

Mathematics Prerequisites

Contents

1	Algebra	2
1.1	Elementary Algebra	2
1.1.1	Equalities	2
1.1.2	Inequalities	2
1.1.3	Infimum and supremum	3
1.1.4	Sums	3
1.1.5	Products	4
1.1.6	Division	4
1.1.7	Sums and products together	4
1.1.8	Quotients	4
1.1.9	Exponentials and Logarithms	5
2	Derivatives	9
3	Integration	10
3.1	Antiderivatives	10
3.2	Regularity Conditions	12
3.3	Fundamental Theorem of Calculus	16
4	Double Integrals	19
5	Linear Algebra	31
5.1	Vectors	31
5.1.1	Special vectors	31
5.1.2	Orthogonality	33
5.2	Matrices	34
5.2.1	Matrix transpose	34
5.2.2	Matrix addition	35
5.2.3	Scalar multiplication	35
5.2.4	Matrix multiplication	36
5.2.5	Matrix-vector multiplication	36
5.3	Special Matrices	36
5.4	Quadratic Forms	39
5.5	Design Matrix	40
6	Vector Calculus	40
7	Additional resources	47
7.1	Calculus	47
7.2	Linear Algebra and Vector Calculus	47
7.3	Numerical Analysis	47
7.4	Real Analysis	47
	References	48

Math is not just a way of calculating numerical answers; it is a way of thinking, using clear definitions for concepts and rigorous logic to organize our thoughts and back up our assertions.

Cheng (2025)

These lecture notes use:

- algebra
- precalculus
- univariate calculus
- linear algebra
- vector calculus

Some key results are listed here.

1 Algebra

1.1 Elementary Algebra

Mastery of Elementary Algebra¹ (a.k.a. “College Algebra”) is a prerequisite for calculus, which is a prerequisite for Epi 202 and Epi 203, which are prerequisites for this course (Epi 204). Nevertheless, each year, some Epi 204 students are still uncomfortable with algebraic manipulations of mathematical formulas. Therefore, I include this section as a quick reference.

1.1.1 Equalities

Theorem 1.1 (Equalities are transitive). *If $a = b$ and $b = c$, then $a = c$*

Theorem 1.2 (Substituting equivalent expressions). *If $a = b$, then for any function $f(x)$, $f(a) = f(b)$*

1.1.2 Inequalities

Theorem 1.3 (Adding to both sides of an inequality). *If $a < b$, then $a + c < b + c$*

Theorem 1.4 (Negating both sides of an inequality). *If $a < b$, then: $-a > -b$*

Theorem 1.5 (Multiplying both sides of an inequality by a nonnegative number). *If $a < b$ and $c \geq 0$, then $ca < cb$.*

Theorem 1.6 (Negation is multiplication by -1).

$$-a = (-1) * a$$

¹https://en.wikipedia.org/wiki/Elementary_algebra

1.1.3 Infimum and supremum

Definition 1.1 (Infimum (greatest lower bound)). The **infimum** of a nonempty set $A \subseteq \mathbb{R}$, written $\inf A$, is the greatest real number m satisfying $m \leq a$ for all $a \in A$:

$$\inf A \stackrel{\text{def}}{=} \max_{t \in \mathbb{R}} \{t : \forall a \in A, a \geq t\}$$

If the infimum belongs to A , it equals the minimum: $\inf A = \min A$.

Exm

Example 1.1 (Numerical examples of infimum).

- $\inf\{1, 2, 3\} = 1$, since 1 is the smallest element.
- $\inf(0.5, 1] = 0.5 = \min[0.5, 1]$: for intervals open below, the infimum equals the minimum of the corresponding closed-below interval, even though $0.5 \notin (0.5, 1]$. More generally, $\inf(c, b] = \min[c, b] = c$ for any $c < b$.
- $\inf\{t \geq 0 : t > 0.5\} = 0.5$, even though 0.5 itself is not in the set.

Definition 1.2 (Supremum (least upper bound)). The **supremum** of a nonempty set $A \subseteq \mathbb{R}$, written $\sup A$, is the smallest real number M satisfying $M \geq a$ for all $a \in A$:

$$\sup A \stackrel{\text{def}}{=} \min_{t \in \mathbb{R}} \{t : \forall a \in A, a \leq t\}$$

If the supremum belongs to A , it equals the maximum: $\sup A = \max A$.

Exm

Example 1.2 (Numerical examples of supremum).

- $\sup\{1, 2, 3\} = 3$, since 3 is the largest element.
- $\sup\{t \geq 0 : t < 0.5\} = 0.5$, even though 0.5 itself is not in the set.

1.1.4 Sums

Theorem 1.7 (adding zero changes nothing).

$$a + 0 = a$$

Theorem 1.8 (Sums are symmetric).

$$a + b = b + a$$

Theorem 1.9 (Sums are associative).

When summing three or more terms, the order in which you sum them does not matter:

$$(a + b) + c = a + (b + c)$$

1.1.5 Products

Theorem 1.10 (Multiplying by 1 changes nothing).

$$a \times 1 = a$$

Theorem 1.11 (Products are symmetric).

$$a \times b = b \times a$$

Theorem 1.12 (Products are associative).

$$(a \times b) \times c = a \times (b \times c)$$

1.1.6 Division

Theorem 1.13 (Division can be written as a product).

$$\frac{a}{b} = a \times \frac{1}{b}$$

1.1.7 Sums and products together

Theorem 1.14 (Multiplication is distributive).

$$a(b + c) = ab + ac$$

1.1.8 Quotients

Definition 1.3 (Quotients, fractions, rates).

A **quotient**, **fraction**, or **rate** is a division of one quantity by another:

$$\frac{a}{b}$$

In epidemiology, rates typically have a quantity involving time or population in the denominator.
c.f. [https://en.wikipedia.org/wiki/Rate_\(mathematics\)](https://en.wikipedia.org/wiki/Rate_(mathematics))

Definition 1.4 (Ratios). A **ratio** is a quotient in which the numerator and denominator are measured using the same unit scales.

c.f. <https://en.wikipedia.org/wiki/Ratio>

Definition 1.5 (Proportion). In statistics, a “proportion” typically means a ratio where the numerator represents a subset of the denominator.

See https://en.wikipedia.org/wiki/Population_proportion.

See also [https://en.wikipedia.org/wiki/Proportion_\(mathematics\)](https://en.wikipedia.org/wiki/Proportion_(mathematics)) for other meanings.

Definition 1.6 (Proportional). Two functions $f(x)$ and $g(x)$ are **proportional** if their ratio $\frac{f(x)}{g(x)}$ does not depend on x . (c.f. [https://en.wikipedia.org/wiki/Proportionality_\(mathematics\)](https://en.wikipedia.org/wiki/Proportionality_(mathematics)))

Additional reference for elementary algebra: https://en.wikipedia.org/wiki/Population_proportion#Mathematical_definition

1.1.9 Exponentials and Logarithms

Theorem 1.15 (logarithm of a product is the sum of the logs of the factors).

$$\log a \cdot b = \log a + \log b$$

Corollary 1.1 (logarithm of a quotient).

The logarithm of a quotient is equal to the difference of the logs of the factors:

$$\log \frac{a}{b} = \log a - \log b$$

Theorem 1.16 (logarithm of an exponential function).

$$\log\{a^b\} = b \cdot \log\{a\}$$

Theorem 1.17 (exponential of a sum).

The exponential of a sum is equal to the product of the exponentials of the addends:

$$\exp\{a + b\} = \exp\{a\} \cdot \exp\{b\}$$

Corollary 1.2 (exponential of a difference).

The exponential of a difference is equal to the quotient of the exponentials of the addends:

$$\exp\{a - b\} = \frac{\exp\{a\}}{\exp\{b\}}$$

Theorem 1.18 (exponential of a product).

$$a^{bc} = (a^b)^c = (a^c)^b$$

Corollary 1.3 (natural exponential of a product).

$$\exp\{ab\} = (\exp\{a\})^b = (\exp\{b\})^a$$

Exercise 1.1. For $a \geq 0$, $b, c \in \mathbb{R}$, When does $(a^b)^c = a^{(b^c)}$?

Solution

Solution 1.1. Short answer: rarely (that's all you need to know for this course).

Long answer:

If $(a^b)^c = a^{(b^c)}$, then since $(a^b)^c = a^{bc}$, we have:

$$\begin{aligned} a^{bc} &= a^{(b^c)} \\ \log\{a^{bc}\} &= \log\{a^{(b^c)}\} \\ bc \cdot \log\{a\} &= b^c \cdot \log\{a\} \end{aligned} \tag{1}$$

Equation 1 holds in each of the following cases:

1. $bc = b^c$ (see Exercise 1.2).
2. $a = 1$ (i.e., $\log\{a\} = 0$).
3. $a = 0$ (i.e., $\log\{a\} = -\infty$) and $\text{sign}\{bc\} = \text{sign}\{b^c\}$.

In particular, when $a = 0$ and $c = 0$, $bc = 0$ and $b^c = 1$ (for any $b \in \mathbb{R}$), so $\text{sign}\{bc\} \neq \text{sign}\{b^c\}$, and $(a^b)^c \neq a^{(b^c)}$:

$$\begin{aligned} (a^b)^c &= (0^b)^0 \\ &= 1 \\ a^{(b^c)} &= 0^{(b^0)} \\ &= 0^1 \\ &= 0 \end{aligned}$$

Exercise 1.2. For $b, c \in \mathbb{R}$, when does $b^c = bc$?

Solution

Solution 1.2. $bc = b^c$ in each of the following cases:

1. $c = 1$.
2. $b = 0$ and $c > 0$.
3. $b = \exp\left\{\frac{\log c}{c-1}\right\}$ (for $c \geq 0$).

See the red contours in Figure 2 for a visualization.

```

`b*c_f` <- function(b, c) b*c
`b^c_f` <- function(b, c) b^c
values_b <- seq(0, 5, by = .01)
values_c <- seq(-.5, 3, by = .01)

`b*c` <- outer(values_b, values_c, `b*c_f`)
`b^c` <- outer(values_b, values_c, `b^c_f`)
`b^c`[is.infinite(`b^c`)] = NA

opacity <- .3
z_min <- min(`b*c`, `b^c`, na.rm = TRUE)
z_max <- 5
plotly::plot_ly(
  x = ~values_b,
  y = ~values_c
) |>
  plotly::add_surface(
    z = ~ t(`b*c`),
    contours = list(
      z = list(
        show = TRUE,
        start = -1,
        end = 1,
        size = .1
      )
    ),
    name = "b*c",
    showscale = FALSE,
    opacity = opacity,
    colorscale = list(c(0, 1), c("green", "green"))
  ) |>
  plotly::add_surface(
    opacity = opacity,
    colorscale = list(c(0, 1), c("red", "red")),
    z = ~ t(`b^c`),
    contours = list(
      z = list(
        show = TRUE,
        start = z_min,
        end = z_max,
        size = .2
      )
    ),
    showscale = FALSE,
    name = "b^c"
  ) |>
  plotly::layout(
    scene = list(
      xaxis = list(
        # type = "log",
        title = "b"
      ),
      yaxis = list(
        # type = "log",
        title = "c"
      ),
      zaxis = list(
        # type = "log",
        range = c(z_min, z_max),
        title = "outcome"
      ),
      camera = list(eye = list(x = -1.25, y = -1.25, z = 0.5)),

```

```

`b^c - b*c_f` <- function(b, c) `b^c_f`(b,c) - `b*c_f`(b,c)

mat1 <- outer(values_b, values_c, `b^c - b*c_f`)
mat1[is.infinite(mat1)] = NA

opacity <- .3
plotly::plot_ly(
  x = ~values_b,
  y = ~values_c
) |>
  plotly::add_surface(
    z = ~ t(mat1),
    contours = list(
      z = list(
        show = TRUE,
        start = 0,
        end = 1,
        size = 1,
        color = "red"
      )
    ),
    name = "b^c - b*c",
    showscale = TRUE,
    opacity = opacity
  ) |>
  plotly::layout(
    scene = list(
      xaxis = list(
        # type = "log",
        title = "b"
      ),
      yaxis = list(
        # type = "log",
        title = "c"
      ),
      zaxis = list(
        title = "outcome"
      ),
      camera = list(eye = list(x = -1.25, y = -1.25, z = 0.5)),
      aspectratio = list(x = .9, y = .8, z = 0.7)
    )
  )

```

Theorem 1.19 ($\exp\{\}$ and $\log\{\}$ are mutual inverses).

$$\exp\{\log\{a\}\} = \log\{\exp\{a\}\} = a$$

2 Derivatives

Theorem 2.1 (Constant rule).

$$\frac{\partial}{\partial x} c = 0$$

Theorem 2.2 (Power rule). *If a is constant with respect to x , then:*

$$\frac{\partial}{\partial x} ay = a \frac{\partial x}{\partial y}$$

Theorem 2.3 (Power rule).

$$\frac{\partial}{\partial x} x^q = qx^{q-1}$$

Theorem 2.4 (Derivative of natural logarithm).

$$\log'\{x\} = \frac{1}{x} = x^{-1}$$

Theorem 2.5 (derivative of exponential).

$$\exp'\{x\} = \exp\{x\}$$

Theorem 2.6 (Product rule).

$$(ab)' = ab' + ba'$$

Theorem 2.7 (Quotient rule).

$$(a/b)' = a'/b - (a/b^2)b'$$

Theorem 2.8 (Chain rule).

$$\begin{aligned} \frac{\partial a}{\partial c} &= \frac{\partial a}{\partial b} \frac{\partial b}{\partial c} \\ &= \frac{\partial b}{\partial c} \frac{\partial a}{\partial b} \end{aligned}$$

or in Euler/Lagrange notation^a:

$$(f(g(x)))' = g'(x)f'(g(x))$$

^ahttps://en.wikipedia.org/wiki/Notation_for_differentiation#Lagrange's_notation

Corollary 2.1 (Chain rule for logarithms).

$$\frac{\partial}{\partial x} \log f(x) = \frac{f'(x)}{f(x)}$$

i Proof

Proof. Apply Theorem 2.8 and Theorem 2.4. □

3 Integration

Integration is the inverse operation of differentiation: it recovers a function from its derivative and accumulates quantities such as areas, totals, and probabilities. We begin with antiderivatives, then state basic integration rules, and conclude with the Fundamental Theorem of Calculus and a worked example from probability.

