

MATH CAMP 2026

Professor: Alan G. Isaac

Day 1

Functions

Matrices

Comparative Statics with Linear Models

Day 2

Differential Calculus: A Brief Review

Basic Optimization

Day 3

Nonlinear Comparative Statics

Multivariate Optimization: Unconstrained

Multivariate Optimization: Constrained

Functions, Equations, and Economic Models

For these notes, see `functionReview.pdf` on Canvas.



By definition: a *function* $f : S \rightarrow T$ must be *source-total* and *target-definite*.

- S is the *source* set (or *domain*).
- T is the *target* set (or *codomain*).
- $S \rightarrow T$ is the *type signature* of f .

The *graph* of f is the following subset of $S \otimes T$.

$$\{ \langle x, f[x] \rangle \mid x \in S \}$$

Note:

An *argument* to a function is a specific input.

We write $f[x]$ to denote the result of *applying* f to the argument x .

In math, $f[x]$ is called the value of f at x .

In computation, $f[x]$ is called the return value of f when applied to x .

Example 1.2 (Production Technology) Let $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ be a production function that transforms nonnegative labor inputs (N) into nonnegative outputs. The set of nonnegative real numbers ($\mathbb{R}_{\geq 0}$) is both the source set and the target set of this production function. The function graph is $\{ \langle N, f[N] \rangle \mid N \in \mathbb{R}_{\geq 0} \}$.



Example 1.5 Let the source set have two elements: $S = \{\perp, \top\}$. Let the target set have the same two elements: $T = \{\perp, \top\}$. Then there are four possible functions. Define one of these by enumeration: $\perp \mapsto \top$ and $\top \mapsto \perp$. This function is the logical *negation* operation. This operation is usually represented by a prefixed negation operator ($\neg$), so that the proposition $\neg p$ has the opposite truth value of p .

Exercise:

Define each of the four possible functions by enumeration.



Which of the following functions is logical negation?

inputs↓	f_1	f_2	f_3	f_4
$\perp$	$\perp$	$\top$	$\perp$	$\top$
$\top$	$\perp$	$\perp$	$\top$	$\top$

```

sS = sT = {False, True}
(* find the possible targets for each source element: *)
choices = Table[s→t, {s,sS}, {t,sT}]
(* make ONE target choice for EACH source element *)
graphs = Tuples@choices
funcs = Association@@@graphs (* define functions by enumeration *)
(* conveniently display all the functions in a column *)
Column[funcs]

```

□

A useful depiction of very simple enumerated functions draws arrows from the elements of the source set to the corresponding elements of the target set. Each arrow links two points, so an arrow represents an ordered pair. The arrow tail is its source point, while the arrow head is its target point. The target point is the *image* of the source point. As an example, Figure 1.3 represents two simple functions where the source set is the first six nonnegative integers and the target set is the last six letters of the English alphabet.

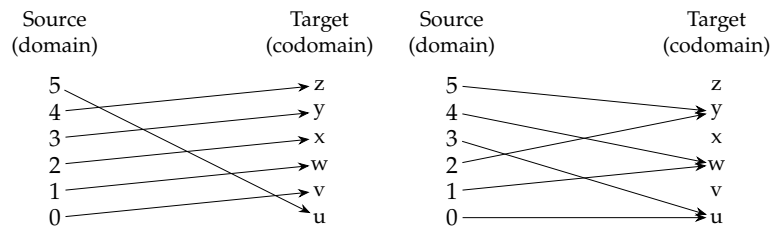


Figure 1.3: Two Simple Functions

The function on the left is source-definite (one-to-one) and is also target-total (onto). The function on the right has neither property.

WL Exercise: use `Association` to implement these functions.

WL note: invert an `Association` with `PositionIndex`.

Exercise:

Figure 1.3 illustrates two functions that share a source set (domain) and a target set (codomain). How many different functions have this source set and this target set?

Answer:

If the source set has m elements and a target set with n elements, each of the m elements in the source set might map to any one of the n elements in the target set. So there are n^m possible functions.



Figure 1.3 effectively defines two simple functions by explicit *enumeration*. Although these functions have a common source set and target set, they differ in important ways. The first function maps some source element to each target element. This is called a mapping *onto* the target set, and such a function is *target-total*.⁶ In contrast, the second simple function has a range that is only part of the target set.

More importantly, every point in the range of the first function is the image of only one source point. Pictorially, this means that only one arrowhead points to any given target point. This is called a *one-to-one* mapping, and such a function is *source-definite*.⁷ Only one source point maps to any one target point. In this sense, the source point of any given target point is definitely determinable. Source-definite functions are special because the mapping can be reversed without ambiguity. For example, if the demand for a good is a source-definite function of its price, then the demand price can be written as a function of quantity.

For a given source set S and a target set T , there are typically many possible functions. Even when the source and target sets are rather small, the number of functions can be very large. For example, if the source and the target set each have half a dozen elements, there are tens of thousands of possible functions.

⁶A target-total function is alternatively called right-total or said to be a surjection.

⁷A source-definite function is sometimes called left-definite or left-unique. It is also said rather unhelpfully to be an *injection*. A function that is target-total as well as source definite it is called a *one-to-one correspondence*, *bijection*, or *isomorphism*.



1.2.2 Rule-Based Definition

A transformation rule describes how to turn any valid input into a unique output. Function definition by means of transformation rules is prevalent in mathematical and computational social science. In contrast to explicit enumeration, function definition by means of transformation rules makes it easy to define functions on very big source sets.

Often a function that would be impossible to define by enumeration can be characterized by a simple transformation rule. This subsection considers transformations that can turn any real input into a real output.

Topics: real function; type signature □

Definition 1.4 (Real Function) The target set of a *real-valued* function is a set of real numbers. The source set of a function *of a real variable* is a set of real numbers. A *real* function is real-valued function of a real variable.

In order to reach conclusions about classes of functions, general properties are more relevant than specific function definitions. When new information about a function follows just from its known general properties, the new information applies to all functions that share those properties.

One salient property of any function is its type signature. The mathematical notation $f: S \rightarrow T$ states that a function named f has source set S and target set T . Equivalently, the *type signature* of f is $S \rightarrow T$. For example, $f: X \rightarrow \mathbb{R}$ concisely indicates that f is a real-valued function. Similarly, $f: \mathbb{R} \rightarrow X$ concisely indicates that f is a function of a real variable. Finally, $f: \mathbb{R} \rightarrow \mathbb{R}$ indicates that f is a real function.

Topics: anonymous functions □

We often name our functions in order to easily discuss them, but we are not forced to do so. A function that is defined but not assigned a name is an *anonymous* function. A complete specification of a function must state its source set and target set, not just its transformation rule. For example, the squaring function $(x \mapsto x * x): \mathbb{R} \rightarrow \mathbb{R}$ differs from $(x \mapsto x * x): \mathbb{Z} \rightarrow \mathbb{Z}$ by having different source and target sets. Here the same transformation rule appears perfectly sensible in both cases. Even though the function signature is a crucial part of a function definition, it is common for the source and target sets to remain implicit. In this case, the signature must be determined from the discussion context.

Example 1.6 (Anonymous Real Functions) The notation $(x \mapsto x): \mathbb{R} \rightarrow \mathbb{R}$ describes an identity function on the entire set of real numbers. The identity function on the real numbers is one of the simplest real functions. It is source-definite and target-total.

The notation $(x \mapsto x^2): \mathbb{R} \rightarrow \mathbb{R}$ describes a squaring function on the real numbers. The real numbers are both the source set and the target set. This squaring function is neither source-definite nor target-total.

WL Exercise: define these anonymous functions with `Function` and apply them.

WL note: As a tradeoff for its overall flexibility, WL offers somewhat limited facilities for type checking the input and outputs of functions. However, it is possible to use pattern matching to do some basic type checking. We will discuss the definition of functions via pattern matching later on.

Topic: Name Binding □

Anonymous functions are often useful, but in many settings it proves helpful to name a function. This facilitates easy reuse of the function. This book uses the equals sign to bind names to objects, including functions. This makes it easy to associate a name with a transformation rule. For example, to bind the name f to the squaring function, write $f = x \mapsto x^2$. (Read this as “ f maps x to x^2 .”) The name of the function is on the left of the equals sign, and the *defining expression* for the function is on the right.⁸

WL Exercise: Contrast ordinary functions with pattern-based functions. Explain the relationship between Set and =.

Example 1.7 (Linear Production Technology) The function $(N \mapsto 100N): \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ may represent a production technology that transforms nonnegative labor inputs (N) into nonnegative outputs. The function graph is $\{ \langle N, 100N \rangle \mid N \in \mathbb{R}_{\geq 0} \}$. Bind the name f to this function as follows: $f = (N \mapsto 100N): \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$. This specifies both the rule and the function’s type signature at the same time. However, the source set and the target set more commonly remain implicit, and we more succinctly write $f = N \mapsto 100N$.

⁸A common alternative function-definition notation is $f[x] = x^2$. (Read this as “ f of x is x^2 .”) This is also meant to bind the name f to the squaring function, but it less cleanly separates function naming from function definition. In addition, the notation is then typically reused for function application, creating an unfortunate semantic ambiguity.

Topics: image, range $\square$

Recall that the notation $f[x]$ denotes the target element associated by the function f with the source element x . This is the *image* of x under f , or equivalently the *value* of f at x . Additionally, for any subset of the source set, $X \subseteq S$, overload this notation by letting $f[X] \stackrel{\text{def}}{=} \{f[x] \mid x \in X\}$; this is the *image* of X under f . (This is substitutive set building.) For example, $f[S]$ is the image of the entire source set, which is the *range* of the function. When f is target-total, $f[S] = T$.

WL illustration of `FunctionRange`:

```
FunctionRange[Sin[x], x, y]
FunctionRange[{(x^2)[x], -1 <= x <= 1}, x, y]
```

Remark 1.3 If f is not target-total, then $f[S] \subset T$. However, one may always restrict the target set to the range of f in order to produce a target-total function: $\tilde{f}: S \rightarrow f[S]$, where $\forall x \in S: \tilde{f}[x] = f[x]$. The new function $\tilde{f}$ is known as a *surjective restriction* of f . (Here, the word *surjective* is synonymous with *target-total*.)



1.2.2.1 Implicit Domain

A full characterization of a function includes a specification of its source set. This is the domain of valid inputs, or argument type, of the function. However, less formal function definitions that do not specify the source set are very common. In this case, the source set is implicitly the largest sensible domain in the context of the definition. In particular, a real function has a *natural* domain that is the set of real numbers x for which its transformation rule can produce a real number. For example, the real squaring function has a natural domain equal to the entire set of real numbers, as is true for any polynomial. In contrast, zero is not in the natural domain of real reciprocal function ($x \mapsto 1/x$), and the negative real numbers ($\mathbb{R}_{<0}$) are not in the natural domain of real square-root function ($x \mapsto \sqrt{x}$).

[WL illustration using FunctionDomain](#)

```
FunctionDomain[1/x, x]  
FunctionDomain[Sqrt[x], x]
```



A function used in social science often has an intended domain that is much smaller than the natural domain. For example, economists refer to a restricted function domain determined by economic considerations as the *economic* domain of the function. Even when $\mathbb{R}$ is the natural domain of a real function, the domain relevant to a social-scientist is often the nonnegative real numbers ($\mathbb{R}_{\geq 0}$). This restriction typically remains implicit, to be understood from the context of the discussion.

Example 1.8 (Economic Domain) When other inputs are constant, an economist might propose a functional dependence of an industry's output on its labor input (N). Theoretical considerations may suggest a specific functional form for the transformation, such as the linear form $N \mapsto aN$ (for some $a > 0$). Aside from anomalies (e.g., power-loom riots), the possible values of the labor input are generally nonnegative. So economists typically treat the economic domain of this production function as the nonnegative real numbers. Using $\mathbb{R}_{\geq 0}$ as the target set as well, this linear production function is a source-definite and target-total.



1.2.3 Computational Functions

A typical computer program decomposes an overall task into smaller subtasks. A single execution of the computer program often performs some of these subtasks repeatedly. It is good practice to implement such subtasks as subroutines of the encompassing program. A subroutine is a reusable chunk of code that we write once in order to use many times.

Code duplication reduces the readability and the maintainability of the source code for a computer program. This observation motivates the so-called *DRY principle* of computer programming: *don't repeat yourself*. In order to avoid code duplication, programmers need convenient facilities for code reuse. Every modern programming language provides these, and *computational* functions are a particularly important example. Defining a named function that consumes a specified type of value in order to compute a needed value proves particularly useful. A particular input value an *argument* to the function, and the resulting output value the *return value* of the function. We can apply the function to any valid input argument in order to produce the needed return value.⁹ Accordingly, functions are ubiquitous in computer programming.

⁹This book generally speaks of a *function application* rather than a *function call*.



Pure functions in computation:

- closed to outside influences
- has no side effects

In a computational discussions, the term *function* often has a very broad use. Although a fully specified mathematical function has a clearly defined source set and target set, a computational function may not. A computational function may even accept no input and yet return a different value each time it runs, as with traditional random number generators.¹⁰ Even when it does require an input argument, a computational function may not restrict the function behavior to ensure that it handles only valid inputs.

Nevertheless, some computational functions behave much like simple mathematical functions. A computational function is *pure* if its behavior does not depend on its execution environment and it does not change this environment. Equivalently, a pure function is closed to outside influences and has no side effects. The output of a pure computational function is completely determined by its explicit inputs.

Reliance on pure functions can make it much easier to read and understand a computer program. The use of pure functions also makes it easier to relate the behavior of computer code to mathematical constructs, since many mathematical arguments involve mathematical functions. It is a good programming practice to work with pure functions whenever doing so proves reasonably convenient.¹¹

¹⁰Another lecture discusses random number generation.

¹¹Nevertheless, impure computational functions can be very useful for their *side effects*. For example, an impure function may change the global program state, or change the state of a program object, or simply print a description of the object state.

Listing 1.1: Computational Function Definition

```
consumption = yd  $\mapsto$  20.0 + 0.8*yd
```

Listing 1.1 introduces the semi-formal style of pseudocode this book uses to define and name computational functions.¹² In this case, the computational definition looks very like its mathematical definition. A name for the function is bound to a function object that is created by the maps-to operator ($\mapsto$).¹³ On the left of the maps-to operator is its named input *parameter*.¹⁴ The expression on the right of the maps-to operator specifies the function's transformation in term of this input parameter. To reduce visual clutter, the pseudocode does not indicate the type of the input argument nor of the output value. The type signature of the function is omitted and remains implicit.

¹²This pseudocode syntax derives almost exactly from the Wolfram Language. Pseudocode is a stylized description of an implementation in computer code. The implementation in your chosen language may look very different. The pseudocode in this book often does not specify the precise types of the values in our computations, relying instead on context. In the present example, the context indicates that applying this function to a numerical input will return a numerical result.

¹³In WL, enter this operator as a three-character sequence: `| - >`.

¹⁴A function parameter is occasionally called a *formal parameter* of the function, in which case the input argument may be called the *actual parameter* of the function. This book instead refers to parameters and arguments.

```
consumption = yd ↦ 20.0 + 0.8*yd
```

In the pseudocode of this book, code outside of a function's definition cannot access its parameters. That is, a function parameter has *local scope*; it is not visible outside the function definition. In this example, *yd* is the only function parameter. Within the function definition, the parameter represents any valid input argument. Notice that the function parameter occurs in the body of the function, which defines the transformation produced by this function. This is natural: the value of total consumption spending is computed based on disposable income. This is the value *returned* by the function. Put more simply, it is the value of the function.

Here are some examples of implementation of Listing 1.1 in different languages. For the purposes of illustration and comparison, even though it is superfluous, each case includes an assignment to a local variable named `spending`. Similarly, each example includes an explicit return statement, even when this is not required by the language. (It is *not* required by Mathematica, Julia, or Rlang.)

Python:

```
def consumption(yd):
    spending = 20.0 + 0.8 * yd
    return spending
```

Mathematica:

```
consumption = Function[yd,
  Module[{spending = 20.0 + 0.8*yd},
    Return[spending]
  ]]
```

Julia:

```
function consumption(yd)
    spending = 20.0 + 0.8 * yd
    return spending
end
```

Golang:

```
func consumption(yd: float64) float64 {
    spending := 20.0 + 0.8 * yd
    return spending
}
```

Rlang:

```
consumption <- function(yd) {
    spending <- 20.0 + 0.8 * yd
    return(spending)
}
```

WL Exercise: Implement the function and apply it to an argument of 50.0.
Hint: WL uses square brackets for function application.



1.2.3.1 Recursive Definition

A function definition is *recursive* if it makes use of the very function being defined. In order for this to work, there must be some way to bring the recursion to a halt. Typically, there is a special argument value for which there is a known result. This special value is the *base case*. The function definition must ensure that recursion will terminate at the base case.

□

The factorial operation serves as a standard first illustration of recursive function definition. Mathematicians often use a post-fixed exclamation mark operator to denote the factorial of a nonnegative integer. When $n \in \mathbb{Z}_{\geq 0}$, define $\text{fac}[n]$ to be the product of the members of the integer interval $[1 .. n]$.

$$\text{fac} = n \mapsto \prod_{i=1}^n i \quad (1.4)$$

(An empty product evaluates to 1). We can also represent this computation recursively.

$$\text{fac} = n \mapsto \begin{cases} 1 & n = 0 \\ n \cdot \text{fac}[n - 1] & \text{otherwise} \end{cases} \quad (1.5)$$

The first case is the base case. The second case is the recursive function application, since it makes use of the very function being defined. In principle, even a very big integer argument must produce recursion all the way to the base case, since the argument is decremented at each recursive step. Remember, the domain of this function is $\mathbb{Z}_{\geq 0}$.

Topic: ternary if □

By relying on a ternary `if` function, we may write this recursive definition even more succinctly.

$$\text{fac} = n \mapsto \text{if}[n == 0, 1, n \cdot \text{fac}[n - 1]] \quad (1.6)$$

The value of this `if` function depends on the value of its first boolean argument. When this is $\top$, the `if` function produces the value of its second argument (1). Otherwise, the `if` function produces the value of its third argument. Many programming languages include such a function, and some include an equivalent ternary operator.

```
fac = n ↦ If [n==0,1,n*fac [n-1]]
```

Remark 1.4 Recursive function are very elegant. However, a computational problem with recursion can arise when it is impossible to return from a function call until its recursive call returns. Every recursive function call then *stacks up* until the base case is reached. The stored function calls must include relevant information from their environment, such as the values and scope of local variables. This means that deep recursion may be computationally expensive. To avoid this expense, programmers often try to write equivalent iterative algorithms, which do the repetitive recalculation without recursion.



1.3 Binary Operations



The number of arguments the function accepts is its *arity*. A unary function consumes one argument and returns a result. A binary function consumes two arguments and returns a result. In various settings, mathematicians often use the term *operation* rather than *function*, but there is no structural conceptual difference.¹⁵ For example, a unary function that consumes one real argument and returns a real result is often simply called a *real function*. However, the very standard transformations of negation or absolute value may well be called operations. Similarly, a binary function that consumes two real arguments and returns a real result is often simply called a *binary real function*. Example 1.9 illustrates this with a modified consumption function.

Example 1.9 (Binary Real Function) Consumption should depend on wealth (F) as well as disposable income (y_d). We may represent this by a binary real function, such as

$$\langle y_d, F \rangle \mapsto 20.0 + 0.8y_d + 0.06F$$

In many settings where the argument types and return type of a binary function are all the same, it is common to speak of a *binary operation*. This is particularly true if there is a operator special syntax for representing the transformation. In Example 1.9, the body of the consumption function uses two common operations in the real numbers: addition and multiplication. Although some programming languages provide named functions (such as `add` and `mul`) for these operations, most popular programming languages provide the familiar infix arithmetic operators. In Example 1.9, note that the multiplication operator is implicit, as is common in mathematical expositions. As we saw in Listing 1.1, however, our implementations in code typically must include the operator explicitly.

¹⁵In computer science, the term *operation* has a much broader usage, which often includes the various side effects that result from code execution.



1.3.1 Arithmetical Operations

Consider the binary operation of addition in the nonnegative integers ($\mathbb{Z}_{\geq 0}$). Each of the operands is a *term*, and the result of the operation is their *sum*. It is conventional to represent this operation by using the plus (+) symbol as an infix operator. This operation is associative, and there is an additive identity ($0 \in \mathbb{Z}_{\geq 0}$). In sum, for any three nonnegative integers (a , b , and c), the following properties hold.

- associativity: $(a + b) + c = a + (b + c)$
- identity: $a + 0 = a$ and $0 + a = a$

Associativity and the existence of an identity are the *monoid* properties. We may say that the nonnegative integers are a monoid under addition. In addition, the addition operation is commutative: the order of the operands does not affect the result.

- commutativity: $a + b = b + a$



Next, consider the entire set of integers ($\mathbb{Z}$). This too is a commutative monoid under addition, but now addition has an additional useful property. Associated with each element $a \in \mathbb{Z}$ is an inverse element, denoted $-a$.

- inverse: $\exists(-a) \in \mathbb{Z}$ such that $a + (-a) = 0$ and $(-a) + a = 0$.

A monoid where every element has an inverse is called a *group*, so addition is a group operation in the integers. Since addition is also commutative, we say that the integers are a commutative group under addition. Addition is similarly a group property in the rational numbers, real numbers, or complex numbers.

**EXERCISE:**

Are the integers a group under multiplication? Explain.

What about the rational numbers?

What about the nonzero rational numbers?

EXERCISE:

Use the associativity and identity properties to prove that if a is an inverse for b and c is also an inverse for b then $a = c$.

SOLUTION: $\square$

Example 1.10 (Unique Inverse) Use the associativity and identity properties to prove that inverse elements are unique. That is, if a is an inverse for b and c is also an inverse for b then $a = c$.

Proof:

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

associativity axiom

$$0 + c = a + 0$$

inverse assumption

$$c = a$$

identity axiom



Multiplication is a second familiar binary operation in the real numbers. Each of the operands is a *factor*, and the result of the operation is their *product*. We typically represent multiplication by a special infix operator notation (using $\cdot$ or $*$ or sometimes $\times$), although this is often suppressed for brevity. For any three real numbers (a , b , and c), the multiplication operation has the following familiar properties.

- associativity: $(a \cdot b) \cdot c = a \cdot (b \cdot c)$
- identity: $a \cdot 1 = a$ and $1 \cdot a = a$
- inverse: $\forall a \neq 0: \exists a^{-1}$ such that $a \cdot a^{-1} = 1$ and $a^{-1} \cdot a = 1$.
- commutativity: $a \cdot b = b \cdot a$

The real numbers are not quite a group under multiplication, because 0 does not have a multiplicative inverse. However, the nonzero real numbers ($\mathbb{R}_{\neq 0}$) are a commutative group under multiplication.



Let $\langle X, +, \cdot \rangle$ be a set with two binary operations, called addition and multiplication. This triple constitutes a *field* if the following properties apply.

- X is a commutative group under addition (call the identity 0).
- $X - \{0\}$ is a commutative group under multiplication (call the identity 1).
- multiplication distributes over addition: $\forall a, b, c \in X$ we have $a \cdot (b + c) = (a \cdot b) + (a \cdot c)$ and $(b + c) \cdot a = (b \cdot a) + (c \cdot a)$

Note that since the operations are commutative, we right distribution follows from left distribution and vice versa. In this case, we can coherently specify only one of these. The distributive law provides a necessary linkage between between addition and multiplication.



Example 1.11 (Null Factor Property) A null factor ensures a zero product: $a \cdot 0 = 0$. Note that $a \cdot 0 = a \cdot (0 + 0) = a \cdot 0 + a \cdot 0$. Add the additive inverse of $a \cdot 0$ to get

$$(a \cdot 0) + (-(a \cdot 0)) = ((a \cdot 0) + (a \cdot 0)) + (-(a \cdot 0))$$

The left side is 0, by definition of the additive inverse.

$$\begin{aligned} 0 &= ((a \cdot 0) + (a \cdot 0)) + (-(a \cdot 0)) \\ &= (a \cdot 0) + ((a \cdot 0) + (-(a \cdot 0))) && \text{distribution} \\ &= (a \cdot 0) + 0 && \text{inverse} \\ &= a \cdot 0 && \text{identity} \end{aligned}$$

Example 1.12 (Null Product Property) A null product must include a zero factor. That is, $a \cdot b = 0$ implies $a = 0$ or $b = 0$.

The proof is by cases. If $a = 0$, we are done. Otherwise, $a \neq 0$, and then since $a \cdot b = 0$,

$$\begin{aligned} a^{-1} \cdot (a \cdot b) &= a^{-1} \cdot 0 && \text{multiply by } a^{-1} \\ (a^{-1} \cdot a) \cdot b &= 0 && \text{associativity; null factor} \\ 1 \cdot b &= 0 && \text{inverse} \\ b &= 0 && \text{identity} \end{aligned}$$



Example 1.13 (Negative One) Multiplication by -1 produces additive inverses. Prove that $-1 \cdot a = -a$ as follows.

$a + (-1 \cdot a) = (1 \cdot a) + (-1 \cdot a)$	identity
$= (1 + (-1)) \cdot a$	distribution
$= 0 \cdot a$	inverse
$= 0$	null factor

Prove $(-a) \cdot (-a) = a \cdot a$.

$(-a) \cdot (-a) = (-1 \cdot a) \cdot (-1 \cdot a)$	negative one
$= -1 \cdot (a \cdot (-1 \cdot a))$	assoc
$= -1 \cdot ((a \cdot -1) \cdot a)$	assoc
$= -1 \cdot ((-1 \cdot a) \cdot a)$	commut
$= -1 \cdot (-1 \cdot (a \cdot a))$	assoc
$= -1 \cdot -(a \cdot a)$	negative one
$= a \cdot a$	negative one

We usually write a^2 for $a \cdot a$, so our result is $(-a)^2 = a^2$.



1.3.1.1 Powers

We typically write x^2 for $x \cdot x$, x^3 for $x \cdot x \cdot x$, and so forth. These are called *powers* of x , and the superscripted numbers are called *exponents*. For a real number x and a nonnegative integer n , consider the computation of x^n . This is a binary operation with two different kinds of inputs: a real number, and an integer. This operation may be defined in terms of repeated multiplication.

$$\text{pow} = \langle x, n \rangle \mapsto \prod_{i=1}^n x \quad (1.7)$$

(Recall that an empty product evaluates to 1.) For this operation, mathematical discussion usually adopt a special superscript notation: x^n . This suggests an obvious computational algorithm: the multiplication of x times itself n times. This is an *iterative* approach. However, the definition also suggests a recursive approach.

$$\text{pow} = \langle x, n \rangle \mapsto \text{if } [n = 0, 1, x \cdot \text{pow}[x, n - 1]] \quad (1.8)$$

Two iterative approaches and a recursive approach, when n is a nonnegative integer.

```
pow01 = {x, n} ↦ Product[x, {i, n}]
pow02 = {x, n} ↦ Nest[a ↦ a*x, 1, n]
pow03 = {x, n} ↦ If [n==0, 1, x*pow03[x, n-1]]
```



1.3.2 Logical and Set Operations

A binary logical operation takes two boolean inputs and produces a boolean output. The most commonly used are and ($\wedge$), or ($\vee$), implies ($\implies$) and iff ($\iff$). The implies operation is often called *material implication*. The iff operation is often called *material bimplication* and is usually pronounced *if and only if*. Table 1.1 defines these operations enumeratively.

Table 1.1: Define Common Logical Operations

p	q	$p \wedge q$	$p \vee q$	$p \implies q$	$p \iff q$
$\top$	$\top$	$\top$	$\top$	$\top$	$\top$
$\top$	$\perp$	$\perp$	$\top$	$\perp$	$\perp$
$\perp$	$\top$	$\perp$	$\top$	$\top$	$\perp$
$\perp$	$\perp$	$\perp$	$\perp$	$\top$	$\top$

An enumerative definition of some common binary logical operations. An entry of $\top$ denotes propositional truth; an entry of $\perp$ denotes falsity. Any possible pair of values for p and q is represented by a single row of the truth table. Each column defines a binary logical operation.

