

Base R

Cheat Sheet

Getting Help

Accessing the help files

?mean
Get help of a particular function.
help.search('weighted mean')
Search the help files for a word or phrase.
help(package = 'dplyr')
Find help for a package.

More about an object

str(iris)
Get a summary of an object's structure.
class(iris)
Find the class an object belongs to.

Using Packages

install.packages('dplyr')
Download and install a package from CRAN.
library(dplyr)
Load the package into the session, making all its functions available to use.
dplyr::select
Use a particular function from a package.
data(iris)
Load a built-in dataset into the environment.

Working Directory

getwd()
Find the current working directory (where inputs are found and outputs are sent).
setwd('C://file/path')
Change the current working directory.
Use projects in RStudio to set the working directory to the folder you are working in.

Vectors

Creating Vectors

c(2, 4, 6)	2 4 6	Join elements into a vector
2:6	2 3 4 5 6	An integer sequence
seq(2, 3, by=0.5)	2.0 2.5 3.0	A complex sequence
rep(1:2, times=3)	1 2 1 2 1 2	Repeat a vector
rep(1:2, each=3)	1 1 1 2 2 2	Repeat elements of a vector

Vector Functions

sort(x) Return x sorted.	rev(x) Return x reversed.
table(x) See counts of values.	unique(x) See unique values.

Selecting Vector Elements

By Position

x[4]	The fourth element.
x[-4]	All but the fourth.
x[2:4]	Elements two to four.
x[-(2:4)]	All elements except two to four.
x[c(1, 5)]	Elements one and five.

By Value

x[x == 10]	Elements which are equal to 10.
x[x < 0]	All elements less than zero.
x[x %in% c(1, 2, 5)]	Elements in the set 1, 2, 5.

Named Vectors

x['apple']	Element with name 'apple'.
-------------------	----------------------------

Programming

For Loop

```
for (variable in sequence){  
  Do something  
}
```

Example

```
for (i in 1:4){  
  j <- i + 10  
  print(j)  
}
```

While Loop

```
while (condition){  
  Do something  
}
```

Example

```
while (i < 5){  
  print(i)  
  i <- i + 1  
}
```

If Statements

```
if (condition){  
  Do something  
} else {  
  Do something different  
}
```

Example

```
if (i > 3){  
  print('Yes')  
} else {  
  print('No')  
}
```

Functions

```
function_name <- function(var){  
  Do something  
  return(new_variable)  
}
```

Example

```
square <- function(x){  
  squared <- x*x  
  return(squared)  
}
```

Reading and Writing Data

Also see the **readr** package.

Input	Ouput	Description
df <- read.table('file.txt')	write.table(df, 'file.txt')	Read and write a delimited text file.
df <- read.csv('file.csv')	write.csv(df, 'file.csv')	Read and write a comma separated value file. This is a special case of read.table/write.table.
load('file.RData')	save(df, file = 'file.Rdata')	Read and write an R data file, a file type special for R.

Conditions	a == b	Are equal	a > b	Greater than	a >= b	Greater than or equal to	is.na(a)	Is missing
	a != b	Not equal	a < b	Less than	a <= b	Less than or equal to	is.null(a)	Is null

Types

Converting between common data types in R. Can always go from a higher value in the table to a lower value.

<code>as.logical</code>	TRUE, FALSE, TRUE	Boolean values (TRUE or FALSE).
<code>as.numeric</code>	1, 0, 1	Integers or floating point numbers.
<code>as.character</code>	'1', '0', '1'	Character strings. Generally preferred to factors.
<code>as.factor</code>	'1', '0', '1', levels: '1', '0'	Character strings with preset levels. Needed for some statistical models.

Maths Functions

<code>log(x)</code>	Natural log.	<code>sum(x)</code>	Sum.
<code>exp(x)</code>	Exponential.	<code>mean(x)</code>	Mean.
<code>max(x)</code>	Largest element.	<code>median(x)</code>	Median.
<code>min(x)</code>	Smallest element.	<code>quantile(x)</code>	Percentage quantiles.
<code>round(x, n)</code>	Round to n decimal places.	<code>rank(x)</code>	Rank of elements.
<code>signif(x, n)</code>	Round to n significant figures.	<code>var(x)</code>	The variance.
<code>cor(x, y)</code>	Correlation.	<code>sd(x)</code>	The standard deviation.

Variable Assignment

```
> a <- 'apple'
> a
[1] 'apple'
```

The Environment

<code>ls()</code>	List all variables in the environment.
<code>rm(x)</code>	Remove x from the environment.
<code>rm(list = ls())</code>	Remove all variables from the environment.

You can use the environment panel in RStudio to browse variables in your environment.

Matrices

```
m <- matrix(x, nrow = 3, ncol = 3)
Create a matrix from x.
```

 <code>m[2,]</code> - Select a row	<code>t(m)</code> Transpose
 <code>m[, 1]</code> - Select a column	<code>m %*% n</code> Matrix Multiplication
 <code>m[2, 3]</code> - Select an element	<code>solve(m, n)</code> Find x in: $m * x = n$

Lists

```
l <- list(x = 1:5, y = c('a', 'b'))
A list is a collection of elements which can be of different types.
```

<code>l[[2]]</code> Second element of l.	<code>l[1]</code> New list with only the first element.	<code>l\$x</code> Element named x.	<code>l['y']</code> New list with only element named y.
---	--	---------------------------------------	--

Also see the **dplyr** package.

Data Frames

```
df <- data.frame(x = 1:3, y = c('a', 'b', 'c'))
A special case of a list where all elements are the same length.
```

x	y
1	a
2	b
3	c

Matrix subsetting

<code>df[, 2]</code>	
<code>df[2,]</code>	
<code>df[2, 2]</code>	

List subsetting

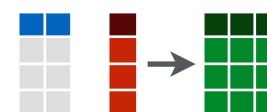
<code>df\$x</code>		<code>df[[2]]</code>	
<i>Understanding a data frame</i>			
<code>View(df)</code>	See the full data frame.		
<code>head(df)</code>	See the first 6 rows.		

`nrow(df)`
Number of rows.

`ncol(df)`
Number of columns.

`dim(df)`
Number of columns and rows.

`cbind` - Bind columns.



`rbind` - Bind rows.



Strings

Also see the **stringr** package.

<code>paste(x, y, sep = ' ')</code>	Join multiple vectors together.
<code>paste(x, collapse = ' ')</code>	Join elements of a vector together.
<code>grep(pattern, x)</code>	Find regular expression matches in x.
<code>gsub(pattern, replace, x)</code>	Replace matches in x with a string.
<code>toupper(x)</code>	Convert to uppercase.
<code>tolower(x)</code>	Convert to lowercase.
<code>nchar(x)</code>	Number of characters in a string.

Factors

<code>factor(x)</code> Turn a vector into a factor. Can set the levels of the factor and the order.	<code>cut(x, breaks = 4)</code> Turn a numeric vector into a factor by 'cutting' into sections.
--	--

Statistics

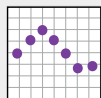
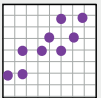
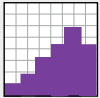
<code>lm(y ~ x, data=df)</code> Linear model.	<code>t.test(x, y)</code> Perform a t-test for difference between means.	<code>prop.test</code> Test for a difference between proportions.
<code>glm(y ~ x, data=df)</code> Generalised linear model.	<code>pairwise.t.test</code> Perform a t-test for paired data.	<code>aov</code> Analysis of variance.
<code>summary</code> Get more detailed information out a model.		

Distributions

	Random Variates	Density Function	Cumulative Distribution	Quantile
Normal	<code>rnorm</code>	<code>dnorm</code>	<code>pnorm</code>	<code>qnorm</code>
Poisson	<code>rpois</code>	<code>dpois</code>	<code>ppois</code>	<code>qpois</code>
Binomial	<code>rbinom</code>	<code>dbinom</code>	<code>pbinom</code>	<code>qbinom</code>
Uniform	<code>runif</code>	<code>dunif</code>	<code>punif</code>	<code>qunif</code>

Plotting

Also see the **ggplot2** package.

 <code>plot(x)</code> Values of x in order.	 <code>plot(x, y)</code> Values of x against y.	 <code>hist(x)</code> Histogram of x.
---	---	---

Dates

See the **lubridate** package.

Advanced R

Cheat Sheet

Created by: Arianne Colton and Sean Chen

Environment Basics

Environment – **Data structure** (with two components below) that powers lexical scoping

```
Create environment: env1<-new.env()
```

1. **Named list** (“Bag of names”) – each name points to an object stored elsewhere in memory.

If an object has no names pointing to it, it gets automatically deleted by the garbage collector.

- Access with: `ls('env1')`
- 2. **Parent environment** – used to implement lexical scoping. If a name is not found in an environment, then R will look in its parent (and so on).
- Access with: `parent.env('env1')`

Four special environments

1. **Empty environment** – ultimate ancestor of all environments
 - Parent: none
 - Access with: `emptyenv()`
2. **Base environment** - environment of the base package
 - Parent: empty environment
 - Access with: `baseenv()`
3. **Global environment** – the interactive workspace that you normally work in
 - Parent: environment of last attached package
 - Access with: `globalenv()`
4. **Current environment** – environment that R is currently working in (may be any of the above and others)
 - Parent: empty environment
 - Access with: `environment()`

Environments

Search Path

Search path – mechanism to look up objects, particularly functions.

- Access with : `search()` – lists all parents of the global environment (see Figure 1)
- Access any environment on the search path:
`as.environment('package:base')`

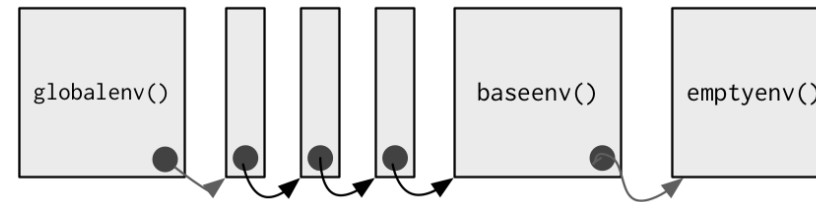


Figure 1 – The Search Path

- Mechanism : always start the search from global environment, then inside the latest attached package environment.
- New package loading with `library()/require()` : new package is attached right after global environment. (See Figure 2)
- Name conflict in two different package : functions with the same name, latest package function will get called.

```
search() :
```

```
'.GlobalEnv' ... 'Autoloads' 'package:base'
```

```
library(reshape2); search()
```

```
'.GlobalEnv' 'package:reshape2' ... 'Autoloads' 'package:base'
```

NOTE: Autoloads : special environment used for saving memory by only loading package objects (like big datasets) when needed

Figure 2 – Package Attachment

Binding Names to Values

Assignment – act of binding (or rebinding) a name to a value in an environment.

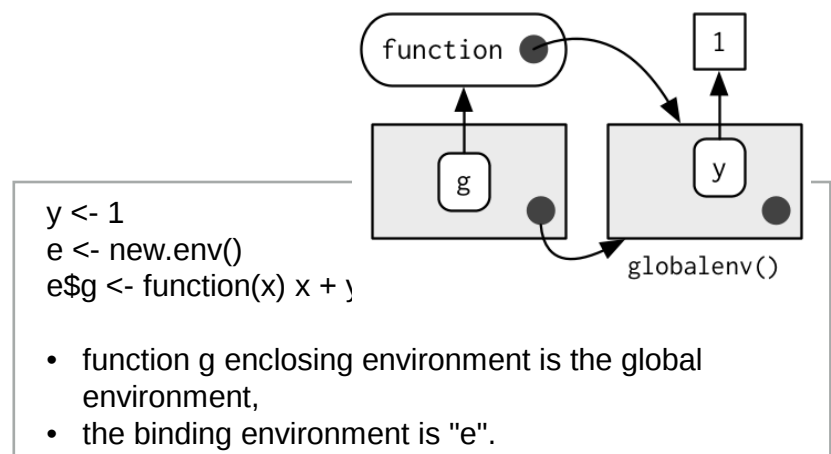
1. `<-` (Regular assignment arrow) – always creates a variable in the current environment
2. `<<-` (Deep assignment arrow) - modifies an existing variable found by walking up the parent environments

Warning: If `<<-` doesn't find an existing variable, it will create one in the global environment.

Function Environments

1. **Enclosing environment** - an environment where the function is created. It determines how function finds value.
 - Enclosing environment never changes, even if the function is moved to a different environment.
 - Access with: `environment('func1')`
2. **Binding environment** - all environments that the function has a binding to. It determines how we find the function.
 - Access with: `pryr::where('func1')`

Example (for enclosing and binding environment):



3. **Execution environment** - new created environments to host a function call execution.
 - Two parents :
 - I. Enclosing environment of the function
 - II. Calling environment of the function
 - Execution environment is thrown away once the function has completed.
4. **Calling environment** - environments where the function was called.
 - Access with: `parent.frame('func1')`
 - Dynamic scoping :
 - About : look up variables in the calling environment rather than in the enclosing environment
 - Usage : most useful for developing functions that aid interactive data analysis

Data Structures

	Homogeneous	Heterogeneous
1d	Atomic vector	List
2d	Matrix	Data frame
nd	Array	

Note: R has no 0-dimensional or scalar types. Individual numbers or strings, are actually vectors of length one, NOT scalars.

Human readable description of any R data structure :

```
str(variable)
```

Every **Object** has a mode and a class

1. **Mode:** represents how an object is stored in memory

- 'type' of the object from R's point of view
- Access with: **typeof()**

2. **Class:** represents the object's abstract type

- 'type' of the object from R's object-oriented programming point of view
- Access with: **class()**

	typeof()	class()
strings or vector of strings	character	character
numbers or vector of numbers	numeric	numeric
list	list	list
data.frame	list	data.frame

Factors

1. Factors are built on top of integer vectors using two attributes :

```
class(x) -> 'factor'
```

```
levels(x) # defines the set of allowed values
```

2. Useful when you know the possible values a variable may take, even if you don't see all values in a given dataset.

Warning on Factor Usage:

1. Factors look and often behave like character vectors, they are actually integers. Be careful when treating them like strings.
2. Most data loading functions automatically convert character vectors to factors. (Use argument `stringAsFactors = FALSE` to suppress this behavior)

Object Oriented (OO) Field Guide

Object Oriented Systems

R has three object oriented systems :

1. **S3** is a very casual system. It has no formal definition of classes. It implements generic function OO.

- **Generic-function OO** - a special type of function called a generic function decides which method to call.

Example:	drawRect(canvas, 'blue')
Language:	R

- **Message-passing OO** - messages (methods) are sent to objects and the object determines which function to call.

Example:	canvas.drawRect('blue')
Language:	Java, C++, and C#

2. **S4** works similarly to S3, but is more formal. Two major differences to S3 :

- **Formal class definitions** - describe the representation and inheritance for each class, and has special helper functions for defining generics and methods.
- **Multiple dispatch** - generic functions can pick methods based on the class of any number of arguments, not just one.

3. **Reference classes** are very different from S3 and S4:

- **Implements message-passing OO** - methods belong to classes, not functions.
- **Notation** - \$ is used to separate objects and methods, so method calls look like **canvas\$drawRect('blue')**.

S3

1. **About S3 :**

- R's first and simplest OO system
- Only OO system used in the base and stats package
- Methods belong to functions, not to objects or classes.

2. **Notation :**

- **generic.class()**

mean.Date()	Date method for the generic - mean()
-------------	--------------------------------------

3. **Useful 'Generic' Operations**

- Get all methods that belong to the 'mean' generic:
 - **Methods('mean')**
- List all generics that have a method for the 'Date' class :
 - **methods(class = 'Date')**

4. **S3 objects** are usually built on top of lists, or atomic vectors with attributes.

- Factor and data frame are S3 class
- Useful operations:

Check if object is an S3 object	is.object(x) & !isS4(x) or pryr::otype()
Check if object inherits from a specific class	inherits(x, 'classname')
Determine class of any object	class(x)

Base Type (C Structure)

R base types - the internal C-level types that underlie the above OO systems.

- **Includes** : atomic vectors, list, functions, environments, etc.
- **Useful operation** : Determine if an object is a base type (Not S3, S4 or RC) **is.object(x)** returns FALSE

- **Internal representation** : C structure (or struct) that includes :

- Contents of the object
- Memory Management Information
- Type
 - Access with: **typeof()**

Functions

Function Basics

Functions – objects in their own right
All R functions have three parts:

body()	code inside the function
formals()	list of arguments which controls how you can call the function
environment()	“map” of the location of the function’s variables (see “Enclosing Environment”)

Every operation is a function call

- +, for, if, [, \$, { ...
- $x + y$ is the same as ``+`(x, y)`

Note: the backtick (```), lets you refer to functions or variables that have otherwise reserved or illegal names.

Lexical Scoping

What is Lexical Scoping?

- Looks up value of a symbol. (see “Enclosing Environment”)
- **findGlobals()** - lists all the external dependencies of a function

```
f <- function() x + 1
codetools::findGlobals(f)
> '+' 'x'
```

```
environment(f) <- emptyenv()
f()
# error in f(): could not find function “+”
```

- R relies on lexical scoping to find everything, even the + operator.

Function Arguments

Arguments – passed by reference and copied on modify

1. Arguments are matched first by exact name (perfect matching), then by prefix matching, and finally by position.

2. Check if an argument was supplied : **missing()**

```
i <- function(a, b) {
  missing(a) -> # return true or false
}
```

3. Lazy evaluation – since x is not used **stop("This is an error!")** never get evaluated.

```
f <- function(x) {
  10
}
f(stop("This is an error!")) -> 10
```

4. Force evaluation

```
f <- function(x) {
  force(x)
  10
}
```

5. Default arguments evaluation

```
f <- function(x = ls()) {
  a <- 1
  x
}
```

<code>f()</code> -> 'a' 'x'	ls() evaluated inside f
<code>f(ls())</code>	ls() evaluated in global environment

Return Values

- **Last expression evaluated or explicit return().**
Only use explicit return() when returning early.
- **Return ONLY single object.**
Workaround is to return a list containing any number of objects.
- **Invisible return object value** - not printed out by default when you call the function.

```
f1 <- function() invisible(1)
```

Primitive Functions

What are Primitive Functions?

1. Call C code directly with **.Primitive()** and contain no R code

```
print(sum) :
> function (..., na.rm = FALSE) .Primitive('sum')
```

2. **formals()**, **body()**, and **environment()** are all NULL
3. Only found in base package
4. More efficient since they operate at a low level

Infix Functions

What are Infix Functions?

1. Function name comes in between its arguments, like + or –
2. All user-created infix functions must start and end with %.

```
`%+%` <- function(a, b) paste0(a, b)
'new' %+% 'string'
```

3. Useful way of providing a default value in case the output of another function is NULL:

```
`%||%` <- function(a, b) if (!is.null(a)) a else b
function_that_might_return_null() %||% default value
```

Replacement Functions

What are Replacement Functions?

1. Act like they modify their arguments in place, and have the special name `xxx <-`
2. Actually create a modified copy. Can use **pryr::address()** to find the memory address of the underlying object

```
`second<-` <- function(x, value) {
  x[2] <- value
  x
}
x <- 1:10
second(x) <- 5L
```

Subsetting

Subsetting returns a copy of the original data, NOT copy-on modified

Simplifying vs. Preserving Subsetting

1. Simplifying subsetting

- Returns the **simplest** possible data structure that can represent the output

2. Preserving subsetting

- Keeps the structure of the output the **same** as the input.
- When you use `drop = FALSE`, it's preserving

	Simplifying*	Preserving
Vector	<code>x[[1]]</code>	<code>x[1]</code>
List	<code>x[[1]]</code>	<code>x[1]</code>
Factor	<code>x[1:4, drop = T]</code>	<code>x[1:4]</code>
Array	<code>x[1,]</code> or <code>x[, 1]</code>	<code>x[1, , drop = F]</code> or <code>x[, 1, drop = F]</code>
Data frame	<code>x[, 1]</code> or <code>x[[1]]</code>	<code>x[, 1, drop = F]</code> or <code>x[1]</code>

Simplifying behavior varies slightly between different data types:

1. Atomic Vector

- `x[[1]]` is the same as `x[1]`

2. List

- `[ ]` always returns a list
- Use `[[ ]]` to get list contents, this returns a single value piece out of a list

3. Factor

- Drops any unused levels but it remains a factor class

4. Matrix or Array

- If any of the dimensions has length 1, that dimension is dropped

5. Data Frame

- If output is a single column, it returns a vector instead of a data frame

Data Frame Subsetting

Data Frame – possesses the **characteristics of both lists and matrices**. If you subset with a single vector, they behave like lists; if you subset with two vectors, they behave like matrices

1. Subset with a single vector : Behave like lists

```
df1[c('col1', 'col2')]
```

2. Subset with two vectors : Behave like matrices

```
df1[, c('col1', 'col2')]
```

The results are the same in the above examples, however, results are different if subsetting with only one column. (see below)

1. Behave like matrices

```
str(df1[, 'col1']) -> int [1:3]
```

- Result: the result is a vector

2. Behave like lists

```
str(df1['col1']) -> 'data.frame'
```

- Result: the result remains a data frame of 1 column

\$ Subsetting Operator

1. About Subsetting Operator

- Useful shorthand for `[[` combined with character subsetting

```
x$y is equivalent to x[['y', exact = FALSE]]
```

2. Difference vs. `[[`

- `$` does partial matching, `[[` does not

```
x <- list(abc = 1)
x$a -> 1      # since "exact = FALSE"
x[['a']] ->    # would be an error
```

3. Common mistake with `$`

- Using it when you have the name of a column stored in a variable

```
var <- 'cyl'
x$var
# doesn't work, translated to x[['var']]
# Instead use x[[var]]
```

Examples

1. Lookup tables (character subsetting)

```
x <- c('m', 'f', 'u', 'f', 'f', 'm', 'm')
lookup <- c(m = 'Male', f = 'Female', u = NA)
lookup[x]
> m f u f f m m
> 'Male' 'Female' NA 'Female' 'Female' 'Male' 'Male'
unname(lookup[x])
> 'Male' 'Female' NA 'Female' 'Female' 'Male' 'Male'
```

2. Matching and merging by hand (integer subsetting)

Lookup table which has multiple columns of information:

```
grades <- c(1, 2, 2, 3, 1)
info <- data.frame(
  grade = 3:1,
  desc = c('Excellent', 'Good', 'Poor'),
  fail = c(F, F, T)
)
```

First Method

```
id <- match(grades, info$grade)
info[id, ]
```

Second Method

```
rownames(info) <- info$grade
info[as.character(grades), ]
```

3. Expanding aggregated counts (integer subsetting)

- Problem:** a data frame where identical rows have been collapsed into one and a count column has been added
- Solution:** `rep()` and integer subsetting make it easy to uncollapse the data by subsetting with a repeated row index: `rep(x, y)` `rep` replicates the values in `x`, `y` times.

```
df1$countCol is c(3, 5, 1)
rep(1:nrow(df1), df1$countCol)
> 1 1 1 2 2 2 2 2 3
```

4. Removing columns from data frames (character subsetting)

There are two ways to remove columns from a data frame:

Set individual columns to NULL	<code>df1\$col3 <- NULL</code>
Subset to return only columns you want	<code>df1[c('col1', 'col2')]</code>

5. Selecting rows based on a condition (logical subsetting)

- This is the most commonly used technique for extracting rows out of a data frame.

```
df1[df1$col1 == 5 & df1$col2 == 4, ]
```

Subsetting continued

Boolean Algebra vs. Sets (Logical and Integer Subsetting)

1. **Using integer subsetting** is more effective when:

- You want to find the first (or last) TRUE.
- You have very few TRUEs and very many FALSEs; a set representation may be faster and require less storage.

2. **which()** - conversion from boolean representation to integer representation

```
which(c(T, F, T F)) -> 1 3
```

- Integer representation length : is always $\leq$ boolean representation length
- Common mistakes :
 - I. Use **x[which(y)]** instead of **x[y]**
 - II. **x[-which(y)]** is not equivalent to **x[!y]**

Recommendation:

Avoid switching from logical to integer subsetting unless you want, for example, the first or last TRUE value

Subsetting with Assignment

1. All subsetting operators can be combined with assignment to modify selected values of the input vector.

```
df1$col1[df1$col1 < 8] <- 0
```

2. Subsetting with nothing in conjunction with assignment :

- Why : Preserve original object class and structure

```
df1[] <- lapply(df1, as.integer)
```

Debugging, Condition Handling and Defensive Programming

Debugging Methods

1. **traceback()** or **RStudio's error inspector**

- Lists the sequence of calls that lead to the error

2. **browser()** or **RStudio's breakpoints tool**

- Opens an interactive debug session at an arbitrary location in the code

3. **options(error = browser)** or **RStudio's "Rerun with Debug" tool**

- Opens an interactive debug session where the error occurred

- Error Options:

options(error = recover)

- Difference vs. 'browser': can enter environment of any of the calls in the stack

options(error = dump_and_quit)

- Equivalent to 'recover' for non-interactive mode
- Creates **last.dump.rda** in the current working directory

In batch R process :

```
dump_and_quit <- function() {  
  # Save debugging info to file  
  last.dump.rda  
  dump.frames(to.file = TRUE)  
  
  # Quit R with error status  
  q(status = 1)  
}  
options(error = dump_and_quit)
```

In a later interactive session :

```
load("last.dump.rda")  
debugger()
```

Condition Handling of Expected Errors

1. **Communicating potential problems to users:**

I. **stop()**

- Action : raise fatal error and force all execution to terminate
- Example usage : when there is no way for a function to continue

II. **warning()**

- Action : generate warnings to display potential problems
- Example usage : when some of elements of a vectorized input are invalid

III. **message()**

- Action : generate messages to give informative output
- Example usage : when you would like to print the steps of a program execution

2. **Handling conditions programmatically:**

I. **try()**

- Action : gives you the ability to continue execution even when an error occurs

II. **tryCatch()**

- Action : lets you specify handler functions that control what happens when a condition is signaled

```
result = tryCatch(code,  
  error = function(c) "error",  
  warning = function(c) "warning",  
  message = function(c) "message"  
)
```

Use `conditionMessage(c)` or `c$message` to extract the message associated with the original error.

Defensive Programming

Basic principle : "fail fast", to raise an error as soon as something goes wrong

1. **stopifnot()** or use 'assertthat' package - check inputs are correct

2. **Avoid subset(), transform() and with()** - these are non-standard evaluation, when they fail, often fail with uninformative error messages.

3. **Avoid [and apply()** - functions that can return different types of output.

- Recommendation : Whenever subsetting a data frame in a function, you should always use **drop = FALSE**

String manipulation with stringr : : CHEAT SHEET

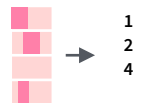


The **stringr** package provides a set of internally consistent tools for working with character strings, i.e. sequences of characters surrounded by quotation marks.

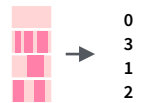
Detect Matches



str_detect(string, **pattern**) Detect the presence of a pattern match in a string.
`str_detect(fruit, "a")`



str_which(string, **pattern**) Find the indexes of strings that contain a pattern match.
`str_which(fruit, "a")`



str_count(string, **pattern**) Count the number of matches in a string.
`str_count(fruit, "a")`

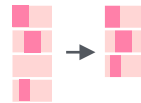


str_locate(string, **pattern**) Locate the positions of pattern matches in a string. Also **str_locate_all**.
`str_locate(fruit, "a")`

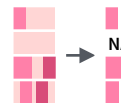
Subset Strings



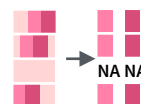
str_sub(string, start = 1L, end = -1L) Extract substrings from a character vector.
`str_sub(fruit, 1, 3); str_sub(fruit, -2)`



str_subset(string, **pattern**) Return only the strings that contain a pattern match.
`str_subset(fruit, "b")`

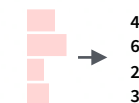


str_extract(string, **pattern**) Return the first pattern match found in each string, as a vector. Also **str_extract_all** to return every pattern match.
`str_extract(fruit, "[aeiou]")`

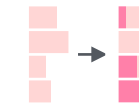


str_match(string, **pattern**) Return the first pattern match found in each string, as a matrix with a column for each () group in pattern. Also **str_match_all**.
`str_match(sentences, "(a|the) ([^ ]+)")`

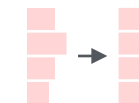
Manage Lengths



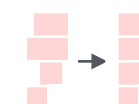
str_length(string) The width of strings (i.e. number of code points, which generally equals the number of characters). `str_length(fruit)`



str_pad(string, width, side = c("left", "right", "both"), pad = " ") Pad strings to constant width. `str_pad(fruit, 17)`



str_trunc(string, width, side = c("right", "left", "center"), ellipsis = "...") Truncate the width of strings, replacing content with ellipsis. `str_trunc(fruit, 3)`

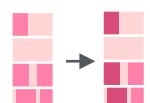


str_trim(string, side = c("both", "left", "right")) Trim whitespace from the start and/or end of a string. `str_trim(fruit)`

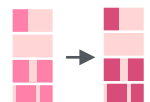
Mutate Strings



str_sub() <- value. Replace substrings by identifying the substrings with `str_sub()` and assigning into the results.
`str_sub(fruit, 1, 3) <- "str"`



str_replace(string, **pattern**, replacement) Replace the first matched pattern in each string. `str_replace(fruit, "a", "-")`



str_replace_all(string, **pattern**, replacement) Replace all matched patterns in each string. `str_replace_all(fruit, "a", "-")`

A STRING
↓
a string

str_to_lower(string, locale = "en")¹ Convert strings to lower case.
`str_to_lower(sentences)`

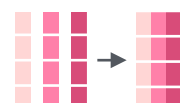
a string
↓
A STRING

str_to_upper(string, locale = "en")¹ Convert strings to upper case.
`str_to_upper(sentences)`

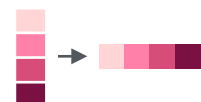
a string
↓
A String

str_to_title(string, locale = "en")¹ Convert strings to title case. `str_to_title(sentences)`

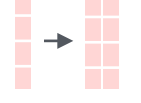
Join and Split



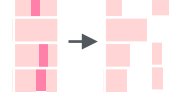
str_c(..., sep = "", collapse = NULL) Join multiple strings into a single string.
`str_c(letters, LETTERS)`



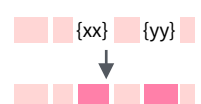
str_c(..., sep = "", collapse = NULL) Collapse a vector of strings into a single string.
`str_c(letters, collapse = "")`



str_dup(string, times) Repeat strings times times. `str_dup(fruit, times = 2)`



str_split_fixed(string, **pattern**, n) Split a vector of strings into a matrix of substrings (splitting at occurrences of a pattern match). Also **str_split** to return a list of substrings.
`str_split_fixed(fruit, " ", n=2)`



str_glue(..., .sep = "", .envir = parent.frame()) Create a string from strings and {expressions} to evaluate. `str_glue("Pi is {pi}")`



str_glue_data(.x, ..., .sep = "", .envir = parent.frame(), .na = "NA") Use a data frame, list, or environment to create a string from strings and {expressions} to evaluate.
`str_glue_data(mtcars, "{rownames(mtcars)} has {hp} hp")`

Order Strings



str_order(x, decreasing = FALSE, na_last = TRUE, locale = "en", numeric = FALSE, ...) ¹ Return the vector of indexes that sorts a character vector. `x[str_order(x)]`



str_sort(x, decreasing = FALSE, na_last = TRUE, locale = "en", numeric = FALSE, ...) ¹ Sort a character vector. `str_sort(x)`

Helpers

apple
banana
pear

str_conv(string, encoding) Override the encoding of a string. `str_conv(fruit, "ISO-8859-1")`

apple
banana
pear

str_view(string, **pattern**, match = NA) View HTML rendering of first regex match in each string. `str_view(fruit, "[aeiou]")`

str_view_all(string, **pattern**, match = NA) View HTML rendering of all regex matches. `str_view_all(fruit, "[aeiou]")`

str_wrap(string, width = 80, indent = 0, exdent = 0) Wrap strings into nicely formatted paragraphs. `str_wrap(sentences, 20)`

Need to Know

Pattern arguments in stringr are interpreted as regular expressions *after any special characters have been parsed*.