3.1 Antiderivatives

Definition 3.1 (Antiderivative). A function F is an **antiderivative** of f on an interval I if:

$$\frac{\partial}{\partial x} F(x) = f(x), \quad \forall x \in I$$

The family of all antiderivatives of f is written as the **indefinite integral**:

$$\int f(x) dx = F(x) + C$$

where C is an arbitrary constant of integration.

(Larson and Edwards 2018, sec. 4.1, pp. 248–249)

Exm

Example 3.1 (Antiderivative of a power function). For $f(x) = x^2$, an antiderivative is $F(x) = \frac{x^3}{3}$, since $\frac{\partial}{\partial x} \frac{x^3}{3} = x^2 = f(x)$.

Adding any constant C gives another antiderivative; for example, with $C = 7$, $F(x) = \frac{x^3}{3} + 7$ also satisfies $F'(x) = x^2$, since adding a constant does not change the derivative. See Figure 3.

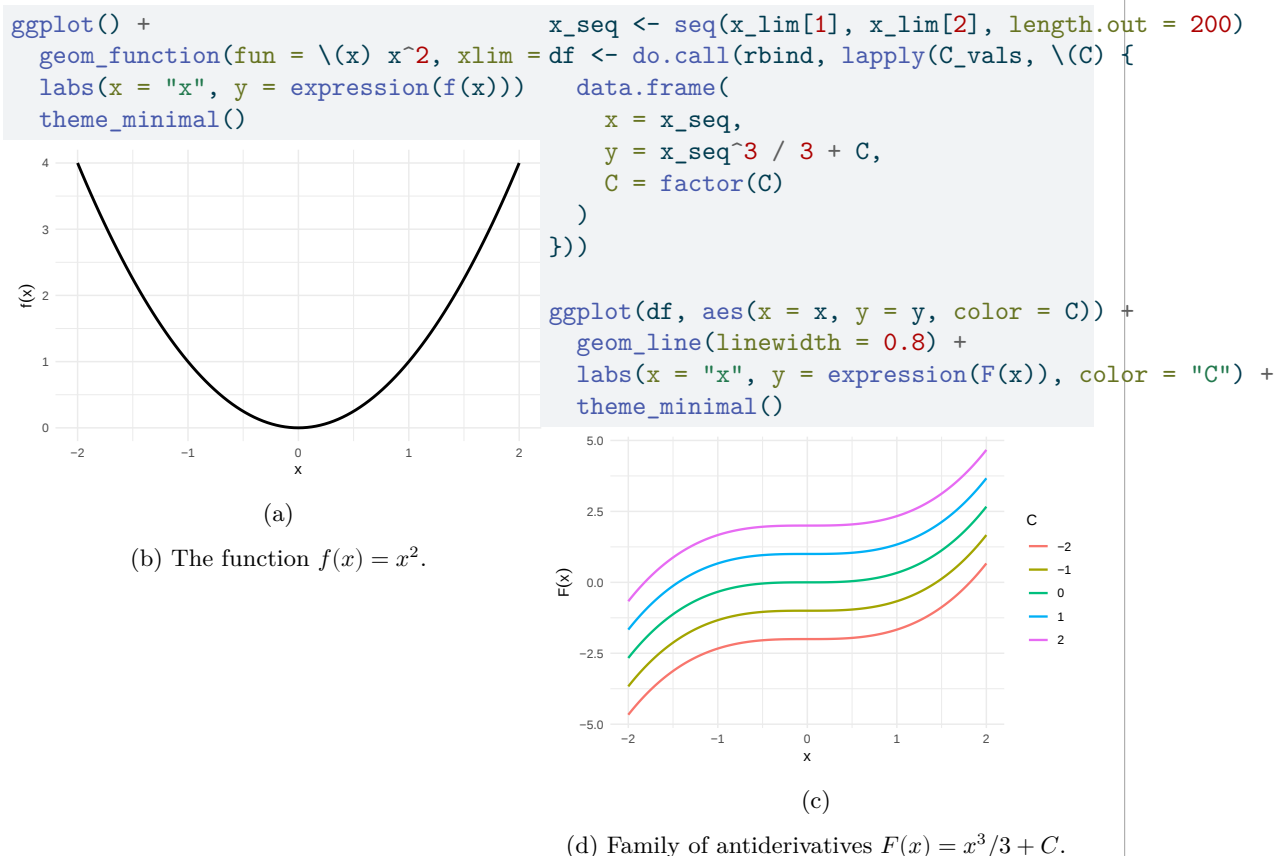


Figure 3: The function $f(x) = x^2$ and five antiderivatives $F(x) = x^3/3 + C$ for $C \in \{-2, -1, 0, 1, 2\}$. Each antiderivative has the same derivative f ; they differ only by a vertical shift.

Theorem 3.1 (Basic integration rules). *Each antiderivative in Theorem 3.1 is defined only up to an arbitrary constant C (see Definition 3.1); the table omits $+C$ from every row for brevity.*

Function $f(x)$	Antiderivative $F(x)$	Condition
c	cx	—
x^n	$\frac{x^{n+1}}{n+1}$	$n \neq -1$
$\frac{1}{x}$	$\ln x $	$x \neq 0$
e^x	e^x	(self-antiderivative)
e^{cx}	$\frac{1}{c}e^{cx}$	$c \neq 0$
$c \cdot f(x)$	$c \cdot F(x)$	—
$f(x) + g(x)$	$F(x) + G(x)$	—

The first two rows and the bottom two rows (linearity) are from (Larson and Edwards 2018, sec. 4.1, p. 250 “Basic Integration Rules”); $1/x$ is from (Larson and Edwards 2018, sec. 5.2, Theorem 5.5, p. 324); e^x and e^{cx} are from (Larson and Edwards 2018, sec. 5.4, Theorem 5.12, p. 346).

Exm

Example 3.2 (Antiderivative of $3x^2 - 1$). By the power rule ($n = 2$) and linearity from Theorem 3.1:

$$\int (3x^2 - 1) dx = 3 \cdot \frac{x^3}{3} - x + C = x^3 - x + C.$$

Verify by differentiating: $\frac{\partial}{\partial x}(x^3 - x + C) = 3x^2 - 1 = f(x)$, as required.

3.2 Regularity Conditions

Definition 3.2 (Differentiable function). A function f is **differentiable at** $x = c$ if the limit

$$f'(c) \stackrel{\text{def}}{=} \lim_{h \rightarrow 0} \frac{f(c+h) - f(c)}{h}$$

exists and is finite. f is **differentiable on** an interval if it is differentiable at every interior point; at a closed endpoint, the appropriate one-sided derivative is used. (Larson and Edwards 2018, sec. 2.1, p. 100)

Definition 3.3 (Continuous function). A function f is **continuous at** $x = c$ if all three conditions hold:

1. $f(c)$ is defined,
2. $\lim_{x \rightarrow c} f(x)$ exists, and
3. $\lim_{x \rightarrow c} f(x) = f(c)$.

f is **continuous on** a closed interval $[a, b]$ if it is continuous at every point of $[a, b]$. (Larson and Edwards 2018, sec. 1.4, p. 73)

Definition 3.4 (Riemann integrable). A bounded function f is **Riemann integrable on** $[a, b]$ if the Riemann integral

$$\int_a^b f(x) dx \stackrel{\text{def}}{=} \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_i^*) \Delta x, \quad \Delta x \stackrel{\text{def}}{=} \frac{b-a}{n},$$

exists and is finite (for equal-width partitions of width $\Delta x = (b-a)/n$), where x_i^* is any point in the i -th subinterval.

(Larson and Edwards 2018, sec. 4.3, p. 272)

Definition 3.5 (General Riemann integrability). More generally, using partitions $\mathcal{P}$ of arbitrary mesh — subintervals of varying widths Δx_i — a bounded function f is Riemann integrable on $[a, b]$ if

$$\int_a^b f(x) dx \stackrel{\text{def}}{=} \lim_{\|\mathcal{P}\| \rightarrow 0} \sum_{i=1}^n f(x_i^*) \Delta x_i$$

exists and is finite, where $\|\mathcal{P}\| = \max_i \Delta x_i$ is the mesh of the partition.

Theorem 3.2 (Equivalence of Riemann sum formulations). *For continuous f on a closed interval $[a, b]$, the equal-width Riemann sum (Definition 3.4) and the arbitrary-mesh Riemann sum (Definition 3.5) give the same value (Rudin 1976, chap. 6). The equal-width form in Definition 3.4 is used throughout Epi 204.*

Before stating the Fundamental Theorem of Calculus, we record two prerequisite results. The FTC requires the integrand f to be continuous, and the following two theorems establish where continuity

comes from (differentiability $\Rightarrow$ continuity) and what it buys us (continuity $\Rightarrow$ integrability).

Theorem 3.3 (Differentiability implies continuity). *If f is differentiable at $x = c$, then f is continuous at $x = c$.*

(Larson and Edwards 2018, Theorem 2.1, p. 106)

Exm

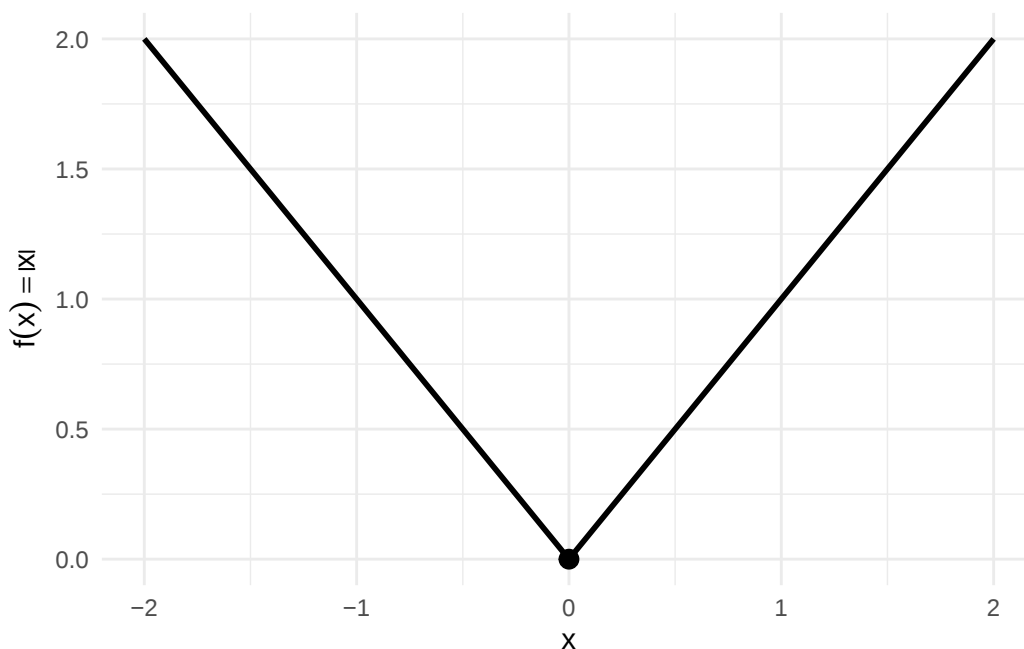
Example 3.3 (Differentiable, hence continuous: $x^3 - x$). $f(x) = x^3 - x$ is differentiable everywhere (with derivative $f'(x) = 3x^2 - 1$), so by Theorem 3.3 it is continuous everywhere.

Exm

Example 3.4 (Continuous but not differentiable: $|x|$). The absolute-value function $f(x) = |x|$ is continuous at $x = 0$ ($\lim_{x \rightarrow 0} |x| = 0 = |0|$), but it is not differentiable at $x = 0$: the left-derivative is -1 and the right-derivative is $+1$.

This counterexample shows that the converse of Theorem 3.3 fails: continuity does not imply differentiability. See Figure 4.

```
ggplot() +  
  geom_function(fun = abs, xlim = c(-2, 2), linewidth = 1) +  
  geom_point(aes(x = 0, y = 0), size = 3) +  
  labs(x = "x", y = expression(f(x) == abs(x))) +  
  theme_minimal()
```



(a)

Figure 4: $f(x) = |x|$ has a sharp corner at $x = 0$ (not differentiable there) but is continuous everywhere: no gaps or jumps.

Theorem 3.4 (Continuity implies integrability). *If f is continuous on the closed interval $[a, b]$, then f is integrable on $[a, b]$ (i.e., the Riemann integral $\int_a^b f(x) dx$ exists and is finite).*

(Larson and Edwards 2018, Theorem 4.4, p. 272)

Exm

Example 3.5 (Continuous, hence integrable: polynomials). Every polynomial is continuous on $\mathbb{R}$, so by Theorem 3.4 every polynomial is integrable on every closed interval $[a, b]$.