□

A binary set operation operates on two sets, say P and Q , in order to produce a new set as a result. A *membership table* can define set operations by relating membership in the operands (here, P and Q) to membership in the result. The number 1 indicates membership in a set, while the number 0 represents nonmembership. Any object under consideration must be associated with exactly one of the four rows of Table 1.2. When classifying an object, pick the row on the left side of the table that represents the memberships of the element in P and Q . This determines its membership in any set that can be constructed from P and Q . The right side of this table defines some common binary set operations, which are represented by infix binary operators: $\cap$ for *intersection*, $\cup$ for *union*, $-$ for *set difference*, and Δ for *symmetric difference*.

Table 1.2: Define Common Set Operations

P	Q	$P \cap Q$	$P \cup Q$	$P - Q$	$P \Delta Q$
1	1	1	1	0	0
1	0	0	1	1	1
0	1	0	1	0	1
0	0	0	0	0	0

An entry of 1 indicates set membership; an entry of 0 indicates nonmembership. Any object is categorized by a single row of the membership table, based on its membership in P and/or Q . Each column defines a binary set operation.



The mathematical expression $\omega \in P$ asserts that an ω is an element of P . This assertion is either true or false. The expression has the value $\top$ if ω is an element of P , and it has the value $\perp$ if ω is not an element of P . Using the elementhood operator ($\in$) along with set-builder notation and the logical operators, we may offer alternative characterizations of set operations. Equation (1.9) does so for intersection and union.

$$\begin{aligned} P \cap Q &\stackrel{\text{def}}{=} \{ \omega \mid \omega \in P \wedge \omega \in Q \} \\ P \cup Q &\stackrel{\text{def}}{=} \{ \omega \mid \omega \in P \vee \omega \in Q \} \end{aligned} \tag{1.9}$$

Recall that the $\wedge$ symbol is the infix *logical and* operator, so the intersection operation selects only the objects that are in both of the sets. Also, the $\vee$ symbol is the infix *logical or* operator, so the union operation selects any object that is in one or the other of the sets.

```
MemberQ[{1, 2, 3}, 1]
1 ∈ Integers
Pi ∈ Reals
```

```
setP = {1,2,3};
setQ = {3,4,5};
setP ∪ setQ
setP ∩ setQ
```

□

1.3.2.1 Predicates

A *predicate* is a boolean-valued function. That is, the target set of a predicate comprises only the two boolean values: $\mathbb{B} = \{\perp, \top\}$. When a predicate p is defined on a set Ω , it effectively partitions it into two subsets: the elements that satisfy predicate, and those that do not. For the moment, let us call these P and $\bar{P}$.

$$\begin{aligned} P &= \{ \omega \in \Omega \mid p[\omega] \} \\ \bar{P} &= \{ \omega \in \Omega \mid \neg p[\omega] \} \end{aligned} \tag{1.10}$$

The set P defined in (1.10) comprises those elements of Ω that satisfy the predicate p . This setbuilder notation provides an example of selective set building. The predicate provides the selection criterion for producing this subset of Ω . The set defined $\bar{P}$ defined in (1.10) comprises those elements of Ω that do not satisfy the predicate p . Naturally, $\bar{P} = \Omega - P$.



Recall that the expression $\omega \in \mathbb{Z}_{\geq 0}$ on the left of the set-builder separator in (1.10) indicates that only nonnegative integers are tested by the predicate. Within the set builder, the variable ω represents any nonnegative integer value whatsoever. Since it occurs to the left of the separator, the variable ω is *bound* by the set-builder expression. That is, it does not refer to anything outside of the set-builder expression. Since ω occurs again to the right of the separator, this set-builder expression tests each nonnegative number with the predicate. Note that since the set builder binds the variable ω , we may (and usually do) simply place the predicate's defining expression to the right of the separator, as in example 1.14.

Example 1.14 (Even and Odd Integers) When defined in the nonnegative integers ($\mathbb{Z}_{\geq 0}$), the *modulo* operation produces the remainder from integer division.^a An infix modulo operator is often written as % or as mod. For integers $m \geq 0$ and $n > 0$, denote the remainder from integer division by $r = m \% n$. Note that $m \% 2$ is zero iff m is even. Therefore, we may use the modulo operation to build the the sets of even and odd numbers, as follows.

$$\begin{aligned} E &= \{ \omega \in \mathbb{Z}_{\geq 0} \mid 0 = \omega \% 2 \} && \text{even numbers} \\ \bar{E} &= \{ \omega \in \mathbb{Z}_{\geq 0} \mid 1 = \omega \% 2 \} && \text{odd numbers} \end{aligned}$$

These two sets form a *partition* of $\mathbb{Z}_{\geq 0}$: they have no overlap, but together they contain all the nonnegative integers. That is, the union $E \cup \bar{E}$ is the set of all nonnegative integers, and the intersection $E \cap \bar{E}$ is empty.

^aRecall that Euclid's division lemma states that, for integers $m \geq 0$ and $n > 0$, there are unique nonnegative integers q (the quotient) and r (the remainder) such that $m = qn + r$ and $r < n$.

□

In WL, use the Mod command for the modulo operation.

```
Table[n~Mod~2, {n, 10}]    (* 0 for even; 1 for odd *)
Table[OddQ[n], {n, 10}]    (* False for even; True for odd *)
```

Selection, Partition, and Categorization

```
n10 = Range[10];
Table[If[EvenQ@n, n, Nothing], {n, n10}]    (* select evens *)
Select[n10, EvenQ]                          (* select evens *)
{evens, odds} = GatherBy[n10, EvenQ]        (* partition *)
evens ∪ odds == n10                         (* True *)
evens ∩ odds == {}                          (* True *)
GroupBy[EvenQ]@n10                         (* categorization *)
```



1.3.3 Comparison Operations

A comparison operation is a binary predicate on some set. In mathematical discussions, the standard comparison operations on the real numbers are often represented by familiar infix binary operators: $<$, $\leq$, $=$, $\geq$, and $>$. In this case, the two operands are real numbers and the result of the comparison is boolean ($\perp$ or $\top$).

Begin the less-than-or-equal operation, where the result is true when the first operand is no greater than the second operand. Represent this operation by an infix weak-precedence operator ($\leq$). For any real numbers a , b , and c , the operation satisfies the following properties.

- reflexive: $(a \leq a)$
- antisymmetric: $(a \leq b) \wedge (b \leq a) \implies (a = b)$
- transitive: $(a \leq b) \wedge (b \leq c) \implies (a \leq c)$
- complete: $(a \leq b) \vee (b \leq a)$

A complete binary relation is often called *strongly connected*. Any comparison operation with these properties is a **complete order**. In addition, the less-than-or-equal operation on the real numbers satisfies the addition rule and the product rule. For any real a , b , and c ,

- addition rule: $(a \leq b) \implies (a + c) \leq (b + c)$
- product rule: $(0 \leq a) \wedge (0 \leq b) \implies (0 \leq a \cdot b)$

□

Define the less-than comparison as follows.

$$(a < b) \iff (a \leq b \wedge a \neq b) \quad (1.11)$$

It follows that for any real a , b , and c , the following are true.

- irreflexive: $\neg(a < a)$
- asymmetric: $(a < b) \implies \neg(b < a)$
- transitive: $(a < b) \wedge (b < c) \implies (a < c)$
- trichotomy: $(a < b) \vee (a = b) \vee (b < a)$

In addition, the less-than operation on the real numbers satisfies the addition rule and the product rule. For any real a , b , and c ,

- addition rule: $(a < b) \implies (a + c) < (b + c)$
- product rule: $(0 < a) \wedge (0 < b) \implies (0 < a \cdot b)$

Prove the addition rule for $<$.

Given $a < b$ we know from the addition rule for $\leq$ that $a + c \leq b + c$.

Note that $a + c = b + c \implies a = b$, but $a < b$, therefore $a + c < b + c$.

Prove the product rule for $<$.

Given $0 < a$ and $0 < b$ we know from the product rule for $\leq$ that $0 \leq a \cdot b$.

Note that $a \cdot b = 0 \implies (a = 0 \vee b = 0)$ (zero product rule). But $0 < a$ and $0 < b$, therefore $0 < a \cdot b$.

□

Show that $0 < 1$.

First prove $a \neq 0 \implies 0 < a \cdot a$.

For $0 < a$ the product rule for $<$ implies $a \cdot 0 < a \cdot a$ so $0 < a \cdot a$.

For $a < 0$ we know $0 < -a$ so the previous result implies $0 < (-a) \cdot (-a) = a \cdot a$.

So $0 < 1 \cdot 1 = 1$.

□ Prove that if $x \leq y$ and $0 \leq z$ then $zx \leq zy$.

$$\begin{aligned}x &\leq y \\0 &\leq y - x \\z \cdot 0 &\leq z \cdot (y - x) \\0 &\leq zy - zx \\zx &\leq zy\end{aligned}$$

Prove that if $x < y$ and $z < 0$ then $zy < zx$.

First prove that $0 < -z$.

$$\begin{aligned}z &< 0 \\z + (-z) &< 0 + (-z) \\0 &< -z\end{aligned}$$

Next apply the product rule using $(-z)$, recalling that $0 < y - x$.

$$\begin{aligned}0 &< (y - x) \\-z \cdot 0 &< -z \cdot (y - x) \\0 &< -z \cdot y + (-z \cdot -x) \\0 &< -(z \cdot y) + z \cdot x \\z \cdot y &< z \cdot x\end{aligned}$$



1.3.3.1 Absolute Value

Thinking once again in terms of a number line, recall the absolute value of a number may be thought of as its distance from the origin. Now, define the unary absolute value operation on the real numbers in terms of the strict-precedence operation. As usual, let $-x$ denote the additive inverse of the real number x , and use a case-based definition of the function.

$$\text{abs} = x \mapsto \begin{cases} -x & x < 0 \\ x & \text{otherwise} \end{cases} \quad (1.12)$$

Programming languages usually find a dedicated function to perform this operation. However, mathematical expositions typically represent the absolute value operation by delimiting vertical bars, where $|x|$ represents the absolute value of x .



The following core properties follow from this definition of the absolute value operation.

- non-negative: $|a| \geq 0$
- separating: $|a| = 0 \iff a = 0$
- product rule: $|ab| = |a||b|$
- triangle inequality: $|a + b| \leq |a| + |b|$



Many additional properties follow from the definition of the absolute value operation. The following are often useful.

- symmetric: $|-a| = |a|$
- reverse triangle inequality: $||a| - |b|| \leq |a - b|$
- absolute bound: $|a| < b$ iff $-b < a < b$
- power rule: $|a^n| = |a|^n$

□

Prove the absolute-bound property by cases: $|a| < b$ iff $-b < a < b$.

Case 1: $a > 0$ so $a > 0 > -b$, and $a = |a| < b$ is given.

Case 2: $a < 0 < b$ so $a < b$ obviously. Note: $|a| < b \implies -|a| > -b$, but in this case $a = -|a|$, so $-b < a$.

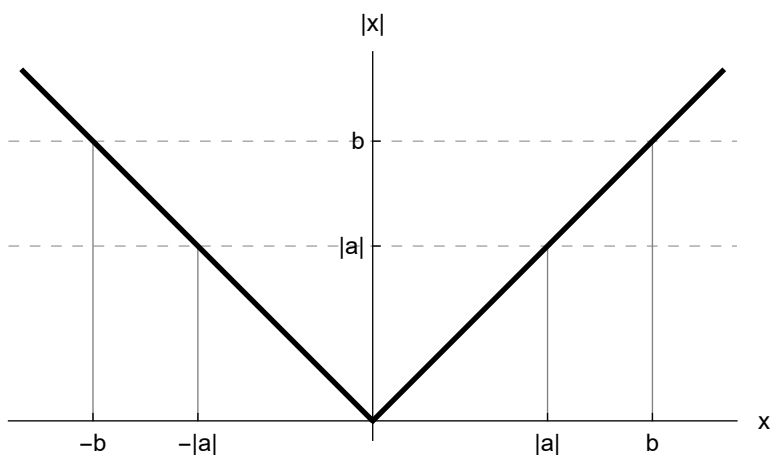


Figure 1.4: Absolute Bound

If $a \geq 0$, then $a = |a|$. If $a < 0$, then $a = -|a|$. Therefore the inequality $|a| < b$ implies $-b < a < b$.

Matrices and Linear Models

For these notes, see `matrix1.pdf` on Canvas.



The Matrix

□ A numerical matrix is a rectangular array of numbers.

Represent the location of a matrix element by a row index and a column index. The convention is to state the row index before the column index. For an arbitrary matrix named $\mathbf{A}$, let $a_{r,k}$ denote the element in row r and column k , as in Figure 1.1.

$$\mathbf{A}_{R \times K} = \begin{bmatrix} a_{1,1} & a_{1,2} & \cdots & a_{1,k} & \cdots & a_{1,K} \\ a_{2,1} & a_{2,2} & \cdots & a_{2,k} & \cdots & a_{2,K} \\ \vdots & \vdots & & \vdots & & \vdots \\ a_{r,1} & a_{r,2} & \cdots & \mathbf{a}_{r,k} & \cdots & a_{r,K} \\ \vdots & \vdots & & \vdots & & \vdots \\ a_{R,1} & a_{R,2} & \cdots & a_{R,k} & \cdots & a_{R,K} \end{bmatrix} \quad \leftarrow \text{row } r$$

$\uparrow$
 column k

Figure 1.1: Matrix Conceptual Layout and Indexing
 The matrix $\mathbf{A}$ is an $R \times K$ matrix: it has R rows and K columns.
 The notation $a_{r,k}$ denotes the element in row r and column k .

Apply WL's `Dimensions` function to each matrix:

```
mA = {{2,3},{5,6},{1,3}};
mB = {{3},{4},{5}};
mC= {{1,2,-5}};
```



Remark 1.1 Matrix *equality* is element-wise: equal matrices have the same shape, and corresponding elements are equal. So for two matrices A and B , the equality $A = B$ means that they have the same shape and in every case $a_{r,k} = b_{r,k}$.

Remark 1.2 A matrix with a single column is often called a *column vector*, and a matrix with a single row is often called a *row vector*.



Example 1.1 (Matrix Shape) Consider the following three matrices. In each case, the matrix has two dimensions.

$$\mathbf{A}_{3 \times 2} = \begin{bmatrix} 2 & 3 \\ 5 & 6 \\ 1 & 3 \end{bmatrix} \quad \mathbf{B}_{3 \times 1} = \begin{bmatrix} 3 \\ 4 \\ 5 \end{bmatrix} \quad \mathbf{C}_{1 \times 3} = [1 \quad 2 \quad -5]$$

Here, the matrix shape is communicated by a subscript appended to the matrix name. Matrix $\mathbf{A}$ is a 3×2 matrix: its shape is $\langle 3, 2 \rangle$. Its row dimension is 3, and its column dimension is 2. Matrix $\mathbf{B}$ is a 3×1 matrix, and matrix $\mathbf{C}$ is a 1×3 matrix. Matrix $\mathbf{B}$ is a column vector, and matrix $\mathbf{C}$ is a row vector.



1.1.1.1 Square Matrices

Definition 1.2 (Square Matrix) A *square* matrix has the same number of rows and columns; this number is the *order* of the square matrix.

Example 1.2 (Some Square Matrices) Each of the following square matrices has order 2. Matrix B is a diagonal matrix. Matrix C is a scalar matrix.

$$A_{2 \times 2} = \begin{bmatrix} 2 & 3 \\ 1 & 3 \end{bmatrix} \quad B_{2 \times 2} = \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix} \quad C_{2 \times 2} = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix}$$

```
mA = {{2,3},{1,3}}; (* square matrix of order 2 *)  
mA // MatrixForm  
mB = {{2,3,4},{5,6,7},{8,9,1}}; (* square matrix of order 3 *)  
mB // MatrixForm  
mB[[2,3]]      (* 7 *)
```



The *main diagonal* (or *principal diagonal*) of a matrix starts with the top, left-most element and sequentially includes all the elements with identical row and column indexes. We may represent the diagonal as an array containing these elements, extracted from the matrix by a `diag` function. For example, if A is an $n \times n$ matrix, then $\text{diag}[A] = \langle a_{i,i} \rangle_{i=1}^n$. Equation (1.1) illustrates this for a 3×3 matrix.

$$A = \begin{bmatrix} \boxed{a} & b & c \\ d & \boxed{e} & f \\ g & h & \boxed{i} \end{bmatrix} \implies \text{diag}[A] = \langle a, e, i \rangle \quad (1.1)$$

```
mA = {{a, b, c},{d, e, f},{g, h, i}}
mA // MatrixForm
Diagonal[mA]
```



Definition 1.3 (Diagonal Matrix) A *diagonal* matrix is a square matrix with zeros for every element that is off the main diagonal. That is, a diagonal matrix is a square matrix with zero elements everywhere except on its principal diagonal (where the row index and column index are equal). A *scalar* matrix is a diagonal matrix with a single value along the principal diagonal.

Example 1.3 (Diagonal Matrices) The matrices A and B are diagonal. The matrix B is also scalar.

$$A = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix} \qquad B = \begin{bmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$$

Creating diagonal matrices by hand:

```
mA = {{1,0,0},{0,2,0},{0,0,3}}; (* diagonal matrix *)
```

Use `DiagonalMatrix` to create a diagonal matrix. (Or possibly even a scalar matrix.)

```
mA = DiagonalMatrix[{1,2,3}]; (* diagonal matrix *)
```

For very large scalar matrices, a better option is often `SparseArray`.

```
mB = SparseArray[{i_ , i_} → 2, {3, 3}] (* 3 x 3 scalar matrix *)  
mB // MatrixForm
```



1.2 Core Matrix Operations

Numerical matrices are rectangular arrays of numbers that we endow with certain operations. This section describes matrix addition, scalar multiplication, and matrix multiplication.



1.2.1 Matrix Addition

Matrix addition is elementwise. Addition is possible only for matrices that have the same shape. Matrices with the same shape are *conformable* for addition. Given two conformable matrices $\mathbf{A}$ and $\mathbf{B}$, define the matrix sum $\mathbf{A} + \mathbf{B}$ as follows.

$$\overbrace{[a_{r,k}]}^{\mathbf{A}} + \overbrace{[b_{r,k}]}^{\mathbf{B}} = \overbrace{[a_{r,k} + b_{r,k}]}^{\mathbf{A+B}} \quad (1.2)$$

This notation means that the $\langle r, k \rangle$ -th element of $\mathbf{A} + \mathbf{B}$ is $a_{r,k} + b_{r,k}$. It is implicit that the three matrices have the same shape: the new matrix has the same shape as the summands. Each element of the result is the sum of the corresponding elements of $\mathbf{A}$ and $\mathbf{B}$.

Example 1.4 (Matrix Addition)

$$\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} + \begin{bmatrix} 7 & 8 & 9 \\ 10 & 11 & 12 \end{bmatrix} = \begin{bmatrix} 1+7 & 2+8 & 3+9 \\ 4+10 & 5+11 & 6+12 \end{bmatrix} = \begin{bmatrix} 8 & 10 & 12 \\ 14 & 16 & 18 \end{bmatrix}$$



Example 1.5 A firm has two factories. Each factory reports a row vector holding its revenues and costs.

$$\mathbf{F}_{1 \times 2} = [200 \ 300] \quad \mathbf{G}_{1 \times 2} = [410 \ 250]$$

In each case, the first number is the factory's revenues, and the second is its costs. Compute the total revenues and costs using matrix addition.

$$\mathbf{F} + \mathbf{G} = [200 + 410 \ 300 + 250] = [610 \ 550]$$

```
mF = {{200, 300}}  
mG = {{410, 250}}  
mF + mG
```



Recall that another lecture explored some properties of common arithmetical operations, including the addition of real numbers. Given our definition of matrix addition, the addition of real matrices can inherit properties directly from the underlying real numbers. Specifically, the real $R \times K$ matrices form a group under the binary operation of matrix addition. The additive identity is the conformable *zero* matrix ($\mathbf{0}_{R \times K}$, where every element is 0), which is often called the *null* matrix. Write the additive inverse of matrix $\mathbf{A}$ as $-\mathbf{A}$, and produce it by replacing each element of $\mathbf{A}$ with its additive inverse. In sum, for any $\mathbf{A}, \mathbf{B}, \mathbf{C}$ within this set of matrices, addition:

- is associative, so that $(\mathbf{A} + \mathbf{B}) + \mathbf{C} = \mathbf{A} + (\mathbf{B} + \mathbf{C})$.
- has an identity $\mathbf{0}$, such that $\mathbf{A} + \mathbf{0} = \mathbf{A}$ and $\mathbf{0} + \mathbf{A} = \mathbf{A}$.
- has an inverse $-\mathbf{A}$, such that $\mathbf{A} + (-\mathbf{A}) = \mathbf{0}$ and $(-\mathbf{A}) + \mathbf{A} = \mathbf{0}$.
- is commutative, so that $\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A}$.

□

Remark 1.3 The additive identity for matrices of a given shape is unique. To see this, suppose A and 0 have the same shape, and A is an additive identity. Then $A = 0 + A$ (since 0 is an additive identity), but also $0 = 0 + A$ (since A is an additive identity). So $A = 0$.

Remark 1.4 Construct the additive inverse of A from the additive inverses of its elements: $-A = [-a_{r,k}]$. This inverse of A is unique. To see this, suppose B and C be inverses for A . Then $B = B + 0 = B + (A + C) = (B + A) + C = 0 + C = C$.

□

Since matrix addition is an elementwise operation, the *associativity* of matrix addition follows from the associativity of addition of the matrix elements. For example, the addition of real numbers is associative, so the addition of real matrices is associative.

$$\begin{aligned}(A + B) + C &= [(a_{r,k} + b_{r,k}) + c_{r,k}] \\ &= [a_{r,k} + (b_{r,k} + c_{r,k})] = A + (B + C)\end{aligned}$$

Similarly, from the commutativity of the addition of real numbers, we know that the addition of real matrices is commutative.

$$A + B = [a_{r,k} + b_{r,k}] = [b_{r,k} + a_{r,k}] = B + A$$

Example 1.6 (Commutativity of Addition) Given $A = \begin{bmatrix} 4 & 7 \\ 3 & 2 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 5 \\ 6 & 8 \end{bmatrix}$, show that $A + B = B + A$.

$$\begin{aligned}A + B &= \begin{bmatrix} 4+1 & 7+5 \\ 3+6 & 2+8 \end{bmatrix} = \begin{bmatrix} 5 & 12 \\ 9 & 10 \end{bmatrix} \\ B + A &= \begin{bmatrix} 1+4 & 5+7 \\ 6+3 & 8+2 \end{bmatrix} = \begin{bmatrix} 5 & 12 \\ 9 & 10 \end{bmatrix}\end{aligned}$$

```
mA = {{4,7},{3,2}}
mB = {{1,5},{6,8}}
mA + mB == mB + mA
```

Remark 1.5 Write $A - B$ instead of $A + (-B)$. This conventional notation aligns with an elementwise understanding of matrix subtraction.

Example 1.7 (Matrix Subtraction) Given $A = \begin{bmatrix} 4 & 7 \\ 3 & 2 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 5 \\ 6 & 8 \end{bmatrix}$, compute $A - B$.

$$A - B = \begin{bmatrix} 4-1 & 7-5 \\ 3-6 & 2-8 \end{bmatrix} = \begin{bmatrix} 3 & 2 \\ -3 & -6 \end{bmatrix}$$

□

1.2.2 Linear Combinations

This subsection defines scalar multiplication and linear combinations for matrices. Any linear combination of matrices of the same shape will be another matrix of that shape. In fact, matrices of a given shape have all the core vector properties: any $R \times K$ real matrix is an element of the vector space of all real matrices with the same shape.

□

Scalar multiplication is elementwise: $\lambda \cdot \mathbf{A}$ multiplies every element of the matrix $\mathbf{A}$ by the scalar λ . That is, the $\langle r, k \rangle$ -th element of $\lambda \cdot \mathbf{A}$ is $\lambda \cdot a_{r,k}$. Compactly summarize this as follows.

$$\lambda \cdot \mathbf{A} = [\lambda \cdot a_{r,k}] \tag{1.3}$$

A scalar for a real matrix is just any real number. It follows that scaling by 1 leaves the matrix unchanged: $1 \cdot \mathbf{A} = \mathbf{A}$. Since $(\alpha\beta)a_{r,k} = \alpha(\beta a_{r,k})$, repeated scaling has a natural property, known as the *associative property of scalar multiplication*.

$$\lambda_2(\lambda_1 \mathbf{A}) = (\lambda_2 \lambda_1) \mathbf{A} \tag{1.4}$$

Remark 1.6 Just as with ordinary multiplication of real numbers, mathematical discussions of scalar multiplication often suppress the operator.

Remark 1.7 Note that $-\mathbf{A} = -1 \cdot \mathbf{A}$, since $-1 \cdot \mathbf{A} = [-a_{r,k}]$.

**Example 1.8 (Scalar Multiplication)**

$$2 \cdot \begin{bmatrix} 1 & 5 \\ 6 & 8 \end{bmatrix} = \begin{bmatrix} 2 \cdot 1 & 2 \cdot 5 \\ 2 \cdot 6 & 2 \cdot 8 \end{bmatrix} = \begin{bmatrix} 2 & 10 \\ 12 & 16 \end{bmatrix}$$



Example 1.9 A firm has a factory that reports a row vector holding the factory's revenues and costs. This year the values reported are $F = [200 \ 300]$. The first number is the factory's revenues, and the second is its costs.

Using scalar multiplication, compute the revenues and costs after these grow by 10% two years in a row. After one year, revenues and costs are

$$1.10 \cdot F = 1.10 \cdot [200 \ 300] = [220 \ 330]$$

After the second year, revenues and costs are

$$1.10 \cdot (1.10 \cdot F) = 1.10 \cdot [220 \ 330] = [242 \ 363]$$

Equivalently, revenues and costs are

$$(1.10 \cdot 1.10) \cdot F = 1.21 \cdot [200 \ 300] = [242 \ 363]$$

```
mF = {{200, 300}}
```

```
1.10 * mF
```

```
1.10^2 * mF
```

□

After defining scalar multiplication and matrix addition, it becomes possible to form any linear combination of matrices that have the same shape.

$$\alpha \mathbf{A} + \beta \mathbf{B} = [\alpha a_{r,k} + \beta b_{r,k}] \quad (1.5)$$

Furthermore, since $(\alpha + \beta)a_{r,k} = \alpha a_{r,k} + \beta a_{r,k}$, the definitions of scalar multiplication and matrix addition imply is a natural distributive law for scalar multiplication. And similarly scalar multiplication distributes across matrix addition. The two scalar distribution laws summarize these observations.

$$\begin{aligned} (\alpha + \beta) \mathbf{A} &= \alpha \mathbf{A} + \beta \mathbf{A} \\ \alpha(\mathbf{A} + \mathbf{B}) &= \alpha \mathbf{A} + \alpha \mathbf{B} \end{aligned} \quad (1.6)$$



Remark 1.8 Scientific programming languages often define elementwise multiplication and division for arrays with the same shape. These are not standard matrix operations.

```
mA = {{4, 7},{3, 2}}  
mB= {{1, 5},{6, 8}}  
mA - mB //MatrixForm
```

Additional elementwise array (but not matrix) operations:

```
mA * mB //MatrixForm  
mA / mB //MatrixForm (* be careful of zeros *)
```



1.2.3 Matrix Multiplication

Matrix multiplication is *not* the elementwise product of two matrices. Instead it is defined in a way that is initially surprising but ultimately proves very useful. Given two matrices $A_{R \times N}$ and $B_{N \times K}$, we define the *matrix product* $A \cdot B$ as follows. If $C = A \cdot B$, then

$$c_{r,k} = \sum_{n=1}^N a_{r,n} b_{n,k} \quad (1.7)$$

See Figure 1.2 for an illustration of this operation. In the product $A \cdot B$, we call A the *premultiplying* matrix and B the *postmultiplying* matrix. To be *conformable* for multiplication, the premultiplying matrix must have the same number of columns as the postmultiplying matrix has rows.