In R, you write regular expressions as *strings*, sequences of characters surrounded by quotes ("" or '') or single quotes('').

Some characters cannot be represented directly in an R string. These must be represented as **special characters**, sequences of characters that have a specific meaning., e.g.

Special Character	Represents
\\	\
\"	"
\n	new line

Run `?\"` to see a complete list

Because of this, whenever a \ appears in a regular expression, you must write it as \\ in the string that represents the regular expression.

Use `writeLines()` to see how R views your string after all special characters have been parsed.

```
writeLines("\\.")
# \.
```

```
writeLines("\\ is a backslash")
# \ is a backslash
```

INTERPRETATION

Patterns in stringr are interpreted as regexs To change this default, wrap the pattern in one of:

regex(pattern, ignore_case = FALSE, multiline = FALSE, comments = FALSE, dotall = FALSE, ...) Modifies a regex to ignore cases, match end of lines as well of end of strings, allow R comments within regex's, and/or to have . match everything including \n. `str_detect("I", regex("i", TRUE))`

fixed() Matches raw bytes but will miss some characters that can be represented in multiple ways (fast). `str_detect("\u0130", fixed("i"))`

coll() Matches raw bytes and will use locale specific collation rules to recognize characters that can be represented in multiple ways (slow). `str_detect("\u0130", coll("i", TRUE, locale = "tr"))`

boundary() Matches boundaries between characters, line_breaks, sentences, or words. `str_split(sentences, boundary("word"))`

Regular Expressions - Regular expressions, or *regexps*, are a concise language for describing patterns in strings.

MATCH CHARACTERS

string (type this)	regex (to mean this)	matches (which matches this)	example
	a (etc.)	a (etc.)	
\\.	\\. (etc.)	.	see("a")
\\!	\\!	!	see("\\.")
\\?	\\?	?	see("\\!")
\\\\	\\\\	\	see("\\\\")
\\(	\\(	(	see("\\\\")
\\)	\\)	)	see("\\\\")
\\{	\\{	{	see("\\{")
\\}	\\}	}	see("\\}")
\\n	\\n	new line (return)	see("\\n")
\\t	\\t	tab	see("\\t")
\\s	\\s	any whitespace (S for <i>non-whitespaces</i>)	see("\\s")
\\d	\\d	any digit (D for <i>non-digits</i>)	see("\\d")
\\w	\\w	any word character (W for <i>non-word chars</i>)	see("\\w")
\\b	\\b	word boundaries	see("\\b")
	[digit:] ¹	digits	see("[digit:]")
	[alpha:] ¹	letters	see("[alpha:]")
	[lower:] ¹	lowercase letters	see("[lower:]")
	[upper:] ¹	uppercase letters	see("[upper:]")
	[alnum:] ¹	letters and numbers	see("[alnum:]")
	[punct:] ¹	punctuation	see("[punct:]")
	[graph:] ¹	letters, numbers, and punctuation	see("[graph:]")
	[space:] ¹	space characters (i.e. \s)	see("[space:]")
	[blank:] ¹	space and tab (but not new line)	see("[blank:]")
	.	every character except a new line	see(".")

¹ Many base R functions require classes to be wrapped in a second set of [], e.g. `[digit:]`

[space:]
new line

[blank:]
space
tab

[graph:]

[punct:]

. , : ; ? ! \ | / ` = * + - ^
_ ~ " ' [] { } () < > @ # \$

[alnum:]

[digit:]

0 1 2 3 4 5 6 7 8 9

[alpha:]

[lower:]

a b c d e f

g h i j k l

m n o p q r

s t u v w x

z

[upper:]

A B C D E F

G H I J K L

M N O P Q R

S T U V W X

Z

ALTERNATES

`alt <- function(rx) str_view_all("abcde", rx)`

regex	matches	example
ab d	or	alt("ab d")
[abe]	one of	alt("[abe]")
[^abe]	anything but	alt("[^abe]")
[a-c]	range	alt("[a-c]")

ANCHORS

`anchor <- function(rx) str_view_all("aaa", rx)`

regex	matches	example
^a	start of string	anchor("^a")
a\$	end of string	anchor("a\$")

LOOK AROUNDS

`look <- function(rx) str_view_all("bacad", rx)`

regex	matches	example
a(=?c)	followed by	look("a(=?c)")
a(?!c)	not followed by	look("a(?!c)")
(?<=b)a	preceded by	look("(?<=b)a")
(?<!=b)a	not preceded by	look("(?<!=b)a")

QUANTIFIERS

`quant <- function(rx) str_view_all("a.aa.aaa", rx)`

regex	matches	example
a?	zero or one	quant("a?")
a*	zero or more	quant("a*")
a+	one or more	quant("a+")
a{n}	exactly n	quant("a{2}")
a{n,}	n or more	quant("a{2,}")
a{n,m}	between n and m	quant("a{2,4}")

GROUPS

`ref <- function(rx) str_view_all("abbaab", rx)`

Use parentheses to set precedent (order of evaluation) and create groups

regex	matches	example
(ab d)e	sets precedence	alt("(ab d)e")

Use an escaped number to refer to and duplicate parentheses groups that occur earlier in a pattern. Refer to each group by its order of appearance

string (type this)	regex (to mean this)	matches (which matches this)	example (the result is the same as ref("abba"))
<code>\\1</code>	<code>\\1 (etc.)</code>	first () group, etc.	<code>ref("(a)(b)\\2\\1")</code>



Apply functions with purrr : : CHEAT SHEET



Apply Functions

Map functions apply a function iteratively to each element of a list or vector.

map(*x*, *fun*, ...) → **fun**(*x*, ...) → **fun**(*x*, ...) → **fun**(*x*, ...)

map(*x*, *fun*, ...) Apply a function to each element of a list or vector. *map(x, is.logical)*

map2(*x*, *y*, *fun*, ...) → **fun**(*x*, *y*, ...) → **fun**(*x*, *y*, ...) → **fun**(*x*, *y*, ...)

map2(*x*, *y*, *fun*, ...) Apply a function to pairs of elements from two lists, vectors. *map2(x, y, sum)*

pmap(*l*, *fun*, ...) → **fun**(*l*, ...) → **fun**(*l*, ...) → **fun**(*l*, ...)

pmap(*l*, *fun*, ...) Apply a function to groups of elements from list of lists, vectors. *pmap(list(x, y, z), sum, na.rm = TRUE)*

invoke_map(*fun*, *x*, ...) → **fun**(*x*, ...) → **fun**(*x*, ...) → **fun**(*x*, ...)

invoke_map(*fun*, *x*, ...) Run each function in a list. Also **invoke**. *l <- list(var, sd); invoke_map(l, x = 1:9)*

lmap(*x*, *fun*, ...) Apply function to each list-element of a list or vector.

imap(*x*, *fun*, ...) Apply *fun* to each element of a list or vector and its index.

OUTPUT

map(), **map2()**, **pmap()**, **imap** and **invoke_map** each return a list. Use a suffixed version to return the results as a specific type of flat vector, e.g. **map2_chr**, **pmap_lgl**, etc.

Use **walk**, **walk2**, and **pwalk** to trigger side effects. Each return its input invisibly.

function	returns
map	list
map_chr	character vector
map_dbl	double (numeric) vector
map_dfc	data frame (column bind)
map_dfr	data frame (row bind)
map_int	integer vector
map_lgl	logical vector
walk	triggers side effects, returns the input invisibly

SHORTCUTS - within a purrr function:

"name" becomes **function(x) x[["name"]]**, e.g. *map(l, "a")* extracts *a* from each element of *l*

~ *x* becomes **function(x) x**, e.g. *map(l, ~ 2 + x)* becomes *map(l, function(x) 2 + x)*

~ *x.y* becomes **function(x, y) x.y**, e.g. *map2(l, p, ~ x + y)* becomes *map2(l, p, function(l, p) l + p)*

~ *..1 ..2* etc becomes **function(..1, ..2, etc) ..1 ..2 etc**, e.g. *pmap(list(a, b, c), ~ ..1 + ..2)* becomes *pmap(list(a, b, c), function(a, b, c) c + a - b)*

Work with Lists

FILTER LISTS

pluck(*x*, ..., *default=NULL*) Select an element by name or index, *pluck(x, "b")*, or its attribute with **attr_getter**. *pluck(x, "b", attr_getter("n"))*

keep(*x*, *p*, ...) Select elements that pass a logical test. *keep(x, is.na)*

discard(*x*, *p*, ...) Select elements that do not pass a logical test. *discard(x, is.na)*

compact(*x*, *p* = identity) Drop empty elements. *compact(x)*

head_while(*x*, *p*, ...) Return head elements until one does not pass. Also **tail_while**. *head_while(x, is.character)*

RESHAPE LISTS

flatten(*x*) Remove a level of indexes from a list. Also **flatten_chr**, **flatten_dbl**, **flatten_dfc**, **flatten_dfr**, **flatten_int**, **flatten_lgl**. *flatten(x)*

transpose(*l*, *names* = NULL) Transposes the index order in a multi-level list. *transpose(x)*

SUMMARISE LISTS

every(*x*, *p*, ...) Do all elements pass a test? *every(x, is.character)*

some(*x*, *p*, ...) Do some elements pass a test? *some(x, is.character)*

has_element(*x*, *y*) Does a list contain an element? *has_element(x, "foo")*

detect(*x*, *fun*, ..., *right=FALSE*, *p*) Find first element to pass. *detect(x, is.character)*

detect_index(*x*, *fun*, ..., *right=FALSE*, *p*) Find index of first element to pass. *detect_index(x, is.character)*

vec_depth(*x*) Return depth (number of levels of indexes). *vec_depth(x)*

JOIN (TO) LISTS

append(*x*, *values*, *after* = length(*x*)) Add to end of list. *append(x, list(d = 1))*

prepend(*x*, *values*, *before* = 1) Add to start of list. *prepend(x, list(d = 1))*

splice(...) Combine objects into a list, storing S3 objects as sub-lists. *splice(x, y, "foo")*

TRANSFORM LISTS

modify(*x*, *fun*, ...) Apply function to each element. Also **map**, **map_chr**, **map_dbl**, **map_dfc**, **map_dfr**, **map_int**, **map_lgl**. *modify(x, ~ + 2)*

modify_at(*x*, *at*, *fun*, ...) Apply function to elements by name or index. Also **map_at**. *modify_at(x, "b", ~ + 2)*

modify_if(*x*, *p*, *fun*, ...) Apply function to elements that pass a test. Also **map_if**. *modify_if(x, is.numeric, ~ + 2)*

modify_depth(*x*, *depth*, *fun*, ...) Apply function to each element at a given level of a list. *modify_depth(x, 1, ~ + 2)*

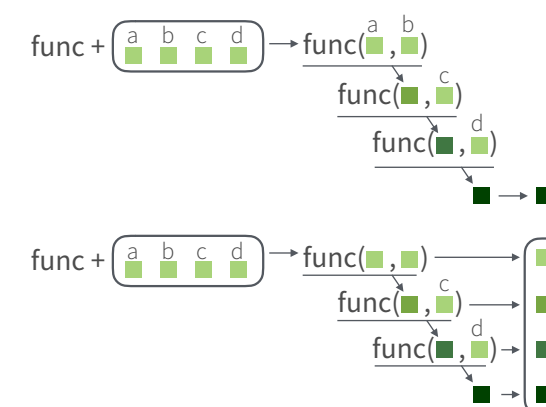
WORK WITH LISTS

array_tree(*array*, *margin* = NULL) Turn array into list. Also **array_branch**. *array_tree(x, margin = 3)*

cross2(*x*, *y*, *filter* = NULL) All combinations of *x* and *y*. Also **cross**, **cross3**, **cross_df**. *cross2(1:3, 4:6)*

set_names(*x*, *nm* = *x*) Set the names of a vector/list directly or with a function. *set_names(x, c("p", "q", "r"))*
set_names(x, tolower)

Reduce Lists



reduce(*x*, *fun*, ..., *init*) Apply function recursively to each element of a list or vector. Also **reduce_right**, **reduce2**, **reduce2_right**. *reduce(x, sum)*

accumulate(*x*, *fun*, ..., *init*) Reduce, but also return intermediate results. Also **accumulate_right**. *accumulate(x, sum)*

Modify function behavior

compose() Compose multiple functions.

lift() Change the type of input a function takes. Also **lift_dl**, **lift_dv**, **lift_ld**, **lift_lv**, **lift_vd**, **lift_vl**.

rerun() Rerun expression *n* times.

negate() Negate a predicate function (a pipe friendly!)

partial() Create a version of a function that has some args preset to values.

safely() Modify func to return list of results and errors.

quietly() Modify function to return list of results, output, messages, warnings.

possibly() Modify function to return default value whenever an error occurs (instead of error).

Nested Data

A **nested data frame** stores individual tables within the cells of a larger, organizing table.

nested data frame

Species	data
setosa	<tibble [50 x 4]>
versicolor	<tibble [50 x 4]>
virginica	<tibble [50 x 4]>

n_iris

"cell" contents

Sepal.L	Sepal.W	Petal.L	Petal.W
5.1	3.5	1.4	0.2
4.9	3.0	1.4	0.2
4.7	3.2	1.3	0.2
4.6	3.1	1.5	0.2
5.0	3.6	1.4	0.2

n_iris\$data[[1]]

Sepal.L	Sepal.W	Petal.L	Petal.W
7.0	3.2	4.7	1.4
6.4	3.2	4.5	1.5
6.9	3.1	4.9	1.5
5.5	2.3	4.0	1.3
6.5	2.8	4.6	1.5

n_iris\$data[[2]]

Sepal.L	Sepal.W	Petal.L	Petal.W
6.3	3.3	6.0	2.5
5.8	2.7	5.1	1.9
7.1	3.0	5.9	2.1
6.3	2.9	5.6	1.8
6.5	3.0	5.8	2.2

n_iris\$data[[3]]

Use a nested data frame to:

- preserve relationships between observations and subsets of data

- manipulate many sub-tables at once with the **purrr** functions **map()**, **map2()**, or **pmap()**.

Use a two step process to create a nested data frame:

1. Group the data frame into groups with **dplyr::group_by()**
2. Use **nest()** to create a nested data frame with one row per group

Species	S.L	S.W	P.L	P.W
setosa	5.1	3.5	1.4	0.2
setosa	4.9	3.0	1.4	0.2
setosa	4.7	3.2	1.3	0.2
setosa	4.6	3.1	1.5	0.2
setosa	5.0	3.6	1.4	0.2
versi	7.0	3.2	4.7	1.4
versi	6.4	3.2	4.5	1.5
versi	6.9	3.1	4.9	1.5
versi	5.5	2.3	4.0	1.3
versi	6.5	2.8	4.6	1.5
virgini	6.3	3.3	6.0	2.5
virgini	5.8	2.7	5.1	1.9
virgini	7.1	3.0	5.9	2.1
virgini	6.3	2.9	5.6	1.8
virgini	6.5	3.0	5.8	2.2

n_iris <- iris %>% **group_by**(Species) %>% **nest**()

tidyr::nest(data, ..., .key = data)

For grouped data, moves groups into cells as data frames.

Unnest a nested data frame with **unnest()**:

n_iris %>% **unnest**()

tidyr::unnest(data, ..., .drop = NA, .id=NULL, .sep=NULL)
Unnests a nested data frame.

Species	data
setos	<tibble [50x4]>
versi	<tibble [50x4]>
virgini	<tibble [50x4]>

Species	S.L	S.W	P.L	P.W
setosa	5.1	3.5	1.4	0.2
setosa	4.9	3.0	1.4	0.2
setosa	4.7	3.2	1.3	0.2
setosa	4.6	3.1	1.5	0.2
setosa	5.0	3.6	1.4	0.2
versi	7.0	3.2	4.7	1.4
versi	6.4	3.2	4.5	1.5
versi	6.9	3.1	4.9	1.5
versi	5.5	2.3	4.0	1.3
versi	6.5	2.8	4.6	1.5
virgini	6.3	3.3	6.0	2.5
virgini	5.8	2.7	5.1	1.9
virgini	7.1	3.0	5.9	2.1
virgini	6.3	2.9	5.6	1.8
virgini	6.5	3.0	5.8	2.2

List Column Workflow

Nested data frames use a **list column**, a list that is stored as a column vector of a data frame. A typical **workflow** for list columns:

1 Make a list column

Species	S.L	S.W	P.L	P.W
setosa	5.1	3.5	1.4	0.2
setosa	4.9	3.0	1.4	0.2
setosa	4.7	3.2	1.3	0.2
setosa	4.6	3.1	1.5	0.2
setosa	5.0	3.6	1.4	0.2
versi	7.0	3.2	4.7	1.4
versi	6.4	3.2	4.5	1.5
versi	6.9	3.1	4.9	1.5
versi	5.5	2.3	4.0	1.3
virgini	6.3	3.3	6.0	2.5
virgini	5.8	2.7	5.1	1.9
virgini	7.1	3.0	5.9	2.1
virgini	6.3	2.9	5.6	1.8

n_iris <- iris %>%
group_by(Species) %>%
nest()

2 Work with list columns

Species	data	model
setosa	<tibble [50x4]>	<S3: lm>
versi	<tibble [50x4]>	<S3: lm>
virgini	<tibble [50x4]>	<S3: lm>

mod_fun <- function(df)
lm(Sepal.Length ~ ., data = df)

m_iris <- n_iris %>%
mutate(model = **map**(data, mod_fun))

3 Simplify the list column

Species	beta
setos	2.35
versi	1.89
virgini	0.69

b_fun <- function(mod)
coefficients(mod)[[1]]

m_iris %>% **transmute**(Species,
beta = **map_dbl**(model, b_fun))

1. MAKE A LIST COLUMN - You can create list columns with functions in the **tibble** and **dplyr** packages, as well as **tidyr**'s **nest()**

tibble::tribble(...)

Makes list column when needed

tribble(~max, ~seq,
3, 1:3,
4, 1:4,
5, 1:5)

max	seq
3	<int [3]>
4	<int [4]>
5	<int [5]>

tibble::tibble(...)

Saves list input as list columns

tibble(max = c(3, 4, 5), seq = list(1:3, 1:4, 1:5))

tibble::enframe(x, name="name", value="value")

Converts multi-level list to tibble with list cols
enframe(list('3'=1:3, '4'=1:4, '5'=1:5), 'max', 'seq')

dplyr::mutate(.data, ...) Also **transmute()**

Returns list col when result returns list.

mtcars %>% **mutate**(seq = **map**(cyl, seq))

dplyr::summarise(.data, ...)

Returns list col when result is wrapped with **list()**

mtcars %>% **group_by**(cyl) %>%
summarise(q = **list**(quantile(mpg)))

2. WORK WITH LIST COLUMNS - Use the **purrr** functions **map()**, **map2()**, and **pmap()** to apply a function that returns a result element-wise to the cells of a list column. **walk()**, **walk2()**, and **pwalk()** work the same way, but return a side effect.

purrr::map(.x, .f, ...)

Apply .f element-wise to .x as .f(.x)

n_iris %>% **mutate**(n = **map**(data, dim))

purrr::map2(.x, .y, .f, ...)

Apply .f element-wise to .x and .y as .f(.x, .y)

m_iris %>% **mutate**(n = **map2**(data, model, list))

purrr::pmap(.l, .f, ...)

Apply .f element-wise to vectors saved in .l

m_iris %>%
mutate(n = **pmap**(list(data, model, data), list))

data	fun	result
<tibble [50x4]>	fun	result 1
<tibble [50x4]>	fun	result 2
<tibble [50x4]>	fun	result 3

data	model	fun	result
<tibble [50x4]>	<S3: lm>	fun	result 1
<tibble [50x4]>	<S3: lm>	fun	result 2
<tibble [50x4]>	<S3: lm>	fun	result 3

data	model	funcs	result
<tibble [50x4]>	<S3: lm>	coef	result 1
<tibble [50x4]>	<S3: lm>	AIC	result 2
<tibble [50x4]>	<S3: lm>	BIC	result 3

3. SIMPLIFY THE LIST COLUMN (into a regular column)

Use the **purrr** functions **map_lgl()**, **map_int()**, **map_dbl()**, **map_chr()**, as well as **tidyr**'s **unnest()** to reduce a list column into a regular column.

purrr::map_lgl(.x, .f, ...)

Apply .f element-wise to .x, return a logical vector
n_iris %>% **transmute**(n = **map_lgl**(data, is.matrix))

purrr::map_int(.x, .f, ...)

Apply .f element-wise to .x, return an integer vector
n_iris %>% **transmute**(n = **map_int**(data, nrow))

purrr::map_dbl(.x, .f, ...)

Apply .f element-wise to .x, return a double vector
n_iris %>% **transmute**(n = **map_dbl**(data, nrow))

purrr::map_chr(.x, .f, ...)

Apply .f element-wise to .x, return a character vector
n_iris %>% **transmute**(n = **map_chr**(data, nrow))

Data Import :: CHEAT SHEET



R's **tidyverse** is built around **tidy data** stored in **tibbles**, which are enhanced data frames.



The front side of this sheet shows how to read text files into R with **readr**.



The reverse side shows how to create tibbles with **tibble** and to layout tidy data with **tidyr**.

OTHER TYPES OF DATA

Try one of the following packages to import other types of files

- **haven** - SPSS, Stata, and SAS files
- **readxl** - excel files (.xls and .xlsx)
- **DBI** - databases
- **jsonlite** - json
- **xml2** - XML
- **httr** - Web APIs
- **rvest** - HTML (Web Scraping)

Save Data

Save **x**, an R object, to **path**, a file path, as:

Comma delimited file

write_csv(x, path, na = "NA", append = FALSE, col_names = !append)

File with arbitrary delimiter

write_delim(x, path, delim = " ", na = "NA", append = FALSE, col_names = !append)

CSV for excel

write_excel_csv(x, path, na = "NA", append = FALSE, col_names = !append)

String to file

write_file(x, path, append = FALSE)

String vector to file, one element per line

write_lines(x, path, na = "NA", append = FALSE)

Object to RDS file

write_rds(x, path, compress = c("none", "gz", "bz2", "xz"), ...)

Tab delimited files

write_tsv(x, path, na = "NA", append = FALSE, col_names = !append)

Read Tabular Data - These functions share the common arguments:

```
read_*(file, col_names = TRUE, col_types = NULL, locale = default_locale(), na = c("", "NA"),
quoted_na = TRUE, comment = "", trim_ws = TRUE, skip = 0, n_max = Inf, guess_max = min(1000,
n_max), progress = interactive())
```

a,b,c
1,2,3
4,5,NA

A	B	C
1	2	3
4	5	NA

Comma Delimited Files

read_csv("file.csv")

To make file.csv run:

write_file(x = "a,b,c\n1,2,3\n4,5,NA", path = "file.csv")

a;b;c
1;2;3
4;5;NA

A	B	C
1	2	3
4	5	NA

Semi-colon Delimited Files

read_csv2("file2.csv")

write_file(x = "a;b;c\n1;2;3\n4;5;NA", path = "file2.csv")

a|b|c
1|2|3
4|5|NA

A	B	C
1	2	3
4	5	NA

Files with Any Delimiter

read_delim("file.txt", delim = "|")

write_file(x = "a|b|c\n1|2|3\n4|5|NA", path = "file.txt")

a b c
1 2 3
4 5 NA

A	B	C
1	2	3
4	5	NA

Fixed Width Files

read_fwf("file.fwf", col_positions = c(1, 3, 5))

write_file(x = "a b c\n1 2 3\n4 5 NA", path = "file.fwf")

Tab Delimited Files

read_tsv("file.tsv") Also **read_table**().

write_file(x = "a\tb\tc\n1\t2\t3\n4\t5\tNA", path = "file.tsv")

USEFUL ARGUMENTS

a,b,c
1,2,3
4,5,NA

Example file

write_file("a,b,c\n1,2,3\n4,5,NA","file.csv")
f <- "file.csv"

1	2	3
4	5	NA

Skip lines

read_csv(f, **skip** = 1)

A	B	C
1	2	3
4	5	NA

No header

read_csv(f, **col_names** = FALSE)

A	B	C
1	2	3

Read in a subset

read_csv(f, **n_max** = 1)

x	y	z
A	B	C
1	2	3
4	5	NA

Provide header

read_csv(f, **col_names** = c("x", "y", "z"))

A	B	C
NA	2	3
4	5	NA

Missing Values

read_csv(f, **na** = c("1", "!"))

Read Non-Tabular Data

Read a file into a single string

read_file(file, locale = default_locale())

Read each line into its own string

read_lines(file, skip = 0, n_max = -1L, na = character(), locale = default_locale(), progress = interactive())

Read Apache style log files

read_log(file, col_names = FALSE, col_types = NULL, skip = 0, n_max = -1, progress = interactive())

Read a file into a raw vector

read_file_raw(file)

Read each line into a raw vector

read_lines_raw(file, skip = 0, n_max = -1L, progress = interactive())

Data types

readr functions guess the types of each column and convert types when appropriate (but will NOT convert strings to factors automatically).

A message shows the type of each column in the result.

```
## Parsed with column specification:
## cols(
##   age = col_integer(),
##   sex = col_character(),
##   earn = col_double()
## )
```

age is an integer

sex is a character

earn is a double (numeric)

1. Use **problems()** to diagnose problems.

x <- **read_csv**("file.csv"); **problems**(x)

2. Use a **col_** function to guide parsing.

- **col_guess()** - the default
- **col_character()**
- **col_double()**, **col_euro_double()**
- **col_datetime**(format = ""), Also **col_date**(format = ""), **col_time**(format = "")
- **col_factor**(levels, ordered = FALSE)
- **col_integer()**
- **col_logical()**
- **col_number()**, **col_numeric()**
- **col_skip()**

x <- **read_csv**("file.csv", **col_types** = **cols**(
A = **col_double**(),
B = **col_logical**(),
C = **col_factor**()))

3. Else, read in as character vectors then parse with a **parse_** function.

- **parse_guess()**
- **parse_character()**
- **parse_datetime()** Also **parse_date()** and **parse_time()**
- **parse_double()**
- **parse_factor()**
- **parse_integer()**
- **parse_logical()**
- **parse_number()**

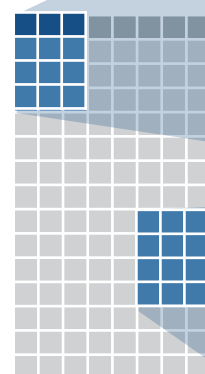
x\$A <- **parse_number**(x\$A)

Tibbles - an enhanced data frame



The **tibble** package provides a new S3 class for storing tabular data, the tibble. Tibbles inherit the data frame class, but improve three behaviors:

- **Subsetting** - `[` always returns a new tibble, `[[` and `$` always return a vector.
- **No partial matching** - You must use full column names when subsetting
- **Display** - When you print a tibble, R provides a concise view of the data that fits on one screen



A large table to display

```
# A tibble: 234 x 6
  manufacturer model displ
  <chr> <chr> <dbl>
1 audi a4 1.8
2 audi a4 1.8
3 audi a4 2.0
4 audi a4 2.0
5 audi a4 2.0
6 audi a4 3.1
7 audi a4 3.1
8 audi a4 3.1
9 audi a4 3.1
10 audi a4 3.1
... with 224 more rows, and
# more variables: year <int>,
# cyl <int>, trans <chr>
```

tibble display

```
156 1999 6 auto(l4)
157 1999 6 auto(l4)
158 2008 6 auto(l4)
159 2008 8 auto(s4)
160 1999 4 manual(m5)
161 1999 4 auto(l4)
162 2008 4 manual(m5)
163 2008 4 manual(m5)
164 2008 4 auto(l4)
165 2008 4 auto(l4)
166 1999 4 auto(l4)
[ reached getOption("max.print")
  -- omitted 68 rows --]
```

data frame display

- Control the default appearance with options:
`options(tibble.print_max = n, tibble.print_min = m, tibble.width = Inf)`
- View full data set with **View()** or **glimpse()**
- Revert to data frame with **as.data.frame()**

CONSTRUCT A TIBBLE IN TWO WAYS

tibble(...)
Construct by columns.
`tibble(x = 1:3, y = c("a", "b", "c"))`

Both make this tibble

tribble(...)
Construct by rows.
`tribble(~x, ~y, 1, "a", 2, "b", 3, "c")`

```
A tibble: 3 x 2
  x y
  <int> <chr>
1 1 a
2 2 b
3 3 c
```

as_tibble(x, ...) Convert data frame to tibble.

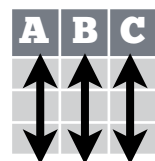
enframe(x, name = "name", value = "value")
Convert named vector to a tibble

is_tibble(x) Test whether x is a tibble.

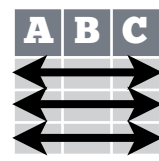
Tidy Data with tidyr

Tidy data is a way to organize tabular data. It provides a consistent data structure across packages.

A table is tidy if:

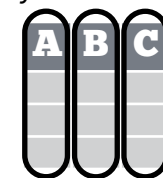


Each **variable** is in its own **column**

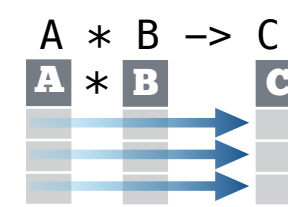


Each **observation**, or **case**, is in its own **row**

Tidy data:



Makes variables easy to access as vectors



Preserves cases during vectorized operations

Reshape Data - change the layout of values in a table

Use **gather()** and **spread()** to reorganize the values of a table into a new layout.

gather(data, key, value, ..., na.rm = FALSE, convert = FALSE, factor_key = FALSE)

gather() moves column names into a **key** column, gathering the column values into a single **value** column.

table4a

country	1999	2000
A	0.7K	2K
B	37K	80K
C	212K	213K



country	year	cases
A	1999	0.7K
B	1999	37K
C	1999	212K
A	2000	2K
B	2000	80K
C	2000	213K

key value

`gather(table4a, `1999`, `2000`,
key = "year", value = "cases")`

spread(data, key, value, fill = NA, convert = FALSE, drop = TRUE, sep = NULL)

spread() moves the unique values of a **key** column into the column names, spreading the values of a **value** column across the new columns.

table2

country	year	type	count
A	1999	cases	0.7K
A	1999	pop	19M
A	2000	cases	2K
A	2000	pop	20M
B	1999	cases	37K
B	1999	pop	172M
B	2000	cases	80K
B	2000	pop	174M
C	1999	cases	212K
C	1999	pop	1T
C	2000	cases	213K
C	2000	pop	1T

key value

`spread(table2, type, count)`

country	year	cases	pop
A	1999	0.7K	19M
A	2000	2K	20M
B	1999	37K	172M
B	2000	80K	174M
C	1999	212K	1T
C	2000	213K	1T

Handle Missing Values

drop_na(data, ...)