Exm

Example 3.6 (Integrable but not continuous: a step function). Let $f(x) = 0$ for $x < \frac{1}{2}$ and $f(x) = 1$ for $x \geq \frac{1}{2}$. Then f is discontinuous at $x = \frac{1}{2}$, but it is integrable on $[0, 1]$:

$$\int_0^1 f(x) dx = \int_0^{1/2} 0 dx + \int_{1/2}^1 1 dx = 0 + \frac{1}{2} = \frac{1}{2}.$$

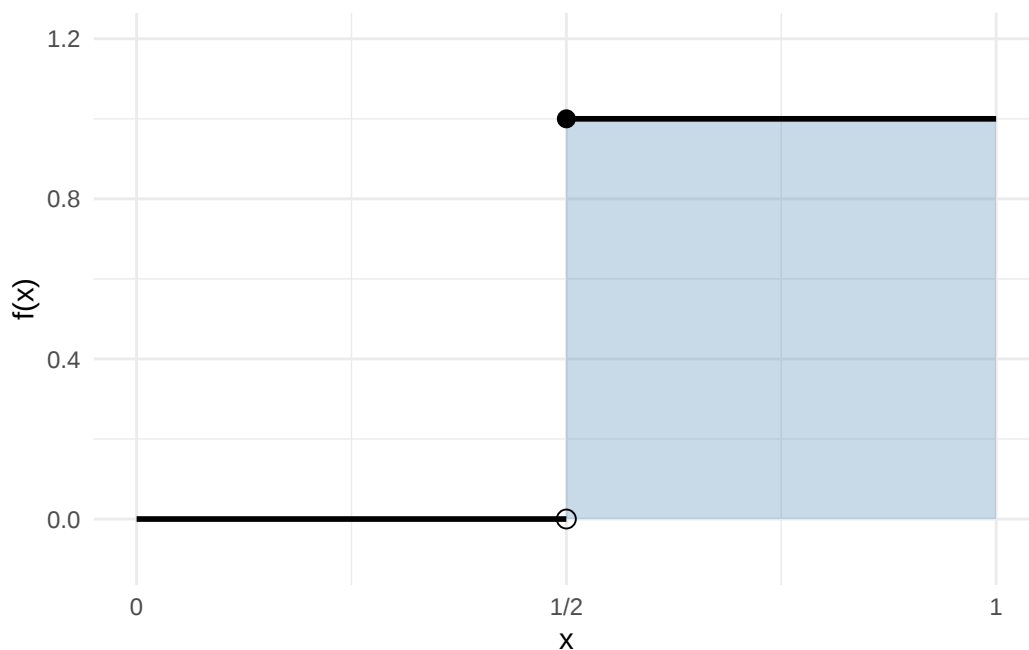
This counterexample shows that the converse of Theorem 3.4 fails: integrability does not imply continuity. See Figure 5.

```

step_df <- data.frame(
  x = c(0, 0.5, 0.5, 1),
  y = c(0, 0, 1, 1),
  segment = c("left", "left", "right", "right")
)

ggplot() +
  geom_rect(
    aes(xmin = 0.5, xmax = 1, ymin = 0, ymax = 1),
    fill = "steelblue", alpha = 0.3
  ) +
  geom_line(
    data = step_df,
    aes(x = x, y = y, group = segment),
    linewidth = 1
  ) +
  geom_point(aes(x = 0.5, y = 0), shape = 1, size = 3) +
  geom_point(aes(x = 0.5, y = 1), shape = 16, size = 3) +
  scale_x_continuous(breaks = c(0, 0.5, 1), labels = c("0", "1/2", "1")) +
  scale_y_continuous(limits = c(-0.1, 1.2)) +
  labs(x = "x", y = "f(x)") +
  theme_minimal()

```



(a)

Figure 5: Step function: $f(x) = 0$ on $[0, \frac{1}{2})$ (open circle at the jump) and $f(x) = 1$ on $[\frac{1}{2}, 1]$ (filled circle). The shaded rectangle has area $\frac{1}{2}$, matching the integral computed in Example 3.6.

Together, Theorem 3.3 and Theorem 3.4 establish the chain:

$$\text{differentiable} \Rightarrow \text{continuous} \Rightarrow \text{integrable}$$

Example 3.4 and Example 3.6 show that neither implication reverses in general.

3.3 Fundamental Theorem of Calculus

Theorem 3.5 (Fundamental Theorem of Calculus). *Let f be a continuous function on a closed interval $[a, b]$.*

Part 1 (Derivative of an integral). *Define $F(x) = \int_a^x f(t) dt$ for $x \in [a, b]$. Then F is differentiable and:*

$$\frac{\partial}{\partial x} \int_a^x f(t) dt = f(x) \quad (2)$$

Note

Continuity on all of $[a, b]$ is a sufficient condition. More generally, Part 1 holds at any individual point x where f is integrable on $[a, b]$ (see Definition 3.4) and continuous at x (see Definition 3.3), even if f has jump discontinuities elsewhere (Rudin 1976, Theorem 6.20, p. 133).

(Larson and Edwards 2018, Theorem 4.11, p. 288)

Part 2 (Evaluation theorem). *The F here may be any antiderivative of f — not just the accumulation function from Part 1. If F is an antiderivative of f on $[a, b]$ (i.e., $\frac{\partial}{\partial x} F(x) = f(x)$ for all $x \in [a, b]$), then:*

$$\int_a^b f(x) dx = F(b) - F(a) \quad (3)$$

Equivalently, with b replaced by a variable upper limit x , integrating the derivative of F recovers the net change in F :

$$\int_a^x F'(t) dt = F(x) - F(a) \quad (4)$$

or equivalently in Leibniz notation:

$$\int_a^x \frac{dF}{dt} dt = F(x) - F(a)$$

(Banner 2007, chap. 18; Larson and Edwards 2018, Theorem 4.9, p. 282)

The two parts of the FTC together express that **differentiation and integration are inverse operations**:

- Part 1: differentiating the integral of f recovers f (Equation 2).
- Part 2: the integral of f over $[a, b]$ equals the difference of any antiderivative's values at the endpoints (Equation 3), which rearranges to “integrating the derivative of F recovers the net change in F ” (Equation 4).

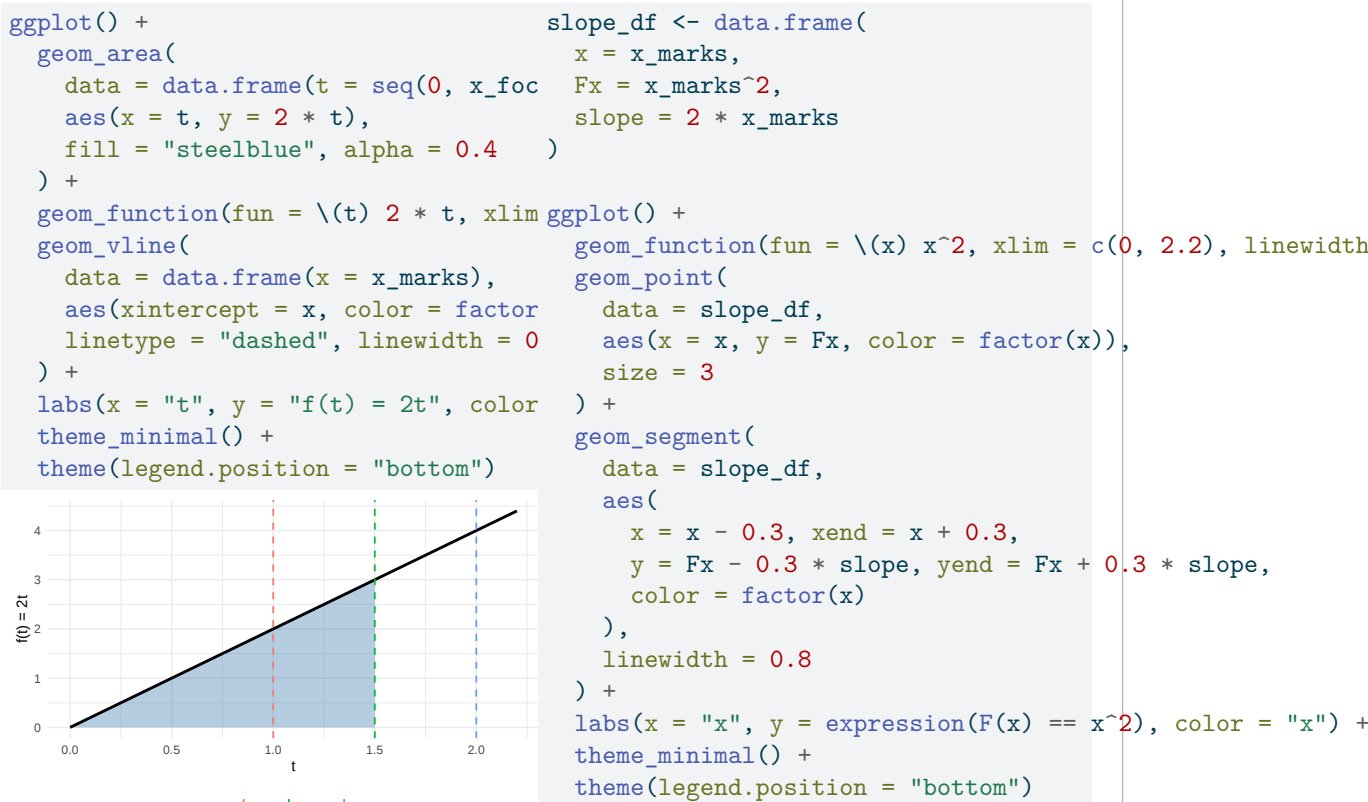
The standard form of the FTC assumes f is continuous on $[a, b]$; continuity is *sufficient* but not strictly necessary (see the callout note inside Theorem 3.5 for the more general statement). Since differentiability implies continuity (Theorem 3.3), the FTC applies in particular whenever f is differentiable — a common situation in Epi 204.

Exm

Example 3.7 (FTC Part 1 visualized: accumulation function for $f(t) = 2t$). Take $f(t) = 2t$ on $[0, 2]$. The accumulation function from 0 is

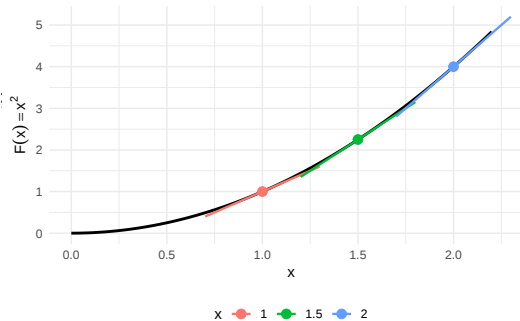
$$F(x) \stackrel{\text{def}}{=} \int_0^x 2t dt = [t^2]_{t=0}^{t=x} = x^2 - 0^2 = x^2,$$

so $F(x) = x^2$, and indeed $F'(x) = 2x = f(x)$, as Theorem 3.5 Part 1 predicts. Figure 6 shows the integrand on the left (shaded area equals $F(x)$ at each x) and the accumulation function $F(x) = x^2$ on the right (its slope at x equals $f(x) = 2x$).



(a)

(b) $f(t) = 2t$; shaded area equals $F(1.5) = 2.25$; vertical lines mark $x \in \{1, 1.5, 2\}$.



(c)

(d) $F(x) = x^2$; tangent slope at each marked x equals $f(x) = 2x$.

Figure 6: Left: $f(t) = 2t$; the shaded area $\int_0^{1.5} 2t \, dt = F(1.5) = 2.25$; vertical lines mark $x \in \{1, 1.5, 2\}$. Right: $F(x) = x^2$; for each marked x , the tangent slope equals $f(x) = 2x$.

Exm

Example 3.8 (CDF and PDF of the exponential distribution). In what follows, f denotes the PDF and F the CDF — the same letters as the antiderivative pair in Definition 3.1, because the FTC will show F is exactly an antiderivative of f .

For the exponential distribution with rate parameter $\lambda > 0$, the probability density function (PDF) is (Kleinbaum and Klein 2012, sec. II, p. 295, “Survival and Hazard Functions for Selected Distributions”):

$$f(t) = \lambda e^{-\lambda t}, \quad t \geq 0$$

FTC Part 2 gives the cumulative distribution function (CDF) from the PDF. Apply the e^{cx}

rule from Theorem 3.1 with $c = -\lambda$ to antidifferentiate the integrand:

$$\begin{aligned}
 F(t) &= \int_0^t \lambda e^{-\lambda u} du \\
 &= \left[\lambda \cdot \frac{1}{-\lambda} e^{-\lambda u} \right]_{u=0}^{u=t} \\
 &= [(-1)e^{-\lambda u}]_{u=0}^{u=t} \\
 &= [-e^{-\lambda u}]_{u=0}^{u=t} \\
 &= -e^{-\lambda t} - (-e^0) \\
 &= -e^{-\lambda t} - (-1) \\
 &= 1 - e^{-\lambda t}
 \end{aligned}$$

FTC Part 1 recovers the PDF from the CDF:

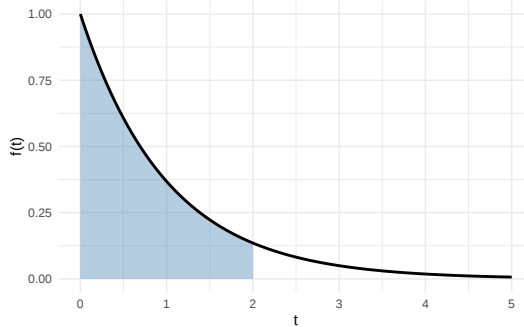
$$\begin{aligned}
 \frac{\partial}{\partial t} F(t) &= \frac{\partial}{\partial t} (1 - e^{-\lambda t}) \\
 &= 0 - (-\lambda)e^{-\lambda t} \\
 &= \lambda e^{-\lambda t} \\
 &= f(t)
 \end{aligned}$$

For a concrete instance: with $\lambda = 1$ (standard exponential), the probability that $T \leq 2$ is:

$$F(2) = 1 - e^{-1 \cdot 2} = 1 - e^{-2} \approx 1 - 0.135 = 0.865$$

See Figure 7.