Remark 1.9 The matrix product $A \cdot B$ is also called a matrix *dot product* of A and B . In mathematical presentations (but not in code), it is very common to suppress the operator and just write AB .²

²Do not assume a programming language will support implicit matrix multiplication, even if it supports implicit multiplication. In **WL**, for example, a space between two matrices is an implicit **Times** command, which produces an elementwise (Hadamard) product, not matrix multiplication. It is generally best to be explicit in your code.

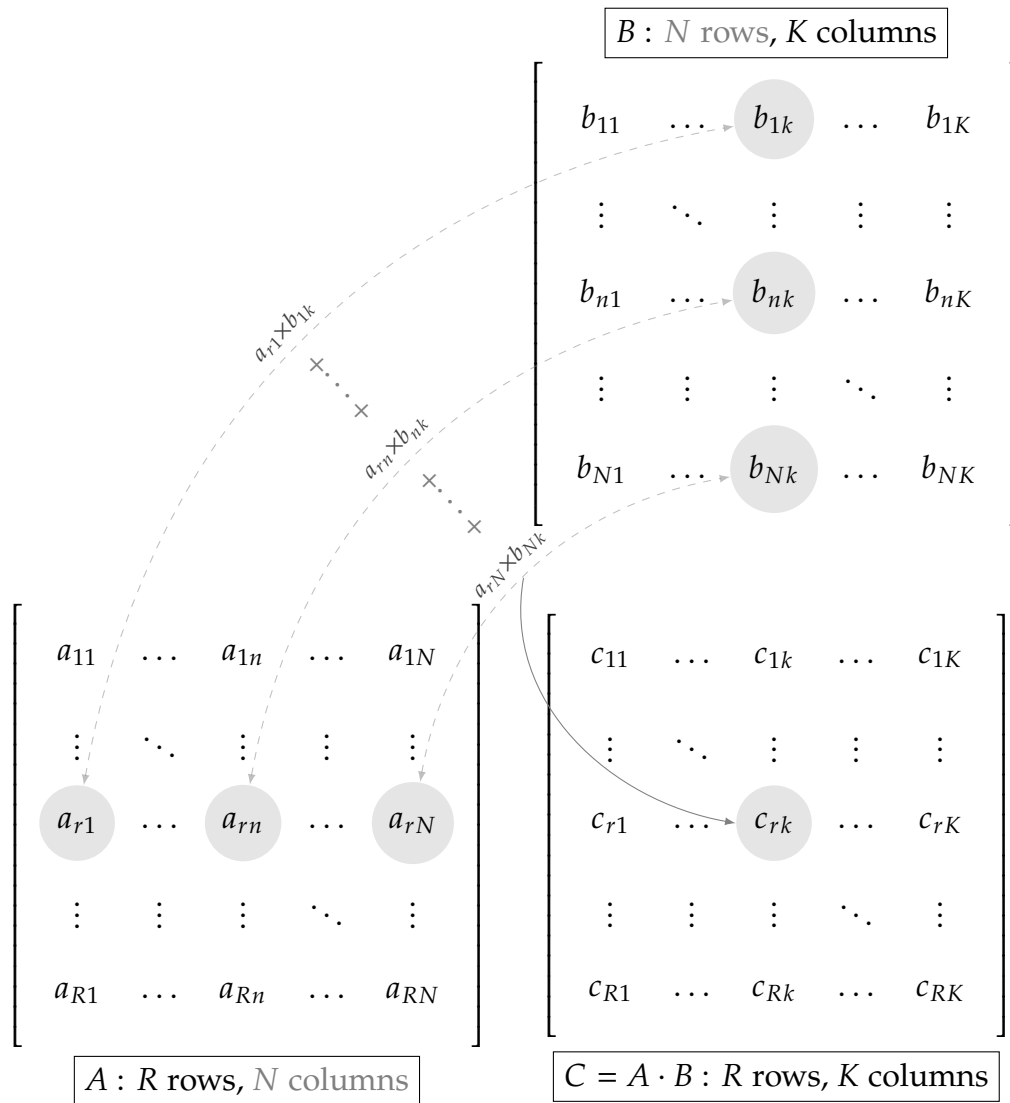


Figure 1.2: Falk's Scheme

Falk's scheme for matrix multiplication illustrates how to produce the r, k -th element of a matrix product from the r -th row of the premultiplying matrix and the k -th column of the postmultiplying matrix.

□

The definition (1.7) implies that each element of $A \cdot B$ is an ordinary dot product of a *row* of the premultiplying matrix and a *column* of the postmultiplying matrix. The r, k -th element of the matrix product $A \cdot B$ is the dot product of the r -th row of A and the k -th column of B . Let $A_{r,:}$ be the vector representing the r -th row of A , and let $B_{:,k}$ be the vector representing the k -th column of B . Then (1.8) expresses this more succinct understanding of the matrix product.

$$C = A \cdot B \implies c_{r,k} = A_{r,:} \cdot B_{:,k} \quad (1.8)$$

Example 1.10 (Matrix Products) Let $A = \begin{bmatrix} 2 & 3 \\ 5 & 6 \end{bmatrix}$, $B = \begin{bmatrix} 8 & 1 \\ 4 & 2 \end{bmatrix}$, and $C = \begin{bmatrix} 8 \\ 4 \end{bmatrix}$. Compute $A \cdot B$ and $A \cdot C$.

$$AB = \begin{bmatrix} 2 & 3 \\ 5 & 6 \end{bmatrix} \begin{bmatrix} 8 & 1 \\ 4 & 2 \end{bmatrix} = \begin{bmatrix} \langle 2, 3 \rangle \cdot \langle 8, 4 \rangle & \langle 2, 3 \rangle \cdot \langle 1, 2 \rangle \\ \langle 5, 6 \rangle \cdot \langle 8, 4 \rangle & \langle 5, 6 \rangle \cdot \langle 1, 2 \rangle \end{bmatrix} = \begin{bmatrix} 28 & 8 \\ 64 & 17 \end{bmatrix}$$

In this first case, a 2×2 matrix premultiplies a 2×2 matrix, yielding a 2×2 matrix. In the next case, a 2×2 matrix premultiplies a 2×1 matrix, yielding a 2×1 matrix.

$$AC = \begin{bmatrix} 2 & 3 \\ 5 & 6 \end{bmatrix} \begin{bmatrix} 8 \\ 4 \end{bmatrix} = \begin{bmatrix} \langle 2, 3 \rangle \cdot \langle 8, 4 \rangle \\ \langle 5, 6 \rangle \cdot \langle 8, 4 \rangle \end{bmatrix} = \begin{bmatrix} 28 \\ 64 \end{bmatrix}$$

Note that because the first (and only) column of C has the same elements as the first column of B , the first (and only) column of AC has the same elements as the first column of AB .

Exercise: Do it in WL!

□ Consider the matrix product $A \cdot B$.

- Is the j -th column of $A \cdot B$ a weighted sum of the columns of A ? Explain.
- Is the i -th row of $A \cdot B$ a weighted sum of the rows of B ? Explain.
- Suppose every element of row i of A is 0. Why is every element of row i of $A \cdot B$ also 0?
- Suppose every element of column j of B is 0. Why is every element of column j of $A \cdot B$ also 0?

A *positive matrix* is a matrix where every element is positive. If A and B are conformable positive matrices, is $A \cdot B$ also positive?

A *nonnegative matrix* is a matrix where every element is nonnegative. If A and B are conformable nonnegative matrices, is $A \cdot B$ also nonnegative? In this case, if neither A or B are positive, can $A \cdot B$ be positive? Example?

Harder:

Suppose A and B are conformable nonnegative matrices, and $A \cdot B$ is positive. Why must $A \cdot (A \cdot B)$ be positive?



1.2.3.1 Identity Matrices

A scalar matrix where the diagonal value is 1 is called an *identity* matrix. To see why, let $\mathbf{I}_N$ be the identity matrix of order N , and show that $\mathbf{I}_R \mathbf{A}_{R \times K} = \mathbf{A}$ and $\mathbf{A}_{R \times K} \mathbf{I}_K = \mathbf{A}$. We say $\mathbf{I}_R$ is a *left identity* for $\mathbf{A}$, and $\mathbf{I}_K$ is a *right identity* for $\mathbf{A}$. If $\mathbf{A}$ is a square matrix of order N , then $\mathbf{I}_N$ is a multiplicative identity for $\mathbf{A}$. That is, it is both a left identity and a right identity. Here are three identity matrices:

$$\mathbf{I}_1 = [1] \qquad \mathbf{I}_2 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \qquad \mathbf{I}_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

`MatrixForm` `/@ IdentityMatrix` `/@ Range`[\[3\]](#)



Example 1.11 (Multiplicative Identity) Consider the real matrix $A = \begin{bmatrix} a & b & c \\ d & e & f \end{bmatrix}$.

$$\begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \cdot A = \begin{bmatrix} 1 \cdot a + 0 \cdot d & 1 \cdot b + 0 \cdot e & 1 \cdot c + 0 \cdot f \\ 0 \cdot a + 1 \cdot d & 0 \cdot b + 1 \cdot e & 0 \cdot c + 1 \cdot f \end{bmatrix} = A$$

$$A \cdot \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} a \cdot 1 + b \cdot 0 + c \cdot 0 & a \cdot 0 + b \cdot 1 + c \cdot 0 & a \cdot 0 + b \cdot 0 + c \cdot 1 \\ d \cdot 1 + b \cdot 0 + c \cdot 0 & a \cdot 0 + e \cdot 1 + c \cdot 0 & a \cdot 0 + b \cdot 0 + f \cdot 1 \end{bmatrix} = A$$

□

1.2.3.2 Associative, Not Commutative

From our work on matrix multiplication, we know how to produce $\langle r, k \rangle$ -th element of $(AB)C$ as the product of the r -th row of AB and the k -th column of C . This is

$$(A_{r,:}B)C_{:,k} = \left[\sum_j \left(\sum_i a_{ri}b_{ij} \right) c_{jk} \right] \quad (1.9)$$

Similarly, produce $\langle r, k \rangle$ -th element of $A(BC)$ as the product of the r -th row of A and the k -th column of BC .

$$A_{r,:}(BC_{:,k}) = \left[\sum_i a_{ri} \left(\sum_j b_{ij}c_{jk} \right) \right] \quad (1.10)$$

The right sides of (1.9) and (1.10) differ only in the order of summation, so they are equal. Therefore

$$(A_{r,:}B)C_{:,k} = A_{r,:}(BC_{:,k}) \quad (1.11)$$

Since this is true for every r and k , we conclude that the matrix product is associative.

$$(AB)C = A(BC) \quad (1.12)$$

The associativity of matrix multiplication means that the order in which we elect to perform adjacent matrix multiplications is irrelevant. As a result, as with the multiplication of numbers, we often simply omit parentheses for matrix products. For example, write $(AB)C$ or $A(BC)$ simply as ABC unless there is a reason to highlight the order of operations.



Example 1.12 Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$, $B = \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix}$, and $C = \begin{bmatrix} 9 & 10 \\ 11 & 12 \end{bmatrix}$. Calculate the element in the second row and first column of ABC without calculating any other elements of this product. Plan to multiply the second row of AB times the first column of C . The second row of AB in turn can be calculated by multiplying the second row of A times B , as follows.

$$\begin{bmatrix} 3 & 4 \end{bmatrix} \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix} = \begin{bmatrix} 43 & 50 \end{bmatrix}$$

Therefore, calculate the desired element as

$$\begin{bmatrix} 43 & 50 \end{bmatrix} \begin{bmatrix} 9 \\ 11 \end{bmatrix} = 937$$

```
{mA,mB,mC} = Reshape[Range[12],{3,2,2}]  
(mA[[2,A11]] . mB)  
% . mC[[A11,1]]
```



Matrix multiplication is *not commutative*: the order of the operands generally affects the matrix product. For one thing, there is an issue of conformity for multiplication: BA may not be defined even if AB is. Recall that the definition of matrix multiplication requires that the premultiplying matrix has the same number of columns as the postmultiplying matrix has rows. Recall that matrices with such compatible shapes are said to be *conformable* for multiplication.

The matrix product has as many rows as the premultiplying matrix and as many columns as the postmultiplying matrix. When the shape of A is $R \times N$ and the shape of B is $N \times K$, then the shape of AB is $R \times K$. Even if $R = K$ so that both AB and BA are defined, they typically yield matrices of different dimensions. And even when both matrices are square, AB will generally differ from BA .



Example 1.13 (Matrix Multiplication) Let $A = \begin{bmatrix} 1 & 2 \end{bmatrix}$ and $B = \begin{bmatrix} 3 \\ 4 \end{bmatrix}$. Form the matrix products AB and BA as follows:

$$AB = \begin{bmatrix} 1 & 2 \end{bmatrix} \begin{bmatrix} 3 \\ 4 \end{bmatrix} = \begin{bmatrix} 1 \cdot 3 + 2 \cdot 4 \end{bmatrix} = \begin{bmatrix} 11 \end{bmatrix}$$
$$BA = \begin{bmatrix} 3 \\ 4 \end{bmatrix} \begin{bmatrix} 1 & 2 \end{bmatrix} = \begin{bmatrix} 3 \cdot 1 & 3 \cdot 2 \\ 4 \cdot 1 & 4 \cdot 2 \end{bmatrix} = \begin{bmatrix} 3 & 6 \\ 4 & 8 \end{bmatrix}$$

Example 1.14 Let $A = \begin{bmatrix} 4 & 7 \\ 3 & 2 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 5 \\ 6 & 8 \end{bmatrix}$. Then AB and BA can be calculated as

$$AB = \begin{bmatrix} 46 & 76 \\ 15 & 31 \end{bmatrix}, \quad BA = \begin{bmatrix} 19 & 17 \\ 48 & 58 \end{bmatrix}$$

□

Some special square matrices do commute under multiplication. Diagonal matrices provide a particularly simple example. Suppose $\mathbf{A}$ and $\mathbf{B}$ are real $N \times N$ diagonal matrices. As we have seen above, the $\langle r, k \rangle$ -th element of $\mathbf{A} \cdot \mathbf{B}$ is $\mathbf{A}_{r,:} \cdot \mathbf{B}_{:,k}$. Since the only nonzero elements are on the diagonals, this product must be zero whenever $r \neq k$. This means that the product is a diagonal matrix, with $a_{i,i}b_{i,i}$ as the i -th diagonal element. Since the multiplication of numbers commutes, the product of diagonal matrices also commutes.

Example 1.15 Let $\mathbf{A} = \begin{bmatrix} a_{11} & 0 \\ 0 & a_{22} \end{bmatrix}$ and $\mathbf{B} = \begin{bmatrix} b_{11} & 0 \\ 0 & b_{22} \end{bmatrix}$.

$$\mathbf{A} \cdot \mathbf{B} = \begin{bmatrix} a_{11}b_{11} & 0 \\ 0 & a_{22}b_{22} \end{bmatrix} = \begin{bmatrix} b_{11}a_{11} & 0 \\ 0 & b_{22}a_{22} \end{bmatrix} = \mathbf{B} \cdot \mathbf{A}$$

□

In general, however, matrix multiplication is not commutative. Correspondingly, there are two separate *distributive laws* for matrices: left distribution (for premultiplication) and right distribution (for postmultiplication).

$$\begin{aligned} A(B + C) &= AB + AC \\ (B + C)A &= BA + CA \end{aligned} \tag{1.13}$$

Here is a brief proof of the left distributive law. Note that $(B + C)_{:,k} = B_{:,k} + C_{:,k}$: the k -th column of the sum is the sum of the k -th columns. Consider the $\langle r, k \rangle$ -th element of $A(B + C)$.

$$\begin{aligned} A_{r,:} \cdot (B_{:,k} + C_{:,k}) &= \sum_{i=1}^I [a_{ri}(b_{ik} + c_{ik})] \\ &= \sum_{i=1}^I (a_{ri}b_{ik} + a_{ri}c_{ik}) \\ &= \sum_{i=1}^I a_{ri}b_{ik} + \sum_{i=1}^I a_{ri}c_{ik} \\ &= A_{r,:} \cdot B_{:,k} + A_{r,:} \cdot C_{:,k} \end{aligned} \tag{1.14}$$



1.3 Multiplicative Inverse

Given a square matrix A , suppose we can find a matrix B of the same order such that $BA = I$ and $AB = I$. We then say that B is a multiplicative inverse for A . If C is also an inverse for A , then $C = C(AB) = (CA)B = IB = B$. That is, any inverse for A is unique. We therefore adopt a special notation for this inverse: A^{-1} .

Definition 1.4 A square matrix A is *invertible* if there is a matrix A^{-1} such that $A^{-1}A = AA^{-1} = I$. Otherwise it is not invertible. When it exists, the matrix A^{-1} is the multiplicative *inverse* of A .

Example 1.16 (Invertibility) A zero matrix is not invertible, For example, whenever $\begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}$ is a factor in a product, the product is a zero matrix. In contrast, any identity matrix is its own inverse. More generally, any diagonal matrix is invertible, as long as its diagonal elements are nonzero: we simply use the reciprocals of the diagonal elements to create a new diagonal matrix. For example, $\begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix}^{-1} = \begin{bmatrix} 1/a & 0 \\ 0 & 1/b \end{bmatrix}$.

Remark 1.10 If A is invertible then so is its inverse, with $(A^{-1})^{-1} = A$. If A and B are invertible then $(AB)^{-1} = B^{-1}A^{-1}$.

Exercise 1.1 Prove that the product of invertible matrices is invertible.

□

Although finding the inverse of larger matrices is more work, finding the inverse of an invertible real 2×2 matrix also follows the very simple rule in (1.15).

$$\mathbf{A} = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \implies \mathbf{A}^{-1} = \frac{1}{ad - bc} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix} \quad (1.15)$$

This rule makes use of scalar multiplication, where the scalar is the fraction $1/(ad - bc)$. The number $ad - bc$ is the *determinant* of this 2×2 matrix, denoted by $\det[\mathbf{A}]$ or $|\mathbf{A}|$. The determinant determines whether we can use this formula to produce an inverse. If $ad = bc$ then $|\mathbf{A}| = 0$ this formula fails. In this case, matrix does not have a multiplicative inverse. (We will return to this claim.) In every other case, the matrix $\mathbf{A}$ is invertible, and the simple rule in (1.15) produces its inverse.

□

Here is one way to see that any 2×2 matrix A with $\det[A] = 0$ does not have an inverse. First show that the determinant of a product of 2×2 matrices is the product of the determinants. (This result is generalized in another lecture.) Let $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ and $B = \begin{bmatrix} e & f \\ g & h \end{bmatrix}$. Then

$$\begin{aligned} |B \cdot A| &= \left| \begin{bmatrix} ea + fc & eb + fd \\ ga + hc & gb + hd \end{bmatrix} \right| \\ &= eahd + fcgb - ebhc - fdga \\ &= (eh - fg)(ad - bc) = |B| |A| \end{aligned} \tag{1.16}$$

Next note that if B is an inverse matrix for A then

$$|B \cdot A| = |I_2| = 1 \tag{1.17}$$

It is impossible to satisfy $|B| |A| = 1$ when $|A| = 0$.

Exercise 1.2 (Determinant of Product is Product of Determinants) Given any numerical matrices $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ and $B = \begin{bmatrix} e & f \\ g & h \end{bmatrix}$, show that $|AB| = |A| |B|$.

We will soon define the determinant of larger matrices and prove this same result.

□

The matrix $\begin{bmatrix} d & -b \\ -c & a \end{bmatrix}$ is the *adjugate* of the matrix A . This is the 2×2 matrix on the right of (1.15). This book uses the notation $A^\#$ or $\text{adj}[A]$ to denote the adjugate of A . With this notation, assuming the determinant is nonzero, produce an inverse matrix from the matrix determinant and the adjugate matrix as follows.

$$A^{-1} = \frac{1}{|A|} A^\# \quad (1.18)$$

Exercise 1.3 Assuming a nonzero determinant, define A and A^{-1} as in (1.15). Show that $A^{-1}A = I_2$ and that $AA^{-1} = I_2$.

$$A^{-1} = \frac{1}{|A|} A^\#$$

```
Clear[a,b,c,d]
mA = {{a,b},{c,d}};
Det[mA]
Adjugate[mA]
Inverse[mA] == (1/Det[mA]) * Adjugate[mA]
```



Example 1.17 If $A = \begin{bmatrix} 2 & 4 \\ 6 & 8 \end{bmatrix}$ then $|A| = -8$ and $A^\# = \begin{bmatrix} 8 & -4 \\ -6 & 2 \end{bmatrix}$. Compute the matrix inverse as $(1/|A|) \cdot A^\#$:

$$A^{-1} = \frac{1}{-8} \begin{bmatrix} 8 & -4 \\ -6 & 2 \end{bmatrix} = \begin{bmatrix} -1.0 & 0.50 \\ 0.75 & -0.25 \end{bmatrix}$$

Check that $A^{-1}A = I_2$ and that $AA^{-1} = I_2$. For example,

$$A^{-1}A = \begin{bmatrix} -2 + 3 & -4 + 4 \\ 1.5 - 1.5 & 3 - 2 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$



1.3.1 Diagonal Inverse

Even when it exists, finding the inverse of a matrix can be a substantial challenge. However, the inverse of a real diagonal matrix follows a simple rule. If there are any zeros on the diagonal, there is no inverse. Otherwise, produce the inverse matrix by replacing each diagonal element with its reciprocal. This rule reflects the observation in section 1.2.3.2 that multiplying two diagonal matrices produces a diagonal matrix whose diagonal is just the elementwise products of the matrix diagonals.

Example 1.18 (Inverse of Diagonal Matrix)

$$\begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 4 \end{bmatrix}^{-1} = \begin{bmatrix} 1/2 & 0 & 0 \\ 0 & 1/3 & 0 \\ 0 & 0 & 1/4 \end{bmatrix}$$

Define the determinant of a diagonal matrix D to be the product of the elements on its diagonal. Write this as $\det[D] = \prod_i d_{i,i}$. (Another lecture returns to this.) This determines whether it is possible to find an inverse of the matrix. If any diagonal element is zero then so is the determinant, and it is not possible. Otherwise, it possible.

Momentarily considering only *diagonal* matrices A and B , note that the determinant of their product is the product of their determinants.

$$\det[A \cdot B] = \prod_i a_{i,i} b_{i,i} = \left(\prod_i a_{i,i} \right) \left(\prod_i b_{i,i} \right) = \det[A] \det[B] \quad (1.19)$$



1.3.2 Cancellation

We have found that some square matrices are invertible but others are not. Let us gather together all of the invertible 2×2 real matrices in a set, which is called the *general linear group* of 2×2 matrices over $\mathbb{R}$ (or $\text{GL}[2, \mathbb{R}]$). We have already seen that the product of any two elements of this set is also in this set: if two matrices are invertible, their product is invertible. If we consider the binary operation of matrix multiplication in this set, we discover some familiar properties. For any members of this set we find that multiplication is associative, I_2 is the multiplicative identity, and by definition A^{-1} is in the set for any A in the set. These are the group properties. In short, this set is a group under the operation of matrix multiplication.

□

The cancellation property is implied by the group properties, and this is particularly useful. For A, B, C in this set, if $CA = CB$ or $AC = BC$, then $A = B$. To see why, consider the first case.

$CA = CB$	given
$C^{-1}(CA) = C^{-1}(CB)$	multiply
$(C^{-1}C)A = (C^{-1}C)B$	associate
$IA = IB$	inverse
$A = B$	identity

The cancellation property is obviously very useful, but notice that we used all the group properties to derive it. This is a clue that we cannot expect cancellation to hold when we are not dealing with invertible matrices, just as $0 \cdot a = 0 \cdot b$ does not imply $a = b$ in the real numbers.



Example 1.19 Consider the following three matrices:

$$A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \quad B = \begin{bmatrix} -1 & 2 \\ 4 & 4 \end{bmatrix} \quad C = \begin{bmatrix} 1 & 2 \\ 2 & 4 \end{bmatrix}$$

We see that

$$CA = \begin{bmatrix} 7 & 10 \\ 14 & 20 \end{bmatrix} \quad CB = \begin{bmatrix} 7 & 10 \\ 14 & 20 \end{bmatrix}$$

Thus we have $CA = CB$ even though $A \neq B$. Note that the matrix C does not have an inverse.



1.3.3 System Solutions

Inverses offer one approach to solving certain systems of equations. Suppose we can set up a system as $A\mathbf{x} = \mathbf{b}$, where A is a known square matrix of order N , $\mathbf{x}$ is a conformable column vector of unknowns, and $\mathbf{b}$ is a known column vector exogenous constants. Then if A is invertible, we can solve for $\mathbf{x} = A^{-1}\mathbf{b}$.

Example 1.20 Consider the famous ball and bat problem (Meyer and Frederick, 2023). A ball and a bat cost \$1.10 in total, and the bat costs \$1.00 more than the ball. What is the price of the ball? Set this up as a matrix equation and then solve with an inverse.

$$\begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} p_{bat} \\ p_{ball} \end{bmatrix} = \begin{bmatrix} 1.10 \\ 1.00 \end{bmatrix}$$

Note that $\begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}^{-1} = \begin{bmatrix} 1/2 & 1/2 \\ 1/2 & -1/2 \end{bmatrix}$ so

$$\begin{bmatrix} p_{bat} \\ p_{ball} \end{bmatrix} = \begin{bmatrix} 1/2 & 1/2 \\ 1/2 & -1/2 \end{bmatrix} \begin{bmatrix} 1.10 \\ 1.00 \end{bmatrix}$$

The bat costs \$1.05 and the ball costs \$0.05.

```
(* basic check of work *)
eq1 = pbat + pball == 1.10;
eq2 = pbat - pball == 1.00;
Solve[eq1 && eq2, {pbat, pball}]
(* matrix approach *)
mA = {{1, 1}, {1, -1}}; vb = {1.10, 1.00};
Inverse[mA] . vb
```

□ Solve for $\vec{x}$ with an inverse.

$$\begin{bmatrix} 1 & 2 \\ -3 & 4 \end{bmatrix} \cdot \vec{x} = \begin{bmatrix} 20 \\ 40 \end{bmatrix}$$



Differential Calculus: Brief Review

For these notes, see `dcalc.pdf` on Canvas.

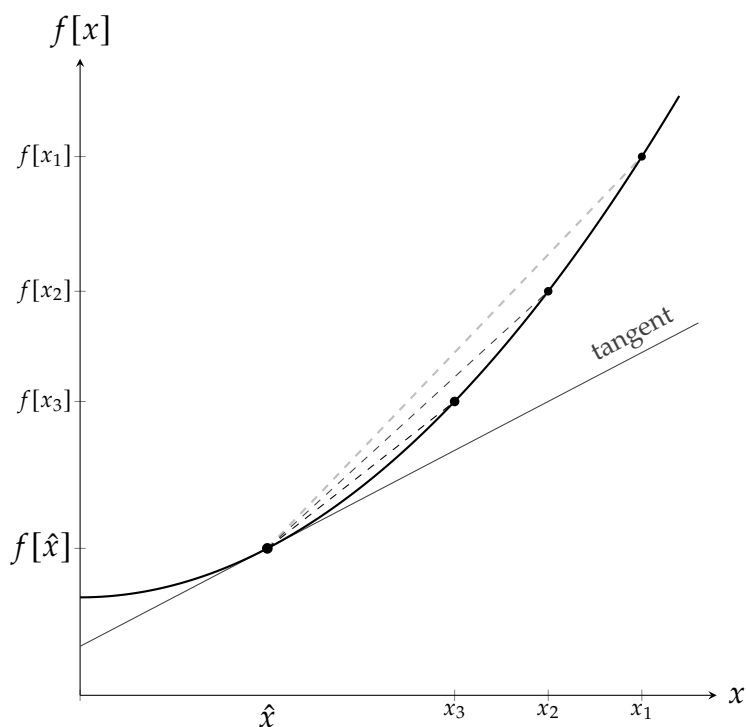


Figure 1.1: Slope at a Point

As the step size $(|x_i - \bar{x}|)$ shrinks, the secant slope approaches the slope of the tangent line at $\hat{x}$.