Drop rows containing NA's in ... columns.

x1	x2
A	1
B	NA
C	NA
D	3
E	NA

`drop_na(x, x2)`

fill(data, ..., .direction = c("down", "up"))

Fill in NA's in ... columns with most recent non-NA values.

x1	x2
A	1
B	NA
C	NA
D	3
E	NA

`fill(x, x2)`

replace_na(data, replace = list(), ...)

Replace NA's by column.

x1	x2
A	1
B	NA
C	NA
D	3
E	NA

`replace_na(x, list(x2 = 2))`

Expand Tables - quickly create tables with combinations of values

complete(data, ..., fill = list())

Adds to the data missing combinations of the values of the variables listed in ...

`complete(mtcars, cyl, gear, carb)`

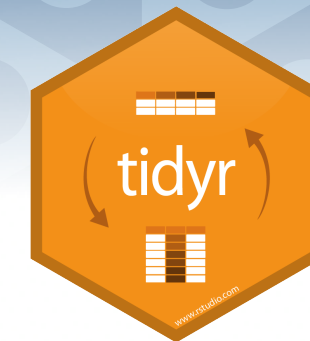
expand(data, ...)

Create new tibble with all possible combinations of the values of the variables listed in ...

`expand(mtcars, cyl, gear, carb)`

Split Cells

Use these functions to split or combine cells into individual, isolated values.



separate(data, col, into, sep = "[^:alnum:]]+", remove = TRUE, convert = FALSE, extra = "warn", fill = "warn", ...)

Separate each cell in a column to make several columns.

table3

country	year	rate
A	1999	0.7K/19M
A	2000	2K/20M
B	1999	37K/172M
B	2000	80K/174M
C	1999	212K/1T
C	2000	213K/1T



country	year	cases	pop
A	1999	0.7K	19M
A	2000	2K	20M
B	1999	37K	172
B	2000	80K	174
C	1999	212K	1T
C	2000	213K	1T

`separate(table3, rate,
into = c("cases", "pop"))`

separate_rows(data, ..., sep = "[^:alnum:]."]
+", convert = FALSE)

Separate each cell in a column to make several rows. Also **separate_rows_()**.

table3

country	year	rate
A	1999	0.7K/19M
A	2000	2K/20M
B	1999	37K/172M
C	1999	212K/1T
C	2000	213K/1T



country	year	rate
A	1999	0.7K
A	1999	19M
A	2000	2K
A	2000	20M
B	1999	37K
B	1999	172M
B	2000	80K
B	2000	174M
C	1999	212K
C	1999	1T
C	2000	213K
C	2000	1T

`separate_rows(table3, rate)`

unite(data, col, ..., sep = "_", remove = TRUE)

Collapse cells across several columns to make a single column.

table5

country	century	year
Afghan	19	99
Afghan	20	0
Brazil	19	99
Brazil	20	0
China	19	99
China	20	0



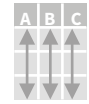
country	year
Afghan	1999
Afghan	2000
Brazil	1999
Brazil	2000
China	1999
China	2000

`unite(table5, century, year,
col = "year", sep = "")`

Data Transformation with dplyr :: CHEAT SHEET



dplyr functions work with pipes and expect **tidy data**. In tidy data:



&



pipes

Each **variable** is in its own **column**

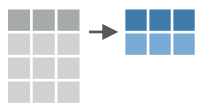
Each **observation**, or **case**, is in its own **row**

$x \%>\% f(y)$ becomes $f(x, y)$

Summarise Cases

These apply **summary functions** to columns to create a new table of summary statistics. Summary functions take vectors as input and return one value (see back).

summary function



summarise(.data, ...) Compute table of summaries.
summarise(mtcars, avg = mean(mpg))



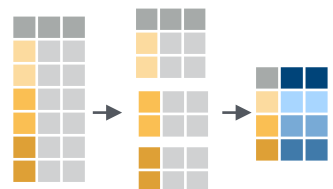
count(x, ..., wt = NULL, sort = FALSE) Count number of rows in each group defined by the variables in ... Also **tally**().
count(iris, Species)

VARIATIONS

summarise_all() - Apply funs to every column.
summarise_at() - Apply funs to specific columns.
summarise_if() - Apply funs to all cols of one type.

Group Cases

Use **group_by()** to create a "grouped" copy of a table. dplyr functions will manipulate each "group" separately and then combine the results.



*mtcars %>%
group_by(cyl) %>%
summarise(avg = mean(mpg))*

group_by(.data, ..., add = FALSE) Returns copy of table grouped by ...
g_iris <- group_by(iris, Species)

ungroup(x, ...) Returns ungrouped copy of table.
ungroup(g_iris)

Manipulate Cases

EXTRACT CASES

Row functions return a subset of rows as a new table.



filter(.data, ...) Extract rows that meet logical criteria. *filter(iris, Sepal.Length > 7)*



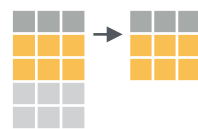
distinct(.data, ..., .keep_all = FALSE) Remove rows with duplicate values.
distinct(iris, Species)



sample_frac(tbl, size = 1, replace = FALSE, weight = NULL, .env = parent.frame()) Randomly select fraction of rows.
sample_frac(iris, 0.5, replace = TRUE)



sample_n(tbl, size, replace = FALSE, weight = NULL, .env = parent.frame()) Randomly select size rows. *sample_n(iris, 10, replace = TRUE)*



slice(.data, ...) Select rows by position.
slice(iris, 10:15)

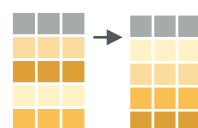
top_n(x, n, wt) Select and order top n entries (by group if grouped data). *top_n(iris, 5, Sepal.Width)*

Logical and boolean operators to use with filter()

<	<=	is.na()	%in%		xor()
>	>=	!is.na()	!	&	

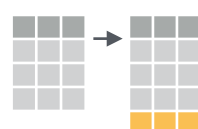
See **?base::logic** and **?Comparison** for help.

ARRANGE CASES



arrange(.data, ...) Order rows by values of a column or columns (low to high), use with **desc()** to order from high to low.
arrange(mtcars, mpg)
arrange(mtcars, desc(mpg))

ADD CASES



add_row(.data, ..., .before = NULL, .after = NULL) Add one or more rows to a table.
add_row(faithful, eruptions = 1, waiting = 1)

Manipulate Variables

EXTRACT VARIABLES

Column functions return a set of columns as a new vector or table.



pull(.data, var = -1) Extract column values as a vector. Choose by name or index.
pull(iris, Sepal.Length)



select(.data, ...) Extract columns as a table. Also **select_if()**.
select(iris, Sepal.Length, Species)

Use these helpers with **select()**, e.g. *select(iris, starts_with("Sepal"))*

contains (match)	num_range (prefix, range)	:, e.g. <i>mpg:cyl</i>
ends_with (match)	one_of (...)	-, e.g. <i>-Species</i>
matches (match)	starts_with (match)	

MAKE NEW VARIABLES

These apply **vectorized functions** to columns. Vectorized funs take vectors as input and return vectors of the same length as output (see back).

vectorized function



mutate(.data, ...) Compute new column(s).
mutate(mtcars, gpm = 1/mpg)



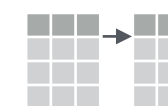
transmute(.data, ...) Compute new column(s), drop others.
transmute(mtcars, gpm = 1/mpg)



mutate_all(.tbl, .funs, ...) Apply funs to every column. Use with **funs()**. Also **mutate_if()**.
mutate_all(faithful, funs(log(.), log2(.)))
mutate_if(iris, is.numeric, funs(log(.)))



mutate_at(.tbl, .cols, .funs, ...) Apply funs to specific columns. Use with **funs()**, **vars()** and the helper functions for **select()**.
mutate_at(iris, vars(-Species), funs(log(.)))



add_column(.data, ..., .before = NULL, .after = NULL) Add new column(s). Also **add_count()**, **add_tally()**. *add_column(mtcars, new = 1:32)*



rename(.data, ...) Rename columns.
rename(iris, Length = Sepal.Length)

Vector Functions

TO USE WITH MUTATE ()

mutate() and **transmute()** apply vectorized functions to columns to create new columns. Vectorized functions take vectors as input and return vectors of the same length as output.

vectorized function

OFFSETS

dplyr::lag() - Offset elements by 1
dplyr::lead() - Offset elements by -1

CUMULATIVE AGGREGATES

dplyr::cumall() - Cumulative all()
dplyr::cumany() - Cumulative any()
cummax() - Cumulative max()
dplyr::cummean() - Cumulative mean()
cummin() - Cumulative min()
cumprod() - Cumulative prod()
cumsum() - Cumulative sum()

RANKINGS

dplyr::cume_dist() - Proportion of all values <=
dplyr::dense_rank() - rank with ties = min, no gaps
dplyr::min_rank() - rank with ties = min
dplyr::ntile() - bins into n bins
dplyr::percent_rank() - min_rank scaled to [0,1]
dplyr::row_number() - rank with ties = "first"

MATH

+, **-**, *****, **/**, **^**, **%/%**, **%%** - arithmetic ops
log(), **log2()**, **log10()** - logs
<, **<=**, **>**, **>=**, **!=**, **==** - logical comparisons
dplyr::between() - x >= left & x <= right
dplyr::near() - safe == for floating point numbers

MISC

dplyr::case_when() - multi-case if_else()
dplyr::coalesce() - first non-NA values by element across a set of vectors
dplyr::if_else() - element-wise if() + else()
dplyr::na_if() - replace specific values with NA
pmax() - element-wise max()
pmin() - element-wise min()
dplyr::recode() - Vectorized switch()
dplyr::recode_factor() - Vectorized switch() for factors

Summary Functions

TO USE WITH SUMMARISE ()

summarise() applies summary functions to columns to create a new table. Summary functions take vectors as input and return single values as output.

summary function

COUNTS

dplyr::n() - number of values/rows
dplyr::n_distinct() - # of uniques
sum(!is.na()) - # of non-NA's

LOCATION

mean() - mean, also **mean(!is.na())**
median() - median

LOGICALS

mean() - Proportion of TRUE's
sum() - # of TRUE's

POSITION/ORDER

dplyr::first() - first value
dplyr::last() - last value
dplyr::nth() - value in nth location of vector

RANK

quantile() - nth quantile
min() - minimum value
max() - maximum value

SPREAD

IQR() - Inter-Quartile Range
mad() - median absolute deviation
sd() - standard deviation
var() - variance

Row Names

Tidy data does not use rownames, which store a variable outside of the columns. To work with the rownames, first move them into a column.

rownames_to_column()
Move row names into col.
`a <- rownames_to_column(iris, var = "C")`

column_to_rownames()
Move col in row names.
`column_to_rownames(a, var = "C")`

Also **has_rownames()**, **remove_rownames()**

Combine Tables

COMBINE VARIABLES

x + **y** = **A B C A B D**

a	t	1		
b	u	2		
c	v	3		

 +

a	t	3		
b	u	2		
d	w	1		

 =

a	t	1	a	t	3
b	u	2	b	u	2
c	v	3	d	w	1

Use **bind_cols()** to paste tables beside each other as they are.

bind_cols(...) Returns tables placed side by side as a single table.
BE SURE THAT ROWS ALIGN.

Use a "Mutating Join" to join one table to columns from another, matching values with the rows that they correspond to. Each join retains a different combination of values from the tables.

left_join(x, y, by = NULL, copy=FALSE, suffix=c(".x",".y"),...)
Join matching values from y to x.

right_join(x, y, by = NULL, copy = FALSE, suffix=c(".x",".y"),...)
Join matching values from x to y.

inner_join(x, y, by = NULL, copy = FALSE, suffix=c(".x",".y"),...)
Join data. Retain only rows with matches.

full_join(x, y, by = NULL, copy=FALSE, suffix=c(".x",".y"),...)
Join data. Retain all values, all rows.

Use **by = c("col1", "col2", ...)** to specify one or more common columns to match on.
`left_join(x, y, by = "A")`

Use a named vector, **by = c("col1" = "col2")**, to match on columns that have different names in each table.
`left_join(x, y, by = c("C" = "D"))`

Use **suffix** to specify the suffix to give to unmatched columns that have the same name in both tables.
`left_join(x, y, by = c("C" = "D"), suffix = c("1", "2"))`

COMBINE CASES

x + **y** = **A B C**

a	t	1
b	u	2
c	v	3

 +

a	t	3
c	v	3
d	w	4

Use **bind_rows()** to paste tables below each other as they are.

bind_rows(..., .id = NULL)
Returns tables one on top of the other as a single table. Set .id to a column name to add a column of the original table names (as pictured)

intersect(x, y, ...)
Rows that appear in both x and y.

setdiff(x, y, ...)
Rows that appear in x but not y.

union(x, y, ...)
Rows that appear in x or y. (Duplicates removed). `union_all()` retains duplicates.

Use **setequal()** to test whether two data sets contain the exact same rows (in any order).

EXTRACT ROWS

x + **y** = **A B C**

a	t	1
b	u	2
c	v	3

 +

a	t	3
b	u	2
d	w	1

Use a "Filtering Join" to filter one table against the rows of another.

semi_join(x, y, by = NULL, ...)
Return rows of x that have a match in y. USEFUL TO SEE WHAT WILL BE JOINED.

anti_join(x, y, by = NULL, ...)
Return rows of x that do not have a match in y. USEFUL TO SEE WHAT WILL NOT BE JOINED.



R Markdown :: CHEAT SHEET

What is R Markdown?

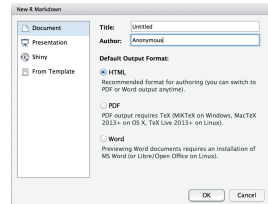


.Rmd files • An R Markdown (.Rmd) file is a record of your research. It contains the code that a scientist needs to reproduce your work along with the narration that a reader needs to understand your work.

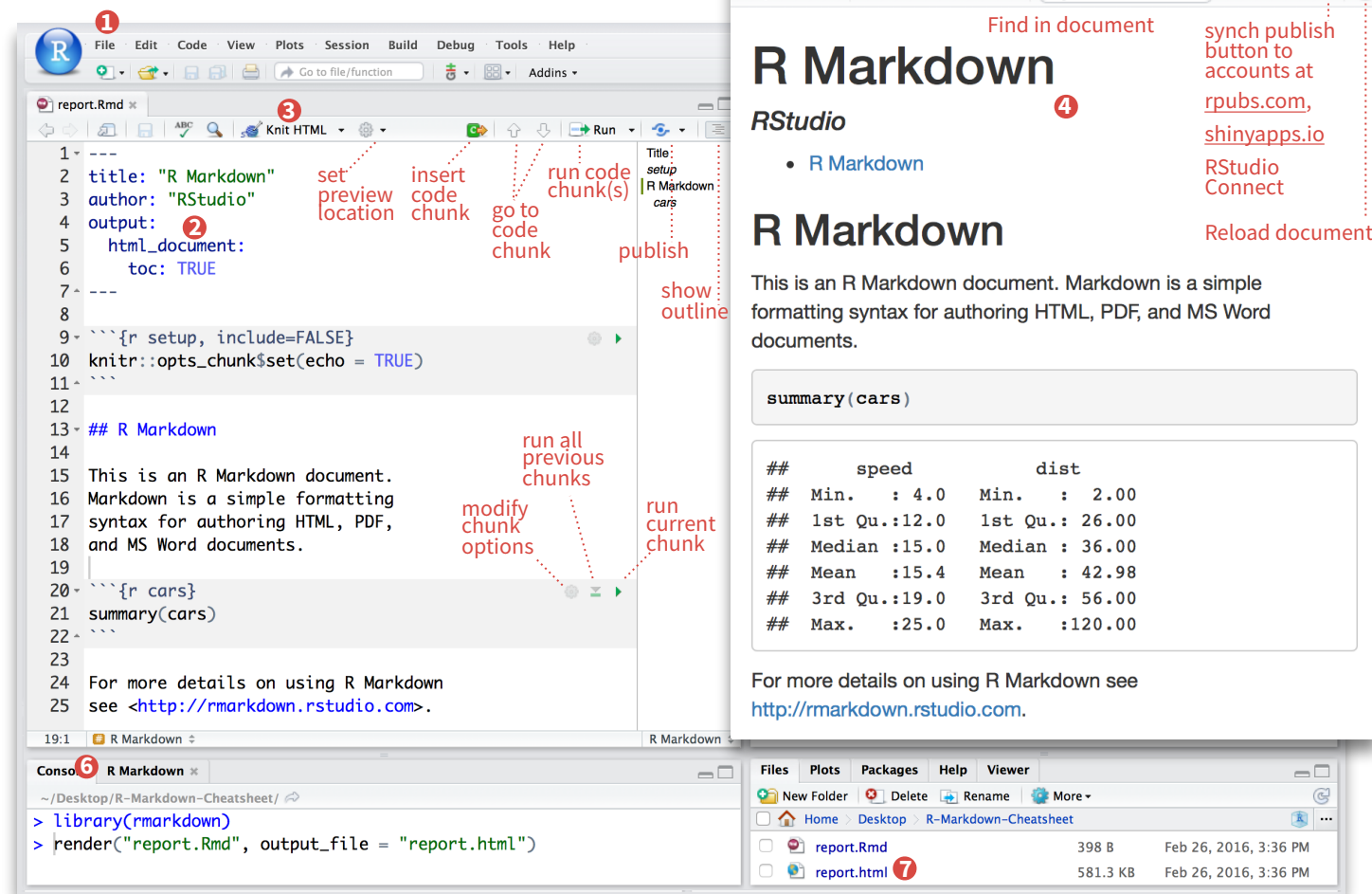
Reproducible Research • At the click of a button, or the type of a command, you can rerun the code in an R Markdown file to reproduce your work and export the results as a finished report.

Dynamic Documents • You can choose to export the finished report in a variety of formats, including html, pdf, MS Word, or RTF documents; html or pdf based slides, Notebooks, and more.

Workflow



- 1 **Open a new .Rmd file** at File ► New File ► R Markdown. Use the wizard that opens to pre-populate the file with a template
- 2 **Write document** by editing template
- 3 **Knit document to create report**; use knit button or `render()` to knit
- 4 **Preview Output** in IDE window
- 5 **Publish** (optional) to web server
- 6 **Examine build log** in R Markdown console
- 7 **Use output file** that is saved along side .Rmd



render

Use `rmarkdown::render()` to render/knit at cmd line. Important args:

input - file to render
output_format

output_options - List of render options (as in YAML)

output_file
output_dir

params - list of params to use

envir - environment to evaluate code chunks in

encoding - of input file

Embed code with knitr syntax

INLINE CODE

Insert with ``r <code>``. Results appear as text without code.

Built with ``r getRversion()`` ➔ Built with 3.2.3

CODE CHUNKS

One or more lines surrounded with ````\{r\}` and `````. Place chunk options within curly braces, after `r`. Insert with

```
```\{r echo=TRUE\}  
getRversion()
```\n
```

```
getRversion()  
## [1] '3.2.3'\n
```

GLOBAL OPTIONS

Set with `knitr::opts_chunk$set()`, e.g.

```
```\{r include=FALSE\}  
knitr::opts_chunk$set(echo = TRUE)
```\n
```

IMPORTANT CHUNK OPTIONS

cache - cache results for future knits (default = FALSE)

cache.path - directory to save cached results in (default = "cache/")

child - file(s) to knit and then include (default = NULL)

collapse - collapse all output into single block (default = FALSE)

comment - prefix for each line of results (default = '##')

dependson - chunk dependencies for caching (default = NULL)

echo - Display code in output document (default = TRUE)

engine - code language used in chunk (default = 'R')

error - Display error messages in doc (TRUE) or stop render when errors occur (FALSE) (default = FALSE)

eval - Run code in chunk (default = TRUE)

fig.align - 'left', 'right', or 'center' (default = 'default')

fig.cap - figure caption as character string (default = NULL)

fig.height, **fig.width** - Dimensions of plots in inches

highlight - highlight source code (default = TRUE)

include - Include chunk in doc after running (default = TRUE)

message - display code messages in document (default = TRUE)

results (default = 'markup')

'asis' - passthrough results

'hide' - do not display results

'hold' - put all results below all code

tidy - tidy code for display (default = FALSE)

warning - display code warnings in document (default = TRUE)

Options not listed above: `R.options`, `aniopts`, `autodep`, `background`, `cache.comments`, `cache.lazy`, `cache.rebuild`, `cache.vars`, `dev`, `dev.args`, `dpi`, `engine.opts`, `engine.path`, `fig.asp`, `fig.env`, `fig.ext`, `fig.keep`, `fig.lp`, `fig.path`, `fig.pos`, `fig.process`, `fig.retina`, `fig.scap`, `fig.show`, `fig.showtext`, `fig.subcap`, `interval`, `out.extra`, `out.height`, `out.width`, `prompt`, `purl`, `ref.label`, `render`, `size`, `split`, `tidy.opts`

.rmd Structure

rmarkdown

YAML Header

Optional section of render (e.g. pandoc) options written as key:value pairs (YAML).

At start of file

Between lines of ---

Text

Narration formatted with markdown, mixed with:

Code Chunks

Chunks of embedded code. Each chunk:

Begins with ````\{r\}`

ends with `````

R Markdown will run the code and append the results to the doc.

It will use the location of the .Rmd file as the **working directory**

Parameters

Parameterize your documents to reuse with different inputs (e.g., data, values, etc.)

1. **Add parameters** • Create and set parameters in the header as sub-values of params

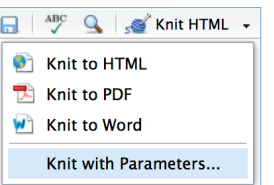
```
---  
params:  
  n: 100  
  d: !r Sys.Date()  
---
```

2. **Call parameters** • Call parameter values in code as `params$<name>`

Today's date is `!r params$d`

3. **Set parameters** • Set values with Knit with parameters or the params argument of `render()`:

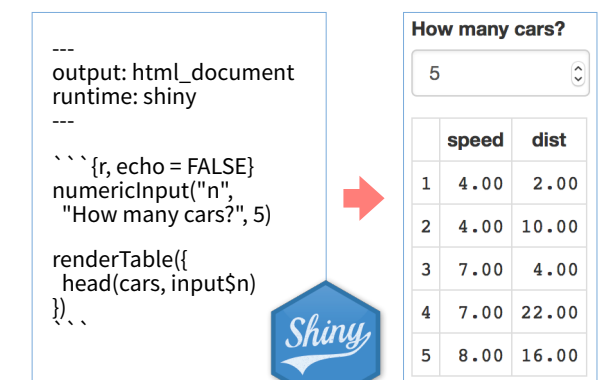
```
render("doc.Rmd", params = list(n = 1,  
  d = as.Date("2015-01-01")))
```



Interactive Documents

Turn your report into an interactive Shiny document in 4 steps

1. Add runtime: shiny to the YAML header.
2. Call Shiny input functions to embed input objects.
3. Call Shiny render functions to embed reactive output.
4. Render with `rmarkdown::run` or click Run Document in RStudio IDE



Embed a complete app into your document with `shiny::shinyAppDir()`

NOTE: Your report will be rendered as a Shiny app, which means you must choose an html output format, like **html_document**, and serve it with an active R Session.



Pandoc's Markdown

Write with syntax on the left to create effect on right (after render)

Plain text
End a line with two spaces
to start a new paragraph.
italics and **bold**
`verbatim code`
sub/superscript^2^~2~
~~strikethrough~~
escaped: * _ \\\
endash: --, emdash: ---
equation: \$A = \pi * r^2\$
equation block:

\$\$E = mc^2\$\$

> block quote

Header1 {#anchor}

Header 2 {#css_id}

Header 3 {.css_class}

Header 4

Header 5

Header 6

<!--Text comment-->

\textbf{Text ignored in HTML}
HTML ignored in pdfs

<http://www.rstudio.com>
[link](www.rstudio.com)
Jump to [Header 1](#anchor)
image:

![Caption](smallorb.png)

* unordered list
+ sub-item 1
+ sub-item 2
- sub-sub-item 1

* item 2

Continued (indent 4 spaces)

1. ordered list
2. item 2
i) sub-item 1
A. sub-sub-item 1

(@) A list whose numbering

continues after

(@) an interruption

Term 1

: Definition 1

Right	Left	Default	Center
12	12	12	12
123	123	123	123
1	1	1	1

- slide bullet 1
- slide bullet 2

(>- to have bullets appear on click)

horizontal rule/slide break:

A footnote [^1]

[^1]: Here is the footnote.

Plain text
End a line with two spaces
to start a new paragraph.
italics and **bold**
`verbatim code`
sub/superscript^2^2
strikethrough
escaped: * _ \
endash: --, emdash: ---
equation: $A = \pi * r^2$
equation block:

block quote

block quote

Header1

Header 2

Header 3

Header 4

Header 5

Header 6

HTML ignored in pdfs

<http://www.rstudio.com>

Jump to [Header 1](#)

image:

- unordered list
 - sub-item 1
 - sub-item 2
 - sub-sub-item 1
- item 2

Continued (indent 4 spaces)

1. ordered list
2. item 2
i. sub-item 1
A. sub-sub-item 1

1. A list whose numbering

continues after

2. an interruption

Term 1

Definition 1

Right	Left	Default	Center
12	12	12	12
123	123	123	123
1	1	1	1

- slide bullet 1
- slide bullet 2

(>- to have bullets appear on click)

horizontal rule/slide break:

A footnote ¹

1. Here is the footnote. ➔

Set render options with YAML

When you render, R Markdown

1. runs the R code, embeds results and text into .md file with knitr
2. then converts the .md file into the finished format with pandoc



Set a document's
default output format
in the YAML header:

```
---  
output: html_document  
---  
# Body
```

output value

html_document

pdf_document

word_document

odt_document

rtf_document

md_document

github_document

ioslides_presentation

slidy_presentation

beamer_presentation

creates

html

pdf (requires Tex)

Microsoft Word (.docx)

OpenDocument Text

Rich Text Format

Markdown

Github compatible markdown

ioslides HTML slides

slidy HTML slides

Beamer pdf slides (requires Tex)

Customize output with
sub-options (listed to
the right):

```
---  
output: html_document:  
  code_folding: hide  
  toc_float: TRUE  
---  
# Body
```

Indent 2
spaces

Indent 4
spaces

html tabsets

Use tablet css class to place sub-headers into tabs

```
# Tabset {.tabset .tabset-fade .tabset-pills}  
## Tab 1  
text 1  
## Tab 2  
text 2  
### End tabset
```

Tabset

Tab 1

Tab 2

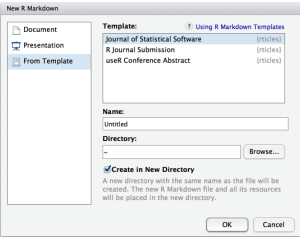
text 1

End tabset

Create a Reusable Template

1. **Create a new package** with a inst/rmarkdown/templates directory
2. In the directory, **Place a folder** that contains:
template.yaml (see below)
skeleton.Rmd (contents of the template)
any supporting files
3. **Install the package**
4. **Access template** in wizard at File ► New File ► R Markdown
template.yaml

```
---  
name: My Template  
---
```



sub-option	description	html	pdf	word	odt	rtf	md	github	ioslides	slidy	beamer
citation_package	The LaTeX package to process citations, natbib, biblatex or none		X					X			X
code_folding	Let readers to toggle the display of R code, "none", "hide", or "show"	X									
colortheme	Beamer color theme to use										X
css	CSS file to use to style document	X								X	X
dev	Graphics device to use for figure output (e.g. "png")	X	X				X	X	X	X	X
duration	Add a countdown timer (in minutes) to footer of slides									X	
fig_caption	Should figures be rendered with captions?	X	X	X	X				X	X	X
fig_height, fig_width	Default figure height and width (in inches) for document	X	X	X	X	X	X	X	X	X	X
highlight	Syntax highlighting: "tango", "pygments", "kate", "zenburn", "textmate"	X	X	X						X	X
includes	File of content to place in document (in_header, before_body, after_body)	X	X		X		X	X	X	X	X
incremental	Should bullets appear one at a time (on presenter mouse clicks)?								X	X	X
keep_md	Save a copy of .md file that contains knitr output	X		X	X	X			X	X	
keep_tex	Save a copy of .tex file that contains knitr output	X									X
latex_engine	Engine to render latex, "pdflatex", "xelatex", or "lualatex"	X									X
lib_dir	Directory of dependency files to use (Bootstrap, MathJax, etc.)	X							X	X	
mathjax	Set to local or a URL to use a local/URL version of MathJax to render equations	X							X	X	
md_extensions	Markdown extensions to add to default definition or R Markdown	X	X	X	X	X	X	X	X	X	X
number_sections	Add section numbering to headers	X	X								
pandoc_args	Additional arguments to pass to Pandoc	X	X	X	X	X	X	X	X	X	X
preserve_yaml	Preserve YAML front matter in final document?						X				
reference_docx	docx file whose styles should be copied when producing docx output			X							
self_contained	Embed dependencies into the doc	X							X	X	
slide_level	The lowest heading level that defines individual slides										X
smaller	Use the smaller font size in the presentation?								X		
smart	Convert straight quotes to curly, dashes to em-dashes, ... to ellipses, etc.	X							X	X	
template	Pandoc template to use when rendering file quarterly_report.html).	X	X		X				X	X	
theme	Bootswatch or Beamer theme to use for page	X									X
toc	Add a table of contents at start of document	X	X	X		X	X	X			X
toc_depth	The lowest level of headings to add to table of contents	X	X	X		X	X	X			
toc_float	Float the table of contents to the left of the main content	X									

Table Suggestions

Several functions format R data into tables

```
Table with kable  
eruptions waiting  
1 3.600 79  
2 1.800 54  
3 3.333 74  
4 2.283 62
```

```
data <- faithful[1:4,]
```

```
`{r results = 'asis'}
```

```
knitr::kable(data, caption = "Table with kable")
```

```
`{r results = "asis"}
```

```
print(xtable::xtable(data, caption = "Table with xtable"),  
      type = "html", html.table.attributes = "border=0")
```

```
`{r results = "asis"}
```

```
stargazer::stargazer(data, type = "html", title = "Table  
with stargazer")
```

```
---
```

```
Table with stargazer  
eruptionswaiting  
1 3.600 79  
2 1.800 54  
3 3.333 74  
4 2.283 62
```

```
`{r results = "asis"}
```

```
stargazer::stargazer(data, type = "html", title = "Table  
with stargazer")
```

```
---
```

Citations and Bibliographies

Create citations with .bib, .bibtex, .copac, .enl, .json, .medline, .mods, .ris, .wos, and .xml files

1. **Set bibliography file** and CSL 1.0
Style file (optional) in the YAML header
2. **Use citation keys in text**

```
---  
bibliography: refs.bib  
cs1: style.csl  
---
```

Smith cited [@smith04].
Smith cited without author [-@smith04].
@smith04 cited in line.