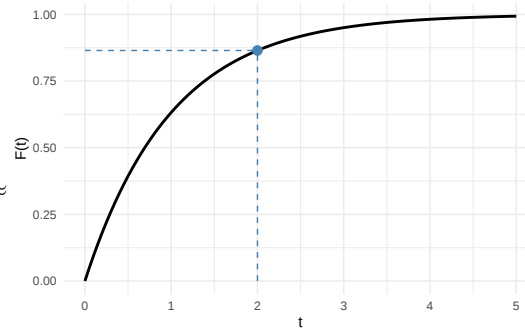
```
ggplot() +
  geom_area(
    data = data.frame(t = seq(0, t_foc
    aes(x = t, y = lambda * exp(-lambda *
    fill = "steelblue", alpha = 0.4
  ) +
  geom_function(
    fun = \(t) lambda * exp(-lambda *
    xlim = c(0, t_max), linewidth = 1
  ) +
  labs(x = "t", y = "f(t)") +
  theme_minimal()
```



(a)

(b) PDF with $\lambda = 1$; shaded area equals $F(2) \approx 0.865$.

```
ggplot() +
  geom_function(
    fun = \(t) 1 - exp(-lambda * t),
    xlim = c(0, t_max), linewidth = 1
  ) +
  geom_point(
    aes(x = t_focus, y = F_at_focus),
    size = 3, color = "steelblue"
  ) +
  geom_segment(
    aes(x = t_focus, xend = t_focus, y = 0, yend = F_at_foc
    linetype = "dashed", color = "steelblue"
  ) +
  geom_segment(
    aes(x = 0, xend = t_focus, y = F_at_focus, yend = F_at_
    linetype = "dashed", color = "steelblue"
  ) +
  labs(x = "t", y = "F(t)") +
  theme_minimal()
```



(c)

(d) CDF with $\lambda = 1$; point marks $F(2) \approx 0.865$.

Figure 7: Exponential distribution with $\lambda = 1$. Left: the PDF $f(t) = \lambda e^{-\lambda t}$; the shaded area under the curve from 0 to 2 equals $F(2) \approx 0.865$. Right: the CDF $F(t) = 1 - e^{-\lambda t}$; the dashed lines mark the value $F(2)$ computed via FTC Part 2.

4 Double Integrals

The **Fubini–Tonelli theorem** states conditions under which the order of integration in a double integral can be exchanged. We state two versions: the Riemann version (Theorem 4.1) is what Epi 204 directly uses for double integrals of continuous functions on simple regions; the σ -finite measure-theoretic version (Theorem 4.2) is included to make the joint-distribution form² corollary in the probability chapter follow from a stated theorem rather than from an aside.

Theorem 4.1 (Fubini’s theorem (Riemann version)). *Let f be **continuous** on a plane region $R \subseteq \mathbb{R}^2$.*

1. **Vertically simple region.** *If R is defined by $a \leq x \leq b$ and $g_1(x) \leq y \leq g_2(x)$, where g_1 and g_2 are continuous on $[a, b]$, then*

$$\iint_R f(x, y) dA = \int_a^b \int_{g_1(x)}^{g_2(x)} f(x, y) dy dx.$$

²[probability.qmd#cor-fubini-joint](#)

2. **Horizontally simple region.** If R is defined by $c \leq y \leq d$ and $h_1(y) \leq x \leq h_2(y)$, where h_1 and h_2 are continuous on $[c, d]$, then

$$\iint_R f(x, y) dA = \int_c^d \int_{h_1(y)}^{h_2(y)} f(x, y) dx dy.$$

When R can be described both ways, the two iterated integrals are equal — so the order of integration can be exchanged.

(Larson and Edwards 2018, Theorem 14.2, p. 982)

Exm

Example 4.1 (Changing the order of integration for a non-rectangular region). Adapted from (Larson and Edwards 2018, sec. 14.2, Example 4, pp. 984–985).

Let X and Y be independent $\text{Uniform}(0, 1)$ random variables, with joint density $f(x, y) = 1$ on the unit square $[0, 1]^2$. Define the function $g(x, y) = e^{-x^2} \mathbb{1}(y \leq x)$, and compute its expectation $E[g(X, Y)]$.

Because the joint density equals 1 on $[0, 1]^2$, this expectation is the double integral of g over the unit square:

$$E[g(X, Y)] = \iint_{[0, 1]^2} g(x, y) dA.$$

The indicator factor $\mathbb{1}(y \leq x)$ equals 1 on the triangular region where $y \leq x$ and 0 elsewhere, so only that region, namely $D = \{(x, y) : x \in [0, 1], y \in [0, x]\}$ (Figure 8), contributes, and there $g(x, y) = e^{-x^2}$:

$$E[g(X, Y)] = \iint_D e^{-x^2} dA.$$

```

region <- data.frame(x = c(0, 1, 1), y = c(0, 0, 1))
ggplot2::ggplot(region, ggplot2::aes(x = x, y = y)) +
  ggplot2::geom_polygon(
    fill = "steelblue", alpha = 0.4, color = "black", linewidth = 0.7
  ) +
  ggplot2::annotate(
    "text", x = 0.65, y = 0.28, label = "D", size = 8, fontface = "italic"
  ) +
  ggplot2::annotate(
    "text", x = 0.36, y = 0.52, label = "y == x",
    parse = TRUE, angle = 45
  ) +
  ggplot2::annotate(
    "text", x = 1.08, y = 0.5, label = "x == 1",
    parse = TRUE, angle = 90, hjust = 0.5
  ) +
  ggplot2::annotate(
    "text", x = 0.5, y = -0.1, label = "y == 0", parse = TRUE
  ) +
  ggplot2::scale_x_continuous(breaks = c(0, 0.5, 1), limits = c(-0.1, 1.3)) +
  ggplot2::scale_y_continuous(breaks = c(0, 0.5, 1), limits = c(-0.2, 1.2)) +
  ggplot2::coord_equal() +
  ggplot2::labs(x = "x", y = "y") +
  ggplot2::theme_minimal()

```

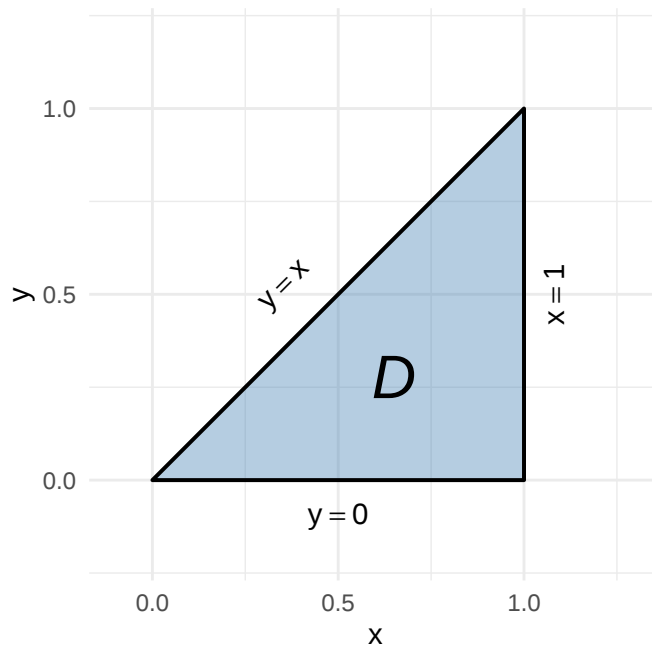


Figure 8: Triangular integration region $D = \{(x, y) : x \in [0, 1], y \in [0, x]\}$, bounded below by $y = 0$, above-left by $y = x$, and on the right by $x = 1$.

Order $dx dy$ is intractable. Re-describing D as $D = \{(x, y) : y \in [0, 1], x \in [y, 1]\}$, the inner integral is

$$\int_y^1 e^{-x^2} dx,$$

which has no elementary antiderivative.

Order $dy\,dx$ works. Applying Theorem 4.1 Part 1 (e^{-x^2} is continuous and D is the vertically simple region $x \in [0, 1]$, $y \in [0, x]$):

$$\begin{aligned}
 \iint_D e^{-x^2} dA &= \int_0^1 \int_0^x e^{-x^2} dy\,dx \\
 &= \int_0^1 e^{-x^2} \left(\int_0^x dy \right) dx \\
 &= \int_0^1 x e^{-x^2} dx \\
 &= \left[-\frac{1}{2} e^{-x^2} \right]_0^1 \\
 &= -\frac{1}{2} (e^{-1} - 1) \\
 &= \frac{e - 1}{2e} \\
 &\approx 0.316
 \end{aligned}$$

The solid whose volume equals this integral is shown in Figure 9.

```

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 <- exp(-x^2)
  z[y > x] <- 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 = "z = exp(-x^2)"),
    camera = list(eye = list(x = 1.6, y = -1.6, z = 0.8))
  ))

```

Example 4.2 (When conditions fail: a counterexample). The conditions in Theorem 4.1 are not merely technical — when they fail, iterated integrals can exist yet disagree.

Let

$$f(x, y) = \frac{x^2 - y^2}{(x^2 + y^2)^2}$$

on the unit square $R = [0, 1] \times [0, 1]$. Strictly, f is defined on $R \setminus \{(0, 0)\}$: the denominator vanishes at the origin, so f is undefined there (we return to this point in the condition check).

Integrating y first, then x :

Using $\frac{\partial}{\partial y} \frac{y}{x^2 + y^2} = \frac{x^2 - y^2}{(x^2 + y^2)^2}$:

$$\begin{aligned} \int_0^1 \int_0^1 f(x, y) dy dx &= \int_0^1 \left[\frac{y}{x^2 + y^2} \right]_{y=0}^{y=1} dx \\ &= \int_0^1 \frac{1}{x^2 + 1} dx \\ &= [\arctan(x)]_0^1 \\ &= \frac{\pi}{4} \end{aligned}$$

Integrating x first, then y :

Using $\frac{\partial}{\partial x} \left(-\frac{x}{x^2 + y^2} \right) = \frac{x^2 - y^2}{(x^2 + y^2)^2}$:

$$\begin{aligned} \int_0^1 \int_0^1 f(x, y) dx dy &= \int_0^1 \left[-\frac{x}{x^2 + y^2} \right]_{x=0}^{x=1} dy \\ &= \int_0^1 \left(-\frac{1}{1 + y^2} \right) dy \\ &= -[\arctan(y)]_0^1 \\ &= -\frac{\pi}{4} \end{aligned}$$

Conclusion: $\frac{\pi}{4} \neq -\frac{\pi}{4}$, so the two iterated integrals are unequal. Theorem 4.1 does not apply here.

Why Theorem 4.1's condition fails: Theorem 4.1 requires f to be **continuous** on R . The denominator $(x^2 + y^2)^2$ vanishes at the origin $(0, 0) \in R$, so f is *not even defined* there — let alone continuous — and the theorem does not apply.

(Wikipedia contributors 2024)

The surface, and the singularity at the origin responsible for the failure, are shown in Figure 10.

```

n_grid <- 81
eps <- 0.04
x_seq <- seq(eps, 1, length.out = n_grid)
y_seq <- seq(eps, 1, length.out = n_grid)

z_mat <- outer(x_seq, y_seq, function(x, y) (x^2 - y^2) / (x^2 + y^2)^2)
z_clip <- 50
z_mat[z_mat > z_clip] <- z_clip
z_mat[z_mat < -z_clip] <- -z_clip

plotly::plot_ly(x = ~x_seq, y = ~y_seq, z = ~t(z_mat)) |>
  plotly::add_surface(
    showscale = FALSE,
    colorscale = list(
      list(0, "#3b4cc0"),
      list(0.5, "#ddddd"),
      list(1, "#b40426")
    )
  ) |>
  plotly::layout(scene = list(
    xaxis = list(title = "x"),
    yaxis = list(title = "y"),
    zaxis = list(title = "f(x, y)", range = c(-z_clip, z_clip)),
    camera = list(eye = list(x = 1.6, y = 1.6, z = 0.6))
  ))

```

Corollary 4.1 (Continuous functions on a rectangle (corollary of Theorem 4.1)). *If $f : [a, b] \times [c, d] \rightarrow \mathbb{R}$ is **continuous** on the closed bounded rectangle $[a, b] \times [c, d]$, then:*

$$\begin{aligned} \int_a^b \left(\int_c^d f(x, y) dy \right) dx &= \int_c^d \left(\int_a^b f(x, y) dx \right) dy \\ &= \iint_{[a, b] \times [c, d]} f(x, y) dx dy. \end{aligned}$$

(Larson and Edwards 2018, Theorem 14.2, p. 982)

i Proof

Proof. A closed bounded rectangle $[a, b] \times [c, d]$ is both vertically simple (with $g_1 \equiv c$, $g_2 \equiv d$) and horizontally simple (with $h_1 \equiv a$, $h_2 \equiv b$). Applying both parts of Theorem 4.1 to f on this rectangle gives the two iterated forms shown. $\square$

Exm

Example 4.3 (Evaluating a double integral on a rectangle). Structure adapted from (Larson and Edwards 2018, sec. 14.2, Example 2, pp. 982–983); the integrand $x^2 + y^2$ is original, chosen so the integral equals $E[g(X, Y)]$ for $g(x, y) = x^2 + y^2$.

Let X and Y be independent $\text{Uniform}(0, 1)$ random variables, with joint density $f(x, y) = 1$ on the unit square $R = \{(x, y) : x \in [0, 1], y \in [0, 1]\}$ (Figure 11). Define the function $g(x, y) = x^2 + y^2$, and compute its expectation $E[g(X, Y)]$.