Average rates of change vs instantaneous rate of change.



Definition 1.1 The real function f defined on $(a \dots b)$ is *differentiable* at $x_0 \in (a \dots b)$ iff the following limit exists.

$$\lim_{x \rightarrow x_0} \frac{f[x] - f[x_0]}{x - x_0}$$

This limiting value is the *derivative* of f at x_0 , typically written as $f'[x_0]$.

$f = x \mapsto x * x$

$x_0 = 2$

Limit $[(f[x] - f[x_0]) / (x - x_0), x \rightarrow x_0]$

$f'[x_0]$

D $[f[x], x] \ /. \ \{x \rightarrow x_0\}$



1.1.2.1 Difference Quotient

Define the stepsize $h = x - x_0$, and substitute in (1.1) in order to rewrite the secant-slope at x_0 as a function of h .

$$h \mapsto \frac{f[x_0 + h] - f[x_0]}{h} \quad (1.2)$$

This is known as the difference quotient at x_0 . Any particular value of h is the *stepsize* for the difference quotient. A positive stepsize produces a *forward* difference quotient, while a negative stepsize produces a *backward* difference quotient.

A function f is *differentiable* at x_0 if the limit exists as the stepsize approaches zero. Wherever f is a differentiable function, define the *derivative function* f' as follows.

$$f' = x \mapsto \lim_{h \rightarrow 0} \frac{f[x + h] - f[x]}{h} \quad (1.3)$$

Remark 1.2 Another common expression for the difference quotient uses dx in place of h . The derivative of f at x is then written as follows.

$$f' = x \mapsto \lim_{dx \rightarrow 0} \frac{f[x + dx] - f[x]}{dx}$$

The limit operation binds its variable, so both notations (h or dx) are completely equivalent. This book uses both notations.



The `DifferenceQuotient` command works with an *expression*, not with a function.

```
Clear[f, x, h]
"difference quotient (general):"
dq = DifferenceQuotient[f[x], {x, h}]
"specialize to squaring fn:"
f = x  $\mapsto$  x^2 (* global *)
dq
"also specify x"
Block[{x = 0.5},
  dq
]
"also specify h"
sslope = Block[{x = 0.5, h = 1.0},
  dq
]
" plot fn and secant line:"
Plot[{f[x], f[x0] + sslope (x - x0)}, {x, 0, 2}]
```



Remark 1.3 (Differentiability from the Right and Left) It is sometimes useful to define the left derivative and the right derivative by changing the limit in (1.3) to consider only backward differences or only forward differences. When f is differentiable at x , both the left derivative and right derivative equal $f'[x]$.

Example 1.1 Consider the identity function $x \mapsto x$ on the closed real interval $[0 .. 1]$. This function has a right derivative at 0 and a left derivative at 1. (Both equal 1.)

Example 1.2 The absolute value function on the real numbers is differentiable everywhere except at 0, where it has a kink. At 0 the left derivative is -1 and the right derivative is 1 .

Example 1.3 Recall from another lecture that the cumulative distribution function for the standard uniform distribution is

$$\text{cdf} = x \mapsto \begin{cases} 0 & x < 0 \\ x & 0 \leq x \leq 1 \\ 1 & 1 < x \end{cases}$$

This is differentiable everywhere except at 0 and 1, where the function is kinked. At 0 the left derivative is 0 and the right derivative is 1. At 1 the left derivative is 1 and the right derivative is 0.



1.1.3 Continuity and More

Differentiable functions have a nice feature: there is a clear way to determine a unique tangent line at a point. Many functions are not differentiable everywhere. Consider for example point x_1 in Figure 1.2, since the function graph has a jump discontinuity. Here, there is no unique concept of a tangent line. Reflecting this, note that at x_1 the value of a backwards secant slope ($h < 0$ or equivalently $x < x_1$) does not approach the value of a forward secant slope ($h > 0$ or equivalently $x > x_1$). This is true no matter how close we insist on being to x_1 . Since the limit from the below does not equal the limit from above, this function does not have a limit at x_1 .

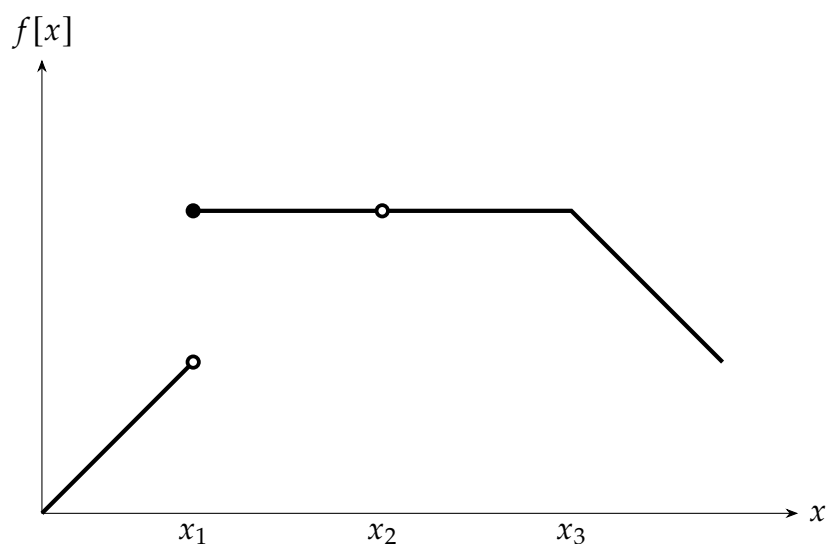


Figure 1.2: Tangency Problems

There is no concept of a tangency to the function graph at a jump or even at a removable discontinuity. Furthermore, there is also no unique tangent to the curve at a kink. Thus the function f is not differentiable at the points x_1 , x_2 , or x_3 .

A difficulty also arises at point x_2 of Figure 1.2, where the function has a removable discontinuity. Although the function has a limit at x_2 , it is not defined at x_2 . This discontinuity breaks differentiability, even though the function's slope is zero on either

side of x_2 . In fact, we cannot even define a secant slope at x_2 . It seems that differentiability requires continuity.



Theorem 1.1 (Differentiability Implies Continuity) Consider a real-valued function f defined on an open interval $(a .. b)$. If f is differentiable at $x \in (a .. b)$, then f must be continuous at x .

Proof: Work with the difference quotient $(f[x + h] - f[x])/h$, which is only defined if $f[x]$ is defined. Show that if the derivative is defined at x then the function f has the limit $f[x]$ at x . By the product rule of limits (and a well-chosen zero),

$$\begin{aligned}\lim_{h \rightarrow 0} f[x + h] &= \lim_{h \rightarrow 0} \frac{f[x + h] - f[x]}{h} h + f[x] \\ &= f'[x] \cdot 0 + f[x] = f[x]\end{aligned}$$

The function is differentiable at x , so $f'[x]$ is well defined. Therefore the function f is continuous at x , since $\lim_{h \rightarrow 0} f[x + h] = f[x]$. ■

The converse of theorem 1.1 is not true: continuity does not imply differentiability. Consider the kink at point x_3 in Figure 1.2. This means that the derivative from the left differs from the derivative from the right, so the function is not differentiable at x_3 . For the same reason, the absolute value function is not differentiable at 0, although it is continuous there.

□

So a to be differentiable, a function must be continuous. However, continuity is not sufficient for differentiability. Consider the point x_3 in Figure 1.2. Now the function is defined and continuous but has a kink. However, at x_3 , a forward difference is negative, while a small backward difference is zero. So at x_3 , no matter how small the step size becomes, the value of the forward difference quotient and backward difference quotient are always very different. At x_3 , the difference quotient (as a function of the step size) has a left limit and a right limit, but even though the function is continuous at x_3 , these two limits are not equal. The difference quotient does not have a limit at x_3 , so the function f is not differentiable at x_3 .



Differentiability is a major analytical convenience. It greatly increases the tractability of many problems. Also, differentiability is often a plausible approximation. Nevertheless, kinks and breaks can have real-world significance.

Example 1.4 (Tax Brackets) The US income tax code changes the tax rate at discrete income levels. The marginal tax rate as a function of income is discontinuous. The average tax paid by an income recipient is called the *effective* tax rate. The discrete changes in the marginal tax rate imply kinks in the effective tax function, so the relationship between income and taxes is not differentiable at those levels. Figure 1.3 uses US data to illustrate these two tax functions.

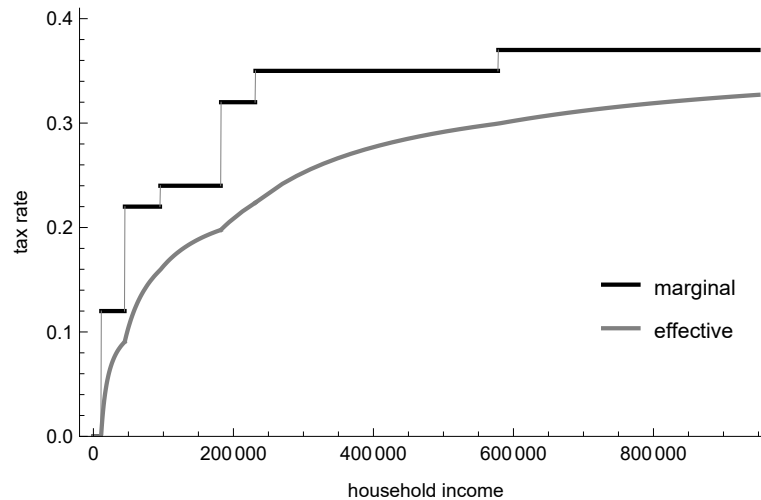


Figure 1.3: Marginal vs Effective Tax Rates

The marginal tax rate undergoes discrete changes; it is discontinuous at the points of change. The implied effective tax rate is continuous, but it has kinks; it is not differentiable at those points. Data Source: IRS 2023 marginal tax rate schedule for individuals.



1.1.3.1 L'Hôpital's Rule

When naively taking a function limit (at a point) would produce $0/0$, L'Hôpital's limit rule for indeterminate forms may be able to help.¹ The limiting expression $0/0$ is considered *indeterminant* in this case because it does not determine whether or not the limit exists. In fact, by replacing the numerator and denominator with their derivatives (at the problematic point), we may be able to produce the desired limit.

Theorem 1.2 (L'Hôpital's Rule) Let $x_t \rightarrow c$ be a convergent real sequence and let f and g be real functions that are continuously differentiable on an open interval containing c . Suppose $\lim_{x \rightarrow c} f[x] = 0$ and $\lim_{x \rightarrow c} g[x] = 0$, so that the limit of the ratio is indeterminate at c . If $g'[c] \neq 0$, then

$$\lim_{x \rightarrow c} \frac{f[x]}{g[x]} = \lim_{x \rightarrow c} \frac{f'[x]}{g'[x]}$$

Proof: The proof here is for a simple version of the rule. Recalling $f[c] = g[c] = 0$,

$$\begin{aligned} \lim_{x \rightarrow c} \frac{f[x]}{g[x]} &= \lim_{x \rightarrow c} \frac{f[x] - f[c]}{g[x] - g[c]} \\ &= \lim_{x \rightarrow c} \frac{f[x] - f[c]}{x - c} \bigg/ \frac{g[x] - g[c]}{x - c} \\ &= \frac{f'[c]}{g'[c]} \end{aligned}$$

By continuity of the derivatives, $f'[c]/g'[c] = \lim_{x \rightarrow c} f'[x]/g'[x]$. ■

² L'Hôpital's rule also works for evaluating limits at infinity or when it is the reciprocals

¹This rule is named after Guillaume de L'Hôpital, who published it in 1696 in a famous first published calculus textbook (L'Hôpital, 1696). Following Stigler's law of eponymy, it was apparently discovered by Johann Bernoulli (Boyer and Merzbach, 2011), who contributed it (and many other results) to the textbook.

²This rule is named after Guillaume de L'Hôpital, who published it in 1696 in a famous first published calculus textbook (L'Hôpital, 1696). Following Stigler's law of eponymy, it was apparently discovered by Johann Bernoulli (Boyer and Merzbach, 2011), who contributed it (and many other results) to the textbook.

of the functions that have zero limits. However, it does not apply more generally. (Look at the proof to see why.) For example, $\lim_{x \rightarrow 1} x/(1+x) = 1/2$, but the ratio of the derivative of the numerator to the derivative of the denominator is 1.

Example 1.5 Let $f = x \mapsto (6x - 6)/(x^2 - 1)$. Find the function limit at 1. Since simple substitution produces $0/0$, try the limit rule for indeterminate forms.

$$\lim_{x \rightarrow 1} f[x] = \lim_{x \rightarrow 1} \frac{6}{2x} = 3$$

Example 1.6 Recall that the monthly payment on an N -month fixed-rate loan of size P at simple *monthly* interest rate of r is

$$\text{pmt} = r \mapsto \begin{cases} \frac{rP}{1-(1+r)^{-N}} & r \neq 0 \\ P/N & r = 0 \end{cases}$$

Show these are consistent by computing the limit as r approaches 0. Since the numerator and denominator both approach 0, apply L'Hôpital's rule:

$$\lim_{r \rightarrow 0} \text{pmt}[r] = \lim_{r \rightarrow 0} \frac{P}{N(1+r)^{-(N+1)}} = P/N$$

Example 1.7 Suppose $f = x \mapsto a + bx$ and $g = x \mapsto e^{gx}$, with b and g both positive, creating problems for evaluating $\lim_{x \rightarrow \infty} f[x]/g[x]$. But $f'[x] = b$ and $g'[x] = ge^{gx}$, so the limit is zero. This is true no matter how big b is or how small g is.



1.1.4 Monotone Functions

A function f is monotone if it is always increasing or always decreasing, and it is strictly monotone if it is always strictly increasing or always strictly decreasing.

- *increasing*: $y \geq x \implies f[y] \geq f[x]$
- *strictly increasing*: $y > x \implies f[y] > f[x]$
- *decreasing*: $y \geq x \implies f[y] \leq f[x]$
- *strictly decreasing*: $y > x \implies f[y] < f[x]$

A function is monotone if it is either increasing or decreasing. A function may be monotone on a real interval even if it is not monotone on its entire domain.



It follows that if a function's derivative has an unchanging sign on an interval, it is monotone on that interval. Consider a real function that has a positive derivative everywhere on an interval. A derivative is a limit of difference quotients, so a positive derivative at a point means that the difference quotients are always positive once we are near enough to the point. That is, the difference quotient $(f[x+h] - f[x])/h$ must be positive for any x and $x+h$ in the interval. Therefore the function is strictly increasing on the interval. Naturally, a nonnegative derivative everywhere on an interval similarly implies that the function is increasing on the interval; a negative derivative everywhere implies it decreasing; and a nonpositive derivative everywhere implies that it is decreasing.

A near converse to these observations also applies: Given $x \in (a \dots b)$ and $h \neq 0$, if the difference quotient is always nonnegative (whenever $x+h$ lies in the interval), then the derivative is always nonnegative. The function is increasing in that interval. Or if this difference quotient is always nonpositive, then the derivative is always nonpositive. The function is decreasing in that interval.

It may therefore seem surprising that a function that is strictly increasing on an interval may nevertheless have a zero derivative in that interval. And symmetrically, a function that is strictly decreasing on an interval may nevertheless have a zero derivative in that interval. One way this can happen is at a boundary point of a closed interval: for example, a function may start out with a zero slope which immediately increases, or its otherwise positive slope may decrease until it just reaches zero at an endpoint. For example, at 0 consider the slope of the squaring function defined on the unit interval. Alternatively, the derivative of a strictly increasing function may decline momentarily to zero and then rise again. For example, at 0 consider the slope of the cubing function defined on $\mathbb{R}$.



1.1.4.1 Affine Functions

The concept of the derivative at a point captures our intuitions about the slope of a function at a point. Many familiar real functions are differentiable and have derivative functions that are simple to characterize. For example, affine functions have a constant slope, so we correctly expect that f' is a constant function whenever f is an affine function. Correspondingly, every affine function is montone, and every non-constant affine function is strictly montone.

Example 1.8 Find the derivative function when $f = x \mapsto a_0 + a_1x$. Start with the function's difference quotient at any given point x :

$$\frac{f[x + h] - f[x]}{h} = \frac{a_0 + a_1(x + h) - (a_0 + a_1x)}{h} = a_1$$

Next, find the limit as h shrinks to zero. This is particularly simple: $\lim_{h \rightarrow 0} a_1 = a_1$. The derivative function is therefore $f' = x \mapsto a_1$. At any point, the derivative of an affine function is its constant slope.

The result in example 1.8 is sometimes called the *affine-function rule* for differentiation. The special case when $a_0 = 0$ is sometimes called the *scalar rule* for differentiation. The special case when $a_1 = 0$ is sometimes called the *constant-function rule* for differentiation. The scalar rule and the constant-function rule are special cases of the affine-function rule.

Example 1.9 (Total Cost and Marginal Cost) Let the total cost of production be an affine function: $tc = x \mapsto c_0 + c_1x$, where x represents the level of production and the function parameters are positive. By the affine function rule, $tc' = x \mapsto c_1$. The rate of change of total cost is called *marginal cost*, so in this case the marginal cost of production is constant.

Economically, this function is sensible only for nonnegative production levels, which implies at 0 it has only a right derivative.



1.1.4.2 Elasticity and Semi-Elasticity

Although social scientists often work with purely numerical functions, many of these implicitly assume units for both the function arguments and its values. These units show up in a function's derivatives. The derivative of the function at a point is the slope of a function at that point. This is the instantaneous rate of change of the function value with respect to the input value. The derivative therefore depends on the units used to measure both the input and the output of the function. For example, 30 kilometers per hour and 500 meters per minute express the same speed in different units.

A semi-elasticity removes the dependence on the output units: it reports the *proportional* rate of change in the output given a change in the input. Define the proportional change in the output when the input changes from x to $x + h$ to be $(f[x + h] - f[x])/f[x]$. This does not depend on the units of the function value. Next, scale this proportional change by the size of the input change (h) to produce the simple discrete growth rate of f per unit change in the argument.

$$\frac{1}{h} \frac{f[x + h] - f[x]}{f[x]} = \frac{f[x + h] - f[x]}{h} \frac{1}{f[x]} \quad (1.4)$$

The limit as $h \rightarrow 0$ of this expression is $f'[x]/f[x]$, which is the semi-elasticity of f at x .

Exercise 1.1 Textbooks often represent market demand (q^d) for a commodity as affine (at nonnegative prices and quantities): $q^d = p \mapsto a - bp$. Demonstrate that the price elasticity of demand is then -1 at the midpoint of the demand curve.



An elasticity removes the dependence on the argument units as well as the function units: it reports the proportional change rate of change of the function value given a *proportional* change in the input. This time, deflate the proportional change in the value of f by the *proprtional* input change (h/x). This produces the simple discrete growth rate of f per percentage change in the argument.

$$\frac{x}{h} \frac{f[x+h] - f[x]}{f[x]} = \frac{f[x+h] - f[x]}{h} \frac{x}{f[x]} \quad (1.5)$$

Letting $h \rightarrow 0$ produces $f'[x]x/f[x]$, which is the point elasticity of f at x , often written as $\varepsilon_f[x]$.



Example 1.10 The parameterized real function $f = x \mapsto ax$ has the derivative function $f' = x \mapsto a$. The semi-elasticity at point x is therefore $f'[x]/f[x] = a/(ax) = 1/x$, while the elasticity at point x is $f'[x]x/f[x] = ax/(ax) = 1$.

Example 1.11 (Arc Elasticity) Rewrite the point elasticity of the real function f as $f'[x]/(f[x]/x)$. This is the ratio of the marginal change to the average value. Allen and Lerner (1934, p.228) argue that a discrete measure of elasticity should compute the average at the midpoint of the secant segment rather than at the initial point. The *arc elasticity* of f at x thereby becomes

$$\frac{f[x+h] - f[x]}{h} \bigg/ \frac{(f[x+h] + f[x])/2}{((x+h) + x)/2}$$

For discrete changes, arc elasticity has the advantage that it depends only on the two input values, not on the direction of change. The limiting value of this expression produces the same measure of point elasticity as before.

Exercise 1.2 From the demand function $q^d = p \mapsto 1000 - p/2$, compute the price elasticity of demand at any price p . Also, compute the corresponding arc elasticity when p rises from 900 to 1100.

```
Clear[p1, p2]
qd = p ↦ 1000 - p/2
(* point elasticity: *)
(qd'[p]*p/qd[p]) // Simplify
(* arc elasticity: *)
arcPctP = (p2 - p1)/((p1 + p2)/2)
arcPctQ = (qd[p2] - qd[p1])/((qd[p1] + qd[p2])/2)
arcE = arcPctP/arcPctQ // Simplify
arcE /. {p1 → 900, p2 → 1100}
(* arc elasticity equivalent: *)
arcMarginal = (qd[p2] - qd[p1])/(p2 - p1) // Simplify
arcAverage = ((qd[p2] + qd[p1])/2)/((p1 + p2)/2) // Simplify
arcE = arcMarginal/arcAverage // Simplify
arcE /. {p1 → 900, p2 → 1100}
```



1.1.5 Quadratic Functions

It is a bit harder to intuit the slope of a quadratic function, since it is different everywhere. However, applying the definition of a derivative once again produces a simple expression. Consider the scaled squaring function, $f = x \mapsto a_2x^2$. The difference quotient at any given x is

$$\frac{f[x+h] - f[x]}{h} = \frac{a_2(x+h)^2 - a_2x^2}{h} = a_2(2x+h)$$

The derivative of f at x is the following limit of this difference quotient.

$$\lim_{h \rightarrow 0} a_2(2x+h) = 2a_2x$$

At every x in the domain of f , this limit is well defined, so f is a differentiable function on $\mathbb{R}$, and its derivative function is $f' = x \mapsto 2a_2x$. For example, if $a_2 = 1$ then $f'[3] = 6$.

Example 1.12 (Total Cost and Marginal Cost) Let the total cost of production be $tc = x \mapsto c_2x^2$, where x represents the level of production and the function parameter is positive. By the quadratic function result, $tc' = x \mapsto 2c_2x$. The rate of change of total cost is called *marginal cost*, so in this case the marginal cost of production is increasing in the level of production.

Economically, this function is sensible only for nonnegative production levels, which implies at 0 it has only a right derivative.



Let $f = x \mapsto ax^2 + bx + c$ and find f' .

```
Clear[x,a,b,c,test]
f = x  $\mapsto$  a*x^2 + b*x + c
f'
Derivative[1][f] (* full form *)
f'[x] (* symbolic differentiation *)
f'[test] (* symbolic differentiation *)
f'[3] (* derivative at a point *)
f'[3] /. {a $\rightarrow$ 1,b $\rightarrow$ 0,c $\rightarrow$ 0} (* specialize to squaring *)
```



1.1.5.1 Marginal Revenue

Recall that the demand-price function (p_d) for a market produces the maximum price at which a market quantity (q) can be sold. It is the inverse of the market demand function. The total revenue function (tr) produces the product of the demand price and the quantity: $\text{tr} = q \mapsto p_d[q] q$. When total revenue is differentiable, economists commonly use the term *marginal revenue* to denote the derivative of the total revenue function: $\text{mr} = \text{tr}'$. Like total revenue and the demand price, marginal revenue is a function of the quantity produced. Figure 1.4 illustrates a total revenue function and its implied marginal revenue function.

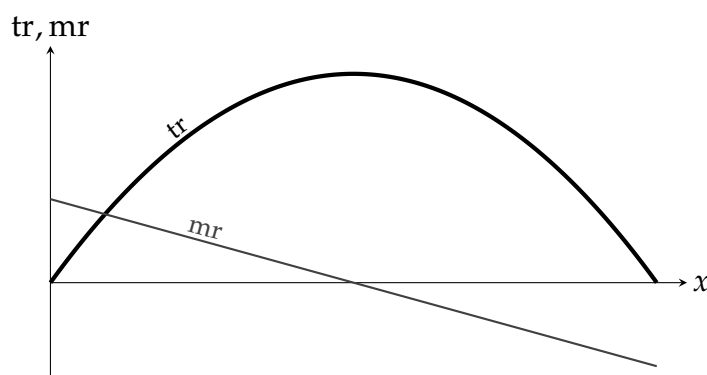


Figure 1.4: Total Revenue and Marginal Revenue

At any value of output (q), total revenue ($\text{tr}[q]$) is changing at the rate $\text{tr}'[q]$, which is the height of the marginal revenue curve ($\text{mr}[q]$). At its maximum, total revenue is neither increasing or decreasing, so the corresponding marginal revenue is zero.

[See demo.](#)

Example 1.13 (Marginal Revenue with Linear Demand) Let the demand function for a monopolist be $p \mapsto a - bp$, where a and b are positive demand function parameters. Its inverse is the demand-price function, $p_d = q \mapsto (a - q)/b$. The total revenue function is therefore

$$\text{tr} = q \mapsto (aq - q^2)/b$$

At any given q , produce a difference quotient:

$$\begin{aligned} \frac{\text{tr}[q+h] - \text{tr}[q]}{h} &= \frac{(a(q+h) - (q+h)^2)/b - (aq - q^2)/b}{h} \\ &= a/b - 2q/b - h/b \end{aligned}$$

Produce the derivative of tr at q by computing the limit (as $h \rightarrow 0$) of this difference quotient: $\text{tr}'[q] = a/b - 2q/b$. Economists call this the marginal revenue (mr) at q . Since tr is everywhere differentiable, the associated marginal revenue function is

$$\text{mr} = q \mapsto a/b - 2q/b$$

Note that with a linear demand curve, the marginal revenue curve is twice as steep as the underlying demand-price curve.

Exercise:

If a firm's demand-price function is $p = q \mapsto 3 - 0.3q$, find the total revenue and marginal revenue.



Example 1.14 (Kinked Demand) A popular model of oligopolistic competition from the 1930s implies that each oligopolistic firm faces a kinked demand curve. The implied total revenue curve is correspondingly kinked, and the marginal revenue curve therefore has a break in it. Figure 1.5 illustrates this possibility.

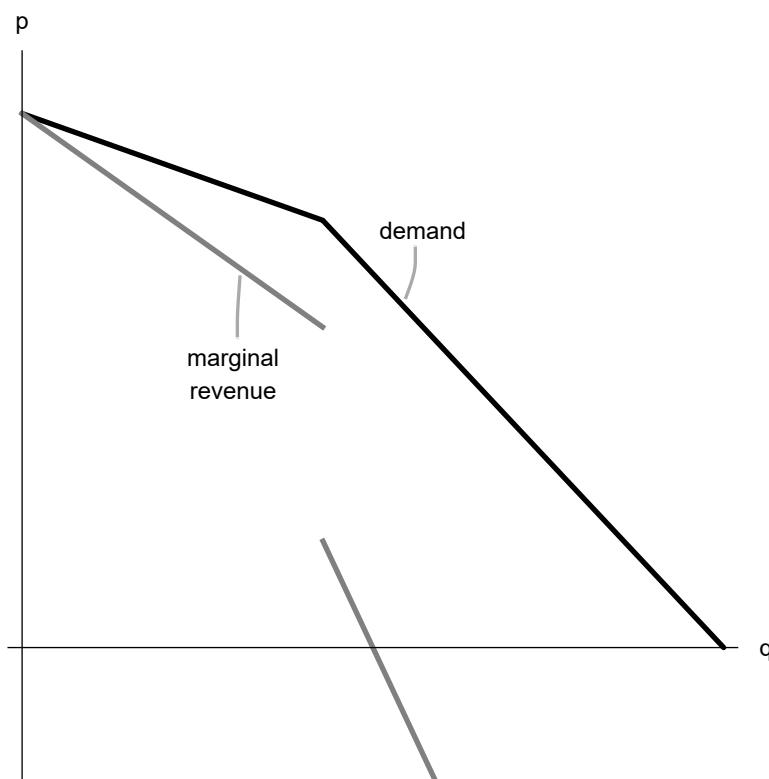


Figure 1.5: Kinked Demand and Marginal Revenue

If a monopolistic competitor faces a kinked demand curve, its marginal revenue function is discontinuous. (Total revenue is not differentiable at the kink.)



