure 9).

```

library(dplyr)
dist =
  expand_grid(1:6) |>
  rowwise() |>
  mutate(total = sum(c_across(everything())) |>
ungroup() |>
count(total) |>
mutate(`p(X=x)` = n/sum(n))

library(ggplot2)

dist |>
  ggplot() +

```

```

aes(x = total, y = `p(X=x)`) +
geom_col() +
xlab("sum of dice (x)") +
ylab("Probability of outcome, Pr(X=x)") +
expand_limits(y = 0)

```

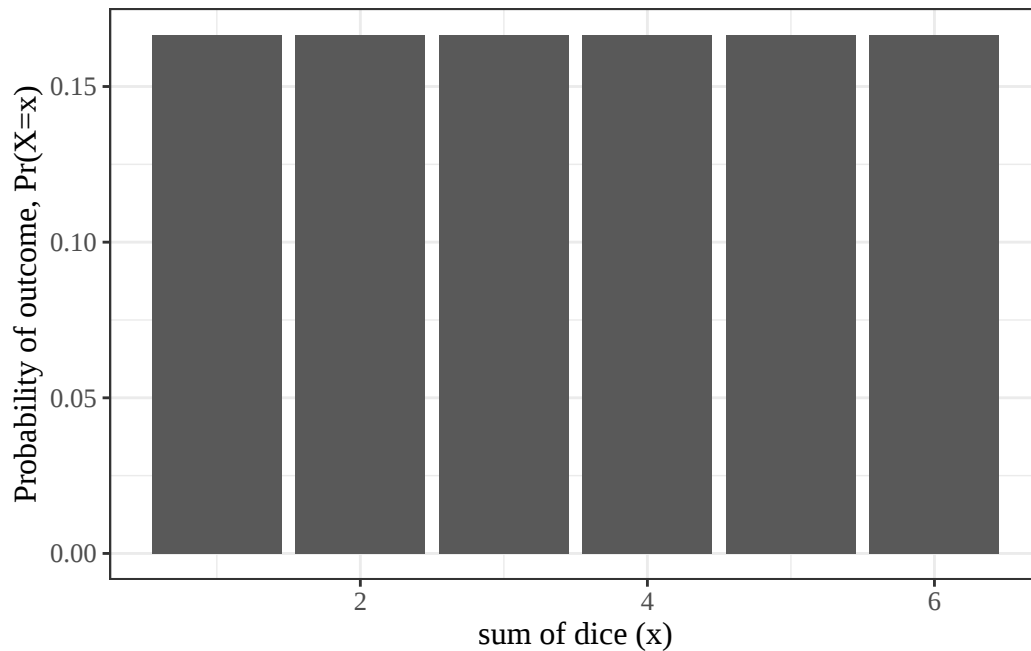


Figure 9: Distribution of the outcome of one die

What distribution does a single die have?

Answer: discrete uniform on 1:6.

4 Additional resources

- Miller (2017)

References

- Billingsley, Patrick. 1995. *Probability and Measure*. 3rd ed. Wiley Series in Probability and Mathematical Statistics. Wiley.
- Casella, George, and Roger Berger. 2002. *Statistical Inference*. 2nd ed. Cengage Learning. <https://www.cengage.com/c/statistical-inference-2e-casella-berger/9780534243128/>.
- Dobson, Annette J, and Adrian G Barnett. 2018. *An Introduction to Generalized Linear Models*. 4th ed. CRC press. <https://doi.org/10.1201/9781315182780>.
- Gut, Allan. 2013. *Probability: A Graduate Course*. 2nd ed. Springer Texts in Statistics. Springer. <https://doi.org/10.1007/978-1-4614-4708-5>.
- Kalbfleisch, John D, and Ross L Prentice. 2011. *The Statistical Analysis of Failure Time Data*.

John Wiley & Sons.

- Klein, John P, and Melvin L Moeschberger. 2003. *Survival Analysis: Techniques for Censored and Truncated Data*. Vol. 1230. Springer. <https://link.springer.com/book/10.1007/b97377>.
- Kleinbaum, David G, and Mitchel Klein. 2012. *Survival Analysis: A Self-Learning Text*. 3rd ed. Springer. <https://link.springer.com/book/10.1007/978-1-4419-6646-9>.
- Miller, Steven J. 2017. *The Probability Lifesaver : All the Tools You Need to Understand Chance*. A Princeton Lifesaver Study Guide. Princeton University Press. <https://press.princeton.edu/books/hardcover/9780691149547/the-probability-lifesaver>.
- Rothman, Kenneth J., Timothy L. Lash, Tyler J. VanderWeele, and Sebastien Haneuse. 2021. *Modern Epidemiology*. Fourth edition. Wolters Kluwer.
- Rudin, Walter. 1976. *Principles of Mathematical Analysis*. 3rd ed. International Series in Pure and Applied Mathematics. McGraw-Hill.
- Soch, Joram. 2020. *Proof: Expected Value of a Non-Negative Random Variable*. The Book of Statistical Proofs. <https://doi.org/10.5281/zenodo.4305949>.
- Vittinghoff, Eric, David V Glidden, Stephen C Shiboski, and Charles E McCulloch. 2012. *Regression Methods in Biostatistics: Linear, Logistic, Survival, and Repeated Measures Models*. 2nd ed. Springer. <https://doi.org/10.1007/978-1-4614-1353-0>.
- Wikipedia contributors. 2026. *Law of the Unconscious Statistician — Wikipedia, the Free Encyclopedia*. https://en.wikipedia.org/wiki/Law_of_the_unconscious_statistician.