Because the joint density equals 1 on R , this expectation is the double integral of g over R :

$$E[g(X, Y)] = \iint_R (x^2 + y^2) dA.$$

```

ggplot2::ggplot() +
  ggplot2::annotate(
    "rect", xmin = 0, xmax = 1, ymin = 0, ymax = 1,
    fill = "steelblue", alpha = 0.4, color = "black", linewidth = 0.7
  ) +
  ggplot2::annotate(
    "text", x = 0.5, y = 0.5, label = "R", size = 8, fontface = "italic"
  ) +
  ggplot2::scale_x_continuous(breaks = c(0, 1), limits = c(-0.15, 1.25)) +
  ggplot2::scale_y_continuous(breaks = c(0, 1), limits = c(-0.15, 1.25)) +
  ggplot2::coord_equal() +
  ggplot2::labs(x = "x", y = "y") +
  ggplot2::theme_minimal()

```

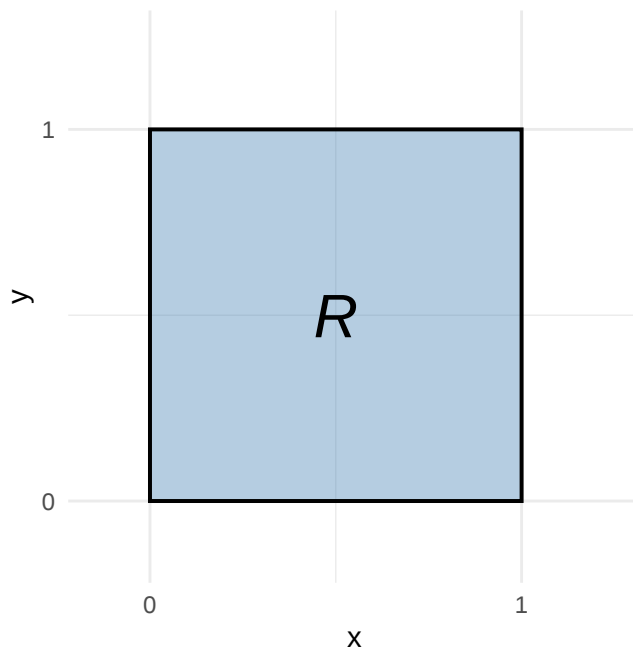


Figure 11: Integration region $R = [0, 1]^2$, the unit square.

The integrand is continuous on R , so Corollary 4.1 applies and either order of integration yields the same value.

Integrating y first, then x :

$$\begin{aligned}
 E[g(X, Y)] &= \int_0^1 \int_0^1 (x^2 + y^2) dy dx \\
 &= \int_0^1 \left[x^2 y + \frac{y^3}{3} \right]_0^1 dx \\
 &= \int_0^1 \left(x^2 + \frac{1}{3} \right) dx \\
 &= \left[\frac{x^3}{3} + \frac{x}{3} \right]_0^1 \\
 &= \frac{2}{3}
 \end{aligned}$$

Integrating x first, then y (verifying the order can be swapped):

$$\begin{aligned}
\int_0^1 \int_0^1 (x^2 + y^2) dx dy &= \int_0^1 \left[\frac{x^3}{3} + y^2 x \right]_0^1 dy \\
&= \int_0^1 \left(\frac{1}{3} + y^2 \right) dy \\
&= \left[\frac{y}{3} + \frac{y^3}{3} \right]_0^1 \\
&= \frac{2}{3}
\end{aligned}$$

Both orders give $\frac{2}{3}$, as Corollary 4.1 guarantees.

As a cross-check, linearity of expectation gives the same value: since $E[X^2] = \int_0^1 x^2 dx = \frac{1}{3}$ for $X \sim \text{Uniform}(0, 1)$ (and likewise for Y),

$$E[g(X, Y)] = E[X^2 + Y^2] = E[X^2] + E[Y^2] = \frac{1}{3} + \frac{1}{3} = \frac{2}{3}.$$

The solid whose volume equals this integral is shown in Figure 12.

```

n_grid <- 41
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) x^2 + y^2)

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 = "z", range = c(0, 2))
  ))

```

Figure 12: Surface $z = x^2 + y^2$ over the unit square $[0, 1]^2$. The double integral $\frac{2}{3}$ is the volume between this surface and the xy -plane, and equals $E[X^2 + Y^2]$.

Theorem 4.2 (Fubini–Tonelli theorem (measure-theoretic form)). Let $(\Omega_1, \mathcal{F}_1, \mu_1)$ and $(\Omega_2, \mathcal{F}_2, \mu_2)$ be **-finite** measure spaces, and let $f : \Omega_1 \times \Omega_2 \rightarrow \mathbb{R}$ be measurable with respect to the product σ -algebra $\mathcal{F}_1 \otimes \mathcal{F}_2$. If either

(a) $f \geq 0$ almost everywhere with respect to $\mu_1 \otimes \mu_2$ (**Tonelli’s theorem**), or

(b) $\int_{\Omega_1 \times \Omega_2} |f| d(\mu_1 \otimes \mu_2) < \infty$ (**Fubini’s theorem**),

then both iterated integrals exist, agree with the double integral, and equal each other:

$$\begin{aligned} \int_{\Omega_1 \times \Omega_2} f d(\mu_1 \otimes \mu_2) &= \int_{\Omega_1} \left(\int_{\Omega_2} f(\omega_1, \omega_2) d\mu_2(\omega_2) \right) d\mu_1(\omega_1) \\ &= \int_{\Omega_2} \left(\int_{\Omega_1} f(\omega_1, \omega_2) d\mu_1(\omega_1) \right) d\mu_2(\omega_2). \end{aligned}$$

(Billingsley 1995, Theorem 18.3; Gut 2013, Theorem 9.1, p. 65; Fubini 1907; Wikipedia contributors 2024)

The measure-theoretic generalization is not required for Epi 204 itself, but it is what justifies the joint-distribution form³ corollary used later in the probability chapter: probability measures are finite (hence σ -finite), so the σ -finiteness condition is automatic. The integrability conditions (nonnegativity or absolute integrability) still need to be verified in each application.

Exm

Example 4.4 (Positive application of Theorem 4.2). Let X and Y be independent Exponential(1) random variables, with joint density $f(x, y) = e^{-(x+y)}$ for $x, y \geq 0$. Their distributions are probability measures on $[0, \infty)$, and probability measures are finite (hence σ -finite), so the product measure on $[0, \infty)^2$ satisfies the σ -finiteness condition of Theorem 4.2. Since $f(x, y) = e^{-(x+y)} \geq 0$, condition (a) (Tonelli theorem, nonnegativity) is also satisfied. By Theorem 4.2, both iterated integrals exist and agree:

$$\begin{aligned} P(X \leq 1, Y \leq 1) &= \int_0^1 \int_0^1 e^{-(x+y)} dy dx \\ &= \int_0^1 e^{-x} \left(\int_0^1 e^{-y} dy \right) dx \\ &= \int_0^1 e^{-x} (1 - e^{-1}) dx \\ &= (1 - e^{-1}) [-e^{-x}]_0^1 \\ &= (1 - e^{-1})^2. \end{aligned}$$

The joint probability calculation also equals $\left(\int_0^1 e^{-x} dx \right)^2 = (1 - e^{-1})^2$, since $f(x, y) = e^{-x} \cdot e^{-y}$ factors as a product of independent densities. Both iterated integrals agree, as Theorem 4.2 guarantees when condition (a) holds.

Exm

Example 4.5 (When neither Fubini–Tonelli condition is satisfied). The same function $f(x, y) = (x^2 - y^2)/(x^2 + y^2)^2$ from Example 4.2 illustrates a case where neither condition of Theorem 4.2 is satisfied.

Why Theorem 4.2’s conditions fail: $\iint_{\mathbb{R}^2} |f| dA = \infty$, which violates condition (b). Switching to polar coordinates (r, θ) near the origin, the integrand satisfies $|f(x, y)| = |x^2 - y^2|/(x^2 + y^2)^2 = |\cos 2\theta|/r^2$, so

³[probability.qmd#cor-fubini-joint](#)

$$\begin{aligned}
\iint_R |f| dA &\geq \int_0^{\pi/2} \int_0^\epsilon \frac{|\cos 2\theta|}{r^2} r dr d\theta \\
&= \left(\int_0^{\pi/2} |\cos 2\theta| d\theta \right) \int_0^\epsilon \frac{dr}{r} \\
&= +\infty,
\end{aligned}$$

since $\int_0^\epsilon dr/r$ diverges. Therefore $\iint_R |f| dA = \infty$, and condition (b) of Theorem 4.2 is not satisfied. (Condition (a) also fails: f takes both positive and negative values, so it is not nonnegative a.e.) The unequal iterated integrals from Example 4.2 are thus consistent with Theorem 4.2: the theorem simply does not apply. (Wikipedia contributors 2024)

5 Linear Algebra

5.1 Vectors

Definition 5.1 (Column vector). A **column vector** of length p is an ordered list of p numbers, written vertically:

$$\tilde{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_p \end{bmatrix}$$

Column vectors are the default convention in these notes and in most statistics textbooks. They are also called $p \times 1$ *matrices*.

Definition 5.2 (Transpose). The **transpose** of a column vector $\tilde{x}$ is the **row vector** with the same sequence of entries, written horizontally:

$$\tilde{x}^\top \equiv \tilde{x}' \equiv [x_1, x_2, \dots, x_p]$$

The transpose operation converts a column vector to a row vector, or more generally, swaps the rows and columns of a matrix (Definition 5.10).

5.1.1 Special vectors

Definition 5.3 (Zero vector). The **zero vector** $\tilde{0}$ of length p has all entries equal to zero:

$$\tilde{0} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}$$

The zero vector is the additive identity for vector addition: $\tilde{x} + \tilde{0} = \tilde{x}$ for any vector $\tilde{x}$ of the same length.

Definition 5.4 (Ones vector). The **ones vector** $\tilde{\mathbf{1}}$ of length p has all entries equal to one:

$$\tilde{\mathbf{1}} = \begin{bmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{bmatrix}$$

The dot product $\tilde{\mathbf{1}}^\top \tilde{\mathbf{x}} = \tilde{\mathbf{1}} \cdot \tilde{\mathbf{x}} = \sum_{i=1}^p x_i$ is the sum of all entries of $\tilde{\mathbf{x}}$.

Definition 5.5 (Indicator vector / standard basis vector). The j -th **indicator vector** (or **standard basis vector**) $\tilde{\mathbf{e}}_j$ of length p has a 1 in position j and 0s elsewhere:

$$(\tilde{\mathbf{e}}_j)_i = \begin{cases} 1 & \text{if } i = j \\ 0 & \text{if } i \neq j \end{cases} \quad \tilde{\mathbf{e}}_j = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} \leftarrow \text{position } j$$

They are also called **unit vectors** or **standard basis vectors**.

Theorem 5.1 (Indicator vectors select entries). For any vector $\tilde{\mathbf{x}}$ of length p and any $j \in \{1, \dots, p\}$:

$$\tilde{\mathbf{e}}_j^\top \tilde{\mathbf{x}} = x_j$$

i Proof

Proof. Writing the product componentwise:

$$\begin{aligned} \tilde{\mathbf{e}}_j^\top \tilde{\mathbf{x}} &= \sum_{i=1}^p (\tilde{\mathbf{e}}_j)_i x_i \\ &= \sum_{i=1}^p \begin{cases} 1 \cdot x_i & \text{if } i = j \\ 0 \cdot x_i & \text{if } i \neq j \end{cases} \\ &= x_j \end{aligned}$$

□

Definition 5.6 (Dot product/linear combination/inner product). For any two real-valued vectors $\tilde{\mathbf{x}} = (x_1, \dots, x_n)$ and $\tilde{\mathbf{y}} = (y_1, \dots, y_n)$, the **dot-product**, **linear combination**, or **inner product** of $\tilde{\mathbf{x}}$ and $\tilde{\mathbf{y}}$ is:

$$\tilde{\mathbf{x}} \cdot \tilde{\mathbf{y}} = \tilde{\mathbf{x}}^\top \tilde{\mathbf{y}} \stackrel{\text{def}}{=} \sum_{i=1}^n x_i y_i$$

i Note

See also the definitions in

- Dobson and Barnett (2018), §1.3 (equation 1.1, page 7)
- Kaplan (2022), here^a.
- wikipedia^b

“Linear combination” can also refer to weighted sums of vectors, or in other words matrix-vector multiplication.

The dot-product has a different generalization for two matrices; see wikipedia^c for more.

^a<https://www.mosaic-web.org/MOSAIC-Calculus/Textbook/Linear-combinations/28-Vectors.html#geometry-arithmetic>

^bhttps://en.wikipedia.org/wiki/Linear_combination

^chttps://en.wikipedia.org/wiki/Dot_product#Dyadics_and_matrices

Theorem 5.2 (Dot product is symmetric). *The dot product is symmetric:*

$$\tilde{x} \cdot \tilde{y} = \tilde{y} \cdot \tilde{x}$$

i Proof

Proof. Apply:

- Definition 5.6
- symmetry of scalar multiplication
- Definition 5.6 again

□

Exm

Example 5.1 (Dot product as matrix multiplication). The dot product of two column vectors $\tilde{x}$ and $\tilde{\beta}$ can be written as a matrix product of the row vector $\tilde{x}^\top$ with the column vector $\tilde{\beta}$:

$$\begin{aligned}\tilde{x} \cdot \tilde{\beta} &= \tilde{x}^\top \tilde{\beta} \\ &= [x_1, x_2, \dots, x_p] \begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_p \end{bmatrix} \\ &= x_1\beta_1 + x_2\beta_2 + \dots + x_p\beta_p\end{aligned}$$

5.1.2 Orthogonality

Definition 5.7 (Orthogonal vectors). Two vectors $\tilde{x}$ and $\tilde{y}$ of the same length are **orthogonal** (written $\tilde{x} \perp \tilde{y}$) if their dot product is zero:

$$\tilde{x} \perp \tilde{y} \iff \tilde{x}^\top \tilde{y} = 0$$

Orthogonality generalizes the geometric notion of perpendicularity to arbitrary dimensions.

Definition 5.8 (Orthonormal vectors). A set of vectors $\{\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_k\}$ is **orthonormal** if the vectors are mutually orthogonal and each has unit length:

$$\tilde{x}_i^\top \tilde{x}_j = \begin{cases} 1 & \text{if } i = j \\ 0 & \text{if } i \neq j \end{cases}$$

The indicator vectors $\tilde{e}_1, \tilde{e}_2, \dots, \tilde{e}_p$ (Definition 5.5) form an orthonormal set.