1.2 Rules to Differentiate By

When introducing the differential calculus, it is traditional to state a few simple rules for differentiating unary real functions. Table 1.1 lists a few common differentiation rules for differentiable functions. Many of these rules are easily proved by applying the definition of the derivative. (See section 1.2.2.) Table 1.2 lists a few common differentiation rules for special functional forms.

**Table 1.1:** Some Differentiation Rules

Function	Derivative	Rule Name
$\alpha f + \beta g$	$\alpha f' + \beta g'$	Linear Combination Rule
$g \circ f$	$x \mapsto g'[f[x]] \cdot f'[x]$	Chain Rule
f^{-1}	$y \mapsto 1/f'[f^{-1}[y]]$	Inverse Function Rule
$f \cdot g$	$f' \cdot g + g' \cdot f$	Product Rule
$1/f$	$-f'/f^2$	Reciprocal Rule
f/g	$(f'g - g'f)/g^2$	Quotient Rule

The real functions f and g are differentiable. The reciprocal and quotient rules apply only where the denominator is not zero.

Table 1.2: Special Functional Forms

Function	Derivative	Rule Name
$x \mapsto x^n$	$(x \mapsto nx^{n-1})$	Power Function Rule
exp	exp	Exponential Rule
ln	$x \mapsto 1/x$	Logarithmic Rule
sin	cos	Sine Rule for Differentiation
cos	$-\sin$	Cosine Rule for Differentiation

□



1.2.1 Example Uses of Differentiation Rules

This subsection briefly illustrates the application of some differentiation rules.

1.2.1.1 Linear Combination Rule

Differentiation is a linear operator on the set of differentiable functions. This is a way of saying that it satisfies the *linear combination rule*: differentiating a linear combination of functions produces a linear combination of derivatives, with the same weights. For differentiable real functions f and g and constants α and β , we may summarize this property as follows.

$$(\alpha f + \beta g)' = \alpha f' + \beta g' \quad (1.6)$$

We can decompose the linear combination rule into two subrules, which are often called the *scalar rule*, and the *sum rule*. The scalar rule states that scaling a function scales its derivative: $(\alpha f)' = \alpha f'$. The sum rule states that the derivative of a sum of functions is the sum of their derivatives: $(f + g)' = f' + g'$.

Illustrating linearity of the differentiation operation:

```
Clear[f, g, h, x]
h = x  $\mapsto$  alpha * f[x] + beta * g[x]
h'[x]
```



The *sum rule* of differentiation captures the additivity of the derivative. Recall that by definition, adding two real functions just adds their respective values at each input: $(f + g)[x] = f[x] + g[x]$. Intuitively, at any x , the slope of $f + g$ is the sum of the slopes of f and g . Therefore, given two differentiable functions f and g , the sum of the functions is also differentiable, with $(f + g)' = f' + g'$.

Example 1.15 (Marginal Cost) A firm has total costs of production equal to fixed costs plus variable costs. Variable costs (vc) depend on the quantity of output produced (q), but fixed costs do not. Define the total cost function (tc) as the following sum:

$$tc = fc + vc$$

Since the fixed cost function is a constant function, its derivative is constant at zero. (Recall the affine-function rule.) Therefore by the sum rule,

$$tc' = vc'$$

Since the marginal cost of production does not depend on fixed costs, it is the same as the marginal variable cost of production.



Example 1.16 (Competitive Profits) A competitive firm faces an exogenous price p for its output q . Its profits are total revenue less total cost: $\Pi = q \mapsto pq - tc[q]$. By the linear combination rule, marginal profits are therefore the difference between marginal revenue (p) and marginal cost (tc').

$$\Pi' = q \mapsto p - tc'[q]$$

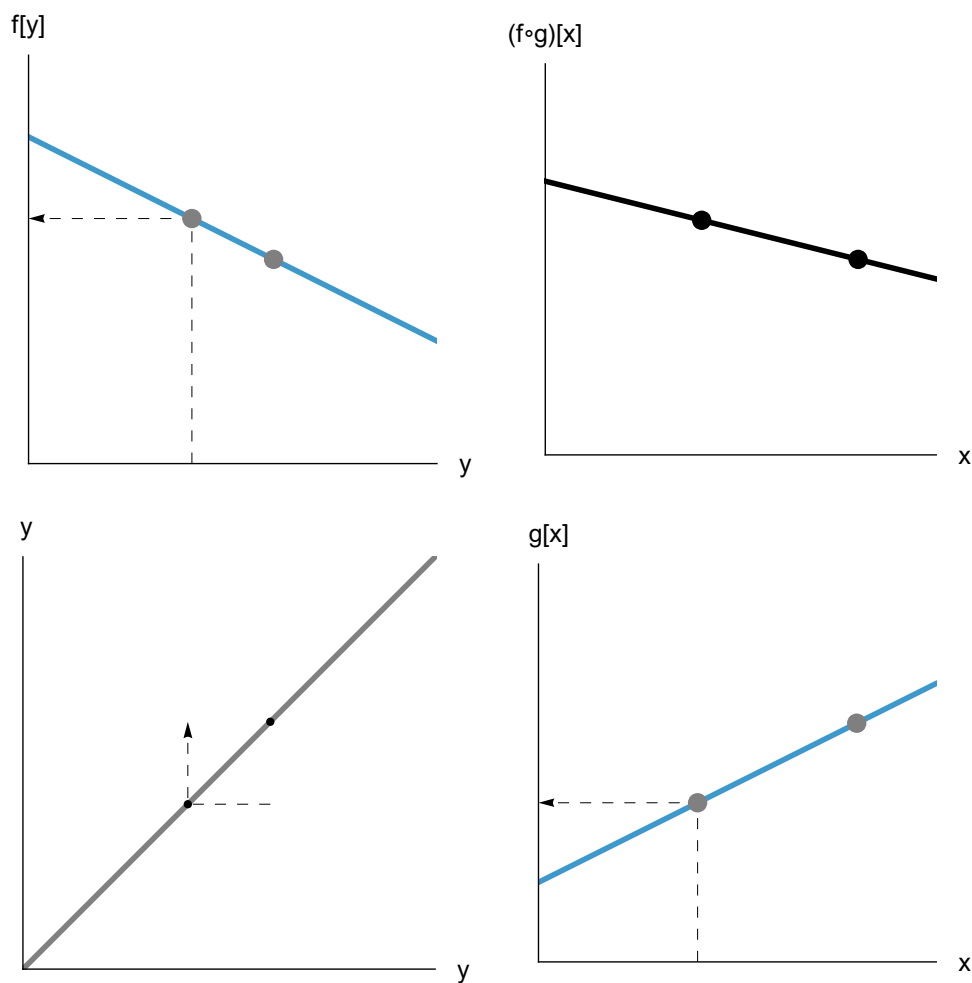
Note that if price is greater than marginal cost and some quantity, the derivative is positive, so the firm can increase profits by increasing the quantity produced.

```
Clear[q]
profit = q ↦ p*q - tc[q]
profit'[q]
```



1.2.1.2 Chain Rule

Recall function composition:



□

The composition of two real functions (f and g) is defined as $(f \circ g) \stackrel{\text{def}}{=} x \mapsto f[g[x]]$. If f and g are differentiable, then so is the composite function $(f \circ g)$. The *chain rule* of differentiation for composite functions states that the derivative function is

$$(f \circ g)' = x \mapsto f'[g[x]] \cdot g'[x] \quad (1.7)$$

Therefore, if f is strictly increasing, then $(f \circ g)'$ must have the same sign as g' . In this case, $(f \circ g)[x_2] > (f \circ g)[x_1] \iff g[x_2] > g[x_1]$. A strictly increasing transformation of a real function is *order preserving*. As shown in a later chapter, sometimes a strictly increasing transformation of a function can aid in optimization.

```
Clear[f, g, x]
(f@g)'[x]
```

□

Example 1.17 Suppose $f = \square \mapsto 3\square^2$ and $g = \square \mapsto 5 + 2\square$. Find the derivative of $f \circ g$ by applying the chain rule. To improve readability, set $y = g[x]$ when useful. Then $(f \circ g)'[x] = f'[y]g'[x] = (6y)(2) = ([6 \cdot (5 + 2x)])(2) = 60 + 24x$. Of course, it is also possible in this case to show that $(f \circ g) = x \mapsto 75 + 60x + 12x^2$ and then differentiate this expression using the sum rule and earlier results.

Illustrating the chain rule of differentiation:

```
Clear[f, g, h, x]
h = f @* g
h'[x]
f = y  $\mapsto$  3*y^2
g = x  $\mapsto$  2*x + 5
h'[x] // Expand
```



1.2.1.3 Inverse Function Rule

If the differentiable real function f has an inverse f^{-1} , then the slope of the inverse is the inverse of the slope. More precisely, $(f^{-1})' = y \mapsto 1/f'[f^{-1}[y]]$. Applying this derivative function to the value $f[x]$ produces an expression of the rule that is perhaps more readable.

$$(f^{-1})'[f[x]] = 1/f'[x] \tag{1.8}$$

The inverse function rule is rather intuitive: since the graph of inverse function just swaps the first and second coordinate of the original function, rise of run for the inverse function is the corresponding run over rise of the original function. As a result, the sign of $(f^{-1})'$ must match the sign of f' .



Example 1.18 The law of demand states that (for positive prices) the quantity demand of a good is a strictly decreasing function of its price: $q'[p] < 0$. The demand price function is the inverse of the ordinary demand function. As illustrated in Figure 1.6, its slope is also negative, since it is the reciprocal of the demand-function slope.

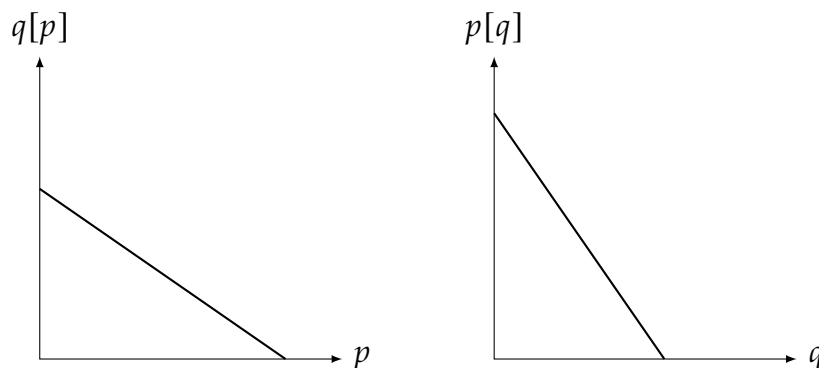


Figure 1.6: Quantity Demanded vs Demand Price

The law of demands states that $q'[p] < 0$. The inverse function rule therefore implies that $p'[q] < 0$ as well. Intuitively, this is because the graph of the inverse function simply reflects the function graph through the 45° line.

□

Example 1.19 In the short run, a firm can only vary its labor input (L), and for nonnegative labor input, its production function is $L \mapsto kL$. For example, when $L = 1$ the quantity produced (q) is k . From the affine-function rule, the derivative of this function is k .

The inverse function represents the labor required to produce any target quantity, q . Find the inverse function on the nonnegative real numbers by solving $q = kL$ for $L = q/k$. Equivalently, the function $q \mapsto q/k$ returns the labor requirement for attaining any q . From the affine-function rule, the derivative of this function is $1/k$, as expected from the inverse function rule.



Example 1.20 A student finds that the study effort (ϵ) required to produce a percentage grade (γ) is described by the grading function $f = \gamma \mapsto \gamma^2$ on the unit interval. (So 100% effort is required to produce a score of 100%.) Our previous work with the squaring function tells us that $f' = \gamma \mapsto 2\gamma$.

The squaring function has an inverse on the unit interval: $f^{-1} = \epsilon \mapsto \sqrt{\epsilon}$. This function shows the grade that results from each effort level. Using the inverse function rule, the slope of the inverse function is $(f^{-1})'[\gamma^2] = 1/(2\gamma)$, or equivalently $(f^{-1})'[\epsilon] = 1/(2\sqrt{\epsilon})$. For example, the slope of square-root function at 0.25 is 1.



1.2.1.4 Product Rule

Given two differentiable functions f and g , their product $(f \cdot g)$ is also differentiable, with

$$(f \cdot g)' = f' \cdot g + g' \cdot f \quad (1.9)$$

The derivative of a function product is a weighted sum of the function derivatives, where the weights are the function values. This is the **product rule** of differentiation. Importantly, this is *not* simply the product of the derivatives.

Example 1.21 (Monopolist Revenue) A monopolist faces the demand-price function p , where $p'[q] < 0$. Total revenue is $\text{tr} = q \mapsto p[q] \cdot q$. This is the product of two functions: the identity function and the demand-price function p . The derivative of the identity function is 1. Therefore the product rule says that at any q , marginal revenue is a weighted sum of $p'[q]$ and 1, where the weights are q and $p[q]$.

$$\text{tr}' = q \mapsto p'[q] \cdot q + 1 \cdot p[q]$$



1.3 Differentials

Recall that if f is an affine function, its difference quotient is constant: it does not depend on x or dx . This is because the change in the function value ($\Delta_{x,dx}f$) is proportional to the step size (dx). That is, given an affine function $f = x \mapsto a_0 + a_1x$, we have

$$\underbrace{f[x + dx] - f[x]}_{\Delta f} = a_1 dx \quad (1.33)$$

Here the factor of proportionality is a_1 , the constant value of the difference quotient of f . To simplify notation, it is common to write $f[x + dx] - f[x]$ more simply as Δf , relying on context to supply the lost arguments.

Use of `DifferenceDelta` requires more arguments:

```
Clear[f, x, dx]
DifferenceDelta[f[x], {x, 1, dx}]
f = x  $\mapsto$  a0 + a1*x;
DifferenceDelta[f[x], {x, 1, dx}]
```

□

Since a nonlinear function does not have a constant slope, characterizing how its value changes is more complicated. However, the slope at any point of a differentiable function is given by its derivative function. Even when a function does not have a constant linear response to changes in its input, for small changes in the function input, the value of the derivative at x can serve a similar role: we can approximate the actual change (Δf) in the function value by $f'[x] dx$. This is the differential approximation of the function change. The sense in which this is a good approximation is clarified by defining the following remainder term.

$$r = \langle x, dx \rangle \mapsto \underbrace{f[x + dx] - f[x]}_{\Delta f} - \underbrace{f'[x] dx}_{df} \quad (1.34)$$

The product $f'[x] dx$ is called the *differential* of f at x , and the notation df is often used to denote this differential.

□

Using this definition of the remainder, break the actual change in the function value (Δf) into the differential (df) and the corresponding remainder (r) as follows.

$$\underbrace{f[x + dx] - f[x]}_{\Delta f} = \underbrace{f'[x] dx}_{df} + r[x, dx] \quad (1.35)$$

Divide both sides by dx and find the limit as it approaches 0.

$$\lim_{dx \rightarrow 0} \frac{f[x + dx] - f[x]}{dx} = f'[x] + \lim_{dx \rightarrow 0} \frac{r[x, dx]}{dx} \quad (1.36)$$

The limit on the left evaluates to $f'[x]$ whenever f is differentiable at x , so for such a function

$$\lim_{dx \rightarrow 0} \frac{r[x, dx]}{dx} = 0 \quad (1.37)$$

Here is one way to summarize this observation: wherever a function is differentiable, the remainder term approaches zero faster than dx does. In this sense, the differential provides a linear approximation of the actual change in the value of the function. Even though the function f is not linear, this linear approximation ($f'[x] dx$) becomes a very good approximation to the actual change ($\Delta_{dx} f[x]$) when the change in the function input (dx) is small.

□

1.3.1 Little- $\mathbf{o}$ Notation

SKIP

Consider two real functions f and g defined near a point a . If

$$\lim_{x \rightarrow a} \frac{f[x]}{g[x]} = 0 \quad (1.38)$$

the ratio $f[x]/g[x]$ approaches 0 as x approaches a . (This is a function limit, so we allow that $f[a]/g[a]$ may not be defined.) In this case we say that f is **little- $\mathbf{o}$** of g near a , which we write as $f \in \mathbf{o}[g; a]$. Intuitively, near a the value of f is very small relative to the value of g .

We have found that for a differentiable function the remainder function disappears as the difference-quotient stepsize vanishes.

$$\lim_{dx \rightarrow 0} \frac{r[x, dx]}{dx} = 0 \quad (1.39)$$

So we may write $r[x_0, \cdot] \in \mathbf{o}[\mathbf{i}; 0]$, where $r[x_0, \cdot]$ is the partially applied remainder function and $\mathbf{i}$ is the identity function. However, the someone puzzling conventional practice is to write $r[x_0, dx] \in \mathbf{o}[dx]$ as $(dx \rightarrow 0)$.

This leads to a third way of defining the derivative. If there exists a value ℓ such that

$$f[x + dx] - f[x] - \ell dx \in \mathbf{o}[dx] \quad (dx \rightarrow 0) \quad (1.40)$$

then f is differentiable at x and ℓ is its derivative there. That is, $f'[x] = \ell$.



1.3.2 Differentials and Tangent Lines

For a differentiable function f , the effect of a small change in the change is well approximated by the differential $f'[x]dx$. Figure 1.8 shows one way to visualize this. When the change in the function argument is small, then the change along the tangent line is a good approximation of the actual change in the value of the function. For a differentiable real function f , the tangent line at $(x_0, f[x_0])$ can be described by the following function, which produces a differential approximation to the value of the function at each x .

$$x \mapsto f[x_0] + f'[x_0](x - x_0) \quad (1.41)$$

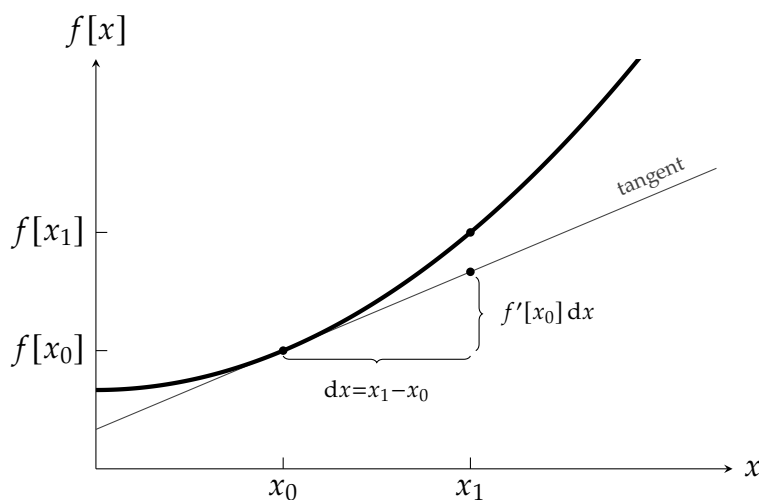


Figure 1.8: Differential and Tangent Line
The change in the height of the tangent line is the function differential $f'[x_0]dx$. This approximates the actual change in the function value, which is $f[x_1] - f[x_0]$.



Example 1.33 Suppose $f = x \mapsto 5x^2 + 2$ and its input argument changes from x to $x + dx$. The actual change in the function value is

$$\begin{aligned}\Delta f &= (5(x + dx)^2 + 2) - (5x^2 + 2) \\ &= 10x dx + 5dx^2\end{aligned}$$

Compute the derivative f' at x as follows:

$$\lim_{dx \rightarrow 0} \frac{10x dx + 5dx^2}{dx} = \lim_{dx \rightarrow 0} 10x + 5dx = 10x$$

The differential is $f'[x] dx = 10x dx$, and the remainder is $5 dx^2$. Suppose x rises from 1.0 to 1.1, so that $dx = 0.1$. Then $f[x]$ rises from 7.0 to 8.05, so that $f[x + dx] - f[x] = 1.05$. The differential approximation to this change is $f'[x]dx = 10 \cdot 0.1 = 1.00$, which is close to the true change. Examine the remainder relative to the step size to confirm that it shrinks to zero.

$$\lim_{dx \rightarrow 0} \frac{\Delta_{dx} f[x] - f'[x] dx}{dx} = \lim_{dx \rightarrow 0} \frac{5 dx^2}{dx} = \lim_{dx \rightarrow 0} 5 dx = 0$$

```
ClearAll[f]
f = x ↦ 5*x^2 + 2
With[{x0=1.0},
  Plot[{f[x], f[x0] + (x - x0)*f'[x0]}, {x, 0, 2}]
]
(* Compare change to differential approximation: *)
f[1.1] - f[1.0] (* actual change *)
f'[1] * (1.1 - 1.0) (* approximate change *)
```

Univariate Optimization: Overview



1.1 Curvature and More

This lecture focuses on differentiable functions. Recall a differentiable real function f has an associated derivative function, which characterizes the slope at each point of f . The differentiable function f is sometimes called the *primitive* function.

The derivative function (f') is again a function, so it may in turn be differentiable. If so, refer to f' as the **first-order** derivative of f , and refer its derivative as the **second-order** derivative of f . Write this second-order derivative as f'' or $f^{(2)}$. (Another common notation is $d^2f[x]/dx^2$.)

Example 1.1 The following table lists several functions along with their first-order and second-order derivative functions.

f	f'	f''
$x \mapsto x^2$	$x \mapsto 2x$	$x \mapsto 2$
$x \mapsto \exp[x]$	$x \mapsto \exp[x]$	$x \mapsto \exp[x]$
$x \mapsto \ln[x]$	$x \mapsto 1/x$	$x \mapsto -1/x^2$
$x \mapsto \sin[x]$	$x \mapsto \cos[x]$	$x \mapsto -\sin[x]$



1.1.1 Concavity and Convexity Redux

The derivative function specifies the rate of change of the primitive function. For example, if $f'[x_0] > 0$, then at x_0 the value of the function increases as its input increases. As a more specific example, if $f'[x_0] = 2$, then at x_0 the value of the function increases twice as fast its input increases. With this in mind, assume the derivative function f' is also differentiable and consider the second-order derivative function, f'' . At any input value x , the value of $f''[x]$ is the rate of change of the value of the derivative function at x . So the second-order derivative yields the rate of change of the slope of f . If $f''[x] < 0$, then the slope is falling in value, while if $f''[x] > 0$, then the slope is rising in value. In this way, the second-order derivative provides information about the *curvature* of the primitive function.

A second-order derivative is the slope's slope: $f'' = (f')'$. This is the *curvature* of f .

□

For example, suppose $f''[x_0] > 0$. This means that at x_0 , the *value* of the *derivative* function increases when the input argument increases. This in turn means that the *slope* of the *primitive* function is increasing. So a positive second-order derivative at x_0 implies that the function is *strictly convex* at x_0 . Similarly, a negative second-order derivative at x_0 implies that the function is *strictly concave* at x_0 .

Example 1.2 (Decreasing Marginal Utility) Macroeconomic models often assume that a representative consumer has a utility function ($u: \mathbb{R}_{>0} \rightarrow \mathbb{R}$) characterized by positive but declining marginal utility. That is, utility is an increasing, concave function of the level of total consumption. So a given utility function u should at any consumption level c be characterized by $u'[c] > 0$ and $u''[c] < 0$. Here are some common functional forms, given as expressions in c .

$u[c]$	$u'[c]$	$u''[c]$
$\sqrt{c}$	$c^{-1/2}/2$	$-c^{-3/2}/4$
$\ln c$	$1/c$	$-1/c^2$
$-(\bar{c} - c)^2/2$	$u'[c] = \bar{c} - c$	$u''[c] = -1$

The all display declining marginal utility, but only the first two ensure marginal utility is positive. This last example only meets the requirement of positive marginal utility when $\bar{c} > c$. Think of $\bar{c}$ as a *bliss* level of consumption: nothing better is possible. Attainment of the bliss point does not appear to characterize actual economies.

Exercise 1.1 Consider the general form of a quadratic polynomial: $f = x \mapsto a_0 + a_1x + a_2x^2$, with $a_2 \neq 0$. Is a quadratic polynomial concave or convex?

If a real function is twice continuously differentiable on some interval, there is a nice relationship between the primary function, its derivative, and its second derivative. The following development repeatedly uses the limit rule for indeterminate forms (L'Hôpital's rule).

$$\begin{aligned}
 \lim_{h \rightarrow 0} \frac{f[x+h] - f[x]}{h^2} + \frac{f[x-h] - f[x]}{h^2} &= \lim_{h \rightarrow 0} \frac{f'[x+h] - f'[x-h]}{2h} \\
 &= \lim_{h \rightarrow 0} \frac{f''[x+h] + f''[x-h]}{2} \\
 &= f''[x]
 \end{aligned} \tag{1.1}$$

This nicely captures the idea that a positive second derivative reflects an increasing slope.



1.1.2 Second-Order Approximation

Given knowledge of the function value $f[x_0]$, the simplest approximation to another value ($f[x]$) is simply the known function value ($f[x_0]$). Call this a zero-th order approximation. As x varies, this approximation is completely accurate only if f is a constant function.

The next simplest approximation matches not just the function's value at x_0 but also its slope at that point. This is the differential approximation, given by the function $x \mapsto f[x_0] + f'[x_0](x - x_0)$. Call this a first-order approximation, because it makes use of the first-order derivative of f . By doing so, it takes advantage of information about the slope of the function. This new approximation function has two desirable properties: it matches the value of f at x_0 , and its derivative it matches the value of f' at x_0 . Nevertheless, at points away from x_0 , the first-order approximation is completely accurate only if f is an affine function.

□

If f is twice differentiable, one may extend this reasoning to produce a second-order (or quadratic) approximation. This uses the second-order derivative at x_0 to add information about the curvature of the function, as illustrated by the following approximation function.

$$x \mapsto f[x_0] + f'[x_0](x - x_0) + f''[x_0] \frac{(x - x_0)^2}{2} \quad (1.2)$$

This second-order approximation is a quadratic polynomial in x . At x_0 its value is just $f[x_0]$, and in addition its first-order derivative equals $f'[x_0]$. Furthermore, its second-order derivative matches $f''[x_0]$. This curvature information can improve the approximation. Nevertheless, at points away from x_0 , the second-order approximation is completely accurate only if f is a quadratic function.

Example 1.3 Consider the function defined by $f = x \mapsto x^2$, with $x_0 = 2$ and $dx = 1$. The true change in the function is $f[3] - f[2] = 9 - 4 = 5$. Since $f'[x] = 2x$ and $f''[x] = 2$, the quadratic approximation to the change is $f'[2] \cdot 1 + f''[2]/2 \cdot 1 = 4 + 1 = 5$.

Consider the function defined by $f = x \mapsto x^3$. Since $f'[x] = 3x^2$ and $f''[x] = 6x$, the function is increasing and convex at $x = 2$. When x rises from 2 to 3, the true change in the function value is $f[3] - f[2] = 27 - 8 = 19$. The quadratic approximation to this change is $f'[2] \cdot 1 + f''[2]/2 \cdot 1 = 12 + 6 = 18$.

[Try it in Mma!](#)



Example 1.4 Consider the function defined by $f = x \mapsto 15 + x + (x-3)^3/2$. Let the function input rise from 2 to 3. Then $\Delta f = f[3] - f[2] = 18 - 16.5 = 1.5$. Since $f'[x] = 1 + (3/2)(x-3)^2$, the linear approximation to this change is $1 - (-3/2) = 2.5$. Since $f''[x] = 3(x-3)$, the quadratic approximation to this change is $2.5 + (-3)/2 = 1.0$. Figure 1.1 illustrates these approximations.