3. **Render.** Bibliography will be
added to end of document

Smith cited (Joe Smith 2004).
Smith cited without author (2004).
Joe Smith (2004) cited in line.

Learn more in
the **stargazer**,
xtable, and **knitr**
packages.

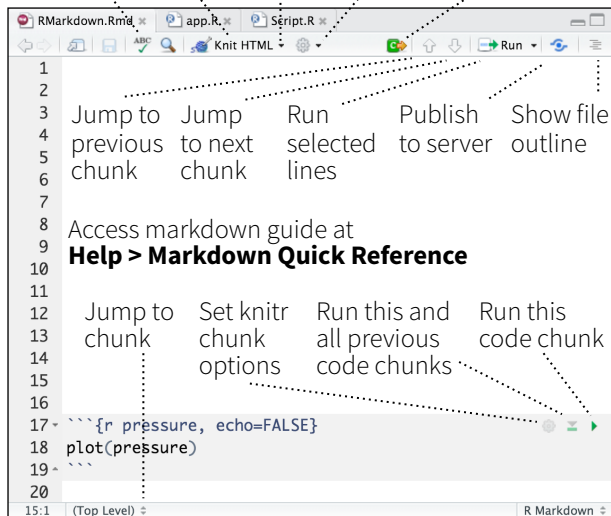
RStudio IDE :: CHEAT SHEET



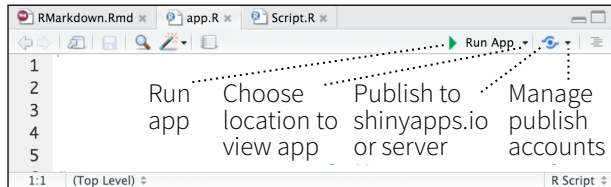
Documents and Apps

Open Shiny, R Markdown, knitr, Sweave, LaTeX, .Rd files and more in Source Pane

Check spelling Render output Choose output format Choose output location Insert code chunk

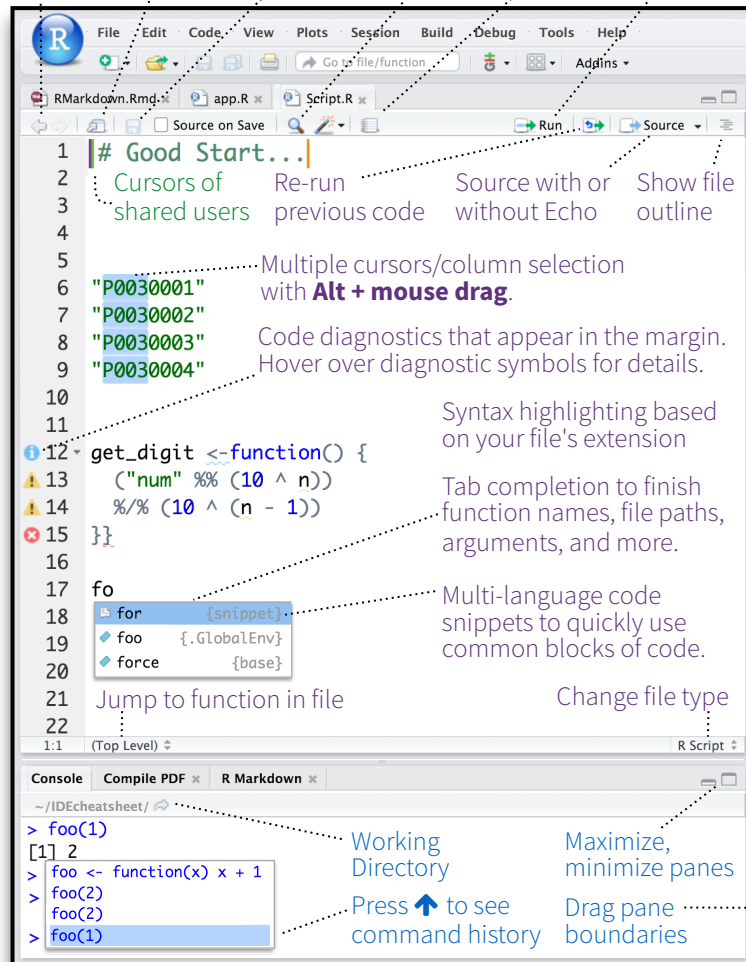


RStudio recognizes that files named **app.R**, **server.R**, **ui.R**, and **global.R** belong to a shiny app



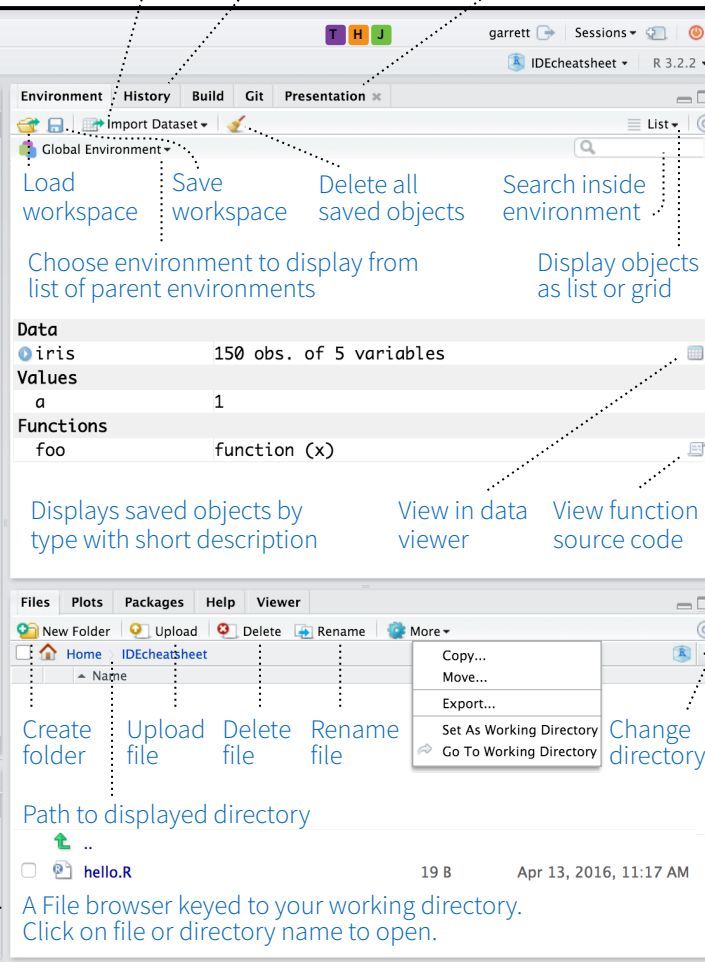
Write Code

Navigate tabs Open in new window Save Find and replace Compile as notebook Run selected code



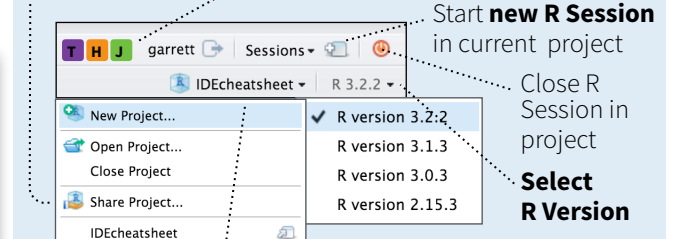
R Support

Import data with wizard History of past commands to run/copy Display .RPres slideshows **File > New File > R Presentation**



Pro Features

Share Project with Collaborators Active shared collaborators



PROJECT SYSTEM

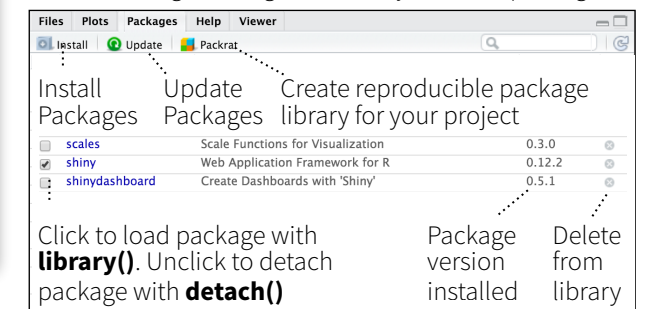
File > New Project

RStudio saves the call history, workspace, and working directory associated with a project. It reloads each when you re-open a project.

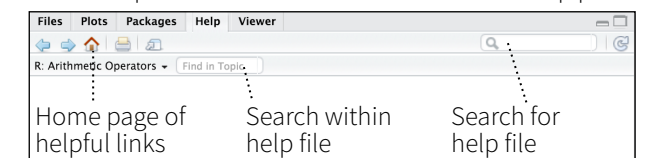
RStudio opens plots in a dedicated Plots pane



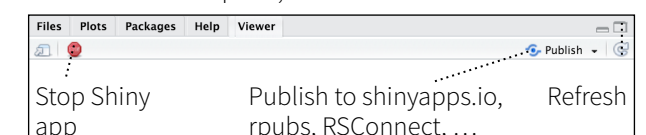
GUI Package manager lists every installed package



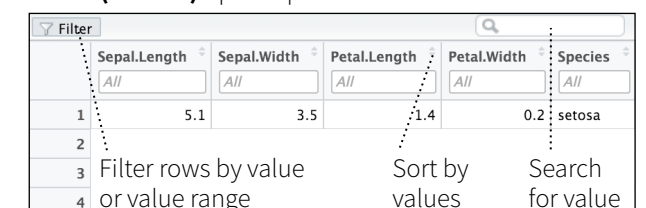
RStudio opens documentation in a dedicated Help pane



Viewer Pane displays HTML content, such as Shiny apps, RMarkdown reports, and interactive visualizations

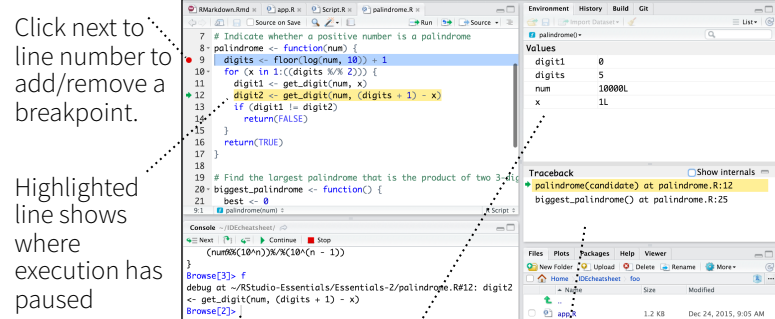


View(<data>) opens spreadsheet like view of data set



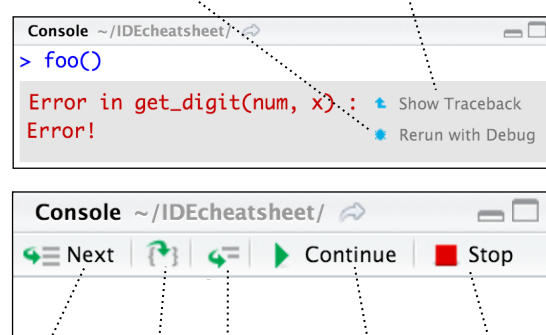
Debug Mode

Open with **debug()**, **browser()**, or a breakpoint. RStudio will open the debugger mode when it encounters a breakpoint while executing code.



Run commands in environment where execution has paused Examine variables in executing environment Select function in traceback to debug

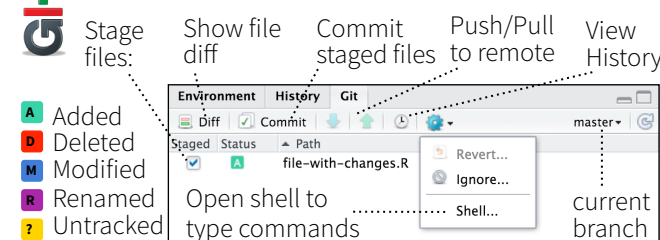
Launch debugger from origin of error Open traceback to examine the functions that R called before the error occurred



Step through code one line at a time Step into and out of functions to run Resume execution mode Quit debug

Version Control with Git or SVN

Turn on at **Tools > Project Options > Git/SVN**

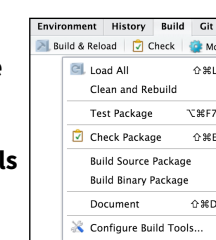


Package Writing

File > New Project > New Directory > R Package

Turn project into package, Enable roxygen documentation with **Tools > Project Options > Build Tools**

Roxygen guide at **Help > Roxygen Quick Reference**





1 LAYOUT

Move focus to Source Editor
Move focus to Console
Move focus to Help
Show History
Show Files
Show Plots
Show Packages
Show Environment
Show Git/SVN
Show Build

Windows/Linux Mac

Ctrl+1
Ctrl+2
Ctrl+3
Ctrl+4
Ctrl+5
Ctrl+6
Ctrl+7
Ctrl+8
Ctrl+9
Ctrl+0

2 RUN CODE

Search command history

Navigate command history
Move cursor to start of line
Move cursor to end of line
Change working directory

Interrupt current command

Clear console

Quit Session (desktop only)

Restart R Session

Run current line/selection

Run current (retain cursor)
Run from current to end
Run the current function
Source a file

Source the current file

Source with echo

Windows/Linux Mac

Ctrl+↑
↑/↓
Home
End
Ctrl+Shift+H
Esc
Ctrl+L
Ctrl+Q
Ctrl+Shift+F10
Ctrl+Enter
Alt+Enter
Ctrl+Alt+E
Ctrl+Alt+F
Ctrl+Alt+G
Ctrl+Shift+S
Ctrl+Shift+Enter

3 NAVIGATE CODE

Goto File/Function

Fold Selected
Unfold Selected
Fold All
Unfold All
Go to line
Jump to
Switch to tab
Previous tab
Next tab
First tab
Last tab
Navigate back
Navigate forward
Jump to Brace
Select within Braces
Use Selection for Find
Find in Files
Find Next
Find Previous
Jump to Word
Jump to Start/End
Toggle Outline

Windows /Linux

Ctrl+.
Alt+L
Shift+Alt+L
Alt+O
Shift+Alt+O
Shift+Alt+G
Shift+Alt+J
Ctrl+Shift+.
Ctrl+F11
Ctrl+F12
Ctrl+Shift+F11
Ctrl+Shift+F12
Ctrl+F9
Ctrl+F10
Ctrl+P
Ctrl+Shift+Alt+E
Ctrl+F3
Ctrl+Shift+F
Win: F3, Linux: Ctrl+G
W: Shift+F3, L:
Ctrl+↔
Ctrl+↑/↓
Ctrl+Shift+O

Mac

Ctrl+.
Cmd+Option+L
Cmd+Shift+Option+L
Cmd+Option+O
Cmd+Shift+Option+O
Cmd+Shift+Option+G
Cmd+Shift+Option+J
Ctrl+Shift+.
Ctrl+F11
Ctrl+F12
Ctrl+Shift+F11
Ctrl+Shift+F12
Cmd+F9
Cmd+F10
Ctrl+P
Ctrl+Shift+Option+E
Cmd+E
Cmd+Shift+F
Cmd+G
Cmd+Shift+G
Option+↔
Cmd+↑/↓
Cmd+Shift+O

4 WRITE CODE

Attempt completion

Navigate candidates
Accept candidate
Dismiss candidates
Undo
Redo
Cut
Copy
Paste
Select All
Delete Line
Select
Select Word
Select to Line Start
Select to Line End
Select Page Up/Down
Select to Start/End
Delete Word Left
Delete Word Right
Delete to Line End
Delete to Line Start
Indent
Outdent
Yank line up to cursor
Yank line after cursor
Insert yanked text

Insert <-

Insert %>%

Show help for function
Show source code
New document
New document (Chrome)
Open document
Save document
Close document
Close document (Chrome)
Close all documents
Extract function
Extract variable
Reindent lines

(Un)Comment lines

Reflow Comment
Reformat Selection
Select within braces
Show Diagnostics
Transpose Letters
Move Lines Up/Down
Copy Lines Up/Down
Add New Cursor Above
Add New Cursor Below
Move Active Cursor Up
Move Active Cursor Down
Find and Replace
Use Selection for Find
Replace and Find

Windows /Linux

Tab or Ctrl+Space

↑/↓
Enter, Tab, or →
Esc
Ctrl+Z
Ctrl+Shift+Z
Ctrl+X
Ctrl+C
Ctrl+V
Ctrl+A
Ctrl+D
Shift+[Arrow]
Ctrl+Shift+↔
Alt+Shift+↔
Alt+Shift+→
Shift+PageUp/Down
Shift+Alt+↑/↓
Ctrl+Backspace

Tab (at start of line)

Shift+Tab

Ctrl+U

Ctrl+K

Ctrl+Y

Alt+-

Ctrl+Shift+M

F1
F2
Ctrl+Shift+N
Ctrl+Alt+Shift+N
Ctrl+O
Ctrl+S
Ctrl+W
Ctrl+Alt+W
Ctrl+Shift+W
Ctrl+Alt+X
Ctrl+Alt+V
Ctrl+I
Ctrl+Shift+C
Ctrl+Shift+/
Ctrl+Shift+A
Ctrl+Shift+E
Ctrl+Shift+Alt+P

Alt+↑/↓

Shift+Alt+↑/↓

Ctrl+Alt+Up

Ctrl+Alt+Down

Ctrl+Alt+Shift+Up

Ctrl+Alt+Shift+Down

Ctrl+F

Ctrl+F3

Ctrl+Shift+J

Mac

Tab or Cmd+Space

↑/↓
Enter, Tab, or →
Esc
Cmd+Z
Cmd+Shift+Z
Cmd+X
Cmd+C
Cmd+V
Cmd+A
Cmd+D
Shift+[Arrow]
Option+Shift+↔
Cmd+Shift+↔
Cmd+Shift+→
Shift+PageUp/Down
Cmd+Shift+↑/↓
Ctrl+Opt+Backspace
Option+Delete
Ctrl+K

Option+Backspace

Tab (at start of line)

Shift+Tab

Ctrl+U

Ctrl+K

Ctrl+Y

Option+-

Cmd+Shift+M

F1
F2
Cmd+Shift+N
Cmd+Shift+Opt+N
Cmd+O
Cmd+S
Cmd+W
Cmd+Option+W
Cmd+Shift+W
Cmd+Option+X
Cmd+Option+V
Cmd+I
Cmd+Shift+C
Cmd+Shift+/
Cmd+Shift+A
Ctrl+Shift+E
Cmd+Shift+Opt+P
Ctrl+T
Option+↑/↓
Cmd+Option+↑/↓
Ctrl+Option+Up
Ctrl+Option+Down
Ctrl+Option+Shift+Up
Ctrl+Opt+Shift+Down
Cmd+F
Cmd+E
Cmd+Shift+J

WHY RSTUDIO SERVER PRO?

RSP extends the the open source server with a commercial license, support, and more:

- open and run multiple R sessions at once
- tune your resources to improve performance
- edit the same project at the same time as others
- see what you and others are doing on your server
- switch easily from one version of R to a different version
- integrate with your authentication, authorization, and audit practices

Download a free 45 day evaluation at

www.rstudio.com/products/rstudio-server-pro/

5 DEBUG CODE

Toggle Breakpoint
Execute Next Line
Step Into Function
Finish Function/Loop
Continue
Stop Debugging

Windows/Linux Mac

Shift+F9
F10
Shift+F4
Shift+F6
Shift+F5
Shift+F8
Shift+F9
F10
Shift+F4
Shift+F6
Shift+F5
Shift+F8

6 VERSION CONTROL

Show diff
Commit changes
Scroll diff view
Stage/Unstage (Git)
Stage/Unstage and move to next

Windows/Linux Mac

Ctrl+Alt+D
Ctrl+Alt+M
Ctrl+↑/↓
Spacebar
Enter
Ctrl+Option+D
Ctrl+Option+M
Ctrl+↑/↓
Spacebar
Enter

7 MAKE PACKAGES

Build and Reload

Load All (devtools)

Test Package (Desktop)

Test Package (Web)

Check Package

Document Package

Windows/Linux Mac

Ctrl+Shift+B
Ctrl+Shift+L
Ctrl+Shift+T
Ctrl+Alt+F7
Ctrl+Shift+E
Ctrl+Shift+D
Cmd+Shift+B
Cmd+Shift+L
Cmd+Shift+T
Cmd+Opt+F7
Cmd+Shift+E
Cmd+Shift+D

8 DOCUMENTS AND APPS

Preview HTML (Markdown, etc.)

Knit Document (knitr)

Compile Notebook

Compile PDF (TeX and Sweave)

Insert chunk (Sweave and Knitr)

Insert code section

Re-run previous region

Run current document

Run from start to current line

Run the current code section

Run previous Sweave/Rmd code

Run the current chunk

Run the next chunk

Sync Editor & PDF Preview

Previous plot

Next plot

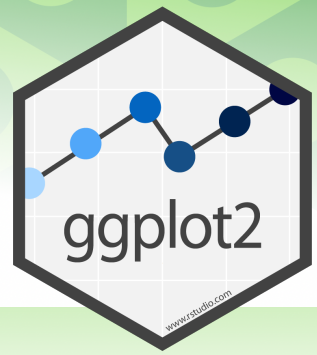
Show Keyboard Shortcuts

Windows/Linux Mac

Ctrl+Shift+K
Ctrl+Shift+K
Ctrl+Shift+K
Ctrl+Shift+K
Ctrl+Alt+I
Ctrl+Shift+R
Ctrl+Shift+P
Ctrl+Alt+R
Ctrl+Alt+B
Ctrl+Alt+T
Ctrl+Alt+P
Ctrl+Alt+C
Ctrl+Alt+N
Ctrl+F8
Ctrl+Alt+F11
Ctrl+Alt+F12
Alt+Shift+K
Cmd+Shift+K
Cmd+Shift+K
Cmd+Shift+K
Cmd+Option+I
Cmd+Shift+R
Cmd+Shift+P
Cmd+Option+R
Cmd+Option+B
Cmd+Option+T
Cmd+Option+P
Cmd+Option+C
Cmd+Option+N
Cmd+F8
Cmd+Option+F11
Cmd+Option+F12
Option+Shift+K



Data Visualization with ggplot2 : : CHEAT SHEET

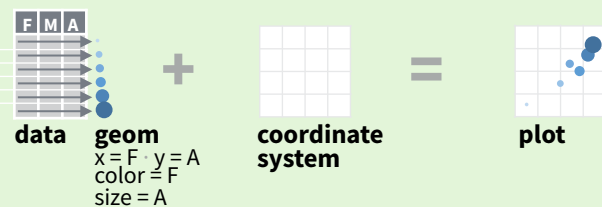


Basics

ggplot2 is based on the **grammar of graphics**, the idea that you can build every graph from the same components: a **data** set, a **coordinate system**, and **geoms**—visual marks that represent data points.



To display values, map variables in the data to visual properties of the geom (**aesthetics**) like **size**, **color**, and **x** and **y** locations.



Complete the template below to build a graph.

```
ggplot (data = <DATA>) +  
<GEOM_FUNCTION> (mapping = aes(<MAPPINGS>),  
  stat = <STAT>, position = <POSITION>) +  
<COORDINATE_FUNCTION> +  
<FACET_FUNCTION> +  
<SCALE_FUNCTION> +  
<THEME_FUNCTION>
```

required

Not required, sensible defaults supplied

ggplot(data = mpg, aes(x = cty, y = hwy)) Begins a plot that you finish by adding layers to. Add one geom function per layer.

qplot(x = cty, y = hwy, data = mpg, geom = "point") Creates a complete plot with given data, geom, and mappings. Supplies many useful defaults.

last_plot() Returns the last plot

ggsave("plot.png", width = 5, height = 5) Saves last plot as 5' x 5' file named "plot.png" in working directory. Matches file type to file extension.

Geoms

Use a geom function to represent data points, use the geom's aesthetic properties to represent variables. Each function returns a layer.

GRAPHICAL PRIMITIVES

```
a <- ggplot(economics, aes(date, unemploy))  
b <- ggplot(seals, aes(x = long, y = lat))
```

- a + geom_blank()**
(Useful for expanding limits)
- b + geom_curve**(aes(yend = lat + 1, xend = long + 1, curvature = z)) - x, xend, y, yend, alpha, angle, color, curvature, linetype, size
- a + geom_path**(lineend = "butt", linejoin = "round", linemitre = 1) - x, y, alpha, color, group, linetype, size
- a + geom_polygon**(aes(group = group)) - x, y, alpha, color, fill, group, linetype, size
- b + geom_rect**(aes(xmin = long, ymin = lat, xmax = long + 1, ymax = lat + 1)) - xmin, xmax, ymin, ymax, alpha, color, fill, linetype, size
- a + geom_ribbon**(aes(ymin = unemploy - 900, ymax = unemploy + 900)) - x, ymin, ymax, alpha, color, fill, group, linetype, size

LINE SEGMENTS

common aesthetics: x, y, alpha, color, linetype, size

- b + geom_abline**(aes(intercept = 0, slope = 1))
- b + geom_hline**(aes(yintercept = lat))
- b + geom_vline**(aes(xintercept = long))
- b + geom_segment**(aes(yend = lat + 1, xend = long + 1))
- b + geom_spoke**(aes(angle = 1:155, radius = 1))

ONE VARIABLE continuous

- ```
c <- ggplot(mpg, aes(hwy)); c2 <- ggplot(mpg)
```
- c + geom\_area**(stat = "bin") - x, y, alpha, color, fill, linetype, size
  - c + geom\_density**(kernel = "gaussian") - x, y, alpha, color, fill, group, linetype, size, weight
  - c + geom\_dotplot**() - x, y, alpha, color, fill
  - c + geom\_freqpoly**() - x, y, alpha, color, group, linetype, size
  - c + geom\_histogram**(binwidth = 5) - x, y, alpha, color, fill, linetype, size, weight
  - c2 + geom\_qq**(aes(sample = hwy)) - x, y, alpha, color, fill, linetype, size, weight

### discrete

- ```
d <- ggplot(mpg, aes(fl))
```
- d + geom_bar**() - x, alpha, color, fill, linetype, size, weight

TWO VARIABLES

continuous x, continuous y

- ```
e <- ggplot(mpg, aes(cty, hwy))
```
- e + geom\_label**(aes(label = cty), nudge\_x = 1, nudge\_y = 1, check\_overlap = TRUE) - x, y, label, alpha, angle, color, family, fontface, hjust, lineheight, size, vjust
  - e + geom\_jitter**(height = 2, width = 2) - x, y, alpha, color, fill, shape, size
  - e + geom\_point**() - x, y, alpha, color, fill, shape, size, stroke
  - e + geom\_quantile**() - x, y, alpha, color, group, linetype, size, weight
  - e + geom\_rug**(sides = "bl") - x, y, alpha, color, linetype, size
  - e + geom\_smooth**(method = lm) - x, y, alpha, color, fill, group, linetype, size, weight
  - e + geom\_text**(aes(label = cty), nudge\_x = 1, nudge\_y = 1, check\_overlap = TRUE) - x, y, label, alpha, angle, color, family, fontface, hjust, lineheight, size, vjust

#### discrete x, continuous y

- ```
f <- ggplot(mpg, aes(class, hwy))
```
- f + geom_col**() - x, y, alpha, color, fill, group, linetype, size
 - f + geom_boxplot**() - x, y, lower, middle, upper, ymax, ymin, alpha, color, fill, group, linetype, shape, size, weight
 - f + geom_dotplot**(binaxis = "y", stackdir = "center") - x, y, alpha, color, fill, group
 - f + geom_violin**(scale = "area") - x, y, alpha, color, fill, group, linetype, size, weight

discrete x, discrete y

- ```
g <- ggplot(diamonds, aes(cut, color))
```
- g + geom\_count**() - x, y, alpha, color, fill, shape, size, stroke

### THREE VARIABLES

```
seals$z <- with(seals, sqrt(delta_long^2 + delta_lat^2))
l <- ggplot(seals, aes(long, lat))
```

- l + geom\_contour**(aes(z = z)) - x, y, z, alpha, colour, group, linetype, size, weight

#### continuous bivariate distribution

- ```
h <- ggplot(diamonds, aes(carat, price))
```
- h + geom_bin2d**(binwidth = c(0.25, 500)) - x, y, alpha, color, fill, linetype, size, weight
 - h + geom_density2d**() - x, y, alpha, colour, group, linetype, size
 - h + geom_hex**() - x, y, alpha, colour, fill, size

continuous function

- ```
i <- ggplot(economics, aes(date, unemploy))
```
- i + geom\_area**() - x, y, alpha, color, fill, linetype, size
  - i + geom\_line**() - x, y, alpha, color, group, linetype, size
  - i + geom\_step**(direction = "hv") - x, y, alpha, color, group, linetype, size

#### visualizing error

- ```
df <- data.frame(grp = c("A", "B"), fit = 4:5, se = 1:2)  
j <- ggplot(df, aes(grp, fit, ymin = fit-se, ymax = fit+se))
```
- j + geom_crossbar**(fatten = 2) - x, y, ymax, ymin, alpha, color, fill, group, linetype, size
 - j + geom_errorbar**() - x, ymax, ymin, alpha, color, group, linetype, size, width (also **geom_errorbarh**())
 - j + geom_linerange**() - x, ymin, ymax, alpha, color, group, linetype, size
 - j + geom_pointrange**() - x, y, ymin, ymax, alpha, color, fill, group, linetype, shape, size

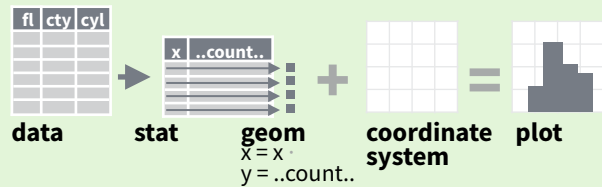
maps

- ```
data <- data.frame(murder = USArrests$Murder,
 state = tolower(rownames(USArrests)))
map <- map_data("state")
k <- ggplot(data, aes(fill = murder))
```
- k + geom\_map**(aes(map\_id = state), map = map) + **expand\_limits**(x = map\$long, y = map\$lat), map\_id, alpha, color, fill, linetype, size

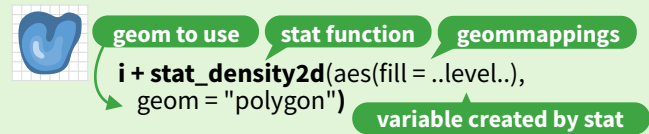
# Stats

An alternative way to build a layer

A stat builds new variables to plot (e.g., count, prop).