5.2 Matrices

Definition 5.9 (Matrix). A **matrix** of dimensions $m \times n$ is a rectangular array of $m \cdot n$ numbers, arranged in m rows and n columns:

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{bmatrix}$$

The entry in row i and column j is denoted a_{ij} or $(\mathbf{A})_{ij}$. A column vector of length p is a special case: a $p \times 1$ matrix. A row vector of length p is a $1 \times p$ matrix.

5.2.1 Matrix transpose

Definition 5.10 (Matrix transpose). The **transpose** of an $m \times n$ matrix $\mathbf{A}$ is the $n \times m$ matrix $\mathbf{A}^\top$ obtained by swapping the rows and columns of $\mathbf{A}$:

$$(\mathbf{A}^\top)_{ij} = a_{ji}$$

Theorem 5.3 (Transpose of a sum).

$$(\mathbf{A} + \mathbf{B})^\top = \mathbf{A}^\top + \mathbf{B}^\top$$

In particular, for column vectors $\tilde{x}$ and $\tilde{y}$:

$$(\tilde{x} + \tilde{y})^\top = \tilde{x}^\top + \tilde{y}^\top$$

Theorem 5.4 (Transpose of a product). *For compatible matrices $\mathbf{A}$ and $\mathbf{B}$:*

$$(\mathbf{AB})^\top = \mathbf{B}^\top \mathbf{A}^\top$$

The order of the factors reverses when transposing a product.

5.2.2 Matrix addition

Definition 5.11 (Zero matrix). The $m \times n$ **zero matrix** $\mathbf{0}_{m \times n}$ (or $\mathbf{0}$ when dimensions are clear from context) has all entries equal to zero:

$$\mathbf{0}_{m \times n} = \begin{bmatrix} 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 0 \end{bmatrix}$$

Definition 5.12 (Matrix addition). Two matrices $\mathbf{A}$ and $\mathbf{B}$ of the same dimensions $m \times n$ can be added element-wise:

$$(\mathbf{A} + \mathbf{B})_{ij} = a_{ij} + b_{ij}$$

Theorem 5.5 (Matrix addition is commutative).

$$\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A}$$

Theorem 5.6 (Matrix addition is associative).

$$(\mathbf{A} + \mathbf{B}) + \mathbf{C} = \mathbf{A} + (\mathbf{B} + \mathbf{C})$$

Theorem 5.7 (Zero matrix is the additive identity).

$$\mathbf{A} + \mathbf{0} = \mathbf{A}$$

Theorem 5.8 (Additive inverse). For any matrix $\mathbf{A}$, the matrix $-\mathbf{A}$ (defined by $(-\mathbf{A})_{ij} = -a_{ij}$) satisfies:

$$\mathbf{A} + (-\mathbf{A}) = \mathbf{0}$$

5.2.3 Scalar multiplication

Definition 5.13 (Scalar multiplication). A matrix $\mathbf{A}$ can be multiplied by a scalar c :

$$(c\mathbf{A})_{ij} = c \cdot a_{ij}$$

5.2.4 Matrix multiplication

Definition 5.14 (Matrix multiplication). The **product** of an $m \times k$ matrix $\mathbf{A}$ and a $k \times n$ matrix $\mathbf{B}$ is the $m \times n$ matrix $\mathbf{C} = \mathbf{AB}$ with entries:

$$c_{ij} = \sum_{s=1}^k a_{is} b_{sj}$$

Matrix multiplication is only defined when the number of columns in $\mathbf{A}$ equals the number of rows in $\mathbf{B}$.

Matrix multiplication is **not** commutative in general: $\mathbf{AB} \neq \mathbf{BA}$.

Theorem 5.9 (Matrix multiplication is associative).

$$(\mathbf{AB})\mathbf{C} = \mathbf{A}(\mathbf{BC})$$

Theorem 5.10 (Matrix multiplication is distributive over addition).

$$\mathbf{A}(\mathbf{B} + \mathbf{C}) = \mathbf{AB} + \mathbf{AC}$$

$$(\mathbf{A} + \mathbf{B})\mathbf{C} = \mathbf{AC} + \mathbf{BC}$$

5.2.5 Matrix-vector multiplication

Definition 5.15 (Matrix-vector multiplication). The product of an $m \times p$ matrix $\mathbf{A}$ and a $p \times 1$ column vector $\tilde{x}$ is the $m \times 1$ column vector $\mathbf{A}\tilde{x}$ with entries:

$$(\mathbf{A}\tilde{x})_i = \sum_{j=1}^p a_{ij} x_j$$

Matrix-vector multiplication is a generalization of the dot product. Each entry of the result is a dot product of a row of $\mathbf{A}$ with the vector $\tilde{x}$.

5.3 Special Matrices

See also Definition 5.11 for the zero matrix.

Definition 5.16 (Square matrix). A matrix is **square** if it has the same number of rows as columns. The number of rows (= columns) is the **order** of the matrix.

Definition 5.17 (Matrix power). For a square matrix $\mathbf{A}$ of order p and a positive integer k , the k -th **power** of $\mathbf{A}$ is:

$$\mathbf{A}^k = \underbrace{\mathbf{A} \mathbf{A} \cdots \mathbf{A}}_{k \text{ copies}}$$

In particular, $\mathbf{A}^2 = \mathbf{A}\mathbf{A}$.

Definition 5.18 (Identity matrix). The $p \times p$ **identity matrix** $\mathbf{I}_p$ (or $\mathbf{I}$ when the size is clear from context) has ones on the main diagonal and zeros elsewhere:

$$(\mathbf{I}_p)_{ij} = \begin{cases} 1 & \text{if } i = j \\ 0 & \text{if } i \neq j \end{cases} \quad \mathbf{I}_p = \begin{bmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \end{bmatrix}$$

Theorem 5.11 (Identity matrix is a multiplicative identity). For any $m \times p$ matrix $\mathbf{A}$:

$$\mathbf{A} \mathbf{I}_p = \mathbf{A}$$

$$\mathbf{I}_m \mathbf{A} = \mathbf{A}$$

Definition 5.19 (Symmetric matrix). A square matrix $\mathbf{A}$ is **symmetric** if $\mathbf{A}^\top = \mathbf{A}$, i.e., $a_{ij} = a_{ji}$ for all i and j .

Covariance matrices and information matrices are symmetric.

Definition 5.20 (Diagonal matrix). A square matrix $\mathbf{D}$ is a **diagonal matrix** if all off-diagonal entries are zero: $d_{ij} = 0$ whenever $i \neq j$:

$$\mathbf{D} = \begin{bmatrix} d_1 & 0 & \cdots & 0 \\ 0 & d_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & d_p \end{bmatrix}$$

Diagonal matrices are denoted $\mathbf{D} = \text{diag}(d_1, d_2, \dots, d_p)$, where $d_1, \dots, d_p$ are the diagonal entries.

Definition 5.21 (Matrix inverse). For a square $p \times p$ matrix $\mathbf{A}$, the **inverse** $\mathbf{A}^{-1}$ (if it exists) is the unique matrix satisfying:

$$\mathbf{A} \mathbf{A}^{-1} = \mathbf{A}^{-1} \mathbf{A} = \mathbf{I}_p$$

A matrix that has an inverse is called **invertible** or **non-singular**.

Theorem 5.12 (Inverse of a product). For invertible matrices $\mathbf{A}$ and $\mathbf{B}$:

$$(\mathbf{AB})^{-1} = \mathbf{B}^{-1} \mathbf{A}^{-1}$$

Definition 5.22 (Idempotent matrix). A square matrix $\mathbf{A}$ is **idempotent** if

$$\mathbf{A}^2 = \mathbf{A}$$

Definition 5.23 (Projection matrix). A square matrix $\mathbf{P}$ is a **projection matrix** (also called an **orthogonal projector**) if it is both symmetric and idempotent:

$$\mathbf{P}^\top = \mathbf{P} \quad \text{and} \quad \mathbf{P}^2 = \mathbf{P}$$

Theorem 5.13 (Complement of a projection matrix). *If $\mathbf{P}$ is a projection matrix, then $\mathbf{I} - \mathbf{P}$ is also a projection matrix.*

i Proof

Proof. We verify symmetry and idempotency.

Symmetry:

$$(\mathbf{I} - \mathbf{P})^\top = \mathbf{I}^\top - \mathbf{P}^\top = \mathbf{I} - \mathbf{P}$$

Idempotency:

$$\begin{aligned} (\mathbf{I} - \mathbf{P})^2 &= (\mathbf{I} - \mathbf{P})(\mathbf{I} - \mathbf{P}) \\ &= \mathbf{I} - \mathbf{P} - \mathbf{P} + \mathbf{P}^2 \\ &= \mathbf{I} - \mathbf{P} - \mathbf{P} + \mathbf{P} \\ &= \mathbf{I} - \mathbf{P} \end{aligned}$$

□

Theorem 5.14 (Hat matrix is a projection matrix). *In a linear regression model with full-rank design matrix $\mathbf{X}$, the **hat matrix***

$$\mathbf{H} = \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$$

is a projection matrix.

i Proof

Proof. We verify symmetry and idempotency.

Symmetry:

$$\begin{aligned} \mathbf{H}^\top &= (\mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top)^\top \\ &= (\mathbf{X}^\top)^\top \cdot ((\mathbf{X}^\top \mathbf{X})^{-1})^\top \cdot \mathbf{X}^\top \\ &= \mathbf{X} \cdot (\mathbf{X}^\top \mathbf{X})^{-1} \cdot \mathbf{X}^\top \\ &= \mathbf{H} \end{aligned}$$

where the third line uses $(\mathbf{X}^\top)^\top = \mathbf{X}$ and the fact that $\mathbf{X}^\top \mathbf{X}$ is symmetric, so its inverse is also symmetric $((\mathbf{X}^\top \mathbf{X})^{-1})^\top = (\mathbf{X}^\top \mathbf{X})^{-1}$.

Idempotency:

$$\begin{aligned}\mathbf{H}^2 &= \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \cdot \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \\ &= \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} (\mathbf{X}^\top \mathbf{X}) (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \\ &= \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \\ &= \mathbf{H}\end{aligned}$$

□

The hat matrix appears in the formula for fitted values in linear regression: $\hat{\mathbf{y}} = \mathbf{X}\hat{\boldsymbol{\beta}} = \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \tilde{\mathbf{y}} = \mathbf{H}\tilde{\mathbf{y}}$. It “puts a hat” on $\tilde{\mathbf{y}}$ — hence the name.

Theorem 5.15 (Projection matrices produce orthogonal decompositions). *If $\mathbf{P}$ is a projection matrix and $\tilde{\mathbf{v}}$ is any vector of compatible dimension, then the two components of the decomposition*

$$\tilde{\mathbf{v}} = \underbrace{\mathbf{P}\tilde{\mathbf{v}}}_{\text{projected}} + \underbrace{(\mathbf{I} - \mathbf{P})\tilde{\mathbf{v}}}_{\text{residual}}$$

are orthogonal:

$$\mathbf{P}\tilde{\mathbf{v}} \perp (\mathbf{I} - \mathbf{P})\tilde{\mathbf{v}}$$

i Proof

Proof.

$$\begin{aligned}(\mathbf{P}\tilde{\mathbf{v}})^\top (\mathbf{I} - \mathbf{P})\tilde{\mathbf{v}} &= \tilde{\mathbf{v}}^\top \mathbf{P}^\top (\mathbf{I} - \mathbf{P})\tilde{\mathbf{v}} \\ &= \tilde{\mathbf{v}}^\top \mathbf{P} (\mathbf{I} - \mathbf{P})\tilde{\mathbf{v}} \\ &= \tilde{\mathbf{v}}^\top (\mathbf{P} - \mathbf{P}^2)\tilde{\mathbf{v}} \\ &= \tilde{\mathbf{v}}^\top (\mathbf{P} - \mathbf{P})\tilde{\mathbf{v}} \\ &= \tilde{\mathbf{v}}^\top \mathbf{0} \tilde{\mathbf{v}} \\ &= 0\end{aligned}$$

where the second line uses symmetry ($\mathbf{P}^\top = \mathbf{P}$) and the fourth line uses idempotency ($\mathbf{P}^2 = \mathbf{P}$). □

5.4 Quadratic Forms

Definition 5.24 (Quadratic form). A **quadratic form** is a mathematical expression of the structure

$$\tilde{\mathbf{x}}^\top \mathbf{S} \tilde{\mathbf{x}}$$

where $\tilde{\mathbf{x}}$ is a $p \times 1$ vector and $\mathbf{S}$ is a $p \times p$ matrix.

Quadratic forms are the matrix generalizations of the scalar expression cx^2 . They occur frequently in statistics:

- The residual sum of squares in linear regression (Section 6) is a quadratic form.
- The variance of a linear combination of estimates (see Inference about Gaussian Linear Regression Models⁴) is a quadratic form: $\text{Var}\left(\tilde{\mathbf{x}}^\top \hat{\boldsymbol{\beta}}\right) = \tilde{\mathbf{x}}^\top \text{Var}\left(\hat{\boldsymbol{\beta}}\right) \tilde{\mathbf{x}}$.

⁴[Linear-models-overview.qmd#sec-infer-LMs](#)

Theorem 5.16 (Symmetric part of a quadratic form). *If $\mathbf{S}$ is a square matrix, then*

$$\tilde{x}^\top \mathbf{S} \tilde{x} = \tilde{x}^\top \left(\frac{1}{2} (\mathbf{S} + \mathbf{S}^\top) \right) \tilde{x}.$$

So the value of a quadratic form depends only on the symmetric part of $\mathbf{S}$.