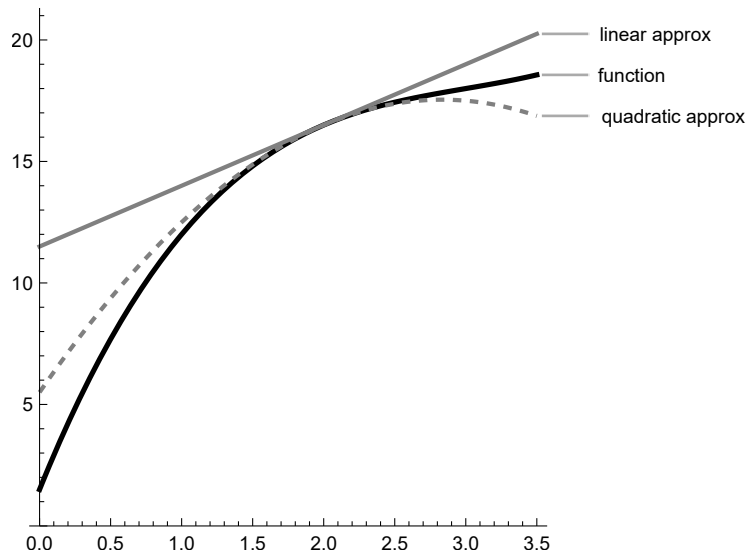


Figure 1.1: Linear and Quadratic Approximations

□

1.1.2.1 Approximating Expected Utility

Let $u[w]$ represent the utility of end of period wealth. More wealth is always beneficial: $u'[w] > 0$. Wealth is uncertain, with mean $\bar{w} = \mathcal{E}[w]$. Consider the case of two outcomes, good and bad, so that $\bar{w} = p_b w_b + p_g w_g$. (As always, the probabilities are positive and sum to one.) Relative to mean wealth, there is a second-order approximation to utility in each outcome.

$$\begin{aligned} u[w_b] &\approx u[\bar{w}] + u'[\bar{w}](w_b - \bar{w}) + u''[\bar{w}]\frac{(w_b - \bar{w})^2}{2} \\ u[w_g] &\approx u[\bar{w}] + u'[\bar{w}](w_g - \bar{w}) + u''[\bar{w}]\frac{(w_g - \bar{w})^2}{2} \end{aligned} \quad (1.3)$$

Use these to approximate expected utility as

$$\begin{aligned} \mathcal{E}[u] &= p_b \left(u[\bar{w}] + u'[\bar{w}](w_b - \bar{w}) + u''[\bar{w}]\frac{(w_b - \bar{w})^2}{2} \right) \\ &\quad + p_g \left(u[\bar{w}] + u'[\bar{w}](w_g - \bar{w}) + u''[\bar{w}]\frac{(w_g - \bar{w})^2}{2} \right) \\ &= u[\bar{w}] + u''[\bar{w}]\frac{\text{var}[w]}{2} \end{aligned} \quad (1.4)$$

Two features of this formulation are salient. First of all, when the second-order derivative is negative, expected utility is less than the utility of expected wealth. This is the case of *risk aversion*: u is concave, so wealth has diminishing marginal utility ($u''[w] < 0$). As a result, $\mathcal{E}[u[w]] < u[\bar{w}]$. Additionally, this second-order approximation depends only on the mean and variance of wealth. Sometimes this observation is treated as suggesting that actions to maximize expected utility may be approximated as actions to trade off expected wealth against its risk.



1.1.3 Higher-Order Approximation

Naturally, second-order derivative functions may themselves be differentiable. Use the notation $f^{(n)}$ to denote the n -th order derivative.

Example 1.5 Find the first, second, third, and fourth-order derivative functions for the following functions.

f	f'	f''	$f^{(3)}$	$f^{(4)}$
$x \mapsto x^4$	$x \mapsto 4x^3$	$x \mapsto 12x^2$	$x \mapsto 24x$	$x \mapsto 24$
$x \mapsto \exp[x]$	$x \mapsto \exp[x]$	$x \mapsto \exp[x]$	$x \mapsto \exp[x]$	$x \mapsto \exp[x]$
$x \mapsto \ln[x]$	$x \mapsto x^{-1}$	$x \mapsto -x^{-2}$	$x \mapsto 2x^{-3}$	$x \mapsto -6x^{-4}$
$x \mapsto \sin[x]$	$x \mapsto \cos[x]$	$x \mapsto -\sin[x]$	$x \mapsto -\cos[x]$	$x \mapsto \sin[x]$

Use Table to produce the elements of this table.

□

Assuming the first n derivatives of f exist, define the n -th order polynomial approximation of f at x_0 as follows. (Remember that $0! = 1$.)

$$\begin{aligned} P_{x_0}^n = x &\mapsto \sum_{k=0}^n f^{(k)}[x_0] \frac{(x - x_0)^k}{k!} \\ &= x \mapsto f[x_0] \frac{1}{0!} + f'[x_0] \frac{(x - x_0)}{1!} + f''[x_0] \frac{(x - x_0)^2}{2!} + \cdots + f^{(n)}[x_0] \frac{(x - x_0)^n}{n!} \end{aligned} \tag{1.5}$$

Exercise 1.2 Show that the first k derivatives of P_k at $h = 0$ equal the first k derivatives of f at $x = a$.

□

A polynomial approximation to $f[x]$ is based on the value of f and its derivatives at the point x_0 . Taylor's theorem tells us something about the quality of this approximation.

Theorem 1.1 (Taylor's Theorem (univariate)) Suppose the real-valued function f is defined on the open real interval I , where it is $n + 1$ times differentiable. Then for distinct points $x_0, x \in I$, there exists z between x and x_0 such that

$$f[x] - P_{x_0}^n[x] = f^{(n+1)}[z] \frac{(x - x_0)^{n+1}}{(n + 1)!}$$

Proof: See for example [https://math.libretexts.org/Bookshelves/Calculus/Calculus_\(Guichard\)/11%3A_Sequences_and_Series/11.12%3A_Taylor's_Theorem](https://math.libretexts.org/Bookshelves/Calculus/Calculus_(Guichard)/11%3A_Sequences_and_Series/11.12%3A_Taylor's_Theorem). ■



In this book, the main import of Taylor's theorem is that the remainders of higher-order approximations shrink very rapidly to zero when x is close to x_0 . However, the theorem is much stronger than such a statement, because it applies to *any* x and x_0 in the interval. Since factorial grows much faster than any power function, Taylor's theorem ensures that as long as all the derivatives exist *and* are bounded, polynomial approximations to a function can be made as good as we like by increasing the polynomial degree. Such nicely behaved functions are *analytic*, and for any x_0 , an analytic real function f has the following *Taylor series* representation:¹

$$f[x] = \sum_{n=0}^{\infty} f^{(n)}[x_0] \frac{(x - x_0)^n}{n!} \quad (1.6)$$

Example 1.6 Recall that the exponential function $x \mapsto e^x$ is its own derivative. Therefore derivatives of every order exist, and they all evaluate to 1 at 0. We can therefore write (since the series is summable) for any $x \in (-\infty .. \infty)$,

$$e^x = \sum_{n=0}^{\infty} \frac{x^n}{n!}$$

¹When $x_0 = 0$, this is also called a *Maclaurin series* representation.



For some functions, we can find an interval where the function equals its Taylor series, although this is not true for the whole real line.

Example 1.7 Expand $f = x \mapsto 1/(1-x)$ around 0. By the power rule, $f^{(n)} = n!(1-x)^{-(n+1)}$, which always equals $n!$ at $x = 0$. If $|x| < 1$, we can therefore write

$$\frac{1}{1-x} = \sum_{n=0}^{\infty} x^n \quad (1.7)$$

However, this is not true outside of $(-1 .. 1)$, because outside of this interval the series does not converge.

Example 1.8 Expand $f = x \mapsto \ln[x]$ around 1. Recall the logarithm rule: on the positive real numbers, $\ln' = x \mapsto 1/x$. Thereafter use the power rule to conclude $f^{(n)}[x] = (-1)^{n-1}(n-1)!x^{-n}$. If $|x-1| < 1$, we can therefore write (after noting $\ln[1] = 0$),

$$\ln[x] = \sum_{n=1}^{\infty} (-1)^n x^n \quad (1.8)$$

However, this is not true outside of $(0 .. 2)$. For larger numbers, the series does not converge.



1.2 Extrema

The following analysis of the search for extrema could focus equally well on either maximization or minimization. After all, for a real-valued function f , finding a largest value of f is equivalent to finding a smallest value of $-f$. Mathematicians and engineers tend to emphasize the minimization perspective. However, social scientist more often spotlight the maximization perspective, and that is the approach of this section.

Maximization of a real-valued function on a set X may be represented as follows.

$$\max_{\vec{x} \in X} f[\vec{x}] \tag{1.9}$$

The real-valued function f is the *objective function*. Maximization involves the search for a *maximizer*, which is a value of $\vec{x}$ that makes $f[\vec{x}]$ as big as possible. Of course not all functions have extrema: consider the real function $x \mapsto a_0 + a_1x$, with $a_1 \neq 0$. However, another lecture found that any continuous real function on a compact set will reach a minimum and a maximum. And examples 1.9 and 1.10 make it clear that finding a maximizer is sometimes very simple and intuitive.



Example 1.9 Define the function $f = x \mapsto x^2$ on the unit interval $[0 .. 1]$. The maximizer is 1, which produces the maximum function value $f[1] = 1$. The minimizer is 0, which produces the minimum function value $f[0] = 0$.

Example 1.10 Fermat (1679) considered a maximization problem based on the figure below. Within the segment $\overline{AC}$, where should we place the point E in order to make the product of the subsegment lengths as large as possible? (The problem is often called the isoperimetric problem because it is equivalent to finding the largest-area rectangle among those with a given perimeter.)



Let ℓ represent half the length $\overline{AC}$, and let $\ell + x$ and $\ell - x$ be the lengths of the two subsegments. Then the product of the lengths is $\ell^2 - x^2$. This expression is maximized by $x = 0$, since $x^2 \geq 0$. That is, to maximize the product of the two lengths, E should lie at the midpoint of $\overline{AC}$. Equivalently, a square is the largest-area rectangle of a given perimeter.

Examples 1.9 and 1.10 each illustrate a unique global maximum, but they involve very different problem types. In example 1.9 the maximizer is at the boundary of the function domain, while in 1.10 the maximizer is an interior point. This section focuses on interior extrema. When the objective function is differentiable, interior extrema are at points where the function is (locally) neither increasing nor decreasing, which leads to a search for values of $\vec{x}$ where the function $f'[\vec{x}] = \vec{0}$. Any such $\vec{x}$ is called a *stationary point* of the function.

Definition 1.1 A point $\hat{\vec{x}}$ is a *stationary point* of the real function f if $f'[\hat{\vec{x}}] = \vec{0}$. A point $\vec{x}$ is a *critical point* of f if it is a stationary point of f or if $f'[\vec{x}]$ is not defined.

□

A core project of the present lecture is to characterize the relationship between stationary points and extrema. Before proceeding with the analysis of stationary points, consider Figure 1.2, which shows the values of a real function f over its entire domain. This illustrates some problems with any analytical association of stationary points with extrema. Points B and E are local maxima. (Some authors prefer the term *relative* maximum.) Point B is also a global maximum. (Some authors prefer the term *absolute* maximum.) However the function is not differentiable at that point. Point C is a local minimum, but point A (at the edge of the feasible set) is the global minimum. Point D is a flat spot. It might be considered either a weak local maximum or a weak local minimum, but globally it is neither. And finally, the function falls from point E to point F and briefly flattens out before falling further.

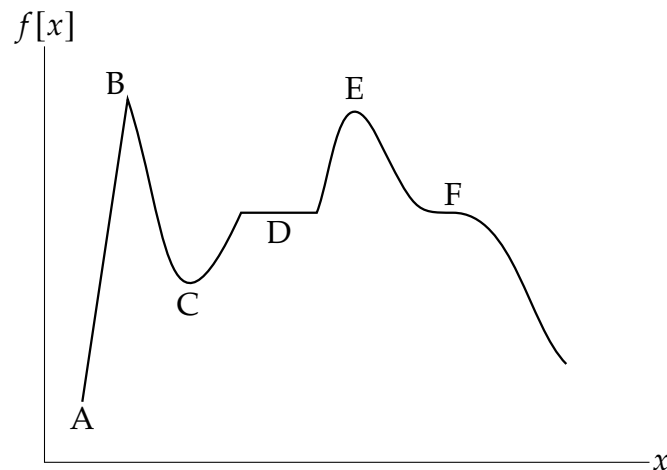


Figure 1.2: Extrema: Local and Global

This figure illustrates several concepts. First, a local extremum need not be a global extremum. In fact, many functions with local extrema lack a global maximum or minimum. (An example is the real function $x \mapsto x^3 - x$.) Second, a function may be stationary at points other than minima and maxima, so a stationary point is not necessarily an extremum. And finally, for an extremum to be associated with a zero first derivative, the

function must be differentiable at the extreme point. This means the function must not be discontinuous or kinked. To handle this, define a *critical point* of a real function to be a stationary point or any point where the function is not differentiable. Any local or global extremum of a continuous function is at a critical point.



1.2.1 Stationary Points

Figure 1.2 suggests that at an interior extremum, the slope of of a differentiable real function must be zero. This is true for both maxima and minima. The zero first-order derivative at a point is a necessary but not a sufficient condition for an extremum. Since this condition involves the first derivative of the function, it is called the *first-order necessary condition (FONC)* for an extremum. The next theorem develops this condition in a bit more detail.

□

Theorem 1.2 (Fermat's Theorem for Stationary Points) Consider the behavior of a continuous real-valued function f over an open interval. If $\hat{x}$ is a local extremum of f and f is differentiable at $\hat{x}$, then $f'[\hat{x}] = 0$.

Proof: Suppose $\hat{x}$ is a local maximizer of f . That is, there is a neighborhood $\mathcal{N}[\hat{x}]$ such that

$$f[\hat{x} + dx] - f[\hat{x}] \leq 0 \quad \forall (\hat{x} + dx) \in \mathcal{N}[\hat{x}] \quad (1.10)$$

In this neighborhood, the difference quotient anchored at $\hat{x}$ must always nonnegative for $dx < 0$. Therefore

$$\lim_{dx \rightarrow 0^-} \frac{f[\hat{x} + dx] - f[\hat{x}]}{dx} \geq 0 \quad (1.11)$$

In this neighborhood, the difference quotient anchored at $\hat{x}$ must always nonpositive for $dx > 0$. Therefore

$$\lim_{dx \rightarrow 0^+} \frac{f[\hat{x} + dx] - f[\hat{x}]}{dx} \leq 0 \quad (1.12)$$

But f is differentiable, so these two limits must be equal. This is possible only if $f'[\hat{x}] = 0$.

(Prove the case of a local minimizer in the same fashion.) ■

Fermat's theorem is a very intuitive result. For example, at a point where the function has a positive slope, there are always nearby inputs where (to the left) the function has a smaller value and where (to the right) the function has a larger value. And at a point where the function has a negative slope, there are always nearby inputs where (to the left) the function has a larger value and where (to the right) the function has a smaller value. (See Figure 1.4 for an illustration.)



Example 1.11 Let tc be the total cost function of a firm, so that $ac = q \mapsto tc[q]/q$ is the average cost function. Assuming differentiability, tc' is the marginal cost function. Suppose the firm minimizes the average cost of production at $\hat{q} > 0$. Then by Fermat's theorem $ac'[\hat{q}] = 0$, which by the quotient rule implies that $tc'[\hat{q}]/\hat{q} - tc[\hat{q}]/\hat{q}^2 = 0$. Multiply by $\hat{q}$ and rearrange to produce $tc'[\hat{q}] = tc[\hat{q}]/\hat{q}$. At the minimum average cost level of production, marginal and average cost are equal: $tc'[\hat{q}] = ac[\hat{q}]$.

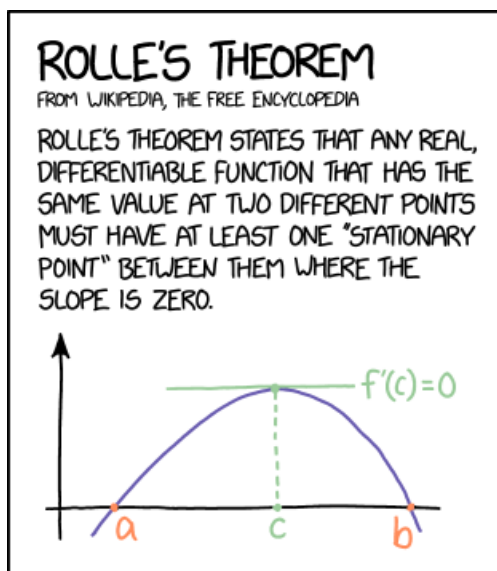


1.2.2 Mean-Value Theorems

Rolle's theorem roughly says that if a differentiable function takes on the same value at two points then it must have a stationary point somewhere between them. That is, if a differentiable function has a zero secant slope across some interval then it also has a zero slope somewhere in that interval.

Theorem 1.3 (Rolle's Theorem) Suppose the real function f is differentiable on the real interval $(a \dots b)$ and is also continuous at the interval endpoints. If $f[a] = f[b]$, then there exists a point $\bar{x} \in (a \dots b)$ such that $f'[\bar{x}] = 0$.

Proof: Weierstrass's extreme value theorem ensures that $[a \dots b]$ contains a minimizer ($x_{\min}$) and a maximizer ($x_{\max}$) of f . If $x_{\min} = x_{\max}$, the function is constant, and the derivative is null everywhere. Otherwise, since the endpoints values of f are equal, either the minimizer or maximizer is interior (or both are). Since f is differentiable, Fermat's theorem for stationary points ensures a zero derivative at an interior extreme point. ■



EVERY NOW AND THEN, I FEEL LIKE THE MATH EQUIVALENT OF THE CLUELESS ART MUSEUM VISITOR SQUINTING AT A PAINTING AND SAYING "C'MON, MY KID COULD MAKE THAT."



The mean-value theorem generalizes Rolle's theorem (and uses it in the proof). It says that for differentiable functions, the secant slope across an interval is also the function's slope at some point in that interval.

□

Theorem 1.4 (Lagrange's Mean-Value Theorem) Suppose f is differentiable on the open real interval $(a .. b)$ and additionally is continuous at the interval endpoints. Then there exists a point $\xi \in (a .. b)$ where $f'[\xi] = (f[b] - f[a])/(b - a)$.

Proof: Define $m = (f[b] - f[a])/(b - a)$, the secant slope of the function between a and b . Use this to define a new function g , as follows.

$$g = x \mapsto f[x] - m(x - a)$$

This definition requires that $g[a] = f[a]$ and $g[b] = f[b]$, so $g[a] = g[b]$. The function g shares the properties of f : it is differentiable on $(a .. b)$ and continuous at the endpoints. Therefore by Rolle's theorem g has a stationary point $\xi \in [a .. b]$. Since $g'[\xi] = 0$ and the definition of g implies that $g'[\xi] = f'[\xi] - m$, it follows that the slope at ξ is the slope of the secant line.

$$f'[\xi] = \frac{f[b] - f[a]}{b - a}$$

■

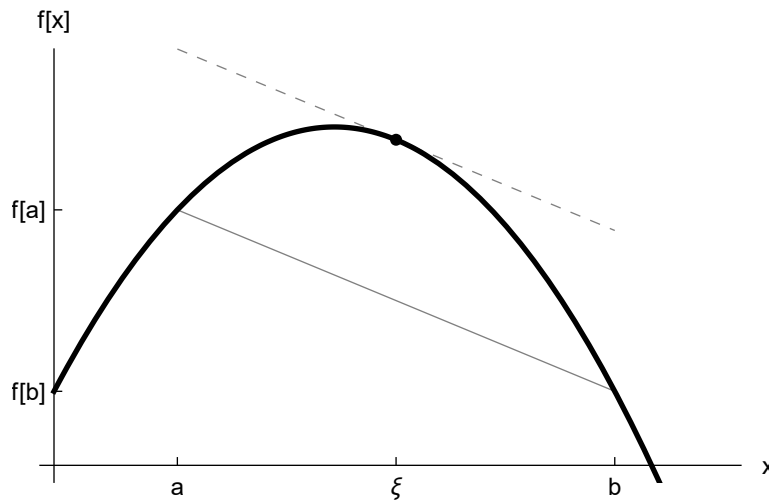


Figure 1.3: Illustrating the Mean-Value Theorem
The slope of the secant segment from $\langle a, f[a] \rangle$ to $\langle b, f[b] \rangle$ can be matched by $f'[\xi]$, the derivative of the function at an input of ξ .

SKIP to x^4 example

□

1.2.2.1 Fixed Points Redux

Theorem 1.5 Consider a differentiable real function f that maps some nonempty closed real interval $\mathcal{I} = \{x \mid a \leq x \leq b\}$ into itself: $f[\mathcal{I}] \subseteq \mathcal{I}$. If there is $0 \leq \lambda < 1$ such that on this interval $|f'[k]| \leq \lambda$, then f is a contraction on this interval.

Proof: Let $x, y \in \mathcal{I}$ be distinct. The MVT says there is $c \in \mathcal{I}$ such that $f'[c] = (f[x] - f[y])/(x - y)$.

Equivalently, $f[x] - f[y] = f'[c](x - y)$. So, $|f[x] - f[y]| = |f'[c]| |x - y|$. Therefore,
 $|f[x] - f[y]| \leq \lambda |x - y|$. ■

Since a contraction mapping ensures a unique fixed point, theorem 1.5 allows us to leverage bounds on the derivative to knowledge about the fixed point existence and uniqueness.

Example 1.12 Consider $f = (x \mapsto \sqrt{1+x}): \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$. Let $\mathcal{I} = [0 .. 3]$ and note that $f[0] = 1$ and $f[3] = 2$. Since that $f' = x \mapsto (1+x)^{-1/2}/2 > 0$, f is strictly increasing, so $f[\mathcal{I}] \subset \mathcal{I}$. Additionally note that $f'' = x \mapsto -(1+x)^{-3/2}/4 < 0$ and $f'[0] = 1/2$. The slope of f on $\mathcal{I}$ starts at $1/2$ and falls to $1/4$. So pick $\lambda = 1/2$, and f is a contraction on $\mathcal{I}$. This function has a unique fixed point.



1.2.3 Curvature and Optimization

At a stationary point $\hat{x}$ of a differentiable real function f , the slope is zero. Its instantaneous rate of change is zero. In this sense, at that point, the value of f tends neither to increase nor to decrease. One possibility is that the function is constant near $\hat{x}$. However, three other prominent cases immediately present themselves.

- on both sides of $\hat{x}$ we have $f[x] < f[\hat{x}]$, so $\hat{x}$ is a strict local maximizer.
- on both sides of $\hat{x}$ we have $f[x] > f[\hat{x}]$, so $\hat{x}$ is a strict local minimizer.
- on one side of $\hat{x}$ we have $f[x] < f[\hat{x}]$, while on the other side of $\hat{x}$ we have $f[x] > f[\hat{x}]$, so $\hat{x}$ is neither a local maximizer nor a local minimizer. In this case $\hat{x}$ must be an *inflection point*.



To understand the behavior of a function near a stationary point, we need to know something about nearby values. Here are three possible sources of this information.

- examine nearby values of the function.
- examine nearby values of f' .
- examine the curvature of the objective function near the stationary point.



An *extremum* of a real function is a maximum value or a minimum value of the function. As previously discussed, points where a differentiable real function has a zero derivative are called stationary points. At a stationary point, the instantaneous rate of change is zero. Stationarity is a necessary condition for the function to reach an extremum in an open interval, since it means the function value is (instantaneously) neither rising nor falling. Now suppose the function is twice differentiable. At stationary points, the second-order derivative may reveal whether we have found a maximum or a minimum by revealing the curvature at that point.

Theorem 1.6 (Second-Order Sufficient Condition) If $f'[\hat{x}] = 0$ and $f''[\hat{x}] < 0$, then $\hat{x}$ is a local maximizer of f . If $f'[\hat{x}] = 0$ and $f''[\hat{x}] > 0$, then $\hat{x}$ is a local minimizer of f .

□

Example 1.13 (Tax-Revenue Curve) Consider the continuous real function $f = t \mapsto \theta(t - t^2)$ defined on the unit interval $[0 .. 1]$. Interpret t as the tax rate and $\theta > 0$ as an economy specific parameter determining total tax revenues at each tax rate. This function is differentiable, with $f'[t] = \theta(1 - 2t)$. Characterize the possible stationary points by

$$1 - 2\hat{t} = 0$$

There is only one: $\hat{t} = 0.5$. Since $f''[t] = -2\theta$, the function is concave. The stationary point is therefore a maximizer of the function, the revenue maximizing tax rate is $1/2$. Figure 1.4 illustrates these considerations.

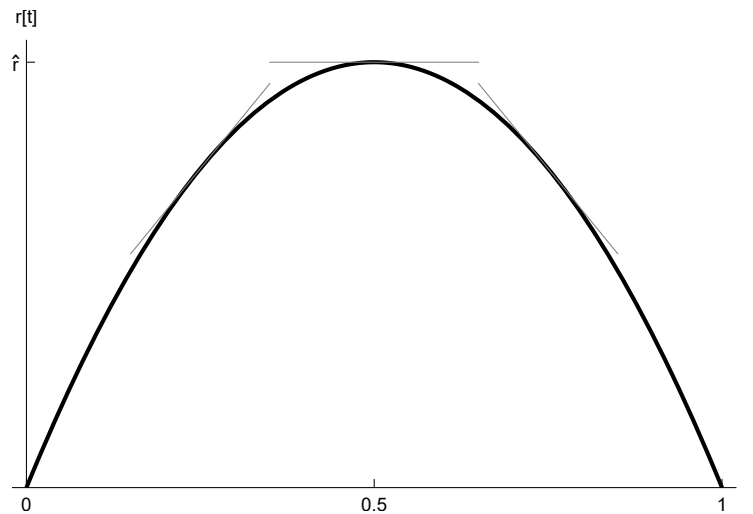


Figure 1.4: Tax Revenue vs Tax Rate

As the tax rate increases, tax revenue rises to a maximum at 0.5 but subsequently falls. The slope of the tax-revenue curve is zero at 0.5, satisfying the first-order necessary condition for an extremum. The slope is also continually decreasing, satisfying the second-order sufficient condition at the stationary point. In fact, the function is globally concave, so 0.5 is a unique maximizer.



1.2.3.1 Inflection Points

Definition 1.2 Consider function $f : X \subseteq \mathfrak{R} \rightarrow \mathfrak{R}$. A point $\hat{x}$ is an *inflection point* of f if f is concave on one side of $\hat{x}$ and convex on the other side. If f is twice differentiable, this means $f''[\hat{x}] = 0$ and f'' changes sign at $\hat{x}$.

A twice differentiable function may have an inflection point, which is an extremum of the slope. If the sign of the second derivative switches from positive to negative or vice versa around x_0 , then x_0 is an inflection point. In this case, the function is concave on one side of x_0 and convex on the other side. In this case, if f is twice continuously differentiable, then $f''[x_0] = 0$. This is just the usual first-order necessary condition for an extremum, but now applied to the derivative of the function f' .

Example 1.14 The function defined by $f = x \mapsto x^3$ has as its second derivative $f''[x] = 6x$, so f is concave to the left of 0 and convex to the right of 0. This means that $x = 0$ is a point of inflection. In more detail, $f'[0] = 0$, so we have a stationary point, but $f''[0] = 0$ too, so it may be an inflection point. That is, it may be an extremum of f' . Since $f'''[0] = 6 > 0$, $x = 0$ minimizes the slope and is indeed a point of inflection.