Visualize a stat by changing the default stat of a geom function, **geom\_bar(stat="count")** or by using a stat function, **stat\_count(geom="bar")**, which calls a default geom to make a layer (equivalent to a geom function). Use **..name..** syntax to map stat variables to aesthetics.



**c + stat\_bin**(binwidth = 1, origin = 10)  
**x, y** | ..count.., ..ncount.., ..density.., ..ndensity..  
**c + stat\_count**(width = 1) **x, y** | ..count.., ..prop..  
**c + stat\_density**(adjust = 1, kernel = "gaussian")  
**x, y** | ..count.., ..density.., ..scaled..

**e + stat\_bin\_2d**(bins = 30, drop = T)  
**x, y, fill** | ..count.., ..density..  
**e + stat\_bin\_hex**(bins=30) **x, y, fill** | ..count.., ..density..  
**e + stat\_density\_2d**(contour = TRUE, n = 100)  
**x, y, color, size** | ..level..  
**e + stat\_ellipse**(level = 0.95, segments = 51, type = "t")

**l + stat\_contour**(aes(z = z)) **x, y, z, order** | ..level..  
**l + stat\_summary\_hex**(aes(z = z), bins = 30, fun = max)  
**x, y, z, fill** | ..value..  
**l + stat\_summary\_2d**(aes(z = z), bins = 30, fun = mean)  
**x, y, z, fill** | ..value..

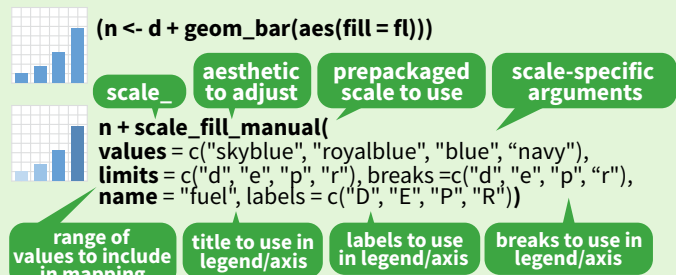
**f + stat\_boxplot**(coef = 1.5) **x, y** | ..lower.., ..middle.., ..upper.., ..width.., ..ymin.., ..ymax..  
**f + stat\_ydensity**(kernel = "gaussian", scale = "area") **x, y** | ..density.., ..scaled.., ..count.., ..n.., ..violinwidth.., ..width..

**e + stat\_ecdf**(n = 40) **x, y** | ..x.., ..y..  
**e + stat\_quantile**(quantiles = c(0.1, 0.9), formula = y ~ log(x), method = "rq") **x, y** | ..quantile..  
**e + stat\_smooth**(method = "lm", formula = y ~ x, se=T, level=0.95) **x, y** | ..se.., ..x.., ..y.., ..ymin.., ..ymax..

**ggplot() + stat\_function**(aes(x = -3:3), n = 99, fun = dnorm, args = list(sd=0.5)) **x** | ..x.., ..y..  
**e + stat\_identity**(na.rm = TRUE)  
**ggplot() + stat\_qq**(aes(sample=1:100), dist = qt, dparam=list(df=5)) **sample, x, y** | ..sample.., ..theoretical..  
**e + stat\_sum**() **x, y, size** | ..n.., ..prop..  
**e + stat\_summary**(fun.data = "mean\_cl\_boot")  
**h + stat\_summary\_bin**(fun.y = "mean", geom = "bar")  
**e + stat\_unique**()

# Scales

**Scales** map data values to the visual values of an aesthetic. To change a mapping, add a new scale.



## GENERAL PURPOSE SCALES

Use with most aesthetics

**scale\_\*\_continuous()** - map cont' values to visual ones  
**scale\_\*\_discrete()** - map discrete values to visual ones  
**scale\_\*\_identity()** - use data values as visual ones  
**scale\_\*\_manual**(values = c()) - map discrete values to manually chosen visual ones  
**scale\_\*\_date**(date\_labels = "%m/%d"), date\_breaks = "2 weeks") - treat data values as dates.  
**scale\_\*\_datetime()** - treat data x values as date times. Use same arguments as scale\_x\_date(). See ?strptime for label formats.

## X & Y LOCATION SCALES

Use with x or y aesthetics (x shown here)

**scale\_x\_log10()** - Plot x on log10 scale  
**scale\_x\_reverse()** - Reverse direction of x axis  
**scale\_x\_sqrt()** - Plot x on square root scale

## COLOR AND FILL SCALES (DISCRETE)

**n <- d + geom\_bar**(aes(fill = fl))  
**n + scale\_fill\_brewer**(palette = "Blues")  
For palette choices:  
RColorBrewer::display.brewer.all()  
**n + scale\_fill\_grey**(start = 0.2, end = 0.8, na.value = "red")

## COLOR AND FILL SCALES (CONTINUOUS)

**o <- c + geom\_dotplot**(aes(fill = ..x..))  
**o + scale\_fill\_distiller**(palette = "Blues")  
**o + scale\_fill\_gradient**(low="red", high="yellow")  
**o + scale\_fill\_gradient2**(low="red", high="blue", mid = "white", midpoint = 25)  
**o + scale\_fill\_gradientn**(colours=topo.colors(6))  
Also: rainbow(), heat.colors(), terrain.colors(), cm.colors(), RColorBrewer::brewer.pal()

## SHAPE AND SIZE SCALES

**p <- e + geom\_point**(aes(shape = fl, size = cyl))  
**p + scale\_shape**() + **scale\_size**()  
**p + scale\_shape\_manual**(values = c(3:7))  
0 1 2 3 4 5 6 7 8 9 10 11 12 13 14 15 16 17 18 19 20 21 22 23 24 25  
**p + scale\_radius**(range = c(1,6))  
**p + scale\_size\_area**(max\_size = 6)

# Coordinate Systems

**r <- d + geom\_bar()**  
**r + coord\_cartesian**(xlim = c(0, 5))  
xlim, ylim  
The default cartesian coordinate system  
**r + coord\_fixed**(ratio = 1/2)  
ratio, xlim, ylim  
Cartesian coordinates with fixed aspect ratio between x and y units  
**r + coord\_flip**()  
xlim, ylim  
Flipped Cartesian coordinates  
**r + coord\_polar**(theta = "x", direction=1)  
theta, start, direction  
Polar coordinates  
**r + coord\_trans**(ytrans = "sqrt")  
xtrans, ytrans, limx, limy  
Transformed cartesian coordinates. Set xtrans and ytrans to the name of a window function.

**π + coord\_quickmap**()  
**π + coord\_map**(projection = "ortho", orientation=c(41, -74, 0))  
projection, orientation, xlim, ylim  
Map projections from the mapproj package (mercator (default), azequalarea, lagrange, etc.)

# Position Adjustments

Position adjustments determine how to arrange geoms that would otherwise occupy the same space.

**s <- ggplot**(mpg, aes(fl, fill = drv))  
**s + geom\_bar**(position = "dodge")  
Arrange elements side by side  
**s + geom\_bar**(position = "fill")  
Stack elements on top of one another, normalize height  
**e + geom\_point**(position = "jitter")  
Add random noise to X and Y position of each element to avoid overplotting  
**e + geom\_label**(position = "nudge")  
Nudge labels away from points  
**s + geom\_bar**(position = "stack")  
Stack elements on top of one another

Each position adjustment can be recast as a function with manual **width** and **height** arguments  
**s + geom\_bar**(position = position\_dodge(width = 1))

# Themes

**r + theme\_bw**()  
White background with grid lines  
**r + theme\_classic**()  
**r + theme\_light**()  
**r + theme\_linedraw**()  
**r + theme\_minimal**()  
Minimal themes  
**r + theme\_dark**()  
dark for contrast  
Empty theme

# Faceting

Facets divide a plot into subplots based on the values of one or more discrete variables.

**t <- ggplot**(mpg, aes(cty, hwy)) + **geom\_point**()  
**t + facet\_grid**(cols = vars(fl))  
facet into columns based on fl  
**t + facet\_grid**(rows = vars(year))  
facet into rows based on year  
**t + facet\_grid**(rows = vars(year), cols = vars(fl))  
facet into both rows and columns  
**t + facet\_wrap**(vars(fl))  
wrap facets into a rectangular layout

Set **scales** to let axis limits vary across facets

**t + facet\_grid**(rows = vars(drv), cols = vars(fl), scales = "free")  
x and y axis limits adjust to individual facets  
"free\_x" - x axis limits adjust  
"free\_y" - y axis limits adjust

Set **labeller** to adjust facet labels

**t + facet\_grid**(cols = vars(fl), labeller = label\_both)  
fl: c fl: d fl: e fl: p fl: r  
**t + facet\_grid**(rows = vars(fl), labeller = label\_bquote(alpha ^ .(fl)))  
 $\alpha^c$   $\alpha^d$   $\alpha^e$   $\alpha^p$   $\alpha^r$

# Labels

**t + labs**( x = "New x axis label", y = "New y axis label", title = "Add a title above the plot", subtitle = "Add a subtitle below title", caption = "Add a caption below plot", <AES> = "New <AES> legend title")  
Use scale functions to update legend labels  
**t + annotate**(geom = "text", x = 8, y = 9, label = "A")  
geom to place manual values for geom's aesthetics

# Legends

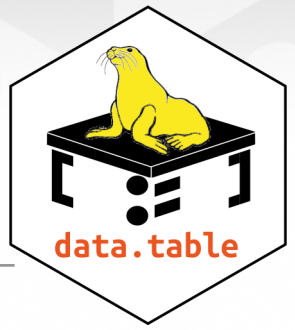
**n + theme**(legend.position = "bottom")  
Place legend at "bottom", "top", "left", or "right"  
**n + guides**(fill = "none")  
Set legend type for each aesthetic: colorbar, legend, or none (no legend)  
**n + scale\_fill\_discrete**(name = "Title", labels = c("A", "B", "C", "D", "E"))  
Set legend title and labels with a scale function.

# Zooming

**Without clipping** (preferred)  
**t + coord\_cartesian**(xlim = c(0, 100), ylim = c(10, 20))  
**With clipping** (removes unseen data points)  
**t + xlim**(0, 100) + **ylim**(10, 20)  
**t + scale\_x\_continuous**(limits = c(0, 100)) + **scale\_y\_continuous**(limits = c(0, 100))



# Data Transformation with data.table :: CHEAT SHEET



## Basics

data.table is an extremely fast and memory efficient package for transforming data in R. It works by converting R's native data frame objects into data.tables with new and enhanced functionality. The basics of working with data.tables are:

**dt[i, j, by]**

Take data.table **dt**,  
subset rows using **i**  
and manipulate columns with **j**,  
grouped according to **by**.

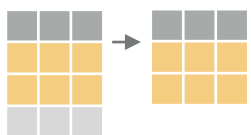
data.tables are also data frames – functions that work with data frames therefore also work with data.tables.

## Create a data.table

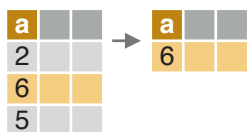
**data.table(a = c(1, 2), b = c("a", "b"))** – create a data.table from scratch. Analogous to data.frame().

**setDT(df)\*** or **as.data.table(df)** – convert a data frame or a list to a data.table.

## Subset rows using i



dt[**1:2**, ] – subset rows based on row numbers.



dt[**a > 5**, ] – subset rows based on values in one or more columns.

### LOGICAL OPERATORS TO USE IN i

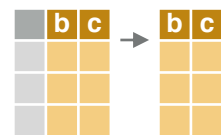
|   |    |          |      |   |           |
|---|----|----------|------|---|-----------|
| < | <= | is.na()  | %in% |   | %like%    |
| > | >= | !is.na() | !    | & | %between% |

## Manipulate columns with j

### EXTRACT



dt[, **c(2)**] – extract columns by number. Prefix column numbers with “-” to drop.



dt[, **.(b, c)**] – extract columns by name.

### SUMMARIZE



dt[, **.(x = sum(a))**] – create a data.table with new columns based on the summarized values of rows.

Summary functions like mean(), median(), min(), max(), etc. can be used to summarize rows.

### COMPUTE COLUMNS\*



dt[, **c := 1 + 2**] – compute a column based on an expression.

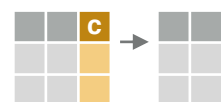


dt[**a == 1**, **c := 1 + 2**] – compute a column based on an expression but only for a subset of rows.



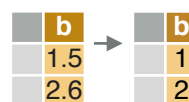
dt[, **`:=`(c = 1, d = 2)**] – compute multiple columns based on separate expressions.

### DELETE COLUMN



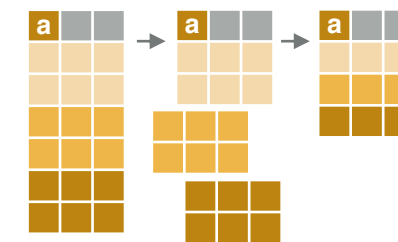
dt[, **c := NULL**] – delete a column.

### CONVERT COLUMN TYPE



dt[, **b := as.integer(b)**] – convert the type of a column using as.integer(), as.numeric(), as.character(), as.Date(), etc..

## Group according to by



dt[, **j, by = .(a)**] – group rows by values in specified columns.

dt[, **j, keyby = .(a)**] – group and simultaneously sort rows by values in specified columns.

### COMMON GROUPED OPERATIONS

dt[, **.(c = sum(b)), by = a**] – summarize rows within groups.

dt[, **c := sum(b), by = a**] – create a new column and compute rows within groups.

dt[, **.SD[1], by = a**] – extract first row of groups.

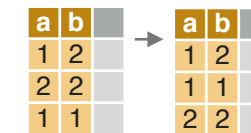
dt[, **.SD[N], by = a**] – extract last row of groups.

## Chaining

dt[...][...] – perform a sequence of data.table operations by chaining multiple “[ ]”.

## Functions for data.tables

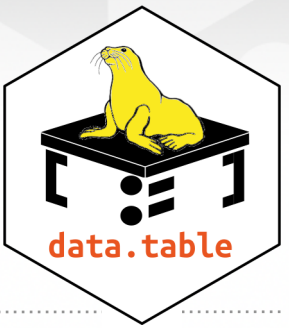
### REORDER



**setorder(dt, a, -b)** – reorder a data.table according to specified columns. Prefix column names with “-” for descending order.

### \* SET FUNCTIONS AND :=

data.table's functions prefixed with “set” and the operator “:=” work without “<-” to alter data without making copies in memory. E.g., the more efficient “setDT(df)” is analogous to “df <- as.data.table(df)”.



## UNIQUE ROWS

| a | b |
|---|---|
| 1 | 2 |
| 2 | 2 |
| 1 | 2 |

**unique**(dt, by = c("a", "b")) – extract unique rows based on columns specified in “by”. Leave out “by” to use all columns.

**uniqueN**(dt, by = c("a", "b")) – count the number of unique rows based on columns specified in “by”.

## RENAME COLUMNS

| a | b |
|---|---|
|   |   |
|   |   |
|   |   |

**setnames**(dt, c("a", "b"), c("x", "y")) – rename columns.

## SET KEYS

**setkey**(dt, a, b) – set keys to enable fast repeated lookup in specified columns using “dt[(value), ]” or for merging without specifying merging columns using “dt\_a[dt\_b]”.

# Combine data.tables

## JOIN

| a | b |
|---|---|
| 1 | c |
| 2 | a |
| 3 | b |

**dt\_a[dt\_b, on = .(b = y)]** – join data.tables on rows with equal values.

| a | b | c |
|---|---|---|
| 1 | c | 7 |
| 2 | a | 5 |
| 3 | b | 6 |

**dt\_a[dt\_b, on = .(b = y, c > z)]** – join data.tables on rows with equal and unequal values.

## ROLLING JOIN

| a | id | date       |
|---|----|------------|
| 1 | A  | 01-01-2010 |
| 2 | A  | 01-01-2012 |
| 3 | A  | 01-01-2014 |
| 1 | B  | 01-01-2010 |
| 2 | B  | 01-01-2012 |

**dt\_a[dt\_b, on = .(id = id, date = date), roll = TRUE]** – join data.tables on matching rows in id columns but only keep the most recent preceding match with the left data.table according to date columns. “roll = -Inf” reverses direction.

## BIND

| a | b |
|---|---|
|   |   |
|   |   |
|   |   |

**rbind**(dt\_a, dt\_b) – combine rows of two data.tables.

| a | b |
|---|---|
|   |   |
|   |   |
|   |   |

**cbind**(dt\_a, dt\_b) – combine columns of two data.tables.

# Reshape a data.table

## RESHAPE TO WIDE FORMAT

| id | y | a | b |
|----|---|---|---|
| A  | x | 1 | 3 |
| A  | z | 2 | 4 |
| B  | x | 1 | 3 |
| B  | z | 2 | 4 |

**dcast**(dt, id ~ y, value.var = c("a", "b"))

Reshape a data.table from long to wide format.

dt A data.table.  
id ~ y Formula with a LHS: ID columns containing IDs for multiple entries. And a RHS: columns with values to spread in column headers.  
value.var Columns containing values to fill into cells.

## RESHAPE TO LONG FORMAT

| id | a | x | a | z | b | x | b | z |
|----|---|---|---|---|---|---|---|---|
| A  | 1 | 2 | 3 | 4 |   |   |   |   |
| B  | 1 | 2 | 3 | 4 |   |   |   |   |

**melt**(dt, id.vars = c("id"), measure.vars = patterns("^a", "^b"), variable.name = "y", value.name = c("a", "b"))

Reshape a data.table from wide to long format.

dt A data.table.  
id.vars ID columns with IDs for multiple entries.  
measure.vars Columns containing values to fill into cells (often in pattern form).  
variable.name, value.name Names of new columns for variables and values derived from old headers.

# Apply function to cols.

## APPLY A FUNCTION TO MULTIPLE COLUMNS

| a | b |
|---|---|
| 1 | 4 |
| 2 | 5 |
| 3 | 6 |

**dt[, lapply(.SD, mean), .SDcols = c("a", "b")]** – apply a function – e.g. mean(), as.character(), which.max() – to columns specified in .SDcols with lapply() and the .SD symbol. Also works with groups.

| a |  | a_m |
|---|--|-----|
| 1 |  | 2   |
| 2 |  | 2   |
| 3 |  | 2   |

**cols <- c("a")**  
**dt[, paste0(cols, "\_m") := lapply(.SD, mean), .SDcols = cols]** – apply a function to specified columns and assign the result with suffixed variable names to the original data.

# Sequential rows

## ROW IDS

| a | b |
|---|---|
| 1 | a |
| 2 | a |
| 3 | b |

**dt[, c := 1:N, by = b]** – within groups, compute a column with sequential row IDs.

## LAG & LEAD

| a | b |
|---|---|
| 1 | a |
| 2 | a |
| 3 | b |
| 4 | b |
| 5 | b |

**dt[, c := shift(a, 1), by = b]** – within groups, duplicate a column with rows lagged by specified amount.

**dt[, c := shift(a, 1, type = "lead"), by = b]** – within groups, duplicate a column with rows leading by specified amount.

# read & write files

## IMPORT

**fread**("file.csv") – read data from a flat file such as .csv or .tsv into R.

**fread**("file.csv", select = c("a", "b")) – read specified columns from a flat file into R.

## EXPORT

**fwrite**(dt, "file.csv") – write data to a flat file from R.

# Machine Learning Modelling in R : : CHEAT SHEET

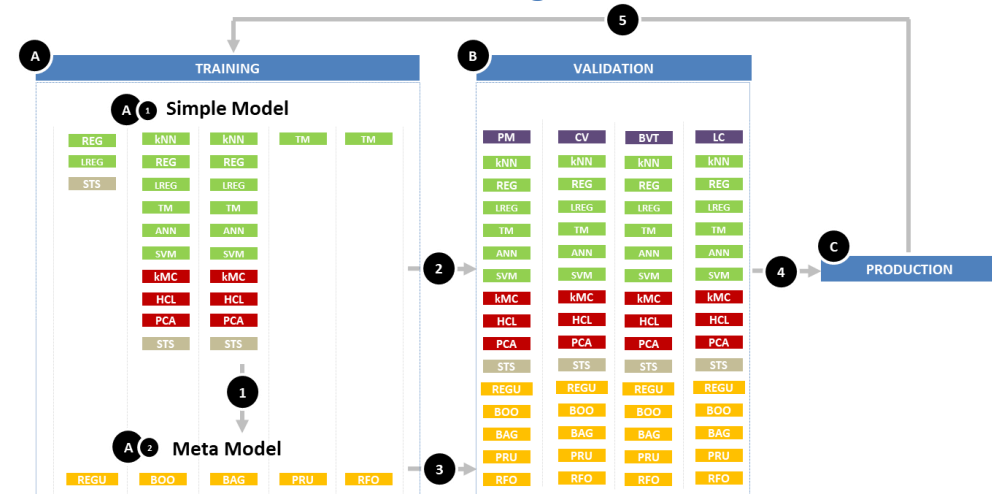
## Supervised & Unsupervised Learning

|                       | ALGORITHM                                  | DESCRIPTION                                                                                                                                                                                                                                                      | R PACKAGE::FUNCTION                                                                | SAMPLE CODE                                                                                                                                                                                                                                                                                        |
|-----------------------|--------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| SUPERVISED LEARNING   | <b>NBC</b><br>Naïve Bayes classifier       | A classification technique based on Bayes' Theorem with an assumption of independence among predictors. In simple terms, a Naïve Bayes classifier assumes that the presence of a particular feature in a class is unrelated to the presence of any other feature | e1071::naiveBayes                                                                  | naiveBayes(class ~., data = x)                                                                                                                                                                                                                                                                     |
|                       | <b>KNN</b><br>k-Nearest Neighbours         | A non-parametric method used for classification and regression. In both cases, the input consists of the k closest training examples in the feature space. The output depends on whether k-NN is used for classification or regression                           | class::knn                                                                         | knn(train, test, cl, k = 1, l = 0, prob = FALSE, use.all = TRUE)                                                                                                                                                                                                                                   |
|                       | <b>REG</b><br>Linear Regression            | Model the linear relationship between a scalar dependant variable Y and one or more explanatory variables (or independent variables) denoted X                                                                                                                   | stats::lm                                                                          | lm(dist ~ speed, data=cars)                                                                                                                                                                                                                                                                        |
|                       | <b>LOG</b><br>Logistic Regression          | Used to predict a binary outcome (1 / 0, Yes / No, True / False) given a set of independent variables.                                                                                                                                                           | stats::glm                                                                         | glm(Y ~., family = binomial (link = 'logit'), data = X)                                                                                                                                                                                                                                            |
|                       | <b>TM</b><br>Tree-Based Models             | The idea is to consecutively divide (branch) the training dataset based on the input features until an assignment criterion with respect to the target variable into a "data bucket" (leaf) is reached                                                           | rpart::rpart                                                                       | rpart(Kyphosis ~ Age + Number + Start, data = kyphosis)                                                                                                                                                                                                                                            |
|                       | <b>ANN</b><br>Artificial Neural Network    | Neural networks are built from units called perceptrons. Perceptrons have one or more inputs, an activation function and an output. An ANN model is built up by combining perceptrons in structured layers.                                                      | neuralnet::neuralnet                                                               | neuralnet(f, data=train, hidden=c(5,3), linear.output=T)                                                                                                                                                                                                                                           |
| UNSUPERVISED LEARNING | <b>SVM</b><br>Support Vector Machine       | A data classification method that separates data using hyperplanes                                                                                                                                                                                               | e1071::svm                                                                         | svm(formula, data = NULL, ..., subset, na.action = na.omit, scale = TRUE)                                                                                                                                                                                                                          |
|                       | <b>PCA</b><br>Principal Component Analysis | A procedure that uses an orthogonal transformation to convert a set of observations of possibly correlated variables into a set of values of linearly uncorrelated variables called principal components.                                                        | stats::prcomp<br>stats::princomp<br>FactoMineR::PCA<br>ade4::dudi.pca<br>amap::acp | stats::prcomp(formula, data = NULL, subset, na.action, ...)<br>stats::princomp(formula, data = NULL, subset, na.action, ...)<br>FactoMineR::PCA(decatlon, quanti.supp = 11:12, quali.supp=13)<br>ade4::dudi.pca(deug\$tab, center = deug\$cent, scale = FALSE, scan = FALSE)<br>amap::acpl(ubisch) |
|                       | <b>KMC</b><br>k-Mean Clustering            | Aims at partitioning n observations into k clusters in which each observation belongs to the cluster with the nearest mean                                                                                                                                       | stats::kmeans                                                                      | kmeans(x, centers, iter.max = 10, nstart = 1, algorithm = c("Hartigan-Wong", "Lloyd", "Forgy", "MacQueen"), trace=FALSE)                                                                                                                                                                           |
|                       | <b>HCL</b><br>Hierarchical Clustering      | An approach which builds a hierarchy from the bottom-up, and doesn't require the number of clusters to be specified beforehand.                                                                                                                                  | stats::hclust                                                                      | hclust(d, method = "complete", members = NULL)                                                                                                                                                                                                                                                     |

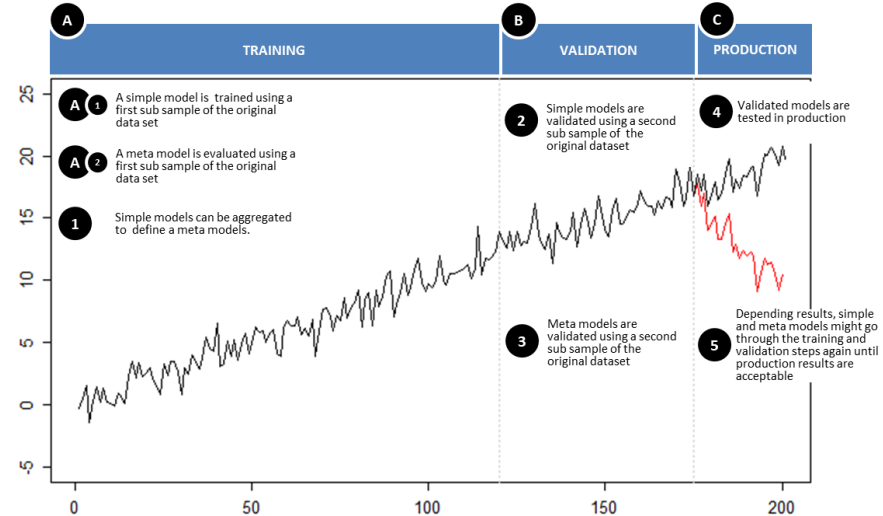
## Meta-Algorithm, Time Series & Model Validation

|                  | ALGORITHM                                                                                                           | DESCRIPTION                                                                                                                                                                                                                                                           | R PACKAGE::FUNCTION                                                                                  | SAMPLE CODE                                                                                                                                                                                                                                      |
|------------------|---------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| META ALGORITHM   | <b>REGU</b><br>Regularisation<br>L1 (Lasso)<br>L2 (Ridge)                                                           | Regularization adds a penalty on the different parameters of a model to reduce the freedom of the model. Hence, the model will be less likely to fit the noise of the training data and will improve the generalization abilities of the model                        | glmnet::glmnet                                                                                       | L1 : glmnet(myMatrixA, myMatrixB, family = "gaussian", alpha = 1)<br>L2 : glmnet(myMatrixA, myMatrixB, family = "gaussian", alpha = 0)                                                                                                           |
|                  | <b>BOO</b><br>Boosting                                                                                              | A process of iteratively refining, e.g. by reweighting, of estimated regression and classification functions (though it has primarily been applied to the latter), in order to improve predictive ability.                                                            | Parametric model - mboost::mboost                                                                    | glmboost(Yen ~., data = cur1[trnids,])                                                                                                                                                                                                           |
|                  | <b>BAG</b><br>Bagging                                                                                               | Bagging is a way to increase the power of a predictive statistical model by taking multiple random samples (with replacement) of the training data set, and using each of them to construct a separate model and separate predictions for the original test set       | All models: foreach<br>Tree models: ipred::bagging                                                   | foreach : d <- data.frame(x=1:10, y=rnorm(10))<br>s <- foreach(d=iter(d, by='row'), combine='bind')<br>%doapar(d)<br>identical(s, d)<br>ipred : bagging(formula, data, subset, na.action=na.rpart, dots)                                         |
|                  | <b>PRU</b><br>Pruning                                                                                               | Pruning is a technique that reduces the size of decision tree by removing sections of the tree that provide little power to classify instances. Pruning reduces the complexity of the final classifier and hence improves predictive accuracy by reducing overfitting | rpart::prune                                                                                         | prune(x, cp = 0.1)                                                                                                                                                                                                                               |
|                  | <b>RFO</b><br>Random Forrest                                                                                        | An ensemble learning method for classification, regression and other tasks, that operate by constructing a multitude of decision trees at training time and outputting the class that is the mode of the classes (classification) or mean prediction (regression)     | randomForest::randomForest                                                                           | randomForest(X ~., data = Y, subset = mySub)                                                                                                                                                                                                     |
| TIME SERIES      | <b>STS</b><br>Lead-lag analysis, Auto-correlation, Spectral analysis, Time series clustering, Seasonality, Trend... | Random sampling of observations for training and testing a model can be an issue when faced with a times dimension. Random sampling may either destroy serial correlation properties in the data which we would like to exploit                                       | stats<br>forecast<br>spectral<br>TTR<br>....                                                         | Auto-correlation: acf(x, lag.max = NULL, type = c("correlation", "covariance", "partial"))<br>Spectral Analysis: specgram(myTs, spans = NULL)<br>Seasonal Decomposition of Time Series: stl(x, s.window = 7, l.window = 50, l.jump = 1)<br>..... |
|                  | <b>PM</b><br>Performance metrics                                                                                    | Depends on the problem:<br>• <b>Regression:</b> squared errors, outliers, error rate...<br>• <b>Classification:</b> Accuracy, precision, recall, F-score...                                                                                                           | Regression: stats::outlierTest, stats::qqPlot...<br>Classification-ROC: Tree: caret::confusionMatrix | Regression: fit <- lm(Y~X, data=myData)<br>outlierTest(fit)<br>qqPlot(fit, main="QQ Plot")                                                                                                                                                       |
| MODEL VALIDATION | <b>BVT</b><br>Bias-Variance Tradeoff                                                                                | Simple models with few parameters are easier to compute but may lead to poorer fits ( <b>high bias</b> ).<br>Complex models may provide more accurate fits but may over-fit the data ( <b>high variance</b> )                                                         | Tailored to the analysis                                                                             | Tailored to the analysis                                                                                                                                                                                                                         |
|                  | <b>CV</b><br>Cross validation                                                                                       | Cross validation compares the test performances of different model realisations with different sets or values of parameters                                                                                                                                           | caret::createDataPartition<br>caret::createFolds                                                     | createDataPartition(classes, p = 0.8, list = FALSE)                                                                                                                                                                                              |
|                  | <b>LC</b><br>Learning Curves                                                                                        | Learning curves plot a model's training and test errors, or the chosen performance metric, depending on the training set size                                                                                                                                         | caret::learning_curve_dat                                                                            | learning_curve_dat(dat, outcome = NULL, proportion = 1:10/10, test_prop = 0, verbose = TRUE, ...)                                                                                                                                                |

## Standard Modelling Workflow



## Time Series View



# Machine Learning with R



mlr

## Introduction

**mlr** offers a unified interface for the basic building blocks of machine learning: tasks, learners, hyperparameters, etc.