5.5 Design Matrix

Definition 5.25 (Design matrix). In a regression model with n observations and p predictors, the **design matrix** (or **model matrix**) $\mathbf{X}$ is the $n \times p$ matrix whose i -th row is the covariate vector $\tilde{x}_i^\top$ for observation i :

$$\mathbf{X} = \begin{bmatrix} \tilde{x}_1^\top \\ \tilde{x}_2^\top \\ \vdots \\ \tilde{x}_n^\top \end{bmatrix} = \begin{bmatrix} x_{11} & x_{12} & \cdots & x_{1p} \\ x_{21} & x_{22} & \cdots & x_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n1} & x_{n2} & \cdots & x_{np} \end{bmatrix}$$

The product $\mathbf{X}\tilde{\beta}$ collects all the linear predictors $\tilde{x}_i^\top \tilde{\beta}$ into a single $n \times 1$ vector:

$$\mathbf{X}\tilde{\beta} = \begin{bmatrix} \tilde{x}_1^\top \tilde{\beta} \\ \vdots \\ \tilde{x}_n^\top \tilde{\beta} \end{bmatrix}$$

The matrix $\mathbf{X}^\top \mathbf{X}$ is a $p \times p$ symmetric matrix that appears in the OLS estimator $\hat{\beta} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \tilde{y}$.

6 Vector Calculus

(adapted from Fieller (2016), §7.2⁵)

This section covers derivatives of functions of vectors and matrices. Linear algebra prerequisites — including vectors, matrices, transpose, dot product, and quadratic forms — are covered in Section 5.

Let $\tilde{x}$ and $\tilde{\beta}$ be column vectors of length p (see Definition 5.1 and Definition 5.6).

Definition 6.1 (Vector derivative). If $f(\tilde{\beta})$ is a function that takes a vector $\tilde{\beta}$ as input, such as $f(\tilde{\beta}) = x' \tilde{\beta}$, then:

$$\frac{\partial}{\partial \tilde{\beta}} f(\tilde{\beta}) = \begin{bmatrix} \frac{\partial}{\partial \beta_1} f(\tilde{\beta}) \\ \frac{\partial}{\partial \beta_2} f(\tilde{\beta}) \\ \vdots \\ \frac{\partial}{\partial \beta_p} f(\tilde{\beta}) \end{bmatrix}$$

Definition 6.2 (Row-vector derivative). If $f(\tilde{\beta})$ is a function that takes a vector $\tilde{\beta}$ as input, such as $f(\tilde{\beta}) = x' \tilde{\beta}$, then:

⁵<https://www.taylorfrancis.com/chapters/mono/10.1201/9781315370200-7/vector-matrix-calculus-nick-fieller?context=ubx&refId=c310b723-786a-4f33-ae56-720a6cccd3a1>

$$\frac{\partial}{\partial \tilde{\beta}^\top} f(\tilde{\beta}) = \begin{bmatrix} \frac{\partial}{\partial \beta_1} f(\tilde{\beta}) & \frac{\partial}{\partial \beta_2} f(\tilde{\beta}) & \cdots & \frac{\partial}{\partial \beta_p} f(\tilde{\beta}) \end{bmatrix}$$

Theorem 6.1 (Row and column derivatives are transposes).

$$\frac{\partial}{\partial \tilde{\beta}^\top} f(\tilde{\beta}) = \left(\frac{\partial}{\partial \tilde{\beta}} f(\tilde{\beta}) \right)^\top$$

$$\frac{\partial}{\partial \tilde{\beta}} f(\tilde{\beta}) = \left(\frac{\partial}{\partial \tilde{\beta}^\top} f(\tilde{\beta}) \right)^\top$$

Definition 6.3 (Constant). $\tilde{x}$ is constant with respect to $\tilde{\beta}$ if

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}} \tilde{x}^\top}_{p \times p} = \mathbf{0}_{p \times p}$$

Exm

Example 6.1 (A constant vector). Let $\tilde{\beta} = (\beta_1, \beta_2)^\top$ and $\tilde{x} = (3, 5)^\top$, so $x_1 = 3$ and $x_2 = 5$ do not depend on $\tilde{\beta}$. Expanding $\frac{\partial}{\partial \tilde{\beta}} \tilde{x}^\top$ into its matrix of scalar partial derivatives (Definition 6.1, applied to each component of the row $\tilde{x}^\top$) and evaluating each entry:

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}} \tilde{x}^\top}_{2 \times 2} = \frac{\partial}{\partial \tilde{\beta}} \begin{bmatrix} x_1 & x_2 \end{bmatrix} = \begin{bmatrix} \frac{\partial}{\partial \beta_1} x_1 & \frac{\partial}{\partial \beta_1} x_2 \\ \frac{\partial}{\partial \beta_2} x_1 & \frac{\partial}{\partial \beta_2} x_2 \end{bmatrix} = \begin{bmatrix} \frac{\partial}{\partial \beta_1} 3 & \frac{\partial}{\partial \beta_1} 5 \\ \frac{\partial}{\partial \beta_2} 3 & \frac{\partial}{\partial \beta_2} 5 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} = \mathbf{0}_{2 \times 2}$$

Every entry is the derivative of a constant, so $\frac{\partial}{\partial \tilde{\beta}} \tilde{x}^\top = \mathbf{0}_{2 \times 2}$ and $\tilde{x}$ is constant with respect to $\tilde{\beta}$ (Definition 6.3).

Theorem 6.2 (Derivative of a dot product). *If $\tilde{x}$ is constant with respect to $\tilde{\beta}$, then:*

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}} (\tilde{x} \cdot \tilde{\beta})}_{p \times 1} = \underbrace{\frac{\partial}{\partial \tilde{\beta}} (\tilde{\beta} \cdot \tilde{x})}_{p \times 1} = \tilde{x}_{p \times 1}$$

i Proof

Proof.

$$\begin{aligned}\frac{\partial}{\partial \tilde{\beta}}(\tilde{x} \cdot \tilde{\beta}) &= \begin{bmatrix} \frac{\partial}{\partial \beta_1}(x_1\beta_1 + x_2\beta_2 + \dots + x_p\beta_p) \\ \frac{\partial}{\partial \beta_2}(x_1\beta_1 + x_2\beta_2 + \dots + x_p\beta_p) \\ \vdots \\ \frac{\partial}{\partial \beta_p}(x_1\beta_1 + x_2\beta_2 + \dots + x_p\beta_p) \end{bmatrix} \\ &= \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_p \end{bmatrix} \\ &= \tilde{x}\end{aligned}$$

□

Exm

Example 6.2 (Derivative of a dot product). Let $\tilde{x} = (3, 5)^\top$ (constant with respect to $\tilde{\beta}$; see Example 6.1) and $\tilde{\beta} = (\beta_1, \beta_2)^\top$. Then $\tilde{x} \cdot \tilde{\beta} = 3\beta_1 + 5\beta_2$, and by Theorem 6.2:

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(\tilde{x} \cdot \tilde{\beta})}_{2 \times 1} = \underbrace{\tilde{x}}_{2 \times 1} = \begin{pmatrix} 3 \\ 5 \end{pmatrix}$$

Verifying entry-wise:

$$\frac{\partial}{\partial \tilde{\beta}}(3\beta_1 + 5\beta_2) = \begin{pmatrix} \frac{\partial}{\partial \beta_1}(3\beta_1 + 5\beta_2) \\ \frac{\partial}{\partial \beta_2}(3\beta_1 + 5\beta_2) \end{pmatrix} = \begin{pmatrix} 3 \\ 5 \end{pmatrix}$$

Both methods agree.

Theorem 6.3 (Product rule for dot-products). *If $\mathbf{a} = \mathbf{a}(\tilde{x})$ and $\mathbf{b} = \mathbf{b}(\tilde{x})$ are differentiable $p \times 1$ vector functions of $\tilde{x}$, then:*

$$\frac{\partial}{\partial \tilde{x}} \underbrace{\tilde{x}}_{p \times 1} \cdot \underbrace{\mathbf{a}}_{p \times 1} \cdot \underbrace{\mathbf{b}}_{p \times 1} = \left(\frac{\partial}{\partial \tilde{x}} \underbrace{\tilde{x}}_{p \times 1} \underbrace{\mathbf{a}^\top}_{1 \times p} \right) \underbrace{\mathbf{b}}_{p \times 1} + \left(\frac{\partial}{\partial \tilde{x}} \underbrace{\tilde{x}}_{p \times 1} \underbrace{\mathbf{b}^\top}_{1 \times p} \right) \underbrace{\mathbf{a}}_{p \times 1}$$

i Proof

Proof. Entry-wise, for $i = 1, \dots, p$:

$$\begin{aligned}\left[\frac{\partial}{\partial \tilde{x}}(\mathbf{a} \cdot \mathbf{b}) \right]_i &= \frac{\partial}{\partial x_i} \sum_{k=1}^p a_k b_k \\ &= \sum_{k=1}^p \left(b_k \frac{\partial}{\partial x_i} a_k + a_k \frac{\partial}{\partial x_i} b_k \right) \\ &= \left[\left(\frac{\partial}{\partial \tilde{x}} \mathbf{a}^\top \right) \mathbf{b} \right]_i + \left[\left(\frac{\partial}{\partial \tilde{x}} \mathbf{b}^\top \right) \mathbf{a} \right]_i\end{aligned}$$

□

Exm

Example 6.3 (Example of the dot-product rule). Let $\tilde{x} = (\beta_1, \beta_2)^\top$, $\mathbf{a}(\tilde{x}) = (\beta_1, \beta_1\beta_2)^\top$, and $\mathbf{b}(\tilde{x}) = (\beta_2, \beta_1)^\top$. Then:

$$\mathbf{a} \cdot \mathbf{b} = \beta_1 \cdot \beta_2 + \beta_1\beta_2 \cdot \beta_1 = \beta_1\beta_2 + \beta_1^2\beta_2$$

By direct calculation:

$$\underbrace{\frac{\partial}{\partial \tilde{x}}(\mathbf{a} \cdot \mathbf{b})}_{2 \times 1} = \frac{\partial}{\partial \tilde{x}}(\beta_1\beta_2 + \beta_1^2\beta_2) = \begin{pmatrix} \beta_2 + 2\beta_1\beta_2 \\ \beta_1 + \beta_1^2 \end{pmatrix}$$

By the product rule (Theorem 6.3), using $\underbrace{\frac{\partial}{\partial \tilde{x}}\mathbf{a}^\top}_{2 \times 2} = \begin{pmatrix} 1 & \beta_2 \\ 0 & \beta_1 \end{pmatrix}$ and $\underbrace{\frac{\partial}{\partial \tilde{x}}\mathbf{b}^\top}_{2 \times 2} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$:

$$\begin{aligned} \underbrace{\frac{\partial}{\partial \tilde{x}}(\mathbf{a} \cdot \mathbf{b})}_{2 \times 1} &= \underbrace{\begin{pmatrix} 1 & \beta_2 \\ 0 & \beta_1 \end{pmatrix}}_{2 \times 2} \underbrace{\begin{pmatrix} \beta_2 \\ \beta_1 \end{pmatrix}}_{2 \times 1} + \underbrace{\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}}_{2 \times 2} \underbrace{\begin{pmatrix} \beta_1 \\ \beta_1\beta_2 \end{pmatrix}}_{2 \times 1} \\ &= \begin{pmatrix} \beta_2 + \beta_1\beta_2 \\ \beta_1^2 \end{pmatrix} + \begin{pmatrix} \beta_1\beta_2 \\ \beta_1 \end{pmatrix} \\ &= \begin{pmatrix} \beta_2 + 2\beta_1\beta_2 \\ \beta_1^2 + \beta_1 \end{pmatrix} \end{aligned}$$

Both methods agree.

Theorem 6.4 (Derivative of a linear map). *If A is an $m \times p$ matrix that is constant with respect to $\tilde{\beta}$, then:*

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(A\tilde{\beta})}_{p \times m} = \underbrace{A^\top}_{p \times m}$$

i Proof

Proof. For entry (i, j) , where row i indexes the denominator $\tilde{\beta}$ (see Definition 6.1) and column j indexes the numerator $A\tilde{\beta}$:

$$\begin{aligned} \left[\frac{\partial}{\partial \tilde{\beta}}(A\tilde{\beta}) \right]_{ij} &= \frac{\partial}{\partial \beta_i}(A\tilde{\beta})_j \\ &= \frac{\partial}{\partial \beta_i} \sum_{k=1}^p a_{jk}\beta_k \\ &= a_{ji} \\ &= [A^\top]_{ij} \end{aligned}$$

□

Exm

Example 6.4 (Derivative of a linear map). Let $A = \begin{pmatrix} 2 & 3 \end{pmatrix}$ (1×2) and $\tilde{\beta} = (\beta_1, \beta_2)^\top$. Then $A\tilde{\beta} = 2\beta_1 + 3\beta_2$, and by Theorem 6.4:

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(A\tilde{\beta})}_{2 \times 1} = \underbrace{A^\top}_{2 \times 1} = \begin{pmatrix} 2 \\ 3 \end{pmatrix}$$

Theorem 6.5 (Vector-derivative of a matrix-vector product). *If A is an $m \times q$ matrix that is constant with respect to $\tilde{\beta}$, and $\tilde{v} = \tilde{v}(\tilde{\beta})$ is a $q \times 1$ vector that depends on the $p \times 1$ vector $\tilde{\beta}$, then:*

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(A\tilde{v})}_{p \times m} = \underbrace{\left(\frac{\partial}{\partial \tilde{\beta}} \tilde{v} \right)}_{p \times q} \underbrace{A^\top}_{q \times m}$$

This result generalizes Theorem 6.4, which is the special case $\tilde{v} = \tilde{\beta}$ (so that $\frac{\partial}{\partial \tilde{\beta}} \tilde{\beta} = \mathbf{I}$ and $\frac{\partial}{\partial \tilde{\beta}}(A\tilde{\beta}) = A^\top$).

i Proof

Proof. For entry (i, j) , where row i indexes the denominator $\tilde{\beta}$ and column j indexes the numerator $A\tilde{v}$:

$$\begin{aligned} \left[\frac{\partial}{\partial \tilde{\beta}}(A\tilde{v}) \right]_{ij} &= \frac{\partial}{\partial \beta_i}(A\tilde{v})_j \\ &= \frac{\partial}{\partial \beta_i} \sum_{k=1}^q a_{jk} v_k \\ &= \sum_{k=1}^q a_{jk} \frac{\partial}{\partial \beta_i} v_k \\ &= \sum_{k=1}^q \left[\frac{\partial}{\partial \tilde{\beta}} \tilde{v} \right]_{ik} [A^\top]_{kj} \\ &= \left[\left(\frac{\partial}{\partial \tilde{\beta}} \tilde{v} \right) A^\top \right]_{ij} \end{aligned}$$