Consider the function $f = x \mapsto x^3 - 3x^2 + 5x + 1$, which has a first-order derivative $f' = x \mapsto 3x^2 - 6x + 5$. Search for an inflection point by looking for the stationary points of f' . The second-order derivative is $f'' = x \mapsto 6x - 6$, so this is a potential inflection point of f . Since $f'''[x] = 6 > 0$, we have indeed found a slope minimizer. The original function is concave to the left of 1 and convex to the right of 1. This time, the function is upward sloping at the inflection point.



1.2.3.2 Sufficient Conditions for an Extremum

Our first order necessary condition describes extrema and horizontal inflection points. To be confident of being at a maximum or a minimum, we need more information. For a strict local maximum, we additionally need the nearby values of the function to be smaller. For a strict local minimum, we additionally need the nearby values of the function to be larger. If we find a critical point, then we can try to verify that near the critical value the function always takes on smaller (or larger) values. This would be *sufficient* information to conclude that we have a local maximum (or minimum).

When a function f is adequately differentiable, the second-order derivative f'' may supply this extra information. Recall that the second-order derivative tells us about the *curvature* of the function. If a function is strictly-concave at a critical point, then we have found a local maximum. If a function is strictly-convex at a critical point, then we have found a local minimum.

Example 1.15 Suppose f is the squaring function ($f = x \mapsto x^2$), so that $f' = x \mapsto 2x$ and $f'' = x \mapsto 2$. We see that a critical point can be found at $x = 0$. Since $f''[0] = 2 > 0$, the function is convex near the critical point and we have therefore found a local minimum. (In fact, the function is globally convex, so we have found a global minimum.)

Suppose f is the cubing function ($f = x \mapsto x^3$), so that $f' = x \mapsto 3x^2$ and $f'' = x \mapsto 6x$. We see that a critical point can be found at $x = 0$. Since $f''[0] = 0$, the second-order condition does not suggest that the function is concave or convex near the critical point. And in fact, we have therefore found an inflection point.



Recall that for a differentiable function, a zero first-order derivative is a necessary condition for an extremum. If this necessary condition is satisfied, we found that a non-zero second-order derivative was sufficient for an extremum. But it is not necessary. We may have a maximum or minimum at a critical point where the second-order derivative is also zero. Such a point is called a [critical point!]degenerate critical point of f . This subsection briefly explores the information provided about a degenerate critical point by higher order derivatives.

Example 1.16 Consider the real function $f = x \mapsto x^4$. Since $f' = x \mapsto 4x^3$, this function is stationary at 0. However, $f'' = x \mapsto 12x^2$, so $f''[0] = 0$. Nevertheless, $x = 0$ is clearly the minimizer of the function: any nonzero input maps to a positive output.

□

Let us use $f^{(n)}[x]$ to denote the n th-order derivative of $f[x]$. Then if $f'[x_0] = 0$ and the first higher-order derivative that is non-zero at x_0 is $f^{(n)}[x_0]$, then we have

- an inflection point if n is odd
- a maximum if n is even and $f^{(n)}[x_0] < 0$
- a minimum if n is even and $f^{(n)}[x_0] > 0$

Generally, this will suffice for us determine whether a critical point is a maximizer, a minimizer, or an inflection point. (However, there are unusual functions with zero derivatives of all orders at an extremum.)

Exercise 1.3 Find the first-order and second-order derivative for the following functions: $x \mapsto x^2$, $x \mapsto -x^2$, $x \mapsto x^3 + 2x^2 + x + 1$, $x \mapsto e^x$, and $x \mapsto \ln x$. Determine whether the following functions are concave, convex, or neither. Characterize any stationary points, making use of the curvature information in the second derivative. Whenever an extremum exists, determine whether it is local or global, maximum or minimum.

[SKIP to constrained optimization](#)



1.2.3.3 Global Extrema

Consider the plot of the real function $f = x \mapsto x^4 - 2x^2$ in Figure 1.5. The first-order derivative function is $f' = x \mapsto 4x^3 - 4x$ which has zeros at $x = -1, 0, 1$. Clearly -1 and 1 are both global minimizers of f , while 0 is the only local maximizer and yet is not a global maximizer. This highlights some of the difficulties in searching for global extrema. However, in special cases things become much easier.

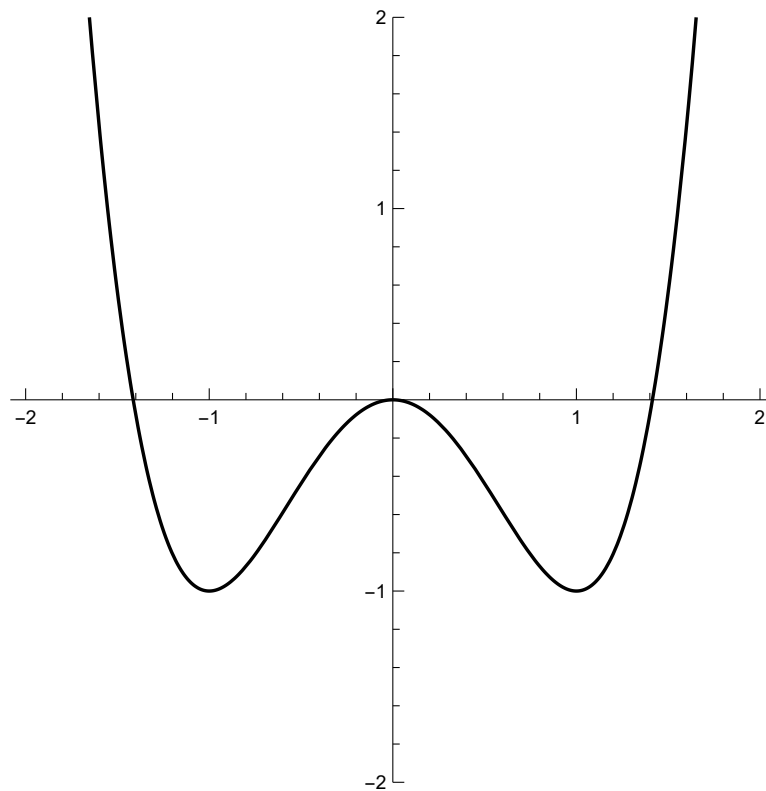


Figure 1.5: Global vs Local Extrema

The function $f = x \mapsto x^4 - 2x^2$ has three critical points, two of which are global minima and one of which is a local maximum.

In our first special case we search for an extremum over any interval. A local extremum that is also the only critical point must be a global extremum on this interval. For example, if $\hat{x}$ is a local maximizer and is the only critical point, then there is a neighborhood of $\hat{x}$ where $f' > 0$ to the left of $\hat{x}$ and $f' < 0$ to the right. The only way to reach a higher

point without a break in the function is for the derivative to change sign outside this neighborhood, which means it would have to become zero somewhere in addition to $\hat{x}$.

For the same reason, globally concave and globally convex functions can have at most one extremum. We continue to focus on twice differentiable functions. Thinking back to our discussion of local extrema, it is easy to guess that a globally concave function will have at most one critical point, which must be a global maximizer (if such exists).

In our second special case we search for an extremum over any compact (closed and bounded) set in the domain. Weierstrass proved that a continuous function always has a global maximum and a global minimum on such a set. This is obvious enough when we consider a differentiable function $f : \mathbb{R} \rightarrow \mathbb{R}$ on a closed interval. Just check the boundary values of the function, and compare them to the values at the critical points, and pick the global extrema.



1.3 Optimization Examples

Example 1.17 The general maxim in economics is that any activity should be undertaken at the level where marginal benefits equal marginal costs. We now apply this to profit maximization, where the marginal benefit of an increase in production is the marginal revenue derived from that production. Profits are the difference between total revenues and total costs.

$$\pi = q \mapsto \text{tr}[q] - \text{tc}[q]$$

The first-order necessary condition for profit maximization, that $\pi' = 0$, is therefore characterized as

$$\text{tr}'[q] - \text{tc}'[q] = 0 \tag{1.13}$$

That is, marginal revenue should equal marginal costs. Graphically, this has two ready interpretations. If we draw the $\text{tr}[q]$ and $\text{tc}[q]$ curves, we are looking for a value of q where they have the same slope. If we draw the mr and mc curves, we are looking for a value of q where they cross.

Of course, this is a necessary not a sufficient condition for a maximum. It also characterizes minima and inflections points. If we also find $\pi'' < 0$, then we are assured of a maximum. In this case

$$\text{tr}''[q] - \text{tc}''[q] < 0$$

at the critical point. This means that the slope of the MR curve should be less than the slope of the MC curve, or equivalently that the MC curve should cut the MR curve from below at the critical point.



SKIP remaining examples; go to Constrained Optimization

Example 1.18 (Revenue Maximizer for a Monopoly) Find the revenue maximizing quantity of production for a monopolist. Let the demand function be $q = p \mapsto a - bp$, where a and b are positive demand function parameters. Recall from example ?? that this demand function implies that the demand-price function is $p = q \mapsto a/b - (1/b)q$, total revenue is therefore $\text{tr} = q \mapsto (a/b)q - (1/b)q^2$, and the associated marginal revenue is $\text{mr} = q \mapsto (a/b) - (2/b)q$. Since $\text{mr} = \text{tr}'$, a stationary point of the total revenue function is a zero of the marginal revenue function.

$$0 = a/b - (2/b)q \implies q = a/2$$

Note that $\text{tr}'' = q \mapsto -2/b$ is a constant function with a negative value, so the stationary point is a revenue maximum. Remember, demand is price elastic when prices are high, so lowering the price below a/b is initially more than offset by the resulting increase in quantity.



Example 1.19 (Profit Maximizer for a Monopoly) Reuse the information in example 1.18, but now add a linear cost function: $tc = cq$. Here c is the marginal cost of production. This monopolist maximizes profits where $mr[q] = mc[q]$, so that

$$(a/b) - (2/b)q = c \implies q = (a/b - c)(b/2) = (a - bc)/2 < a/2$$

The profit maximizing level of production is less than the revenue maximizing level of production. Figure 1.6 illustrates this.

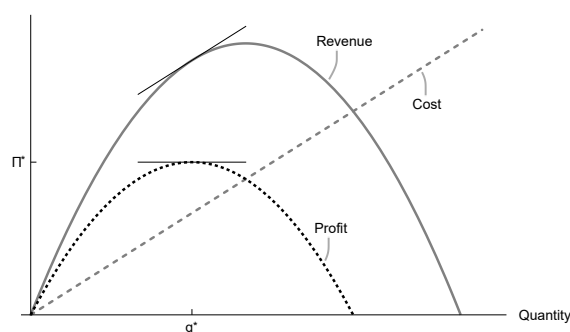


Figure 1.6: Optimum for a Monopolist

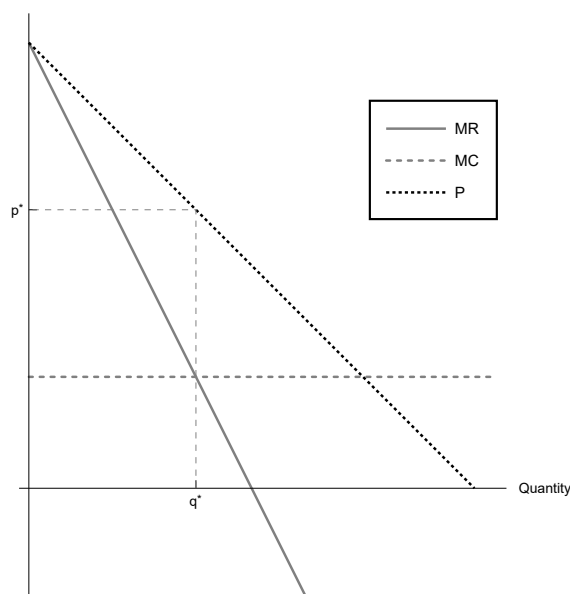


Figure 1.7: Optimum for a Monopolist



Example 1.20 A simple algebraic formulation of the notion of increasing marginal costs of production is the quadratic cost function $tc = q \mapsto \alpha q^2$, with $\alpha > 0$ and q taking on any nonnegative level of production. The total revenue of a price-taking firm is proportion to the market price: $tr = q \mapsto \bar{p}q$. Profits are the difference between total revenue and the cost of production: $\Pi = q \mapsto \bar{p}q - \alpha q^2$. The first order necessary condition that $\Pi'[q] = 0$ naturally equates marginal revenue to marginal cost. That is, $\bar{p} = 2\alpha q$, which describes a single stationary point:

$$q = \bar{p}/(2\alpha)$$

The second-order derivative is a constant function, with $\Pi''[q] = -2\alpha < 0$ at any level of q including the stationary point, so the profit function is concave and the stationary point is a unique maximizer. Note that taken literally this model places no upper bound on revenues: revenue maximization is not profit maximization. Note that taken literally this model places a zero lower bound on costs: in this sense, cost minimization is not profit maximization.

Exercise 1.4 Consider a monopolist with constant marginal costs. Suppose the inverse demand curve is

$$P = Q^{-1/\beta}$$

where $\beta > 0$. Show that β is the elasticity of demand. Find the first order condition for profit maximization. Demonstrate that satisfaction of the second order condition depends on the value of β and suggest why.

Exercise 1.5 Suppose the monetary authority dislikes unemployment and really dislikes high inflation, as represented the the loss function $L = U + \pi^2$. Here U is the difference between the unemployment rate and a target rate of unemployment, and π is the inflation rate. The monetary authority picks the inflation rate to minimize its loss function. We suppose U to be determined by $U = 0.5(\pi - \pi^e)$. (You may think of this as a “Lucas supply curve”.) Here π^e is expected inflation. Again, actual inflation is the choice variable for the monetary authority. If the monetary authority treats expected inflation as a fixed constant, what inflation rate will it set? If the monetary authority treats expected inflation as always equal to actual inflation, what inflation rate will it set?

Exercise 1.6 Anton et al. (2002) offer an example similar to the following. An oil well is located 6km offshore. The refinery is located 8km further down the coast. Laying pipe costs \$1M/km under water and \$0.5M/km over land. Assume pipe is laid in a straight line, whether underwater or along the coast. What is the cheapest way to lay pipe from the well to the refinery?

General Discussion: The Pythagorean theorem implies that the well is 10km from the refinery across the water ($\sqrt{6^2 + 8^2} = 10$). So laying underwater pipe straight to the refinery costs \$10M. The oil well is 6km

offshore, so land can be reached with \$6M in costs, but then it costs another \$4M to lay 8km over land. Both approaches involve a total cost also of \$10M. Is there no cheaper proposal?



Example 1.21 A wine dealer owns a case of fine wine that can be sold for $1,000e^{\sqrt{2}t}$ if she holds onto it for t years. There are no storage costs, but the present value of the wine is discounted at the (continuously compounded) interest rate is 10% per annum. At what date t should the dealer sell her wine? At time t , the current value is $fv[t] = 1,000e^{\sqrt{2}t}$. Discounting back to the present yields

$$pv = t \mapsto e^{-0.1t} * 1000 * e^{\sqrt{2}t} = e^{\sqrt{2}t - 0.1t} * 1000$$

Find the first-order derivative, using the exponential rule and the chain rule.

$$pv' = t \mapsto (\sqrt{2/t} - 0.1)e^{\sqrt{2}t - 0.1t} * 1000$$

To maximize the present value of the wine investment, set to zero the first derivative with respect to t , and solve for t .

$$0 = 1/\sqrt{2t} - 0.1 \implies t = 50$$

Note that this is the point in time where the rate of growth of the future value equals the discount rate.

Let us try an even simpler approach. Any strictly increasing transformation will have the same extrema, so use the natural log to produce the equivalent objective function is $f = t \mapsto \ln 1000 - 0.1t + \sqrt{2}t$. Take the derivative with respect to t and set it equal to zero.

$$0 = -0.1 + 1/\sqrt{2t}$$

which again has the solution $t = 50$. Be sure to check the second order condition!

$$f''[t] = -1/(2\sqrt{2}t^{3/2}) < 0$$

The function is concave, so the stationary point is a maximizer.

You can easily check your work with **WL**:

```
wpv = 1000*Exp[Sqrt[2 t]]*Exp[-0.1 t]
dwpvdt = D[wpv, t]
soln = Solve[dwpvdt == 0, t] // Flatten
D[dwpvdt, t] /. soln (* check 2nd order condition *)
```

Nonlinear Comparative Statics: First Approach

Structural form for a simple IS-LM model:

$$Y = A[i^-, \pi^+, Y^+, F^+]$$

$$m = L[i^-, Y^+]$$

Reduced form for a simple IS-LM model:

$$i = i[m, \pi, F]$$

$$Y = Y[m, \pi, F]$$

Given the structural form and its responses, we want to sign the responses of the reduced form as follows: produce the differential of the *system*, and solve for the change in the endogs in terms of the exogs. Start with the money market.

$$m = L[i, Y]$$

$$dm = dL$$

$$dL = L_i di + L_Y dY$$

$$dm = \overbrace{L_i di + L_Y dY}^{dL}$$

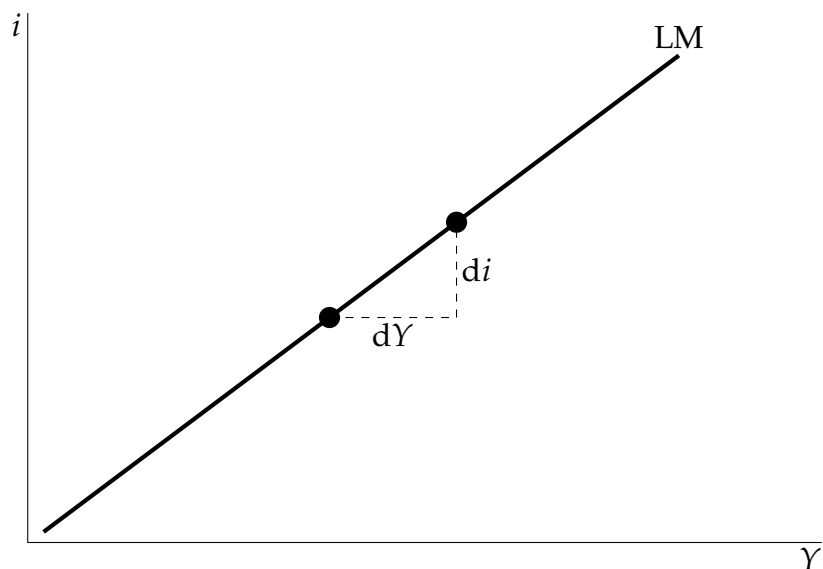


Figure 1.1: Slope of LM Curve (Keynesian Model)
In a simple Keynesian macromodel, the LM curve illustrates the possible money market equilibria given a fixed value of the real money supply.

REVIEW money market: $m = L[\bar{i}, Y]$ (and cover IS-LM homework)

$$dm = dL$$

$$dL = L_i di + L_Y dY$$

$$dm = L_i di + L_Y dY$$

The total differential can be used to find the slope of the LM curve, which shows how i and Y must change together, given m . To represent the constancy of m , set $dm = 0$ in the total differential of the LM equation.

$$0 = L_i di + L_Y dY \quad (1.43)$$

Let this characterize an implicit functional dependence of i and Y , and solve for the implied slope of the LM curve.

$$\left. \frac{di}{dY} \right|_{LM} = -\frac{L_Y}{L_i} > 0 \quad (1.44)$$

The LM curve represents the way i and Y must change together to maintain equilibrium in the money market, *ceteris paribus*. That is, this determines the slope of the “Keynesian” LM curve. Under the standard assumptions that $L_Y > 0$ and $L_i < 0$, the “Keynesian” LM curve has a positive slope.

$$Y = A[i - \pi, Y, F]$$

Next consider the goods market. Recall that the equation $Y = A[i - \pi, Y, F]$ represents equilibrium in the goods market. This must hold both before and after any exogenous changes. That is, we require that we start out in goods market equilibrium, and we also require that we end up in goods market equilibrium. It follows that the changes in real income must equal the changes in real aggregate demand. Looking at the equation for the IS curve, we can see that this means that the change in real income (dY) must equal the change in real aggregate demand (dA).

$$dY = dA \tag{1.45}$$

$$Y = A[i - \pi, Y, F]$$

The change in aggregate demand has three sources: changes in r , changes in Y , and changes in F . We represent these changes as dr , dY , and dF . Of course, the changes in aggregate demand depend not only on the size of the changes in these arguments, but also on how sensitive aggregate demand is to each of these arguments.

$$dA = A_r \overbrace{(di - d\pi)}^{dr} + A_Y dY + A_F dF \quad (1.46)$$

Putting these two pieces together yields the total differential of the IS equation:

$$dY = \overbrace{A_r(di - d\pi) + A_Y dY + A_F dF}^{dA} \quad (1.47)$$

Note that $A[\cdot, \cdot, \cdot]$ has only three arguments. Do not be misled by the fact that we choose to write r as $i - \pi$. This does not change the number of arguments of the aggregate demand function. For example, there is no partial derivative A_i .

Notice that if only i and Y change, we must have

$$dY = A_r di + A_Y dY \quad (1.48)$$

$$\left. \frac{di}{dY} \right|_{IS} = \frac{1 - A_Y}{A_r} < 0 \quad (1.49)$$

This is the way i and Y must change together to maintain equilibrium in the goods market. That is, this determines the slope of the “Keynesian” IS curve. Under the standard assumptions that $0 < A_Y < 1$ and $A_r < 0$, the “Keynesian” IS curve has a negative slope.

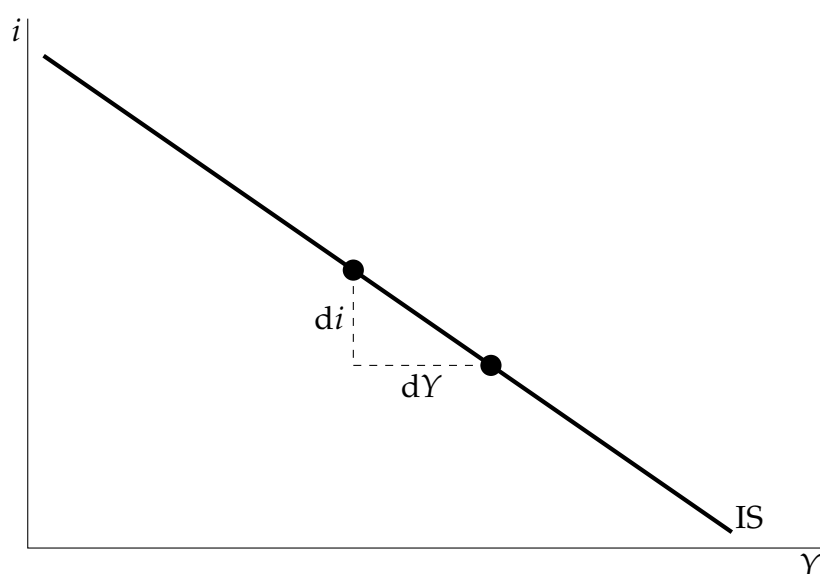


Figure 1.2: Slope of IS Curve (Keynesian Model)
In a simple Keynesian macromodel, the IS curve illustrates the possible goods market equilibria given a fixed value F and π .

The total differential of the IS equation states the requirements for staying on the IS curve. The total differential of the LM equation states the requirements for staying on the LM curve. Make these two statements simultaneously to characterize the requirements for staying on both curves.

$$dY = A_r(di - d\pi) + A_Y dY + A_F dF \quad (1.50)$$

$$dm = L_i di + L_Y dY \quad (1.51)$$

Starting from an initial equilibrium in the goods and money markets, these two total differentials together state the interrelationships required to remain on both the IS and LM curves simultaneously.

A simple Keynesian macromodel focuses on the endogenous determination of i and Y . (Note that this is the first mention of which variables are endogenous.) In terms to the two total differentials, this means that di and dY are the endogenous variables of the system. These are endogenously determined to ensure the simultaneous satisfaction of the IS and LM equations. This is a *linear* system in the two new endogenous variables, di and dY .

It can be helpful to prepare to set up the system as a matrix equation by moving all terms involving the endogenous variables to the left.

$$-A_r di + dY - A_Y dY = -A_r d\pi + A_F dF \quad (1.52)$$

$$L_i di + L_Y dY = dm \quad (1.53)$$

Now rewrite this 2-equation system as a matrix equation in the form $\mathbf{J} \cdot d\vec{x} = d\vec{b}$.

$$\begin{bmatrix} -A_r & (1 - A_Y) \\ L_i & L_Y \end{bmatrix} \begin{bmatrix} di \\ dY \end{bmatrix} = \begin{bmatrix} -A_r d\pi + A_F dF \\ dm \end{bmatrix} \quad (1.54)$$

Next, solve for the reduced form. One approach to this solution is the method of inverses: multiplying both sides by $\mathbf{J}^{-1}$.

$$\begin{aligned} \begin{bmatrix} di \\ dY \end{bmatrix} &= \frac{1}{-A_r L_Y - (1 - A_Y) L_i} \begin{bmatrix} L_Y & -(1 - A_Y) \\ -L_i & -A_r \end{bmatrix} \begin{bmatrix} -A_r d\pi + A_F dF \\ dm \end{bmatrix} \\ &= \frac{1}{A_r L_Y + (1 - A_Y) L_i} \begin{bmatrix} -L_Y & (1 - A_Y) \\ L_i & A_r \end{bmatrix} \begin{bmatrix} -A_r d\pi + A_F dF \\ dm \end{bmatrix} \end{aligned} \quad (1.55)$$

Qualitative comparative statics attempts to determine the qualitative responses of the endogenous variable to the exogenous variables. To achieve this requires some qualitative information about the structural equations, embodied in the assumed signs of their partial derivatives. Let $\Delta = A_r L_Y + (1 - A_Y) L_i$, and notice that our assumptions on the partial derivatives of the structural functions imply that $\Delta < 0$. The reduced form can then be rewritten as

$$\begin{bmatrix} di \\ dY \end{bmatrix} = \frac{1}{\Delta} \begin{bmatrix} -L_Y & (1 - A_Y) \\ L_i & A_r \end{bmatrix} \begin{bmatrix} -A_r d\pi + A_F dF \\ dm \end{bmatrix} \quad (1.56)$$

The implicit function theorem provides the conditions for us to produce this reduced form solution. It ensures that implicit in the original system are functions relating i and Y to the exogenous variables. To characterize the partial derivatives of the implicit reduced form, use this reduced form to consider one exogenous change at a time. It turns out that we have enough information to attach a sign to all of these partial derivatives. These predicted signs are the qualitative comparative statics for this model.

$$\begin{aligned}\begin{bmatrix} \partial i / \partial m \\ \partial Y / \partial m \end{bmatrix} &= \frac{1}{\Delta} \begin{bmatrix} -L_Y & (1 - A_Y) \\ L_i & A_r \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} \\ &= \frac{1}{\Delta} \begin{bmatrix} (1 - A_Y) \\ A_r \end{bmatrix} = \begin{bmatrix} - \\ + \end{bmatrix}\end{aligned}\tag{1.57}$$

$$\begin{aligned}\begin{bmatrix} \partial i / \partial \pi \\ \partial Y / \partial \pi \end{bmatrix} &= \frac{1}{\Delta} \begin{bmatrix} -L_Y & (1 - A_Y) \\ L_i & A_r \end{bmatrix} \begin{bmatrix} -A_r \\ 0 \end{bmatrix} \\ &= \frac{1}{\Delta} \begin{bmatrix} L_Y A_r \\ -L_i A_r \end{bmatrix} = \begin{bmatrix} + \\ + \end{bmatrix}\end{aligned}\tag{1.58}$$

$$\begin{aligned}\begin{bmatrix} \partial i / \partial F \\ \partial Y / \partial F \end{bmatrix} &= \frac{1}{\Delta} \begin{bmatrix} -L_Y & (1 - A_Y) \\ L_i & A_r \end{bmatrix} \begin{bmatrix} A_F \\ 0 \end{bmatrix} \\ &= \frac{1}{\Delta} \begin{bmatrix} -L_Y A_F \\ L_i A_F \end{bmatrix} = \begin{bmatrix} + \\ + \end{bmatrix}\end{aligned}\tag{1.59}$$

Multivariate Optimization: Overview

1.1.1 First-Order Necessary Condition

A *stationary point* $\bar{x}$ of the differentiable real function f is a point where its gradient is null: $\nabla f[\bar{x}] = 0$.¹ An interior local maximum of f must occur at a stationary point, since changing the input a little must not increase the value of the function. Any nonzero partial derivative of f would mean that a small change in the input could increase the value of the function. Similar reasoning applies for a local minimum. So a first-order necessary condition for an extremum is a null gradient.