**Tasks** contain a description of a task (classification, regression, clustering, etc.) and a data set.

**Learners** specify a machine learning algorithm (GLM, SVM, xgboost, etc.) and its parameters.

**Hyperparameters** are learner settings that can be specified directly or tuned. A **parameter set** lists the possible hyperparameters for a given learner.

**Wrapped Models** are learners that have been trained on a task and can be used to make predictions.

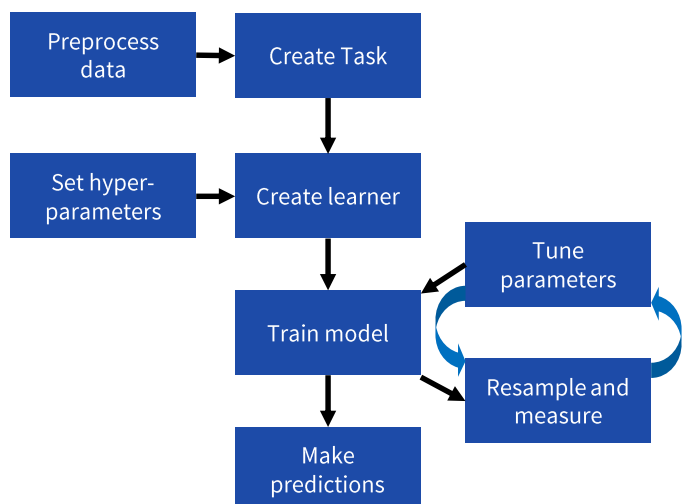
**Predictions** are the results of applying a model to either new data or the original training data.

**Measures** control how learner performance is evaluated, e.g. RMSE, LogLoss, AUC, etc.

**Resampling** estimates generalization performance by separating training data from test data. Common strategies include holdout and cross-validation.

Links: [Tutorial](#) | [CRAN](#) | [Github](#)

## mlr workflow



## Setup

### Preprocessing data

`createDummyFeatures(obj=, target=, method=, cols=)`  
Creates (0,1) flags for each non-numeric variable excluding `target`. Can be applied to entire dataset or only specific `cols`

`normalizeFeatures(obj=, target=, method=, cols=, range=, on.constant=)`  
Normalizes numerical features according to specified `method`:

- "center" (subtract mean)
- "scale" (divide by std. deviation)
- "standardize" (center and scale)
- "range" (linear scale to given range, default `range=c(0,1)`)

`mergeSmallFactorLevels(task=, cols=, min.perc=)`  
Combine infrequent factor levels into a single merged level

`summarizeColumns(obj=)` where `obj` is a data.frame or task.  
Provides type, NA, and distributional data about each column

See also `capLargeValues` `dropFeatures` `removeConstantFeatures` `summarizeLevels`

### Creating a task



`makeClassifTask(data=, target=)`  
Classification of a target variable, with optional positive class `positive`

0 63 100



`makeRegrTask(data=, target=)`  
Regression on a target variable



`makeMultilabelTask(data=, target=)`  
Classification where the target can belong to more than one class per observation



`makeClusterTask(data=)`  
Unsupervised clustering on a data set



`makeSurvTask(data=, target=c("time", "event"))`  
Survival analysis with a survival time column and an event column

`makeCostSensTask(data=, costs=)`  
Cost-sensitive classification where each observation-cost pair has a specified cost

Other arguments that can be passed to a `task`:

- `weights`=Weighting vector to apply to observations
- `blocking`=Factor vector where each level indicates a block of observations that will not be split up in resampling

### Making a learner

`makeLearner(cl=predict.type=, ..., par.vals=)`  
Choose an algorithm class to perform the task and determine what that algorithm will predict

- `cl`=name of algorithm, e.g. "classif.xgboost" "regr.randomForest" "cluster.kmeans"
  - `predict.type`="response" returns a prediction type that matches the source data; "prob" returns a predicted probability for classification problems only; "se" returns the a standard error of the prediction for regression problems only. Only certain learners can return "prob" and "se"
  - `par.vals`=takes a list of hyperparameters and passes them to the learner; parameters can also be passed directly (...)
- You can make multiple learners at once with `makeLearners()`

mlr has integrated over 170 different learning algorithms

- Full list: `View(listLearners())` shows all learners
- Available learners for a task: `View(listLearners(task=))`
- Filtered list: `View(listLearners("classif", properties=c("prob", "factors")))` shows all classification learners "classif" which can predict probabilities "prob" and handle factor inputs "factors"
- See also `getLearnerProperties()`

## Training & Testing

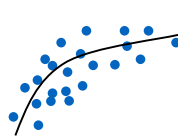
### Setting hyperparameters



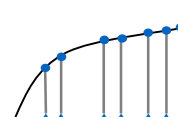
`setHyperPars(learner=, ...)`  
Set the hyperparameters (settings) for each learner, if you don't want to use the defaults. You can also specify hyperparameters in the `makeLearner()` call

`getParamSet(learner=)`  
Show the possible universe of parameters for your learner; can take a learner directly, or a text string such as "classif.qda"

### Train a model and predict



`train(learner=, task=)`  
Train a model (`WrappedModel`) by applying a learner to a task. By default, the model will train on all observations. The underlying model can be extracted with `getLearnerModel()`



`predict(object=, task=, newdata=)`  
Use a trained model to make predictions on a task or dataset. The resulting `pred` object can be viewed with `View(pred)` or accessed by `as.data.frame(pred)`

### Measuring performance

`performance(pred=, measures=)`  
Calculate performance of predictions according to one or more of several measures (use `listMeasures()` for full list):

- `classif` `acc` `auc` `bac` `ber` `brier` `[.scaled]` `f1` `fdr` `fn` `fnr` `fp` `fpr` `gmean` `multiclass` `[.au1u .aunp .aunu .brier]` `npv` `ppv` `qsr` `ssr` `tn` `tnr` `tp` `tpr` `wkappa`
- `regr` `arsq` `expvar` `kendalltau` `mae` `mape` `medae` `medse` `mse` `msle` `rae` `rmse` `rmsle` `rrse` `rsq` `sae` `spearmanrho` `sse`
- `cluster` `db` `dunn` `G1` `G2` `silhouette`
- `multilabel` `multilabel` `[.f1 .subset01 .tpr .ppv .acc .hamloss]`
- `costsens` `mcp` `meancosts`
- `surv` `cindex`
- other `featperc` `timeboth` `timepredict` `timetrain`

For detailed performance data on classification tasks, use:

- `calculateConfusionMatrix(pred=)`
- `calculateROCMeasures(pred=)`

### Resampling a learner

`makeResampleDesc(method=, ..., stratify=)`  
`method` must be one of the following:

- "CV" (cross-validation, for number of folds use `iters=`)
  - "LOO" (leave-one-out cross-validation, for folds use `iters=`)
  - "RepCV" (repeated cross-validation, for number of repetitions use `reps=`, for folds use `folds=`)
  - "Subsample" (aka Monte-Carlo cross-validation, for iterations use `iters=`, for train % use `split=`)
  - "Bootstrap" (out-of-bag bootstrap, uses `iters=`)
  - "Holdout" (for train % use `split=`)
- `stratify` keeps target proportions consistent across samples.

`makeResampleInstance(desc=, task=)` can reduce noise by ensuring the resampling is done identically every time.

`resample(learner=, task=, resampling=, measures=)`  
Train and test model according to specified resampling strategy.

mlr includes several pre-specified resample descriptions: `cv2` (2-fold cross-validation), `cv3`, `cv5`, `cv10`, `hout` (holdout with split 2/3 for training, 1/3 for testing).

Convenience functions also exist to `resample()` with a specific strategy: `crossval()`, `repcv()`, `holdout()`, `subsample()`, `bootstrapOoB()`, `bootstrapB632()`, `bootstrapB632plus()`

## Refining Performance

### Tuning hyperparameters

Set search space using `makeParamSet(make<type>Param())`

- `makeNumericParam(id=, lower=, upper=, trafo=)`
- `makeIntegerParam(id=, lower=, upper=, trafo=)`
- `makeIntegerVectorParam(id=, len=, lower=, upper=, trafo=)`
- `makeDiscreteParam(id=, values=c(...))` (can also be used to test discrete values of numeric or integer parameters)

`trafo` transforms the parameter output using a specified function, e.g. `lower=-2, upper=2, trafo=function(x) 10^x` would test values between 0.01 and 100, scaled exponentially

Other acceptable parameter types include `Logical` `LogicalVector` `CharacterVector` `DiscreteVector`

Set a search algorithm with `makeTuneControl<type>()`

- `Grid(resolution=10L)` Grid of all possible points
- `Random(maxit=100)` Randomly sample search space
- `MBO(budget=)` Use Bayesian model-based optimization
- `Trace(n.instances=)` Iterated racing process
- Other types: `CMAES`, `Design`, `GenSA`

Tune using `tuneParams(learner=, task=, resampling=, measures=, par.set=, control=)`

## Quickstart

### Prepare data for training and testing

```
library(mlbench)
data(Soybean)
soy = createDummyFeatures(Soybean, target="Class")
tsk = makeClassifTask(data=soy, target="Class")
ho = makeResampleInstance("Holdout", tsk)
tsk.train = subsetTask(tsk, ho$train.inds[[1]])
tsk.test = subsetTask(tsk, ho$test.inds[[1]])
```

Convert the factor inputs in the Soybean dataset into (0,1) dummy features which can be used by the XGboost algorithm. Create a task to predict the "Class" column. Create a train set with 2/3 of data and a test set with the remaining 1/3 (default).

### Create learner and evaluate performance

```
lrn = makeLearner("classif.xgboost", nrounds=10)
cv = makeResampleDesc("CV", iters=5)
res = resample(lrn, tsk.train, cv, acc, ps, tc)
```

Create an XGboost learner which will build 10 trees. Then test performance using 5-fold cross-validation. Accuracy should be between 0.90-0.92.

### Tune hyperparameters and retrain model

```
ps = makeParamSet(makeNumericParam("eta", 0, 1),
 makeNumericParam("lambda", 0, 200),
 makeIntegerParam("max_depth", 1, 20))
tc = makeTuneControlMBO(budget=100)
tr = tuneParams(lrn, tsk.train, cv5, acc, ps, tc)
lrn = setHyperPars(lrn, par.vals=tr$x)
```

Tune hyperparameters `eta`, `lambda`, and `max_depth` by defining a search space and using Model Based Optimization (MBO) to control the search. Then perform 100 rounds of 5-fold cross-validation, improving accuracy to ~0.93. Update the XGboost learner with the tuned hyperparameters.

```
mdl = train(lrn, tsk.train)
prd = predict(mdl, tsk.test)
calculateConfusionMatrix(prd)
mdl = train(lrn, tsk)
```

Train the model on the train set and make predictions on the test set. Show performance as a confusion matrix. Finally, re-train model on the full set to use on new data. You are now ready to go out into the real world and make 93% accurate predictions!

Legend for functions (not all parameters shown):  
`function(required_parameters=, optional_parameters=)`

## Configuration

- mlr's default settings can be changed using `configureMlr()`:
- `show.info` Whether to show verbose output by default when training, tuning, resampling, etc. (`TRUE`)
  - `on.learner.error` How to handle a learner error. "stop" halts execution, "warn" returns NAs and displays a warning, "quiet" returns NAs with no warning ("stop")
  - `on.learner.warning` How to handle a learner warning. "warn" displays a warning, "quiet" suppresses it ("warn")
  - `on.par.without.desc` How to handle a parameter with no description. "stop", "warn", "quiet" ("stop")
  - `on.par.out.of.bounds` How to handle a parameter with an out-of-bounds value. "stop", "warn", "quiet" ("stop")
  - `on.measure.not.applicable` How to handle a measure not applicable to a learner. "stop", "warn", "quiet" ("stop")
  - `show.learner.output` Whether to show learner output to the console during training (`TRUE`)
  - `on.error.dump` Whether to create an error dump for crashed learners if `on.learner.error` is not set to "stop" (`TRUE`)

Use `getMlrOptions()` to see current settings

## Parallelization

mlr works with the `parallelMap` package to take advantage of multicore and cluster computing for faster operations. mlr automatically detects which operations are able to run in parallel.

To begin parallel operation use:

- `parallelStart(mode=, cpus=, level=)`
- `mode` determines how the parallelization is performed:
    - "local" no parallelization applied, simply uses `mapply`
    - "multicore" multicore execution on a single machine, uses `parallel::mclapply`. Not available in Windows.
    - "socket" multicore execution in socket mode
    - "mpi" Snow MPI cluster on one or multiple machines using `parallel::makeCluster` and `parallel::clusterMap`
    - "BatchJobs" Batch queuing HPC clusters using `BatchJobs::batchMap`
  - `cpus` determines how many logical cores will be used
  - `level` controls parallelization: "mlr.benchmark" "mlr.resample" "mlr.selectFeatures" "mlr.tuneParams" "mlr.ensemble"

To end parallelization, use `parallelStop()`

## Imputation

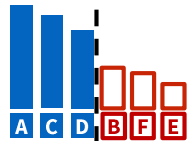
- `impute(obj=, target=, cols=, dummy.cols=, dummy.type=)`  
Applies specified logic to data frame or task containing NAs and returns an imputation description which can be used on new data
- `obj`=data frame or task on which to perform imputation
  - `target`=specify target variable which will not be imputed
  - `cols`=column names and logic for imputation\*
  - `dummy.cols`=column names to create a NA (T/F) column\*
  - `dummy.type`=set to "numeric" to use (0,1) instead of (T/F)
- \*Can also use `classes` and `dummy.classes` in place of `cols`

Imputation logic is passed to `cols` or `classes` via a list, e.g.:  
`cols=list(V1=imputeMean())` where `V1` is the column to which to apply the imputation, and `imputeMean()` is the imputation method. Available imputation methods include:  
`imputeConst(const=)` `imputeMedian()` `imputeMode()` `imputeMin(multiplier=)` `imputeMax(multiplier=)` `imputeNormal(mean=, sd=)` `imputeHist(breaks=, use.mids=)` `imputeLearner(learner=, features=)`  
`impute` returns a list containing the imputed dataset or task as well as an imputation description that can be used to reapply the same imputation to new data using `reimpute`

`reimpute(obj=, desc=)` Imputes missing values on a task or dataset (`obj`) using a description (`desc`) created by `impute`

## Feature Extraction

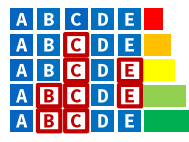
### Feature filtering



`filterFeatures(task=, method=, perc=, abs=, threshold=)`  
Uses a learner-agnostic feature evaluation method to rank feature importance, then includes only features in the top n percent (`perc=`), top n (`abs=`), or which meet a set performance threshold (`threshold=`).

Outputs a task with features that failed the test omitted. `method` defaults to "randomForestSRC.rfsrc", but can be set to: "anova.test" "carscore" "cforest.importance" "chi.squared" "gain.ratio" "information.gain" "kruskal.test" "linear.correlation" "mrmr" "oneR" "permutation.importance" "randomForest.importance" "randomForestSRC.rfsrc" "randomForestSRC.var.select" "rank.correlation" "relief" "symmetrical.uncertainty" "univariate.model.score" "variance"

### Feature selection



`selectFeatures(learner=, task=, resampling=, measures=, control=)`  
Uses a feature selection algorithm (`control`) to resample and build a model repeatedly using different feature sets each time in order to find the best set.

Available controls include:

- `makeFeatSelControlExhaustive(max.features=)` Try every combination of features up to optional `max.features`
- `makeFeatSelControlRandom(maxit=, prob=, max.features=)` Randomly sample features with probability `prob` (default 0.5) until `maxit` (default 100) iterations; return the best one found
- `makeFeatSelControlSequential(method=, maxit=, max.features=, alpha=, beta=)` Perform an iterative search using a `method` from the following: "sfs" forward search, "sbs" backward search, "sffs" floating forward search, "sfbs" floating backward search. `alpha` indicates minimum improvement required to add a feature; `beta` indicates minimum required to remove a feature
- `makeFeatSelControlGA(maxit=, max.features=, mu=, lambda=, crossover.rate=, mutation.rate=)` Genetic algorithm trains on random feature vectors, then uses crossover on the best performers to produce 'offspring', repeated over generations. `mu` is size of parent population, `lambda` is size of children population, `crossover.rate` is probability of choosing a bit from first parent, `mutation.rate` is probability of flipping a bit (on or off)

`selectFeatures` returns a `FeatSelResult` object which contains optimal features and an optimization path. To apply feature selection result (`fsr`) to your task (`tsk`), use:  
`tsk = subsetTask(tsk, features=fsr$x)`

## Benchmarking

`benchmark(learners=, tasks=, resamplings=, measures=)`  
Allows easy comparison of multiple learners on a single task, a single learner on multiple tasks, or multiple learners on multiple tasks. Returns a benchmark result object.

Benchmark results can be accessed with a variety of functions beginning with `getBMR<object>`: `AggrPerformance` `FeatSelResults` `FilteredFeatures` `LearnerIds` `LearnerShortNames` `Learners` `MeasureIds` `Measures` `Models` `Performances` `Predictions` `TaskDescs` `TaskIds` `TuneResults`

mlr contains several toy tasks which are useful for benchmarking: `agri.task` `bc.task` `bh.task` `costiris.task` `iris.task` `lung.task` `mtcars.task` `pid.task` `sonar.task` `wdbc.task` `yeast.task`

## Visualization

### Performance

`generateThreshVsPerfData(obj=, measures=)` Measure performance at different probability cutoffs to determine optimal decision threshold for binary classification problems

- `plotThreshVsPerf(obj)` Plot visual representation of threshold curve(s) from `ThreshVsPerfData`
- `plotROCCurves(obj)` Plot receiver operating characteristic (ROC) curve from `ThreshVsPerfData`. Must set `measures=list(fpr, tpr)`

### Residuals

- `plotResiduals(obj=)` Plots residuals for `Prediction` or `BenchmarkResult`

### Learning curve

`generateLearningCurveData(learners=, task=, resampling=, percs=, measures=)` Measure performance of learner(s) trained on different percentages of task data

- `plotLearningCurve(obj=)` Plot curve showing learner performance vs. proportion of data used, uses `LearningCurveData`

### Feature importance

`generateFilterValuesData(task=, method=)` Get feature importance rankings using specified filter method

- `plotFilterValues(obj=)` Plot bar chart of feature importance based on filter method using `FilterValuesData`

### Hyperparameter tuning

`generateHyperParsEffectData(tune.result=)` Get the impact of different hyperparameter settings on model performance

- `plotHyperParsEffect(hyperpars.effect.t.data=, x=, y=, z=)` Create a plot showing hyperparameter impact on performance using `HyperParsEffectData`

See also:

- `plotOptPath(op=)` Display details of optimization process. Takes `<obj>$opt.path`, where `<obj>` is an object of class `tuneResult` or `FeatSelResult`
- `plotTuneMultiCritResult(res=)` Show pareto front for results of tuning to multiple performance measures

### Partial dependence

`generatePartialDependenceData(obj=, input=)` Get partial dependence of model (`obj`) prediction over each feature of data (`input`)

- `plotPartialDependence(obj=)` Plots partial dependence of model using `PartialDependenceData`

### Benchmarking

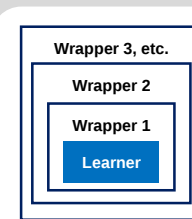
- `plotBMRBoxplots(bmr=)` Distribution of performances
- `plotBMRSummary(bmr=)` Scatterplot of avg. performances
- `plotBMRRankAsBarChart(bmr=)` Rank learners in bar plot

### Other

- `generateCritDifferencesData(bmr=, measure=, p.value=, test=)` Perform critical-differences test using either the Bonferroni-Dunn ("bd") or "Nemenyi" test
- `plotCritDifferences(obj=)`

- `generateCalibrationData(obj=)` Evaluate calibration of probability predictions vs. true incidence
- `plotCalibration(obj=)`

## Wrappers



**Wrappers** fuse a learner with additional functionality. mlr treats a learner with wrappers as a single learner, and hyperparameters of wrappers can be tuned jointly with underlying model parameters. Models trained with wrappers will apply them to new data.

### Preprocessing and imputation

`makeDummyFeaturesWrapper(learner=)`  
`makeImputeWrapper(learner=, classes=, cols=)`  
`makePreprocWrapper(learner=, train=, predict=)`  
`makePreprocWrapperCaret(learner=, ...)`  
`makeRemoveConstantFeaturesWrapper(learner=)`

### Class imbalance

`makeOverBaggingWrapper(learner=)`  
`makeSMOTWrapper(learner=)`  
`makeUndersampleWrapper(learner=)`  
`makeWeightedClassesWrapper(learner=)`

### Cost-sensitive learning

`makeCostSensClassifWrapper(learner=)`  
`makeCostSensRegrWrapper(learner=)`  
`makeCostSensWeightedPairsWrapper(learner=)`

### Multilabel classification

`makeMultilabelBinaryRelevanceWrapper(learner=)`  
`makeMultilabelClassifierChainsWrapper(learner=)`  
`makeMultilabelDBRWrapper(learner=)`  
`makeMultilabelNestedStackingWrapper(learner=)`  
`makeMultilabelStackingWrapper(learner=)`

### Other

`makeBaggingWrapper(learner=)`  
`makeConstantClassWrapper(learner=)`  
`makeDownsampleWrapper(learner=, dw.perc=)`  
`makeFeatSelWrapper(learner=, resampling=, control=)`  
`makeFilterWrapper(learner=, fw.perc=, fw.abs=, fw.threshold=)`  
`makeMultiClassWrapper(learner=)`  
`makeTuneWrapper(learner=, resampling=, par.set=, control=)`

## Nested Resampling

mlr supports **nested resampling** for complex operations such as tuning and feature selection through wrappers. In order to get a good estimate of generalization performance and avoid data leakage, both an outer (for tuning/feature selection) and an inner (for the base model) resampling process are advised.

- Outer resampling can be specified in `resample` or `benchmark`
- Inner resampling can be specified in `makeTuneWrapper`, `makeFeatSelWrapper`, etc.

## Ensembles

`makeStackedLearner(base.learners=, super.learner=, method=)` Combines multiple learners to create an ensemble

- `base.learners`=learners to use for initial predictions
- `super.learner`=learner to use for final prediction
- `method`=how to combine base learner predictions:
  - "average" simple average of all base learners
  - "stack.nocv", "stack.cv" train super learner on results of base learners, with or without cross-validation
  - "hill.climb" search for optimal weighted average
  - "compress" with a neural network for faster performance

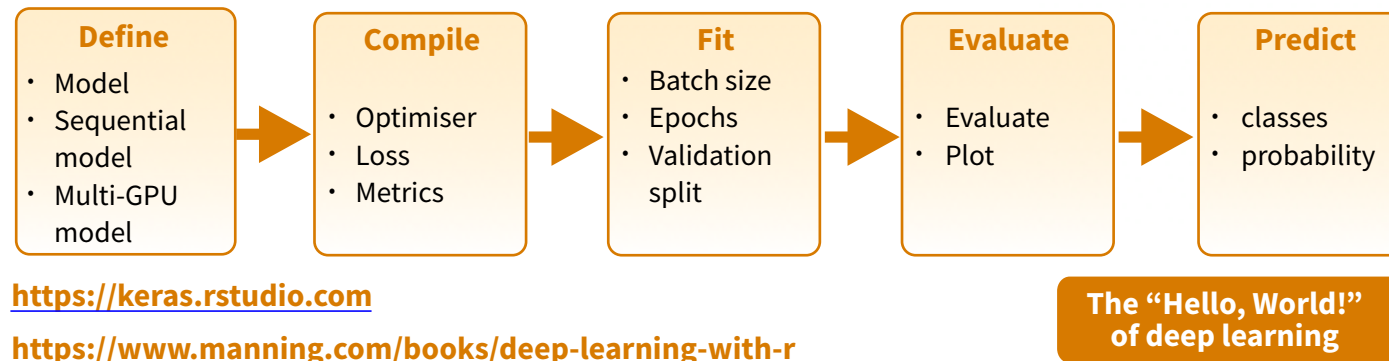
# Deep Learning with Keras :: CHEAT SHEET



## Intro

**Keras** is a high-level neural networks API developed with a focus on enabling fast experimentation. It supports multiple back-ends, including TensorFlow, CNTK and Theano.

TensorFlow is a lower level mathematical library for building deep neural network architectures. The **keras** R package makes it easy to use Keras and TensorFlow in R.



## INSTALLATION

The **keras** R package uses the Python **keras** library. You can install all the prerequisites directly from R.

[https://keras.rstudio.com/reference/install\\_keras.html](https://keras.rstudio.com/reference/install_keras.html)

```
library(keras)
install_keras()
```

See ?install\_keras for GPU instructions

This installs the required libraries in an Anaconda environment or virtual environment 'r-tensorflow'.

## Working with keras models

### DEFINE A MODEL

**keras\_model()** Keras Model

**keras\_model\_sequential()** Keras Model composed of a linear stack of layers

**multi\_gpu\_model()** Replicates a model on different GPUs

### COMPILE A MODEL

**compile(object, optimizer, loss, metrics = NULL)**  
Configure a Keras model for training

### FIT A MODEL

**fit(object, x = NULL, y = NULL, batch\_size = NULL, epochs = 10, verbose = 1, callbacks = NULL, ...)**  
Train a Keras model for a fixed number of epochs (iterations)

**fit\_generator()** Fits the model on data yielded batch-by-batch by a generator

**train\_on\_batch()** **test\_on\_batch()** Single gradient update or model evaluation over one batch of samples

### EVALUATE A MODEL

**evaluate(object, x = NULL, y = NULL, batch\_size = NULL)** Evaluate a Keras model

**evaluate\_generator()** Evaluates the model on a data generator

### PREDICT

**predict()** Generate predictions from a Keras model

**predict\_proba()** and **predict\_classes()**  
Generates probability or class probability predictions for the input samples

**predict\_on\_batch()** Returns predictions for a single batch of samples

**predict\_generator()** Generates predictions for the input samples from a data generator

### OTHER MODEL OPERATIONS

**summary()** Print a summary of a Keras model

**export\_savedmodel()** Export a saved model

**get\_layer()** Retrieves a layer based on either its name (unique) or index

**pop\_layer()** Remove the last layer in a model

**save\_model\_hdf5(); load\_model\_hdf5()** Save/Load models using HDF5 files

**serialize\_model(); unserialize\_model()**  
Serialize a model to an R object

**clone\_model()** Clone a model instance

**freeze\_weights(); unfreeze\_weights()**  
Freeze and unfreeze weights

### CORE LAYERS

**layer\_input()** Input layer

**layer\_dense()** Add a densely-connected NN layer to an output

**layer\_activation()** Apply an activation function to an output

**layer\_dropout()** Applies Dropout to the input

**layer\_reshape()** Reshapes an output to a certain shape

**layer\_permute()** Permute the dimensions of an input according to a given pattern

**layer\_repeat\_vector()** Repeats the input n times

**layer\_lambda(object, f)** Wraps arbitrary expression as a layer

**layer\_activity\_regularization()**  
Layer that applies an update to the cost function based input activity

**layer\_masking()** Masks a sequence by using a mask value to skip timesteps

**layer\_flatten()** Flattens an input

### # input layer: use MNIST images



```
mnist <- dataset_mnist()
x_train <- mnist$train$x; y_train <- mnist$train$y
x_test <- mnist$test$x; y_test <- mnist$test$y
```

### # reshape and rescale

```
x_train <- array_reshape(x_train, c(nrow(x_train), 784))
x_test <- array_reshape(x_test, c(nrow(x_test), 784))
x_train <- x_train / 255; x_test <- x_test / 255
```

```
y_train <- to_categorical(y_train, 10)
y_test <- to_categorical(y_test, 10)
```

### # defining the model and layers

```
model <- keras_model_sequential()
model %>%
 layer_dense(units = 256, activation = 'relu',
 input_shape = c(784)) %>%
 layer_dropout(rate = 0.4) %>%
 layer_dense(units = 128, activation = 'relu') %>%
 layer_dense(units = 10, activation = 'softmax')
```

### # compile (define loss and optimizer)

```
model %>% compile(
 loss = 'categorical_crossentropy',
 optimizer = optimizer_rmsprop(),
 metrics = c('accuracy')
)
```

### # train (fit)

```
model %>% fit(
 x_train, y_train,
 epochs = 30, batch_size = 128,
 validation_split = 0.2
)
model %>% evaluate(x_test, y_test)
model %>% predict_classes(x_test)
```

# More layers

## CONVOLUTIONAL LAYERS



**layer\_conv\_1d()** 1D, e.g. temporal convolution



**layer\_conv\_2d\_transpose()**  
Transposed 2D (deconvolution)

**layer\_conv\_2d()** 2D, e.g. spatial convolution over images



**layer\_conv\_3d\_transpose()**  
Transposed 3D (deconvolution)  
**layer\_conv\_3d()** 3D, e.g. spatial convolution over volumes

**layer\_conv\_lstm\_2d()**  
Convolutional LSTM

**layer\_separable\_conv\_2d()**  
Depthwise separable 2D



**layer\_upsampling\_1d()**  
**layer\_upsampling\_2d()**  
**layer\_upsampling\_3d()**  
Upsampling layer



**layer\_zero\_padding\_1d()**  
**layer\_zero\_padding\_2d()**  
**layer\_zero\_padding\_3d()**  
Zero-padding layer



**layer\_cropping\_1d()**  
**layer\_cropping\_2d()**  
**layer\_cropping\_3d()**  
Cropping layer

## POOLING LAYERS



**layer\_max\_pooling\_1d()**  
**layer\_max\_pooling\_2d()**  
**layer\_max\_pooling\_3d()**  
Maximum pooling for 1D to 3D



**layer\_average\_pooling\_1d()**  
**layer\_average\_pooling\_2d()**  
**layer\_average\_pooling\_3d()**  
Average pooling for 1D to 3D



**layer\_global\_max\_pooling\_1d()**  
**layer\_global\_max\_pooling\_2d()**  
**layer\_global\_max\_pooling\_3d()**  
Global maximum pooling



**layer\_global\_average\_pooling\_1d()**  
**layer\_global\_average\_pooling\_2d()**  
**layer\_global\_average\_pooling\_3d()**  
Global average pooling

## ACTIVATION LAYERS



**layer\_activation(object, activation)**  
Apply an activation function to an output



**layer\_activation\_leaky\_relu()**  
Leaky version of a rectified linear unit



**layer\_activation\_parametric\_relu()**  
Parametric rectified linear unit



**layer\_activation\_thresholded\_relu()**  
Thresholded rectified linear unit



**layer\_activation\_elu()**  
Exponential linear unit

## DROPOUT LAYERS



**layer\_dropout()**  
Applies dropout to the input



**layer\_spatial\_dropout\_1d()**  
**layer\_spatial\_dropout\_2d()**  
**layer\_spatial\_dropout\_3d()**  
Spatial 1D to 3D version of dropout

## RECURRENT LAYERS



**layer\_simple\_rnn()**  
Fully-connected RNN where the output is to be fed back to input

**layer\_gru()**  
Gated recurrent unit - Cho et al

**layer\_cudnn\_gru()**  
Fast GRU implementation backed by CuDNN

**layer\_lstm()**  
Long-Short Term Memory unit - Hochreiter 1997

**layer\_cudnn\_lstm()**  
Fast LSTM implementation backed by CuDNN

## LOCALLY CONNECTED LAYERS

**layer\_locally\_connected\_1d()**  
**layer\_locally\_connected\_2d()**  
Similar to convolution, but weights are not shared, i.e. different filters for each patch

# Preprocessing

## SEQUENCE PREPROCESSING

**pad\_sequences()**  
Pads each sequence to the same length (length of the longest sequence)

**skipgrams()**  
Generates skipgram word pairs

**make\_sampling\_table()**  
Generates word rank-based probabilistic sampling table

## TEXT PREPROCESSING

**text\_tokenizer()** Text tokenization utility

**fit\_text\_tokenizer()** Update tokenizer internal vocabulary

**save\_text\_tokenizer(); load\_text\_tokenizer()**  
Save a text tokenizer to an external file

**texts\_to\_sequences(); texts\_to\_sequences\_generator()**  
Transforms each text in texts to sequence of integers

**texts\_to\_matrix(); sequences\_to\_matrix()**  
Convert a list of sequences into a matrix

**text\_one\_hot()** One-hot encode text to word indices

**text\_hashing\_trick()**  
Converts a text to a sequence of indexes in a fixed-size hashing space

**text\_to\_word\_sequence()**  
Convert text to a sequence of words (or tokens)

## IMAGE PREPROCESSING

**image\_load()** Loads an image into PIL format.

**flow\_images\_from\_data()**  
**flow\_images\_from\_directory()**  
Generates batches of augmented/normalized data from images and labels, or a directory

**image\_data\_generator()** Generate minibatches of image data with real-time data augmentation.

**fit\_image\_data\_generator()** Fit image data generator internal statistics to some sample data

**generator\_next()** Retrieve the next item

**image\_to\_array(); image\_array\_resize()**  
**image\_array\_save()** 3D array representation



# Pre-trained models

Keras applications are deep learning models that are made available alongside pre-trained weights. These models can be used for prediction, feature extraction, and fine-tuning.