□

Exm

Example 6.5 (Vector-derivative of a matrix-vector product). Let $A = \begin{pmatrix} 2 & 3 \end{pmatrix}$ (1×2 , constant) and $\tilde{v}(\tilde{\beta}) = (\beta_1^2, \beta_2^2)^\top$. Then $A\tilde{v} = 2\beta_1^2 + 3\beta_2^2$. By Theorem 6.5:

$$\begin{aligned} \underbrace{\frac{\partial}{\partial \tilde{\beta}}(A\tilde{v})}_{2 \times 1} &= \begin{pmatrix} 2\beta_1 & 0 \\ 0 & 2\beta_2 \end{pmatrix} \begin{pmatrix} 2 \\ 3 \end{pmatrix} \\ &= \begin{pmatrix} 4\beta_1 \\ 6\beta_2 \end{pmatrix} \end{aligned}$$

Corollary 6.1 (Derivative of a dot product, transpose-product form). *If $\tilde{x}$ is constant with respect to $\tilde{\beta}$, then:*

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(\tilde{x}^\top \tilde{\beta})}_{p \times 1} = \tilde{x}$$

This vector derivative formula looks a lot like non-vector calculus, except that you have to transpose the coefficient: in scalar calculus $\frac{\partial}{\partial x}(cx) = c$, but here the coefficient $\tilde{x}^\top$ (a row vector) becomes $\tilde{x}$ (a column vector) in the result.

i Proof

Proof. Using Theorem 6.2:

Since $\tilde{x}^\top \tilde{\beta} = \tilde{x} \cdot \tilde{\beta}$ (Definition 5.6), and $\tilde{x}$ is constant with respect to $\tilde{\beta}$:

$$\frac{\partial}{\partial \tilde{\beta}}(\tilde{x}^\top \tilde{\beta}) = \frac{\partial}{\partial \tilde{\beta}}(\tilde{x} \cdot \tilde{\beta}) = \tilde{x}$$

by Theorem 6.2. □

i Proof

Proof. Using Theorem 6.5:

Since $\tilde{x}$ is constant with respect to $\tilde{\beta}$, $A = \tilde{x}^\top$ is a constant $1 \times p$ matrix. Applying Theorem 6.5 with $\tilde{v} = \tilde{\beta}$ (so $\frac{\partial}{\partial \tilde{\beta}} \tilde{\beta} = \mathbf{I}$):

$$\begin{aligned} \frac{\partial}{\partial \tilde{\beta}}(\tilde{x}^\top \tilde{\beta}) &= \left(\frac{\partial}{\partial \tilde{\beta}} \tilde{\beta} \right) (\tilde{x}^\top)^\top \\ &= \mathbf{I} \cdot \tilde{x} \\ &= \tilde{x} \end{aligned}$$

□

Exm

Example 6.6 (Derivative of a transpose product). Let $\tilde{x} = (3, 5)^\top$ and $\tilde{\beta} = (\beta_1, \beta_2)^\top$. Then $\tilde{x}^\top \tilde{\beta} = 3\beta_1 + 5\beta_2$, and by Corollary 6.1:

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}} \begin{pmatrix} \tilde{x}^\top & \tilde{\beta} \\ 1 \times 2 & 2 \times 1 \end{pmatrix}}_{2 \times 1} = \underbrace{\tilde{x}}_{2 \times 1} = \begin{pmatrix} 3 \\ 5 \end{pmatrix}$$

Theorem 6.6 (Derivative of a quadratic form). *For a quadratic form (Definition 5.24), if S is a symmetric $p \times p$ matrix that is constant with respect to $\tilde{\beta}$, then:*

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top S \tilde{\beta})}_{p \times 1} = \underbrace{2S \tilde{\beta}}_{p \times 1}$$

i Proof

Proof. Expanding entry-wise, $\tilde{\beta}^\top S \tilde{\beta} = \sum_{j=1}^p \sum_{k=1}^p s_{jk} \beta_j \beta_k$. Differentiating component-wise with respect to β_i for $i = 1, \dots, p$:

$$\begin{aligned} \left[\frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top S \tilde{\beta}) \right]_i &= \frac{\partial}{\partial \beta_i} \sum_{j=1}^p \sum_{k=1}^p s_{jk} \beta_j \beta_k \quad (\text{expand quadratic form}) \\ &= \sum_{k=1}^p s_{ik} \beta_k + \sum_{j=1}^p s_{ji} \beta_j \quad (\text{product rule for } \beta_i \beta_k) \\ &= [S \tilde{\beta}]_i + [S^\top \tilde{\beta}]_i \quad (\text{matrix-vector multiplication definition}) \\ &= [(S + S^\top) \tilde{\beta}]_i \quad (\text{linearity of matrix multiplication}) \end{aligned}$$

When S is symmetric ($S = S^\top$), $S + S^\top = 2S$, so $\frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top S \tilde{\beta}) = 2S\tilde{\beta}$. □

This operation is like taking the derivative of cx^2 with respect to x in non-vector calculus.

Exm

Example 6.7 (Derivative of a quadratic form). Let $S = \begin{pmatrix} 3 & 1 \\ 1 & 2 \end{pmatrix}$ (2×2 , symmetric and constant) and $\tilde{\beta} = (\beta_1, \beta_2)^\top$. Then $\tilde{\beta}^\top S \tilde{\beta} = 3\beta_1^2 + 2\beta_1\beta_2 + 2\beta_2^2$. By Theorem 6.6:

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top S \tilde{\beta})}_{2 \times 1} = 2S\tilde{\beta} = 2 \begin{pmatrix} 3 & 1 \\ 1 & 2 \end{pmatrix} \begin{pmatrix} \beta_1 \\ \beta_2 \end{pmatrix} = \begin{pmatrix} 6\beta_1 + 2\beta_2 \\ 2\beta_1 + 4\beta_2 \end{pmatrix}$$

Differentiating component-wise directly:

$$\begin{pmatrix} \frac{\partial}{\partial \beta_1}(3\beta_1^2 + 2\beta_1\beta_2 + 2\beta_2^2) \\ \frac{\partial}{\partial \beta_2}(3\beta_1^2 + 2\beta_1\beta_2 + 2\beta_2^2) \end{pmatrix} = \begin{pmatrix} 6\beta_1 + 2\beta_2 \\ 2\beta_1 + 4\beta_2 \end{pmatrix}$$

Both methods agree.

Corollary 6.2 (Derivative of a simple quadratic form).

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top \tilde{\beta})}_{p \times 1} = \underbrace{2\tilde{\beta}}_{p \times 1}$$

i Proof

Proof. Applying Theorem 6.6 with $S = \mathbf{I}_{p \times p}$ (which is symmetric and constant with respect to $\tilde{\beta}$):

$$\begin{aligned} \frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top \tilde{\beta}) &= \frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top \mathbf{I}_{p \times p} \tilde{\beta}) \quad (\text{rewrite with identity matrix}) \\ &= 2\mathbf{I}_{p \times p} \tilde{\beta} \quad (\text{apply @thm-quadratic-form with } S = \mathbf{I}_{p \times p}) \\ &= 2\tilde{\beta} \quad (\text{identity matrix property}) \end{aligned}$$

□

This vector derivative is like taking the derivative of x^2 .

Exm

Example 6.8 (Derivative of a sum of squares). Let $\tilde{\beta} = (\beta_1, \beta_2)^\top$, so $\tilde{\beta}^\top \tilde{\beta} = \beta_1^2 + \beta_2^2$. By Corollary 6.2:

$$\underbrace{\frac{\partial}{\partial \tilde{\beta}}(\tilde{\beta}^\top \tilde{\beta})}_{2 \times 1} = 2\tilde{\beta} = \begin{pmatrix} 2\beta_1 \\ 2\beta_2 \end{pmatrix}$$

Direct partial differentiation yields the same column vector.

Theorem 6.7 (Vector chain rule).

$$\frac{\partial z}{\partial \tilde{x}} = \frac{\partial y}{\partial \tilde{x}} \frac{\partial z}{\partial y}$$

or in Euler/Lagrange notation:

$$(f(g(\tilde{x})))' = \tilde{g}'(\tilde{x})f'(g(\tilde{x}))$$

See <https://quickfem.com/finite-element-analysis/>, specifically https://quickfem.com/wp-content/uploads/IFEM.AppF_.pdf

See also https://en.wikipedia.org/wiki/Gradient#Relationship_with_Fr%C3%A9chet_derivative

This chain rule is like the univariate chain rule (Theorem 2.8), but the order matters now. The version presented here is for the gradient⁶ (column vector); the total derivative⁷ (row vector) would be the transpose of the gradient⁸.

Corollary 6.3 (Vector chain rule for quadratic forms).

$$\frac{\partial}{\partial \tilde{\beta}} (\tilde{\varepsilon}(\tilde{\beta}) \cdot \tilde{\varepsilon}(\tilde{\beta})) = \left(\frac{\partial}{\partial \tilde{\beta}} \tilde{\varepsilon}(\tilde{\beta}) \right) (2\tilde{\varepsilon}(\tilde{\beta}))$$

7 Additional resources

7.1 Calculus

- Kaplan (2022)
- Khuri (2003)
- Banner (2007)
- Larson and Edwards (2018)
- Miller (2016)
 - <http://www.youtube.com/watch?v=xYzQL0TUtBA>
 - http://www.youtube.com/watch?v=Ps2SBo_WjoE

7.2 Linear Algebra and Vector Calculus

- Fieller (2016)
- Banerjee and Roy (2014)
- Searle and Khuri (2017)

7.3 Numerical Analysis

- Hua Zhou⁹'s lecture notes for “UCLA Biostat 216 - Mathematical Methods for Biostatistics” (2023 Fall)¹⁰

7.4 Real Analysis

- Grinberg (2017)

⁶<https://en.wikipedia.org/wiki/Gradient>

⁷https://en.wikipedia.org/wiki/Total_derivative

⁸https://en.wikipedia.org/wiki/Gradient#Relationship_with_total_derivative

⁹<https://hua-zhou.github.io/>

¹⁰<https://ucla-biostat-216.github.io/2023fall/schedule/schedule.html>

References

- Banerjee, Sudipto, and Anindya Roy. 2014. *Linear Algebra and Matrix Analysis for Statistics*. Vol. 181. Crc Press Boca Raton. <https://www.routledge.com/Linear-Algebra-and-Matrix-Analysis-for-Statistics/Banerjee-Roy/p/book/9781420095388>.
- Banner, Adrian D. 2007. *The Calculus Lifesaver : All the Tools You Need to Excel at Calculus*. A Princeton Lifesaver Study Guide. Princeton University Press. <https://press.princeton.edu/books/paperback/9780691130880/the-calculus-lifesaver>.
- Billingsley, Patrick. 1995. *Probability and Measure*. 3rd ed. Wiley Series in Probability and Mathematical Statistics. Wiley.
- Cheng, Eugenia. 2025. “Opinion | How Math Turned Me from a D.E.I. Skeptic to a Supporter.” *The New York Times*. <https://www.nytimes.com/2025/09/05/opinion/math-dei.html>.
- Dobson, Annette J, and Adrian G Barnett. 2018. *An Introduction to Generalized Linear Models*. 4th ed. CRC press. <https://doi.org/10.1201/9781315182780>.
- Fieller, Nick. 2016. *Basics of Matrix Algebra for Statistics with R*. Chapman; Hall/CRC. <https://doi.org/10.1201/9781315370200>.
- Fubini, Guido. 1907. “Sugli Integrali Multipli.” *Rendiconti Della Reale Accademia Dei Lincei. Classe Di Scienze Fisiche, Matematiche e Naturali* 16: 608–14.
- Grinberg, Raffi. 2017. *The Real Analysis Lifesaver: All the Tools You Need to Understand Proofs*. 1st ed. Princeton Lifesaver Study Guides. Princeton University Press. <https://press.princeton.edu/books/paperback/9780691172934/the-real-analysis-lifesaver>.
- Gut, Allan. 2013. *Probability: A Graduate Course*. 2nd ed. Springer Texts in Statistics. Springer. <https://doi.org/10.1007/978-1-4614-4708-5>.
- Kaplan, Daniel. 2022. *MOSAIC Calculus*. Wwww.mosaic-web.org. www.mosaic-web.org¹¹.
- Khuri, André I. 2003. *Advanced Calculus with Applications in Statistics*. John Wiley & Sons. <https://www.wiley.com/en-us/Advanced+Calculus+with+Applications+in+Statistics%2C+2nd+Edition-p-9780471391043>.
- 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>.
- Larson, Ron, and Bruce H. Edwards. 2018. *Calculus*. 11th ed. Cengage Learning. <https://www.cengage.com/c/calculus-11e-larson/>.
- Miller, Steven J. 2016. *The Probability Lifesaver: Calculus Review Problems*. https://web.williams.edu/Mathematics/sjmiller/public_html/probabilitylifesaver/index.htm#:~:text=http%3A/web.williams.edu/Mathematics/sjmiller/public_html/probabilitylifesaver/supplementalchap_calcreview.pdf.
- Rudin, Walter. 1976. *Principles of Mathematical Analysis*. 3rd ed. International Series in Pure and Applied Mathematics. McGraw-Hill.
- Searle, Shayle R, and Andre I Khuri. 2017. *Matrix Algebra Useful for Statistics*. John Wiley & Sons.
- Wikipedia contributors. 2024. *Fubini’s Theorem — Wikipedia, the Free Encyclopedia*. https://en.wikipedia.org/wiki/Fubini%27s_theorem.

¹¹<https://www.mosaic-web.org>