Theorem 1.1 (First-Order Necessary Condition) Suppose the real function f is differentiable on an open set X . If $\hat{x} \in X$ is a local extreme point of f , then $\nabla f[\hat{x}] = 0$.

Proof: For each x_i , consider the univariate mapping $x_i \mapsto f[\hat{x}_1, \dots, x_i, \dots, \hat{x}_N]$ and apply the familiar univariate reasoning. This implies $\partial f[\hat{x}]/\partial x_i = 0$ for every x_i . ■

Example 1.1 The real function $f = x_1, x_2 \mapsto x_1^2 - 4x_1 + x_2^2$ has the gradient $\nabla f[x_1, x_2] = \langle 2x_1 - 4, 2x_2 \rangle$ so it has a stationary point at $\langle x_1, x_2 \rangle = \langle 2, 0 \rangle$. This is the only stationary point so it is the only possible extreme point. (As we shall see, it is a minimizer.)

```
f = {x1, x2} |-> x1^2 - 4 x1 + x2^2;
D[f[x1, x2], {{x1, x2}}]
Solve[0 == %, {x1, x2}]
```

¹Recall that the vector of first-order partial derivatives is called the *gradient* of f , commonly denoted by ∇f , Df or even more simply by f' . Other fairly common notations include $D_x f$, f_x , $\partial f / \partial \vec{x}^\top$, or $\nabla_x f$.

To explore this necessary condition for an interior extremum, let $f: \mathbb{R}^N \rightarrow \mathbb{R}$ be differentiable at $\vec{x}$. Choose an arbitrary vector $\vec{h} \in \mathbb{R}^N$, define the unary function g by

$$g = \alpha \mapsto f[\vec{x} + \alpha\vec{h}] - f[\vec{x}] \quad (1.1)$$

If f reaches a maximum at $\vec{x}$, then g reaches a maximum of 0 at $\alpha = 0$. Since f is differentiable at $\vec{x}$, the unary function g is differentiable at 0. By the chain rule, the first-order derivative of g is

$$g' = \alpha \mapsto \nabla f[\vec{x} + \alpha\vec{h}] \cdot \vec{h} \quad (1.2)$$

It is necessary that $g'[0] = 0$ at an extremum, regardless of the choice of $\vec{h}$. Therefore ∇f must be zero at the maximizing input $\vec{x}$.

Example 1.2 The binary real function $f = \langle x, y \rangle \mapsto 1 - x^2 - y^2$ is evidently maximized at $\langle 0, 0 \rangle$. Its gradient is $\nabla f = \langle x, y \rangle \mapsto \langle -2x, -2y \rangle$. Let $\vec{h} = \langle h_x, h_y \rangle$ be an arbitrary 2-vector, and for any given $\vec{x}$ define a *unary* real function $g = \alpha \mapsto f[\vec{x} + \alpha\vec{h}] - f[\vec{x}]$. Then

$$\begin{aligned} g[\alpha] &= (1 - (x + \alpha h_x)^2 - (y + \alpha h_y)^2) - (1 - x^2 - y^2) \\ &= -(\alpha^2 h_x^2 + \alpha^2 h_y^2 + 2\alpha x h_x + 2\alpha y h_y) \end{aligned}$$

Find the derivative of g :

$$\begin{aligned} g'[\alpha] &= -(2\alpha h_x^2 + 2\alpha h_y^2 + 2x h_x + 2y h_y) \\ &= \langle -2(x + \alpha h_x), -2(y + \alpha h_y) \rangle \cdot \langle h_x, h_y \rangle \end{aligned}$$

Recall that if f reaches a maximum at $\vec{x}$ then g reaches a maximum at $\alpha = 0$. It follows that $\nabla f[\vec{x}] \cdot \vec{h} = 0$ must hold regardless of the value of $\vec{h}$.

$$\langle -2x, -2y \rangle \cdot \langle h_x, h_y \rangle = \vec{0}$$

This is satisfied only at $\langle x, y \rangle = \vec{0}$.

```
g = alpha |-> f[x + alpha*hx, y + alpha*hy] - f[x, y];
g'
f = {x, y} |-> 1 - (x^2 + y^2)
g'[alpha] // Expand
% == Grad[f[x + alpha*hx, y + alpha*hy], {x, y}] . {hx, hy} // Expand
```

1.1.2 Second-Order Conditions

Theorem 1.1 suggests a strategy for finding the extrem points of differentiable functions: impose $\nabla f[x] = 0$, and find the points that simultaneously satisfy the entire resulting system of equations. Any interior extreme point is a stationary point of the function, so we can restrict consideration to these solutions. But is a stationary $\bar{x}$ a maximizer, or a minimizer, or neither? Just as in the univariate case, pinning this down requires an exploration of the curvature of the function at the stationary point. If the function is locally concave there, it is a local maximizer. If it is locally convex there, it is a local minimizer.

□

1.1.2.1 Second-Order Partial Derivatives

If f is an N -ary differentiable real function, then it has N first-order partial derivatives, each of which are real functions. Recall that the gradient ∇f is the array of these first-order partial derivatives. If all first-order partial derivatives of f are continuous on the set X , we call f *continuously differentiable* on X . The notation $f \in \mathcal{C}_1$ denotes continuous differentiability.

Each first-order partial derivative may in turn be differentiable with respect to each of the N function arguments. The derivative of a first-order partial derivative is a second-order partial derivative. The N -ary function f can therefore have N^2 second-order partial derivatives. The notation $\nabla^2 f$ denotes the matrix of these second-order partial derivatives. The notation $\nabla_{i,j}^2 f$ denotes the i, j -th element of this matrix, which is the i -th partial derivative of the j -th partial derivative of f . This is $\partial_i(\partial_j f)$, also written in compact delta notation as $\partial_{i,j}^2 f$. Verbose delta notation for this second-order derivative is also common:

$$\nabla^2 f = \left[\frac{\partial^2 f}{\partial x_i \partial x_j} \right] \quad (1.3)$$

□

□

Once again economists use a variety of notations to denote this second-order partial derivative, the most common of which is $f_{i,j}$. If all second-order partial derivatives of f are continuous on the open set X , we call f *twice-continuously differentiable* on X . We write this $f \in \mathcal{C}_2$. If $f \in \mathcal{C}_2$, then $f_{i,j} = f_{j,i}$. Economists know this as Young's theorem; it is also called Clairaut's theorem or Schwartz's theorem.

Theorem 1.2 (Young's Theorem) Let $f[x]$ be twice-continuously differentiable. Then

$$\frac{\partial^2 f}{\partial x_i \partial x_j} = \frac{\partial^2 f}{\partial x_j \partial x_i}$$

That is, the order of differentiation does not matter. This is often called the *symmetry* of second-order partial derivatives.

Proof: <https://www.youtube.com/watch?v=wLkXUmVKyi4>

■

□

□

The matrix of second-order partial derivatives is the *hessian* matrix of f . Once again, economists use a variety of notations to designate the hessian matrix, including $\mathbf{H}f$, $D^2f[\mathbf{x}]$, and $f_{\tilde{\mathbf{x}}\tilde{\mathbf{x}}^\top}$ are all common representations. Perhaps the simplest notation is $f''[\mathbf{x}]$, which draws on the unary convention for a second-order derivative. This course generally denotes the value of hessian matrix of the function f at the point $\mathbf{x}$ by $\nabla_f^2[\mathbf{x}]$,

$$\nabla_f^2 = \mathbf{x} \mapsto [f_{i,j}[\mathbf{x}]] = \begin{bmatrix} f_{1,1}[\mathbf{x}] & \dots & f_{1,N}[\mathbf{x}] \\ \vdots & & \vdots \\ f_{N,1}[\mathbf{x}] & \dots & f_{N,N}[\mathbf{x}] \end{bmatrix} \quad (1.4)$$

Example 1.3 Suppose we have the Cobb-Douglas production function $f[K, N] = K^\alpha N^{1-\alpha}$. To find the hessian matrix, begin by finding the gradient. Letting $\mathbf{x} = [K, N]$, the gradient is

$$\frac{\partial f}{\partial \mathbf{x}} = [\alpha(K/N)^{\alpha-1}, (1-\alpha)(K/N)^\alpha]$$

Next, differentiate each element of the gradient with respect to each variable.

$$\begin{aligned} \nabla_f^2[\mathbf{x}] &= \frac{\partial^2 f[\mathbf{x}]}{\partial \mathbf{x} \partial \mathbf{x}} = \begin{bmatrix} \frac{\partial^2 f[\mathbf{x}]}{\partial K \partial K} & \frac{\partial^2 f[\mathbf{x}]}{\partial K \partial N} \\ \frac{\partial^2 f[\mathbf{x}]}{\partial N \partial K} & \frac{\partial^2 f[\mathbf{x}]}{\partial N \partial N} \end{bmatrix} \\ &= \begin{bmatrix} (\alpha-1)\alpha K^{\alpha-2} N^{1-\alpha} & (1-\alpha)\alpha K^{\alpha-1} N^{-\alpha} \\ (1-\alpha)\alpha K^{\alpha-1} N^{-\alpha} & -\alpha(1-\alpha)K^\alpha N^{-(1+\alpha)} \end{bmatrix} \end{aligned}$$

Along the diagonal we find negative elements, reflecting the *diminishing marginal productivity* of each input. In contrast, the cross-partials are positive, since an increase in one factor increases the marginal productivity of the other factor.

The hessian matrix summarizes information about the curvature of the function at a point. The function is strictly concave if the hessian matrix is negative definite. The function is strictly convex if the hessian matrix is positive definite. The function is concave if the hessian matrix is negative semi-definite. The function is convex if the hessian matrix is positive semi-definite.

1.1.2.2 Hessian Matrix

To explore the curvature, let $\nabla^2 f$ produce the matrix of second-order partial derivatives of the twice differentiable function f . When $f: \mathbb{R}^N \rightarrow \mathbb{R}$ is continuously twice differentiable near a point, its hessian at a point provides information about its local curvature. In terms of the representative $\langle i, j \rangle$ -th matrix element,

$$\nabla^2 f = \mathbf{x} \mapsto \left[\frac{\partial^2 f[\mathbf{x}]}{\partial x_i \partial x_j} \right] \quad (1.5)$$

To reduce notation, let $f_{i,j}$ represent the second-order partial derivative with respect to the i -th and j -th arguments.²

$$\nabla^2 f = \mathbf{x} \mapsto [f_{i,j}[\mathbf{x}]] \quad (1.6)$$

This matrix of second-order partials derivatives is the ***hessian*** matrix. If f is continuously twice-differentiable then $f_{i,j} = f_{j,i}$, so the hessian matrix is symmetric.

For the function f , the Hoy book uses two notations for the Hessian: H and $\nabla^2 f$. Typically H denotes the value of the hessian at a specific point (such as a stationary point).

²There are many different notations for the hessian, including $D^2 f$, H_f , and $\frac{\partial^2 f[\hat{\mathbf{x}}]}{\partial \mathbf{x} \partial \mathbf{x}^T}$. (The last treats the input argument as a column vector.)

1.1.2.3 Second-Order Condition

$$\begin{aligned}g &= \alpha \mapsto f[\vec{x} + \alpha \vec{h}] - f[\vec{x}] \\g' &= \alpha \mapsto \nabla f[\vec{x} + \alpha \vec{h}] \cdot \vec{h} \\g'' &= \alpha \mapsto \vec{h} \cdot \nabla^2 f[\vec{x} + \alpha \vec{h}] \cdot \vec{h}\end{aligned}$$

At a local maximum the hessian matrix must be negative semidefinite (sometimes called nonpositive definite). To explore this, once again consider the univariate function defined in (1.1). Our work on the univariate case tells us that if $\hat{x}$ maximizes f then g must be concave at 0.

$$g''[0] = \vec{h} \cdot \nabla^2 f[\hat{x}] \cdot \vec{h} \leq 0 \tag{1.7}$$

Since $\vec{h}$ is arbitrary, the hessian matrix must be negative semidefinite. This is a second-order necessary condition for a maximum. The same argument shows that at a minimum the hessian matrix must be positive semidefinite (sometimes called nonnegative definite).

The continuing parallels with the univariate case raises a natural question: does a negative definite hessian matrix at a stationary point imply local concavity of the objective function, and is it a sufficient condition for stationary point to be a strict local maximizer? The answer to both questions is yes.

Theorem 1.3 Let $\hat{x}$ be a stationary point of the continuously twice-differentiable function $f: X \subseteq \mathbb{R}^N \rightarrow \mathbb{R}$. If $\nabla_f^2[\hat{x}]$ is negative definite, then $\hat{x}$ is a strict local maximizer of f . (Similarly, if $\nabla_f^2[\hat{x}]$ is positive definite, then $\hat{x}$ is a strict local minimizer of f .)

Proof: See Simon and Blume (1994, ch.30). ■

Summary for a *binary* function.

At a stationary point, a necessary 2nd-order condition

- for a strict max is negative elements on the hessian diagonal
- for a strict min is positive elements on the hessian diagonal

Given these, in the 2d case, it is then

sufficient for that the hessian determinant be positive

The matrix $H = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ with $b = c$ (so that it is symmetric) is

- nonpositive definite if $a, d \leq 0$ and $\det[H] \geq 0$.
- negative definite if $a, d < 0$ and $\det[H] > 0$.
- positive definite if $a, d > 0$ and $\det[H] > 0$.
- nonnegative definite if $a, d \geq 0$ and $\det[H] \geq 0$.

Otherwise, it is *indefinite*.

Note: Each case can also be characterized in terms of the eigenvalues of H .

- nonpositive definite iff $\lambda_1, \lambda_2 \leq 0$.
- negative definite iff $\lambda_1, \lambda_2 < 0$.
- positive definite iff $\lambda_1, \lambda_2 > 0$.
- nonnegative definite iff $\lambda_1, \lambda_2 \geq 0$.

Definition 1.1 Let $\vec{x}$ be a stationary point of the continuously twice-differentiable function $f: \mathbb{R}^N \rightarrow \mathbb{R}$. If the hessian matrix $\nabla^2 f[\vec{x}]$ is indefinite, then $\vec{x}$ is called a *saddle point* of f .

Theorem 1.4 A saddle point is not a minimizer or a maximizer.

Proof: See Simon and Blume (1994, ch.30). ■

Example 1.4 Consider the quadratic form $f = \langle x_1, x_2 \rangle \mapsto x_1^2 - x_2^2$. Look for stationary points by setting $\nabla f[\vec{x}] = \vec{0}$.

$$\langle 2x_1, -2x_2 \rangle = \langle 0, 0 \rangle$$

Solve this system of equations to find a stationary point at $\vec{0}$. Next, compute the value of the hessian matrix at this stationary point to find

$$\nabla^2 f[\vec{0}] = \begin{bmatrix} 2 & 0 \\ 0 & -2 \end{bmatrix}$$

which is indefinite. The stationary point is therefore a saddle point. Specifically, changes in x_1 will increase the value of the function, while changes in x_2 will decrease the value of the function.

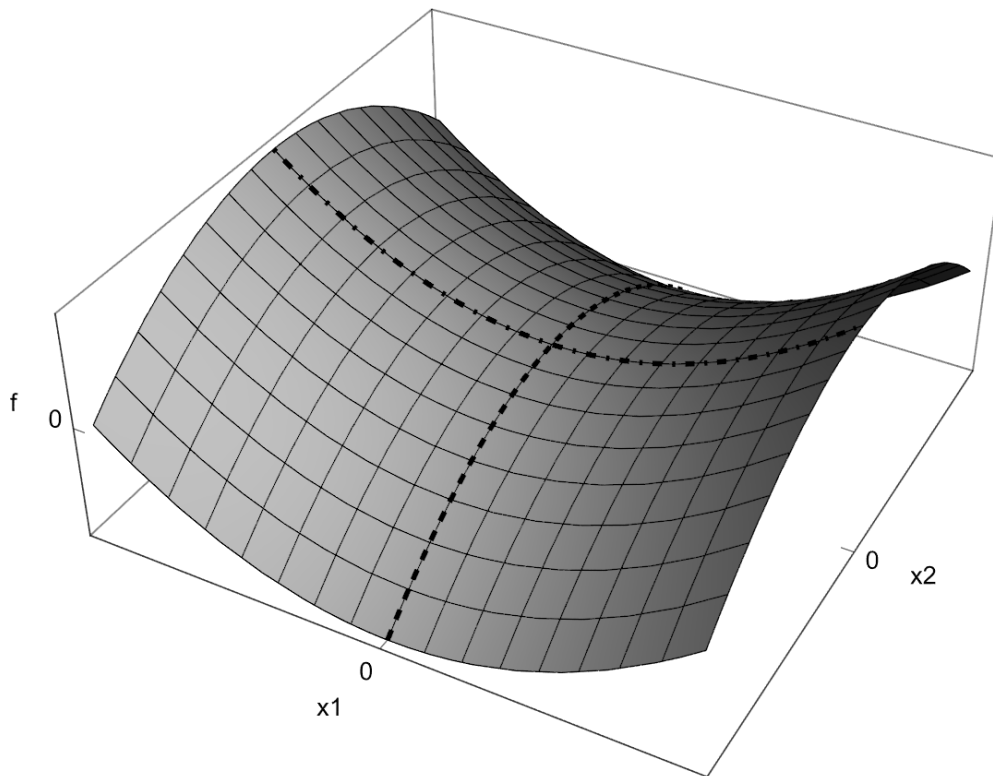


Figure 1.1: Indefinite Quadratic Form

Suppose the true model is

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (1.23)$$

where the $\boldsymbol{\varepsilon}$ are mean-zero random shocks to the underlying relationship. These shocks influence our estimate $\hat{\boldsymbol{\beta}}$.

$$\begin{aligned} \hat{\boldsymbol{\beta}} &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} \\ &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top (\mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}) \\ &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon} \\ &= \boldsymbol{\beta} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon} \end{aligned} \quad (1.24)$$

Consider the expected value of $\hat{\boldsymbol{\beta}}$. However, if $\mathbf{X}$ is uncorrelated with $\boldsymbol{\varepsilon}$, we have an unbiased estimate of the true coefficient parameter.⁵

⁵Correspondingly, write the residual vector as

$$\begin{aligned} \hat{\boldsymbol{\varepsilon}} &= (\mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top) \mathbf{y} \\ &= (\mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top) (\mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}) \\ &= (\mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top) \boldsymbol{\varepsilon} \end{aligned}$$

1.2.3 Omitted Variables

Suppose we obtain the regression estimates

$$Y = b_0 + b_1 X_1 + U \quad (1.25)$$

but X_2 should have been in the regression. That is we should have gotten the estimates

$$Y = b'_0 + b'_1 X_1 + b'_2 X_2 + U' \quad (1.26)$$

How are the two estimates related?

We will explore this by considering the supplementary regression of X_2 on X_1 :

$$X_2 = c_0 + c_1 X_1 + V \quad (1.27)$$

In every case the regression holds exactly on average. (This follows from the OLS regression, which yields a zero average error.) Therefore the deviations of Y from its mean, y , can be written

$$y = b_1 x_1 + U \quad (1.28)$$

or

$$y = b'_1 x_1 + b'_2 x_2 + U' \quad (1.29)$$

Note that the errors are unchanged, since they have a zero mean value.

Similarly, the deviation of X_2 from its mean can be written

$$x_2 = c_1 x_1 + V \quad (1.30)$$

Substitute this into (1.27) to get

$$y = b'_1 x_1 + b'_2 (c_1 x_1 + V) + U' \quad (1.31)$$

Then subtract (1.29) from (1.26) to get

$$\mathbf{0} = (b_1 - b'_1 - b'_2 c_1) x_1 + U - U' - b'_2 V \quad (1.32)$$

The final step is to sum (1.30) over all the observations, recalling the errors U , U' , and V are each zero on average.

$$0 = (b_1 - b'_1 - b'_2 c_1) \bar{x}_1 \quad (1.33)$$

or

$$b_1 = b'_1 + b'_2 c_1 \quad (1.34)$$

So omission of X_2 from the regression biases the coefficient estimate to an extent that depends on the correlation between the omitted and included variables.

Note: Oksanen (1998) notes that these results hold for any regressions yielding zero average errors.

Note: the derivation works even when b'_1 and b'_2 are vectors.

1.2.4 Instrumental Variables

Optional reading: Davidson and McKinnon 7.4, Pindyck and Rubinfeld (1997, ch.7)

But what if X is correlated with the error term? Well perhaps we could generate a proxy for X that is not. For example, suppose we have a collection of variables Z that are not correlated with the error term but that we think can proxy X .

$$X = Z\gamma + v \quad (1.35)$$

Then we can produce our X proxy as

$$\hat{X} = Z(Z'Z)^{-1}Z'X = P_Z X \quad (1.36)$$

Now reformulate our original problem as

$$Y = X\beta + u = P_Z X\beta + u + \hat{v}\beta \quad (1.37)$$

and produce the standard least squares estimate of β :

$$\hat{\beta} = (X'P_Z'P_Z X)^{-1}X'P_Z'Y \quad (1.38)$$

Noting that P_Z is symmetric idempotent, we have

$$\hat{\beta} = (X'P_Z X)^{-1}X'P_Z Y \quad (1.39)$$

This is the instrumental variables estimator, where the instruments are Z .

This estimator is also known as the 2SLS estimator, because it can be easily produced as two linear regressions: first regress X on the instruments, then regress Y on the fitted values of X .

1.3 Constrained Optimization

Constrained optimization adds new technical difficulties. Just as in the univariate case, binding constraints generally invalidate the first-order and second-order necessary conditions that we developed in the unconstrained case. This section focuses on the consequences of binding equality constraints in a multivariate optimization problem. This section considers a single constraint: in addition to an objective function $f[\vec{x}]$, there is a constraint that $g[\vec{x}] = 0$. Both functions are assumed to be twice continuously differentiable ($f, g \in C^2$). This section focuses on constrained maximization of a real valued n -ary function subject to a single equality constraint.

1.3.1 First-Order Necessary Conditions

The *feasible set* F is the set of values that satisfy the constraint.

$$F = \{ \vec{x} \mid g[\vec{x}] = 0 \} \quad (1.40)$$

A constrained maximization problem considers only the values that $f[\vec{x}]$ takes on the specified feasible set. The constrained maximization problem of this section is

$$\max_{\vec{x} \in F} f[\vec{x}] \quad (1.41)$$

A maximizer $\hat{x}$ is a solution to this problem: f achieves an maximum in F at this point.

This section develops a very interesting characterization of the relationship between the objective function f and the constraint function g at a constrained maximum: the gradients of the functions are collinear. That is, for some number λ ,

$$\nabla f[\hat{x}] = \lambda \nabla g[\hat{x}] \quad (1.42)$$

Satisfaction of this relationship is a first order necessary condition for a constrained extremum.

To see why, consider a proof for the special case of a binary objective function f and a single equality constraint in the two choice variables.⁶ Let both functions be continuously differentiable. The optimization problem is to pick x and y so as to maximize $f[x, y]$, subject to the constraint that $g[x, y] = 0$.

Begin by focusing on the binary constraint function, g . Given a differentiable real function ψ , consider the unary function $\gamma = x \mapsto g[x, \psi[x]]$. Apply the chain rule to conclude that

$$\begin{aligned}\gamma' &= x \mapsto g_x[x, \psi[x]] + g_y[x, \psi[x]] \psi'[x] \\ &= \nabla g[x, \psi[x]] \cdot \langle 1, \psi'[x] \rangle\end{aligned}\tag{1.43}$$

⁶The proof strategy draws on Edwards (1973, p.93).

Next, assume that g satisfies the requirements of implicit function theorem, so that satisfaction of the constrained $g[x, y] = 0$ implies that y is implicitly a function of x . This implicit function (ψ) identically satisfies

$$g[x, \psi[x]] \equiv 0 \tag{1.44}$$

That is, for each x , $\psi[x]$ is the value of y that ensure the constraint is satisfied. Therefore changing x has no effect on the value of $g[x, \psi[x]]$. Just as before, making use of the chain rule yields

$$\nabla g[x, \psi[x]] \cdot \langle 1, \psi'[x] \rangle = 0 \tag{1.45}$$

The vector $\langle 1, \psi'[x] \rangle$ is the direction of movement that will not lead to a constraint violation.

Example 1.6 Let $p_x x + p_y y = m$ describe a budget constraint for a two-good consumer choice problem. Define $g = \langle x, y \rangle \mapsto p_x x + p_y y - m$, which has the particularly simple gradient $\nabla g = \langle x, y \rangle \mapsto \langle p_x, p_y \rangle$, as illustrated in Figure 1.2. Solve $g[x, y] = 0$ for the implicit function

$$\psi = x \mapsto (m - p_x x)/p_y$$

and note that $\psi'[x] = -p_x/p_y$, as illustrated in Figure 1.2. Define $\gamma = x \mapsto g[x, \psi[x]]$, so that

$$\gamma[x] = p_x x + p_y((m - p_x x)/p_y) - m$$

This is identically equal to 0, so the function γ has a zero derivative everywhere.

$$\begin{aligned}\gamma'[x] &= g_1 \cdot 1 + g_2 \cdot \psi'[x] \\ &= \langle p_x, p_y \rangle \cdot \langle 1, -p_x/p_y \rangle = 0\end{aligned}$$

Using total differential to find $\psi' \dots$

Impose $p_x x + p_y y - m = 0$ so that (given prices and wealth)

$$\begin{aligned}p_x dx + p_y dy &= 0 \\ dy/dx &= -p_x/p_y\end{aligned}$$

This is ψ' , the slope of the budget constraint. It tells us how much y must fall given a unit increase in x , so that the constraint will continue to be satisfied.

Example 1.7 Let $p_x x + p_y y = m$ describe a budget constraint for a two-good consumer choice problem. Define $g = \langle x, y \rangle \mapsto p_x x + p_y y - m$, which has the particularly simple gradient $\nabla g = \langle x, y \rangle \mapsto \langle p_x, p_y \rangle$, as illustrated in Figure 1.2. Solve $g[x, y] = 0$ for the implicit function

$$\psi = x \mapsto (m - p_x x)/p_y$$

and note that $\psi'[x] = -p_x/p_y$, as illustrated in Figure 1.2. Define $\gamma = x \mapsto g[x, \psi[x]]$, so that

$$\gamma[x] = p_x x + p_y((m - p_x x)/p_y) - m$$

This is identically equal to 0, so the function γ has a zero derivative everywhere.

$$\begin{aligned}\gamma'[x] &= g_1 \cdot 1 + g_2 \cdot \psi'[x] \\ &= \langle p_x, p_y \rangle \cdot \langle 1, -p_x/p_y \rangle = 0\end{aligned}$$

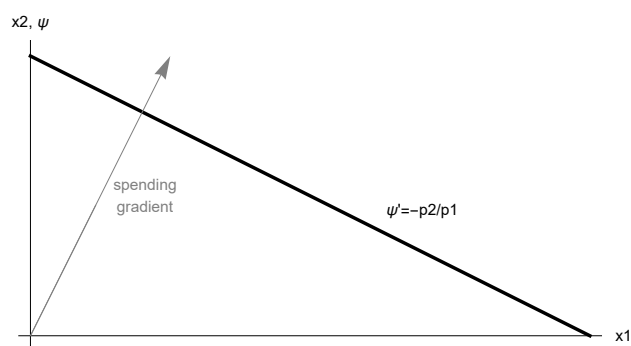


Figure 1.2: Budget Line and Spending Gradient