**application\_xception()**  
**xception\_preprocess\_input()**  
Xception v1 model

**application\_inception\_v3()**  
**inception\_v3\_preprocess\_input()**  
Inception v3 model, with weights pre-trained on ImageNet

**application\_inception\_resnet\_v2()**  
**inception\_resnet\_v2\_preprocess\_input()**  
Inception-ResNet v2 model, with weights trained on ImageNet

**application\_vgg16(); application\_vgg19()**  
VGG16 and VGG19 models

**application\_resnet50()** ResNet50 model

**application\_mobilenet()**  
**mobilenet\_preprocess\_input()**  
**mobilenet\_decode\_predictions()**  
**mobilenet\_load\_model\_hdf5()**  
MobileNet model architecture

## IMAGENET

[ImageNet](https://www.image-net.org/) is a large database of images with labels, extensively used for deep learning

**imagenet\_preprocess\_input()**  
**imagenet\_decode\_predictions()**  
Preprocesses a tensor encoding a batch of images for ImageNet, and decodes predictions

# Callbacks

A callback is a set of functions to be applied at given stages of the training procedure. You can use callbacks to get a view on internal states and statistics of the model during training.

**callback\_early\_stopping()** Stop training when a monitored quantity has stopped improving  
**callback\_learning\_rate\_scheduler()** Learning rate scheduler  
**callback\_tensorboard()** TensorBoard basic visualizations

# Data Science in Spark with Sparklyr : : CHEAT SHEET

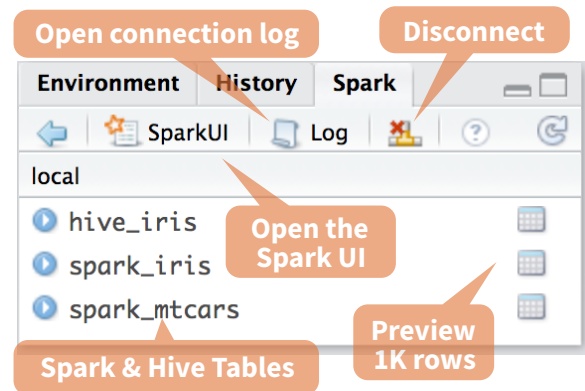


## Intro

**sparklyr** is an R interface for Apache Spark™, it provides a complete **dplyr** backend and the option to query directly using **Spark SQL** statement. With sparklyr, you can orchestrate distributed machine learning using either **Spark's MLlib** or **H2O Sparkling Water**.

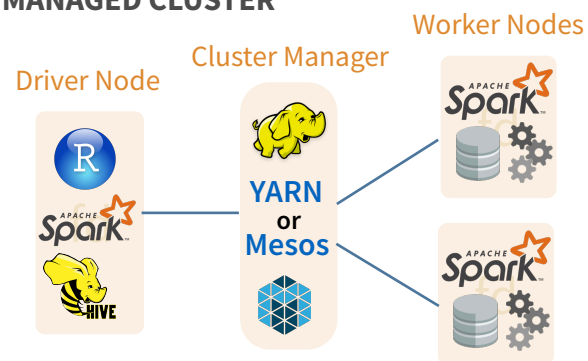
Starting with **version 1.044**, **RStudio Desktop, Server and Pro** include integrated support for the **sparklyr** package. You can create and manage connections to Spark clusters and local Spark instances from inside the IDE.

### RStudio Integrates with sparklyr

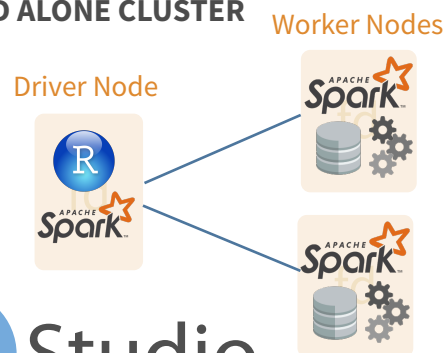


## Cluster Deployment

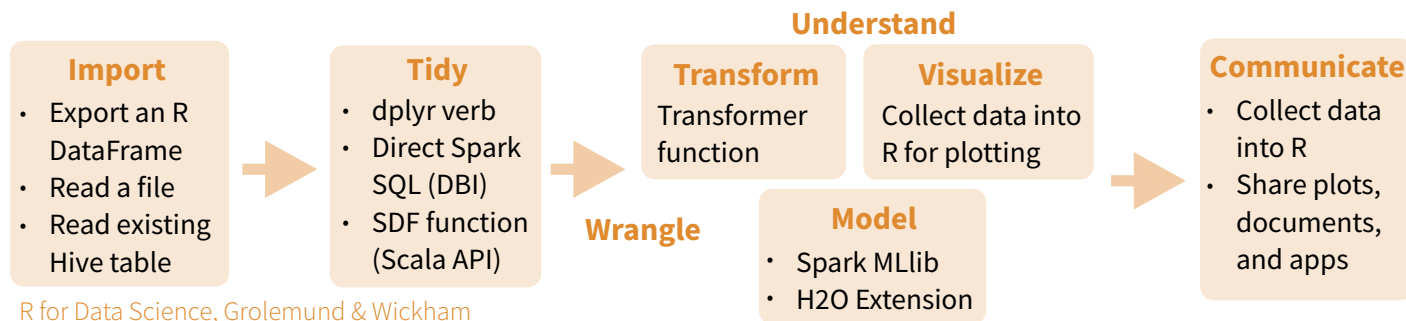
### MANAGED CLUSTER



### STAND ALONE CLUSTER



## Data Science Toolchain with Spark + sparklyr



## Getting Started

### LOCAL MODE (No cluster required)

1. Install a local version of Spark:  
**spark\_install("2.0.1")**
2. Open a connection  
**sc <- spark\_connect(master = "local")**

### ON A YARN MANAGED CLUSTER

1. Install RStudio Server or RStudio Pro on one of the existing nodes, preferably an edge node
2. Locate path to the cluster's Spark Home Directory, it normally is "/usr/lib/spark"
3. Open a connection  
**spark\_connect(master="yarn-client", version = "1.6.2", spark\_home = [Cluster's Spark path])**

### ON A MESOS MANAGED CLUSTER

1. Install RStudio Server or Pro on one of the existing nodes
2. Locate path to the cluster's Spark directory
3. Open a connection  
**spark\_connect(master="[mesos URL]", version = "1.6.2", spark\_home = [Cluster's Spark path])**

### USING LIVY (Experimental)

1. The Livy REST application should be running on the cluster
2. Connect to the cluster  
**sc <- spark\_connect(method = "livy", master = "http://host:port")**

## Tuning Spark

### EXAMPLE CONFIGURATION

```
config <- spark_config()
config$spark.executor.cores <- 2
config$spark.executor.memory <- "4G"
sc <- spark_connect(master="yarn-client",
 config = config, version = "2.0.1")
```

### IMPORTANT TUNING PARAMETERS with defaults

|                                     |                                               |
|-------------------------------------|-----------------------------------------------|
| • spark.yarn.am.cores               | • spark.executor.instances                    |
| • spark.yarn.am.memory <b>512m</b>  | • spark.executor.extraJavaOptions             |
| • spark.network.timeout <b>120s</b> | • spark.executor.heartbeatInterval <b>10s</b> |
| • spark.executor.memory <b>1g</b>   | • sparklyr.shell.executor-memory              |
| • spark.executor.cores <b>1</b>     | • sparklyr.shell.driver-memory                |

## Using sparklyr

A brief example of a data analysis using Apache Spark, R and sparklyr in local mode

```
library(sparklyr); library(dplyr); library(ggplot2);
library(tidyr);
set.seed(100)
```

**Install Spark locally**

```
spark_install("2.0.1")
```

**Connect to local version**

```
sc <- spark_connect(master = "local")
```

```
import_iris <- copy_to(sc, iris, "spark_iris",
 overwrite = TRUE)
```

**Copy data to Spark memory**

```
partition_iris <- sdf_partition(
 import_iris, training=0.5, testing=0.5)
```

**Partition data**

```
sdf_register(partition_iris,
 c("spark_iris_training", "spark_iris_test"))
```

**Create a hive metadata for each partition**

```
tidy_iris <- tbl(sc, "spark_iris_training") %>%
 select(Species, Petal_Length, Petal_Width)
```

**Spark ML Decision Tree Model**

```
model_iris <- tidy_iris %>%
 ml_decision_tree(response="Species",
 features=c("Petal_Length", "Petal_Width"))
```

```
test_iris <- tbl(sc, "spark_iris_test")
```

**Create reference to Spark table**

```
pred_iris <- sdf_predict(
 model_iris, test_iris) %>%
 collect
```

**Bring data back into R memory for plotting**

```
pred_iris %>%
 inner_join(data.frame(prediction=0:2,
 lab=model_iris$model.parameters$labels)) %>%
 ggplot(aes(Petal_Length, Petal_Width, col=lab)) +
 geom_point()
```

```
spark_disconnect(sc)
```

**Disconnect**

# Reactivity

## COPY A DATA FRAME INTO SPARK

```
sdf_copy_to(sc, iris, "spark_iris")
```

```
sdf_copy_to(sc, x, name, memory, repartition,
overwrite)
```

## IMPORT INTO SPARK FROM A FILE

Arguments that apply to all functions:

**sc, name, path, options = list(), repartition = 0, memory = TRUE, overwrite = TRUE**

**CSV** `spark_read_csv( header = TRUE, columns = NULL, infer_schema = TRUE, delimiter = ";", quote = "\"", escape = "\\", charset = "UTF-8", null_value = NULL)`

**JSON** `spark_read_json()`

**PARQUET** `spark_read_parquet()`

## SPARK SQL COMMANDS

```
DBI::dbWriteTable(sc, "spark_iris", iris)
```

```
DBI::dbWriteTable(conn, name, value)
```

## FROM A TABLE IN HIVE

```
my_var <- tbl_cache(sc, name=
"hive_iris")
```

```
tbl_cache(sc, name, force = TRUE)
Loads the table into memory
```

```
my_var <- dplyr::tbl(sc,
name= "hive_iris")
```

```
dplyr::tbl(scr, ...)
```

Creates a reference to the table without loading it into memory

# Wrangle

## SPARK SQL VIA DPLYR VERBS

Translates into Spark SQL statements

```
my_table <- my_var %>%
filter(Species=="setosa") %>%
sample_n(10)
```

## DIRECT SPARK SQL COMMANDS

```
my_table <- DBI::dbGetQuery(sc, "SELECT *
FROM iris LIMIT 10")
```

```
DBI::dbGetQuery(conn, statement)
```

## SCALA API VIA SDF FUNCTIONS

```
sdf_mutate(.data)
```

Works like dplyr mutate function

```
sdf_partition(x, ..., weights = NULL, seed =
sample (.Machine$integer.max, 1))
```

```
sdf_partition(x, training = 0.5, test = 0.5)
```

```
sdf_register(x, name = NULL)
```

Gives a Spark DataFrame a table name

```
sdf_sample(x, fraction = 1, replacement =
TRUE, seed = NULL)
```

```
sdf_sort(x, columns)
```

Sorts by >=1 columns in ascending order

```
sdf_with_unique_id(x, id = "id")
```

```
sdf_predict(object, newdata)
```

Spark DataFrame with predicted values

## ML TRANSFORMERS

```
ft_binarizer(my_table,input.col="Petal_Le
ngth", output.col="petal_large",
threshold=1.2)
```

Arguments that apply to all functions:  
**x, input.col = NULL, output.col = NULL**

```
ft_binarizer(threshold = 0.5)
```

Assigned values based on threshold

```
ft_bucketizer(splits)
```

Numeric column to discretized column

```
ft_discrete_cosine_transform(inverse =
FALSE)
```

Time domain to frequency domain

```
ft_elementwise_product(scaling.col)
```

Element-wise product between 2 cols

```
ft_index_to_string()
```

Index labels back to label as strings

```
ft_one_hot_encoder()
```

Continuous to binary vectors

```
ft_quantile_discretizer(n.buckets=5L)
```

Continuous to binned categorical values

```
ft_sql_transformer(sql)
```

```
ft_string_indexer(params = NULL)
```

Column of labels into a column of label indices.

```
ft_vector_assembler()
```

Combine vectors into single row-vector

# Visualize & Communicate

## DOWNLOAD DATA TO R MEMORY

```
r_table <- collect(my_table)
plot(Petal_Width~Petal_Length, data=r_table)
dplyr::collect(x)
```

Download a Spark DataFrame to an R DataFrame

```
sdf_read_column(x, column)
```

Returns contents of a single column to R

## SAVE FROM SPARK TO FILE SYSTEM

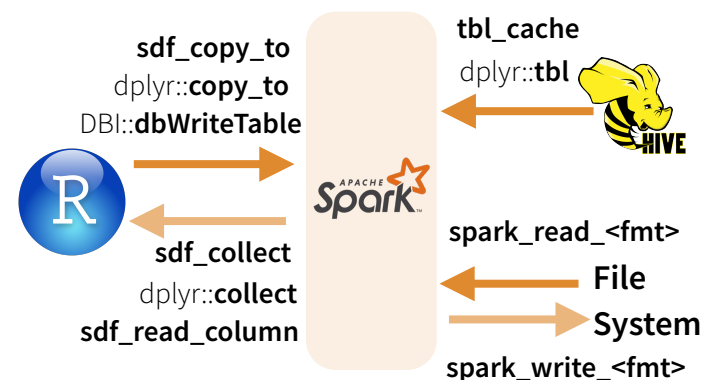
Arguments that apply to all functions: **x, path**

**CSV** `spark_read_csv( header = TRUE, delimiter = ";", quote = "\"", escape = "\\", charset = "UTF-8", null_value = NULL)`

**JSON** `spark_read_json(mode = NULL)`

**PARQUET** `spark_read_parquet(mode = NULL)`

# Reading & Writing from Apache Spark



# Extensions

Create an R package that calls the full Spark API & provide interfaces to Spark packages.

## CORE TYPES

`spark_connection()` Connection between R and the Spark shell process

`spark_jobj()` Instance of a remote Spark object

`spark_dataframe()` Instance of a remote Spark DataFrame object

## CALL SPARK FROM R

`invoke()` Call a method on a Java object

`invoke_new()` Create a new object by invoking a constructor

`invoke_static()` Call a static method on an object

## MACHINE LEARNING EXTENSIONS

`ml_create_dummy_variables()` `ml_options()`

`ml_prepare_dataframe()` `ml_model()`

`ml_prepare_response_features_intercept()`

# Model (MLlib)

```
ml_decision_tree(my_table,
response = "Species", features =
c("Petal_Length", "Petal_Width"))
```

```
ml_als_factorization(x, user.column = "user",
rating.column = "rating", item.column = "item",
rank = 10L, regularization.parameter = 0.1, iter.max = 10L,
ml.options = ml_options())
```

```
ml_decision_tree(x, response, features, max.bins = 32L, max.depth
= 5L, type = c("auto", "regression", "classification"), ml.options =
ml_options()) Same options for: ml_gradient_boosted_trees
```

```
ml_generalized_linear_regression(x, response, features,
intercept = TRUE, family = gaussian(link = "identity"), iter.max =
100L, ml.options = ml_options())
```

```
ml_kmeans(x, centers, iter.max = 100, features = dplyr::tbl_vars(x),
compute.cost = TRUE, tolerance = 1e-04, ml.options = ml_options())
```

```
ml_lda(x, features = dplyr::tbl_vars(x), k = length(features), alpha =
(50/k) + 1, beta = 0.1 + 1, ml.options = ml_options())
```

```
ml_linear_regression(x, response, features, intercept = TRUE,
alpha = 0, lambda = 0, iter.max = 100L, ml.options = ml_options())
Same options for: ml_logistic_regression
```

```
ml_multilayer_perceptron(x, response, features, layers, iter.max =
100, seed = sample(.Machine$integer.max, 1), ml.options =
ml_options())
```

```
ml_naive_bayes(x, response, features, lambda = 0, ml.options =
ml_options())
```

```
ml_one_vs_rest(x, classifier, response, features, ml.options =
ml_options())
```

```
ml_pca(x, features = dplyr::tbl_vars(x), ml.options = ml_options())
```

```
ml_random_forest(x, response, features, max.bins = 32L,
max.depth = 5L, num.trees = 20L, type = c("auto", "regression",
"classification"), ml.options = ml_options())
```

```
ml_survival_regression(x, response, features, intercept =
TRUE, censor = "censor", iter.max = 100L, ml.options = ml_options())
```

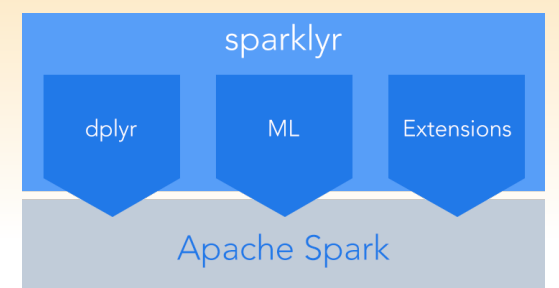
```
ml_binary_classification_eval(predicted_tbl_spark, label, score,
metric = "areaUnderROC")
```

```
ml_classification_eval(predicted_tbl_spark, label, predicted_lbl,
metric = "f1")
```

```
ml_tree_feature_importance(sc, model)
```

sparklyr

is an R  
interface  
for



# Basic Regular Expressions in R

## Cheat Sheet

### Character Classes

|                                               |                                                               |
|-----------------------------------------------|---------------------------------------------------------------|
| <code>[[:digit:]]</code> or <code>\\d</code>  | Digits; [0-9]                                                 |
| <code>\\D</code>                              | Non-digits; [^0-9]                                            |
| <code>[[:lower:]]</code>                      | Lower-case letters; [a-z]                                     |
| <code>[[:upper:]]</code>                      | Upper-case letters; [A-Z]                                     |
| <code>[[:alpha:]]</code>                      | Alphabetic characters; [A-z]                                  |
| <code>[[:alnum:]]</code>                      | Alphanumeric characters [A-z0-9]                              |
| <code>\\w</code>                              | Word characters; [A-z0-9_]                                    |
| <code>\\W</code>                              | Non-word characters                                           |
| <code>[[:xdigit:]]</code> or <code>\\x</code> | Hexadec. digits; [0-9A-Fa-f]                                  |
| <code>[[:blank:]]</code>                      | Space and tab                                                 |
| <code>[[:space:]]</code> or <code>\\s</code>  | Space, tab, vertical tab, newline, form feed, carriage return |
| <code>\\S</code>                              | Not space; [^[:space:]]                                       |
| <code>[[:punct:]]</code>                      | Punctuation characters; !"#%&'()*+,-./:;<=>?@[\\]^_`{ }~      |
| <code>[[:graph:]]</code>                      | Graphical characters; [[:alnum:]][[:punct:]]                  |
| <code>[[:print:]]</code>                      | Printable characters; [[:alnum:]][[:punct:]]\\s               |
| <code>[[:cntrl:]]</code> or <code>\\c</code>  | Control characters; \\n, \\r etc.                             |

### Special Metacharacters

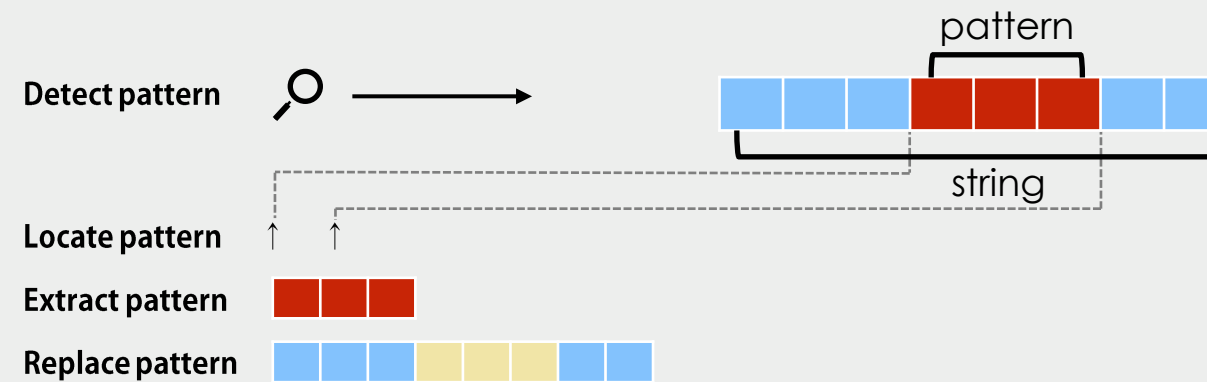
|                  |                 |
|------------------|-----------------|
| <code>\\n</code> | New line        |
| <code>\\r</code> | Carriage return |
| <code>\\t</code> | Tab             |
| <code>\\v</code> | Vertical tab    |
| <code>\\f</code> | Form feed       |

### Lookaraounds and Conditionals\*

|                             |                                                                                        |
|-----------------------------|----------------------------------------------------------------------------------------|
| <code>(?=)</code>           | Lookahead (requires PERL = TRUE), e.g. <code>(?=yx)</code> : position followed by 'xy' |
| <code>(?!)</code>           | Negative lookahead (PERL = TRUE); position NOT followed by pattern                     |
| <code>(?&lt;=)</code>       | Lookbehind (PERL = TRUE), e.g. <code>(?&lt;=yx)</code> : position following 'xy'       |
| <code>(?&lt;!)</code>       | Negative lookbehind (PERL = TRUE); position NOT following pattern                      |
| <code>?(if)then</code>      | If-then-condition (PERL = TRUE); use lookaheads, optional char. etc in if-clause       |
| <code>?(if)then else</code> | If-then-else-condition (PERL = TRUE)                                                   |

\*see, e.g. <http://www.regular-expressions.info/lookaround.html>  
<http://www.regular-expressions.info/conditional.html>

## Functions for Pattern Matching



```
> string <- c("Hiphopotamus", "Rhymenoceros", "time for bottomless lyrics")
> pattern <- "t.m"
```

### Detect Patterns

```
grep(pattern, string)
[1] 1 3

grep(pattern, string, value = TRUE)
[1] "Hiphopotamus"
[2] "time for bottomless lyrics"

grepl(pattern, string)
[1] TRUE FALSE TRUE

stringr::str_detect(string, pattern)
[1] TRUE FALSE TRUE
```

### Split a String using a Pattern

`strsplit(string, pattern)` or `stringr::str_split(string, pattern)`

### Locate Patterns

```
regexpr(pattern, string)
find starting position and length of first match

gregexpr(pattern, string)
find starting position and length of all matches

stringr::str_locate(string, pattern)
find starting and end position of first match

stringr::str_locate_all(string, pattern)
find starting and end position of all matches
```

### Extract Patterns

```
regmatches(string, regexpr(pattern, string))
extract first match [1] "tam" "tim"

regmatches(string, gregexpr(pattern, string))
extract all matches, outputs a list
[[1]] "tam" [[2]] character(0) [[3]] "tim" "tom"

stringr::str_extract(string, pattern)
extract first match [1] "tam" NA "tim"

stringr::str_extract_all(string, pattern)
extract all matches, outputs a list

stringr::str_extract_all(string, pattern, simplify = TRUE)
extract all matches, outputs a matrix

stringr::str_match(string, pattern)
extract first match + individual character groups

stringr::str_match_all(string, pattern)
extract all matches + individual character groups
```

### Replace Patterns

```
sub(pattern, replacement, string)
replace first match

gsub(pattern, replacement, string)
replace all matches

stringr::str_replace(string, pattern, replacement)
replace first match

stringr::str_replace_all(string, pattern, replacement)
replace all matches
```

### Character Classes and Groups

|                     |                                                                    |
|---------------------|--------------------------------------------------------------------|
| <code>.</code>      | Any character except \\n                                           |
| <code> </code>      | Or, e.g. <code>(a b)</code>                                        |
| <code>[...]</code>  | List permitted characters, e.g. <code>[abc]</code>                 |
| <code>[a-z]</code>  | Specify character ranges                                           |
| <code>[^...]</code> | List excluded characters                                           |
| <code>(...)</code>  | Grouping, enables back referencing using \\N where N is an integer |

### Anchors

|                     |                                       |
|---------------------|---------------------------------------|
| <code>^</code>      | Start of the string                   |
| <code>\$</code>     | End of the string                     |
| <code>\\b</code>    | Empty string at either edge of a word |
| <code>\\B</code>    | NOT the edge of a word                |
| <code>\\&lt;</code> | Beginning of a word                   |
| <code>\\&gt;</code> | End of a word                         |

### Quantifiers

|                    |                                         |
|--------------------|-----------------------------------------|
| <code>*</code>     | Matches at least 0 times                |
| <code>+</code>     | Matches at least 1 time                 |
| <code>?</code>     | Matches at most 1 time; optional string |
| <code>{n}</code>   | Matches exactly n times                 |
| <code>{n,}</code>  | Matches at least n times                |
| <code>{n,m}</code> | Matches between n and m times           |

### General Modes

By default R uses *extended* regular expressions. You can switch to *PCRE regular expressions* using `PERL = TRUE` for base or by wrapping patterns with `perl()` for stringr.

All functions can be used with literal searches using `fixed = TRUE` for base or by wrapping patterns with `fixed()` for stringr.

All base functions can be made case insensitive by specifying `ignore.cases = TRUE`.

### Escaping Characters

Metacharacters (`.` `*` `+` etc.) can be used as literal characters by escaping them. Characters can be escaped using `\\` or by enclosing them in `\\Q...\\E`.

### Case Conversions

Regular expressions can be made case insensitive using `(?i)`. In backreferences, the strings can be converted to lower or upper case using `\\L` or `\\U` (e.g. `\\L\\1`). This requires `PERL = TRUE`.

### Greedy Matching

By default the asterisk `*` is greedy, i.e. it always matches the longest possible string. It can be used in lazy mode by adding `?`, i.e. `*?`.

Greedy mode can be turned off using `(?U)`. This switches the syntax, so that `(?U)a*` is lazy and `(?U)a*?` is greedy.

### Note

Regular expressions can conveniently be created using e.g. the packages `rex` or `rebus`.

# caret Package

## Cheat Sheet

### Specifying the Model

Possible syntaxes for specifying the variables in the model:

```
train(y ~ x1 + x2, data = dat, ...)
train(x = predictor_df, y = outcome_vector, ...)
train(recipe_object, data = dat, ...)
```

- `rfe`, `sbf`, `gafs`, and `safs` only have the `x/y` interface.
- The `train` formula method will **always** create dummy variables.
- The `x/y` interface to `train` will not create dummy variables (but the underlying model function might).

**Remember** to:

- Have column names in your data.
- Use factors for a classification outcome (not 0/1 or integers).
- Have valid R names for class levels (not "0"/"1")
- Set the random number seed prior to calling `train` repeatedly to get the same resamples across calls.
- Use the `train` option `na.action = na.pass` if you will be imputing missing data. Also, use this option when predicting new data containing missing values.

To pass options to the underlying model function, you can pass them to `train` via the ellipses:

```
train(y ~ ., data = dat, method = "rf",
 # options to `randomForest`:
 importance = TRUE)
```

### Parallel Processing

The `foreach` package is used to run models in parallel. The `train` code does not change but a “do” package must be called first.

|                                    |                                      |
|------------------------------------|--------------------------------------|
| # on MacOS or Linux                | # on Windows                         |
| <code>library(doMC)</code>         | <code>library(doParallel)</code>     |
| <code>registerDoMC(cores=4)</code> | <code>cl &lt;- makeCluster(2)</code> |
|                                    | <code>registerDoParallel(cl)</code>  |

The function `parallel::detectCores` can help too.

### Preprocessing

Transformations, filters, and other operations can be applied to the *predictors* with the `preProc` option.

```
train(, preProc = c("method1", "method2"), ...)
```

Methods include:

- `"center"`, `"scale"`, and `"range"` to normalize predictors.
- `"BoxCox"`, `"YeoJohnson"`, or `"expoTrans"` to transform predictors.
- `"knnImpute"`, `"bagImpute"`, or `"medianImpute"` to impute.
- `"corr"`, `"nzv"`, `"zv"`, and `"conditionalX"` to filter.
- `"pca"`, `"ica"`, or `"spatialSign"` to transform groups.

`train` determines the order of operations; the order that the methods are declared does not matter.

The `recipes` package has a more extensive list of preprocessing operations.

### Adding Options

Many `train` options can be specified using the `trainControl` function:

```
train(y ~ ., data = dat, method = "cubist",
 trControl = trainControl(<options>))
```

### Resampling Options

`trainControl` is used to choose a resampling method:

```
trainControl(method = <method>, <options>)
```

Methods and options are:

- `"cv"` for K-fold cross-validation (`number` sets the # folds).
- `"repeatedcv"` for repeated cross-validation (`repeats` for # repeats).
- `"boot"` for bootstrap (`number` sets the iterations).
- `"LGOV"` for leave-group-out (`number` and `p` are options).
- `"LOO"` for leave-one-out cross-validation.
- `"oob"` for out-of-bag resampling (only for some models).
- `"timeslice"` for time-series data (options are `initialWindow`, `horizon`, `fixedWindow`, and `skip`).

### Performance Metrics

To choose how to summarize a model, the `trainControl` function is used again.

```
trainControl(summaryFunction = <R function>,
 classProbs = <logical>)
```

Custom R functions can be used but `caret` includes several: `defaultSummary` (for accuracy, RMSE, etc), `twoClassSummary` (for ROC curves), and `prSummary` (for information retrieval). For the last two functions, the option `classProbs` must be set to `TRUE`.

### Grid Search

To let `train` determine the values of the tuning parameter(s), the `tuneLength` option controls how many values **per tuning** parameter to evaluate.

Alternatively, specific values of the tuning parameters can be declared using the `tuneGrid` argument:

```
grid <- expand.grid(alpha = c(0.1, 0.5, 0.9),
 lambda = c(0.001, 0.01))
```

```
train(x = x, y = y, method = "glmnet",
 preProc = c("center", "scale"),
 tuneGrid = grid)
```

### Random Search

For tuning, `train` can also generate random tuning parameter combinations over a wide range. `tuneLength` controls the total number of combinations to evaluate. To use random search:

```
trainControl(search = "random")
```

### Subsampling

With a large class imbalance, `train` can subsample the data to balance the classes them prior to model fitting.

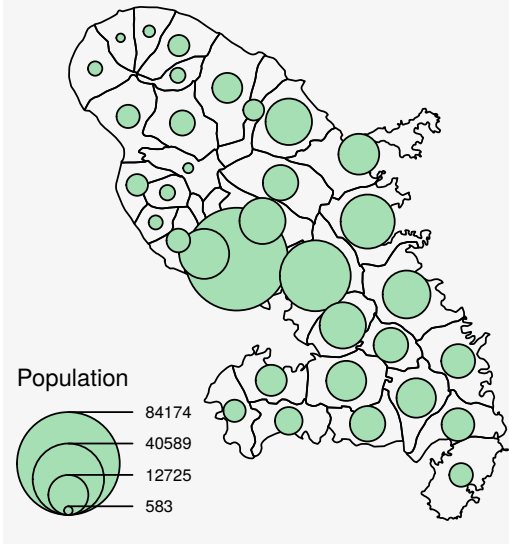
```
trainControl(sampling = "down")
```

Other values are `"up"`, `"smote"`, or `"rose"`. The latter two may require additional package installs.

# Thematic maps with cartography : : CHEAT SHEET

Use cartography with spatial objects from sf or sp packages to create thematic maps.

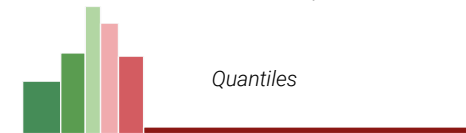
```
library(cartography)
library(sf)
mtq <- st_read("martinique.shp")
plot(st_geometry(mtg))
propSymbolsLayer(x = mtq, var = "P13_POP",
 legend.title.txt = "Population",
 col = "#a7dfb4")
```



## Classification

Available methods are: quantile, equal, q6, fisher-jenks, mean-sd, sd, geometric progression...

```
bks1 <- getBreaks(v = var, nclass = 6,
 method = "quantile")
bks2 <- getBreaks(v = var, nclass = 6,
 method = "fisher-jenks")
pal <- carto.pal("green.pal", 3, "wine.pal", 3)
hist(var, breaks = bks1, col = pal)
```

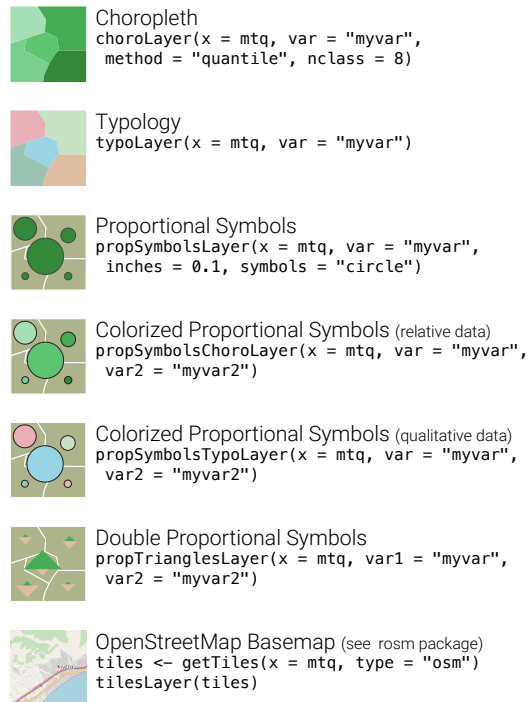


```
hist(var, breaks = bks2, col = pal)
```



## Symbology

In most functions the x argument should be an sf object. sp objects are handled through spdf and df arguments.



**Choropleth**  
choroLayer(x = mtq, var = "myvar",  
method = "quantile", nclass = 8)

**Typology**  
typoLayer(x = mtq, var = "myvar")

**Proportional Symbols**  
propSymbolsLayer(x = mtq, var = "myvar",  
inches = 0.1, symbols = "circle")

**Colorized Proportional Symbols (relative data)**  
propSymbolsChoroLayer(x = mtq, var = "myvar",  
var2 = "myvar2")

**Colorized Proportional Symbols (qualitative data)**  
propSymbolsTypoLayer(x = mtq, var = "myvar",  
var2 = "myvar2")

**Double Proportional Symbols**  
propTrianglesLayer(x = mtq, var1 = "myvar",  
var2 = "myvar2")

**OpenStreetMap Basemap** (see rosm package)  
tiles <- getTiles(x = mtq, type = "osm")  
tilesLayer(tiles)

**Isopleth** (see SpatialPosition package)  
smoothLayer(x = mtq, var = "myvar",  
typefct = "exponential", span = 500,  
beta = 2)

**Discontinuities**  
disclayer(x = mtq.borders, df = mtq,  
var = "myvar", threshold = 0.5)

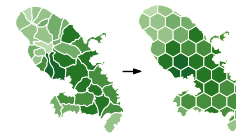
**Flows**  
proLinkLayer(x = mtq\_link, df = mtq\_df,  
var = "fij")

**Dot Density**  
dotDensityLayer(x = mtq, var = "myvar")

**Labels**  
labelLayer(x = mtq, txt = "myvar",  
halo = TRUE, overlap = FALSE)

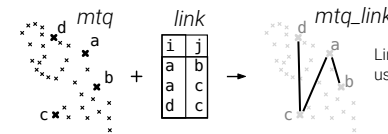
## Transformations

**Polygons to Grid**  
mtq\_grid <- getGridLayer(x = mtq, cellsize = 3.6e+07,  
type = "hexagonal", var = "myvar")



Grids layers can be used by  
choroLayer() or propSymbolsLayer().

**Points to Links**  
mtq\_link <- getLinkLayer(x = mtq, df = link)



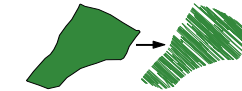
Links layers can be  
used by \*LinkLayer().

**Polygons to Borders**  
mtq\_border <- getBorders(x = mtq)



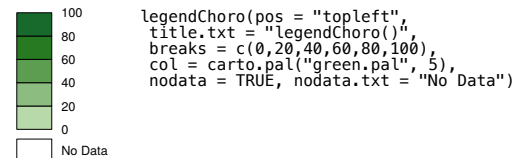
Borders layers can be used by  
disclayer() function

**Polygons to Pencil Lines**  
mtq\_pen <- getPencilLayer(x = mtq)

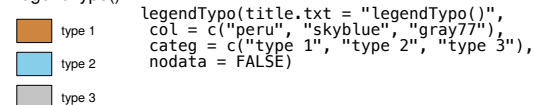


## Legends

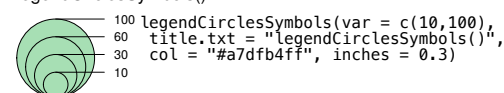
legendChoro()



legendTypo()



legendCirclesSymbols()



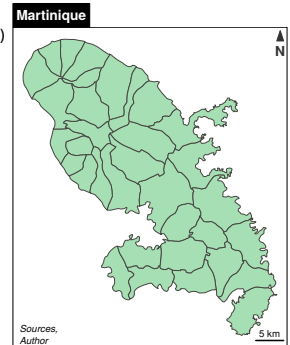
See also legendSquaresSymbols(), legendBarsSymbols(),  
legendGradLines(), legendPropLines() and legendPropTriangles().

## Map Layout

**North Arrow:**  
north(pos = "topright")

**Scale Bar:**  
barscale(size = 5)

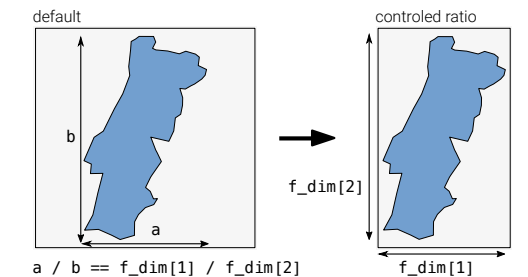
**Full Layout:**  
layoutLayer(  
title = "Martinique",  
tabtitle = TRUE,  
frame = TRUE,  
author = "Author",  
sources = "Sources",  
north = TRUE,  
scale = 5)



**Figure Dimensions**

Get figure dimensions based on the dimension ratio of a spatial object,  
figure margins and output resolution.

```
f_dim <- getFigDim(x = sf_obj, width = 500,
mar = c(0,0,0,0))
png("fig.png", width = 500, height = f_dim[2])
par(mar = c(0,0,0,0))
plot(sf_obj, col = "#729fcf")
dev.off()
```



## Color Palettes

carto.pal(pal1 = "blue.pal", n1 = 5,  
pal2 = sand.pal, n2 = 3)



display.carto.all(n = 8)



## Dataset Operations

### DATA IMPORT / EXPORT

**h2o.uploadFile:** Upload a file into H2O from a client-side path, and parse it.

**h2o.downloadCSV:** Download a H2O dataset to a client-side CSV file.

**h2o.importFile:** Import a file into H2O from a server-side path, and parse it.

**h2o.exportFile:** Export an H2O Data Frame to a server-side file.

**h2o.parseRaw:** Parse a raw data file.

### NATIVE R TO H2O COERCION

**as.h2o:** Convert a R object to an H2O object

### H2O TO NATIVE R COERCION

**as.data.frame:** Check if an object is a data frame, and coerce it if possible.

### DATA GENERATION

**h2o.createFrame:** Creates a data frame in H2O with real-valued, categorical, integer, and binary columns specified by the user, with optional randomization.

**h2o.runif:** Produce a vector of random uniform numbers.

**h2o.interaction:** Create interaction terms between categorical features of an H2O Frame.

**h2o.target\_encode\_apply:** Target encoding map to an H2O Data Frame, which can improve performance of supervised learning models for high cardinality categorical columns.

### DATA SAMPLING / SPLITTING

**h2o.splitFrame:** Split an existing H2O dataset according to user-specified ratios.

### MISSING DATA HANDLING

**h2o.impute:** Impute a column of data using the mean, median, or mode.

**h2o.insertMissingValues:** Replaces a user-specified fraction of entries in an H2O dataset with missing values.

**h2o.na\_omit:** Remove Rows With NAs.

## General Operations

### SUBSCRIPTING

Subscripting example to pull (/push) pieces from (/to) a H2O Parsed Data object.

|                                                    |                                                                                          |
|----------------------------------------------------|------------------------------------------------------------------------------------------|
| <pre>x[j] ## column J x[i, j] x[[i]] x\$name</pre> | <pre>x[i] &lt;- value x[i, j, ...] &lt;- value x[[i]] &lt;- value x\$i &lt;- value</pre> |
| Selection                                          | Value Assignment                                                                         |

### SUBSETTING

**h2o.head, h2o.tail:** Object's Start or End.

### DATA ATTRIBUTES

**h2o.names:** Return column names for an H2O Frame. Also: **h2o.colnames**

**names<-:** Set the row or column names of a H2O Frame. Also: **colnames<-**

**h2o.dim:** Retrieve object dimensions.

**h2o.length:** Length of vector, list or factor.

**h2o.nrow:** Number of H2O Frame rows.

**h2o.ncol:** Number of H2O Frame columns.

**h2o.anyFactor:** Check if an H2O Frame object has any categorical data columns.

**is.factor, is.character, is.numeric:** Check Column's Data Type.

**DATA TYPE COERCION:** Convert to:

**h2o.asfactor, as.factor:** Factor.

**h2o.as\_date, as.Date:** Date.

**h2o.ascharacter, as.character:** Character.

**h2o.asnumeric, as.numeric:** Numeric.

### BASIC DATA MANIPULATION

**c:** Combine Values into a Vector or List.

**h2o.cbind, h2o.rbind:** Combine a sequence of H2O datasets by column (cbind) or rows (rbind).

**h2o.merge:** Merges 2 H2O Frames.

**h2o.arrange:** Sorts an H2O Frame by columns.

### ELEMENT INDEX SELECTION

**h2o.which:** True Condition's Row Numbers

### CONDITIONAL VALUE SELECTION

**h2o.ifelse:** Apply conditional statements to numeric vectors in an H2O Frame.

## Math Operations

(math) vectorized function

### MATH

**h2o.abs:** Compute the absolute value of x.

**h2o.sqrt:** Principal Square Root of x,  $\sqrt{x}$ .

**h2o.ceiling:** Take a single numeric argument x and return a numeric vector containing the smallest integers not less than the corresponding elements of x.

**h2o.floor:** Take a single numeric argument x and return a numeric vector containing the largest integers not greater than the corresponding elements of x.

**h2o.trunc:** Take a single numeric argument x and return a numeric vector containing the integers formed by truncating the values in x toward 0.

**h2o.log:** Compute natural logarithms. See also: **h2o.log10, h2o.log2, h2o.log1p**

**h2o.exp:** Compute the exponential function

**h2o.cos, h2o.cosh, h2o.acos, h2o.sin, h2o.tan, h2o.tanh, Math:** ?groupGeneric

**sign:** Return a vector with the signs of the corresponding elements of x (the sign of a real number is 1, 0, or -1 if the number is positive, zero, or negative, respectively).

**&&** (Vectorized AND), **||** (Vectorized OR), **!x, %in%, Ops:** +, -, \*, /, ^, %%, %/%, ==, !=, <, <=, >=, >, &, |, !

### CUMULATIVE

**h2o.cummax:** Vector of the cumulative maxima of the elements of the argument.

**h2o.cummin:** Vector of the cumulative minima of the elements of the argument.

**h2o.cumprod:** Vector of the cumulative products of the elements of the argument.

**h2o.cumsum:** Vector of the cumulative sums of the elements of the argument.

### PRECISION

**h2o.round:** Round values to the specified number of decimal places. The default is 0.

**h2o.signif:** Round values to the specified number of significant digits.

## Group By Summaries

(group by) summary function

**nrow:** Count the number of rows.

**max:** All input argument's Maximum.

**min:** All input argument's Minimum.

**sum:** All argument values Sum.

**mean:** (Trimmed) arithmetic mean.

**sd:** Calculate the standard deviation of a column of continuous real valued data.

**var:** Compute the variance of x.

## Generic Summaries

### NON-GROUP\_BY SUMMARIES

**h2o.median:** Calculate the median of x.

**h2o.range:** Input argument's Min/Max Vector

**h2o.cor:** Correlation Matrix of H2O Frames.

**h2o.quantile:** Obtain and display quantiles for an H2O Frame Column.

**h2o.hist:** Compute a histogram over a numeric H2O Frame Column.

**h2o.prod:** Product of all arguments values.

**h2o.any:** Given a set of logical vectors, determine if at least one of the values is true.

**h2o.all:** Given a set of logical vectors, determine if all of the values are true.

### NON-GROUP\_BY SUMMARIES: GENERIC

**h2o.summary:** Produce result summaries of the results of various model fitting functions.

## Aggregations

### ROW / COLUMN AGGREGATION

**apply:** Apply a function over an H2O parsed data object (an array) margins.

### GROUP BY AGGREGATION

**h2o.group\_by:** Apply an aggregate function to each group of an H2O dataset.

### TABULATION

**h2o.table:** Use the cross-classifying factors to build a table of counts at each combination of factor levels.

## Data Modeling

### MODEL TRAINING: SUPERVISED LEARNING

**h2o.deeplearning:** Perform Deep Learning Neural Networks on an H2OFrame.

**h2o.gbm:** Build Gradient Boosted Regression Trees or Classification Trees.

**h2o.glm:** Fit a Generalized Linear Model, specified by a response variable, a set of predictors, and the error distribution.

**h2o.naiveBayes:** Compute Naive Bayes classification probabilities on an H2O Frame.

**h2o.randomForest:** Perform Random Forest Classification on an H2O Frame.

**h2o.xgboost:** Build an Extreme Gradient Boosted Model using the XGBoost backend.

**h2o.stackedEnsemble:** Build a stacked ensemble (aka. Super Learner) using the specified H2O base learning algorithms.

**h2o.automl:** Automates the Supervised Machine Learning Model Training Process: Automatically Trains and Cross-validates a set of Models, and trains a Stacked Ensemble.

### MODEL TRAINING: UNSUPERVISED LEARNING

**h2o.prcomp:** Perform Principal Components Analysis on the given H2O Frame.

**h2o.kmeans:** Perform k-means Clustering on the given H2O Frame.

**h2o.anomaly:** Detect anomalies in a H2O Frame using a H2O Deep Learning Model with Auto-Encoding.

**h2o.deepfeatures:** Extract the non-linear features from a H2O Frame using a H2O Deep Learning Model.

**h2o.glrn:** Builds a Generalized Low Rank Decomposition of an H2O Frame.

**h2o.svd:** Singular value decomposition of an H2O Frame using the power method.

**h2o.word2vec:** Trains a word2vec model on a String column of an H2O data frame.

### SURVIVAL MODELS: TIME-TO-EVENT

**h2o.coxph:** Trains a Cox Proportional Hazards Model (CoxPH) on an H2O Frame.

### GRID SEARCH

**h2o.grid:** Efficient method to build multiple models with different hyperparameters.

**h2o.getGrid:** Get a grid object from H2O distributed K/V store.

### MODEL SCORING

**h2o.predict:** Obtain predictions from various fitted H2O model objects.

**h2o.scoreHistory:** Get Model Score History.

### MODEL METRICS

**h2o.make\_metrics:** Given predicted values (target for regression, class-1 probabilities, or binomial or per-class probabilities for multinomial), compute a model metrics object.

### GENERAL MODEL HELPER

**h2o.performance:** Evaluate the predictive performance of a Supervised Learning Regression or Classification Model via various metrics. Set **xval = TRUE** for retrieving the training cross-validation metrics.

### REGRESSION MODEL HELPER

**h2o.mse:** Display the mean squared error calculated from “Predicted Responses” and “Actual (Reference) Responses”. Set **xval = TRUE** for retrieving the cross-validation MSE.

### CLASSIFICATION MODEL HELPERS

**h2o.accuracy:** Get Model Accuracy metric.

**h2o.auc:** Retrieve the AUC (area under ROC curve). Set **xval = TRUE** for retrieving the cross-validation AUC.

**h2o.confusionMatrix:** Display prediction errors for classification data (“Predicted” vs “Reference : Real Values”).

**h2o.hit\_ratio\_table:** Retrieve the Hit Ratios. Set **xval = TRUE** for retrieving the cross-validation Hit Ratio.

### CLUSTERING MODEL HELPER

**h2o.betweenss:** Get the between cluster Sum of Squares.

**h2o.centers:** Retrieve the Model Centers.

### PREDICTOR VARIABLE IMPORTANCE

**h2o.varimp:** Retrieve the variable importance

**h2o.varimp\_plot:** Plot Variable Importances.

## Data Munging

### GENERAL COLUMN MANIPULATION

**is.na:** Display missing elements.

### FACTOR LEVEL MANIPULATIONS

**h2o.levels:** Display a list of the unique values found in a categorical data column.

**h2o.relevel:** Reorders levels of an H2O factor, similarly to standard R's relevel.

**h2o.setLevels:** Set Levels of H2O Factor.

### NUMERIC COLUMN MANIPULATIONS

**h2o.cut:** Convert H2O Numeric Data to Factor by breaking it into Intervals.

### CHARACTER COLUMN MANIPULATIONS

**h2o.strsplit:** “String Split”: Splits the given factor column on the input split.

**h2o.tolower:** Convert the characters of a character vector to lower case.

**h2o.toupper:** Convert the characters of a character vector to upper case.

**h2o.trim:** “Trim spaces”: Remove leading and trailing white space.

**h2o.gsub:** Match a pattern & replace **all** instances (occurrences) of the matched pattern with the replacement string globally.

**h2o.sub:** Match a pattern & replace the **first** instance (occurrence) of the matched pattern with the replacement string.

### DATE MANIPULATIONS

**h2o.month:** Convert Milliseconds to Months in H2O Datasets (Scale: 0 to 11).

**h2o.year:** Convert Milliseconds to Years in H2O Datasets, indexed starting from 1900.

**h2o.day:** Convert Milliseconds to Day of Month in H2O Datasets (Scale: 1 to 31).

**h2o.hour:** Convert Milliseconds to Hour of Day in H2O Datasets (Scale: 0 to 23).

**h2o.dayOfWeek:** Convert Milliseconds to Day of Week in a H2OFrame (Scale: 0 to 6)

### MATRIX OPERATIONS

**%\*%:** Multiply two conformable matrices.

**t:** Returns the transpose of an H2OFrame.

## Cluster Operations

### H2O KEY VALUE STORE ACCESS

**h2o.assign:** Assign H2O hex.keys to R objects.

**h2o.getFrame:** Get H2O dataset Reference.

**h2o.getModel:** Get H2O model reference.

**h2o.ls:** Display a list of object keys in the running instance of H2O.

**h2o.rm:** Remove specified H2O Objects from the H2O server, but not from the R environment.

**h2o.removeAll:** Remove All H2O Objects from the H2O server, but not from the R environment.

### H2O MODEL IMPORT / EXPORT

**h2o.loadModel:** Load H2OModel from disk.

**h2o.saveModel:** Save H2OModel object to disk.

**h2o.download\_pojo:** Download the Scoring POJO (Plain Old Java Object) of an H2O Model.

**h2o.download\_mojo:** Download the model in MOJO format.

### H2O CLUSTER CONNECTION

**h2o.init:** Connect to a running H2O instance using all CPUs on the host.

**h2o.shutdown:** Shut down the specified H2O instance. All data on the server will be lost!

### H2O CLUSTER INFORMATION

**h2o.clusterInfo:** Display the name, version, uptime, total nodes, total memory, total cores and health of a cluster running H2O.

**h2o.clusterStatus:** Retrieve information on the status of the cluster running H2O.

### H2O LOGGING

**h2o.clearLog:** Clear all H2O R command and error response logs from the local disk.

**h2o.downloadAllLogs:** Download all H2O log files to the local disk.

**h2o.logAndEcho:** Write a message to the H2O Java log file and echo it back.

**h2o.openLog:** Open existing logs of H2O R POST commands and error responses on disk.

**h2o.getLogPath:** Get the file path for the H2O R command and error response logs.

**h2o.startLogging:** Begin logging H2O R POST commands and error responses.

**h2o.stopLogging:** Stop logging H2O R POST commands and error responses.

# Data & Variable Transformation

## with sjmisc Cheat Sheet



**sjmisc** complements **dplyr**, and helps with data transformation tasks and recoding variables.

**sjmisc** works together seamlessly with **dplyr** and pipes. All functions are designed to support labelled data.



## Design Philosophy

The design of **sjmisc** functions follows the tidyverse-approach: first argument is always the data (either a *data frame* or *vector*), followed by variable names to be processed by the functions.

The returned object for each function *equals the type of the data-argument*.

### Vector input

- If the data-argument is a *vector*, functions return a *vector*.



```
rec(mtcars$carb, rec = "1,2=1; 3,4=2; else=3")
```

### Data frame input

- If the data-argument is a *data frame*, functions return a *data frame*.



```
rec(mtcars, carb, rec = "1,2=1; 3,4=2; else=3")
```

## The ...-ellipses Argument

Apply functions to a single variable, selected variables or to a complete data frame.

Variable selection is powered by **dplyr**'s **select()**: Separate variables with comma, or use **dplyr**'s select-helpers to select variables, e.g. **?rec**:

```
rec(mtcars, one_of(c("gear", "carb")),
 rec = "min:3=1; 4:max=2")
```

```
rec(mtcars, gear, carb, rec = "min:3=1; 4:max=2")
```

## Descriptives and Summaries

Most of the **sjmisc** functions (including recode-functions) also work on grouped data frames:

```
library(dplyr)
efc %>%
 group_by(e16sex, c172code) %>%
 frq(e42dep)
```

## Frequency Tables

```
frq(x, ..., sort.frq = c("none", "asc", "desc"),
 weight.by = NULL, auto.grp = NULL, ...)
```

Print frequency tables of (labelled) vectors. Uses variable labels as table header.

```
data(efc); frq(efc, e42dep, c161sex)
```

Use this data set in examples!

```
flat_table(data, ..., margin = c("counts",
 "cell", "row", "col"), digits = 2,
 show.values = FALSE)
```

Print contingency tables of (labelled) vectors. Uses value labels.

```
flat_table(efc, e42dep, c172code, e16sex)
```

```
count_na(x, ...)
```

Print frequency table of tagged NA values.

```
library(haven); x <- labelled(c(1:3,
 tagged_na("a", "a", "z")), labels =
 c("Refused" = tagged_na("a"), "N/A" =
 tagged_na("z")))
count_na(x)
```

## Descriptive Summary

```
descr(x, ..., max.length = NULL)
```

Descriptive summary of data frames, including variable labels in output.

```
descr(efc, contains("cop"), max.length = 20)
```

## Finding Variables in a Data Frame

Use **find\_var()** to search for variables by names, value or variable labels. Returns vector/data frame.

```
variables with "cop" in names and variable labels
find_var(efc, pattern = "cop", out = "df")
```

```
variables with "level" in names and value labels
find_var(efc, "level", search = "name_value")
```

## Recode and Transform Variables

Recode functions add a *suffix* to new variables, so original variables are preserved.

By default, original input data frame and new created variables are returned. Use **append = FALSE** to return the recoded variables only.

```
rec(x, ..., rec, as.num = TRUE, var.label =
 NULL, val.labels = NULL, append = TRUE,
 suffix = "_r")
```

Recode values, return result as numeric, character or categorical (factor).

```
rec(mtcars, carb, rec = "1,2=1; 3,4=2; else=3")
```

```
dicho(x, ..., dich.by = "median", as.num =
 FALSE, var.label = NULL, val.labels = NULL,
 append = TRUE, suffix = "_d")
```

Dichotomise variable by median, mean or specific value.

```
dicho(mtcars, disp)
```

```
split_var(x, ..., n, as.num = FALSE,
 val.labels = NULL, var.label = NULL,
 inclusive = FALSE, append = TRUE,
 suffix = "_g")
```

Split variable into equal sized groups. Unlike **dplyr::ntile()**, does not split original categories into different values (see examples in **?split\_var**).

```
split_var(mtcars, mpg, disp, n = 3)
```

```
group_var(x, ..., size = 5, as.num = TRUE,
 right.interval = FALSE, n = 30, append =
 TRUE, suffix = "_gr")
```

Split variable into groups with equal value range, or into a max. # of groups (value range per group is adjusted to match # of groups).

```
group_var(mtcars, mpg, disp, size = 5)
```

```
group_var(mtcars, mpg, size = "auto", n = 4)
```

```
std(x, ..., robust = "sd", include.fac = FALSE,
 append = TRUE, suffix = "_z")
```

Z-standardise variables. Also **center()**.

```
std(efc, e17age, c160age)
```

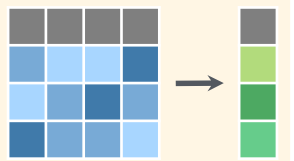
```
recode_to(x, ..., lowest = 0, highest = -1,
 append = TRUE, suffix = "_r0")
```

Shift ("renumber") categories or values.

```
recode_to(mtcars$gear)
```

## Summarise Variables and Cases

The summary functions mostly mimic base R equivalents, but are designed to work together with pipes and **dplyr**.



```
row_sums(x, ..., na.rm = TRUE, var =
 "rowsums", append = FALSE)
```

Row sums of data frames.

```
row_sums(efc, c82cop1:c90cop9)
```

```
row_means(x, ..., n, var = "rowmeans",
 append = FALSE)
```

Row means, for at least *n* valid (non-NA) values.

```
row_means(efc, c82cop1:c90cop9, n = 7)
```

```
row_count(x, ..., count, var = "rowcount",
 append = FALSE)
```

Row-wise count # of values in data frames.

Also **col\_count()**.

```
row_count(efc, c82cop1:c90cop9, count = 2)
```

## Other Useful Functions

- add\_columns()** and **replace\_columns()** to combine data frames, but either replace or preserve existing columns.
- set\_na()** and **replace\_na()** to convert regular into missing values, or vice versa. **replace\_na()** also replaces specific tagged NA values only.
- remove\_var()** and **var\_rename()** to remove variables from data frames, or rename variables.
- group\_str()** to group similar string values. Useful for variables with similar, but not identically spelled string values that should be "merged".
- merge\_df()** to full join data frames and preserve value and variable labels.
- to\_long()** to gather multiple columns in data frames from wide into long format.

## Use with %>% and dplyr

```
use sjmisc-functions in pipes
mtcars %>% select(gear, carb) %>%
 rec(rec = "min:3=1; 4:max=2")
```

```
use sjmisc-function inside mutate
mtcars %>% select(gear, carb) %>% mutate(
 carb2 = rec(carb, rec = "1,2=0;3:8=1"),
 gear2 = rec(gear, rec = "3=1;4:max=2"))
```

# randomizr: : CHEAT SHEET

## Two Arm Trials

**Simple** random assignment is like flipping coins for each unit separately.

```
simple_ra(N = 100, prob = 0.5)
```

**Complete** random assignment allocates a fixed number of units to each condition.

```
complete_ra(N = 100, m = 50)
complete_ra(N = 100, prob = 0.5)
```

**Block** random assignment conducts complete random assignment separately for groups of units.

```
blocks <- rep(c("A", "B", "C"),
 c(50, 100, 200))

defaults to half of each block
block_ra(blocks = blocks)

can change with block_m
block_ra(blocks = blocks,
 block_m = c(20, 30, 40))
```

**Cluster** random assignment allocates whole groups of units to conditions together.

```
clusters <- rep(letters, times = 1:26)
cluster_ra(clusters = clusters)
```

**Block and cluster** random assignment conducts cluster random assignment separately for groups of clusters.

```
clusters <- rep(letters, times = 1:26)
blocks <- rep(paste0("block_", 1:5),
 c(15, 40, 65, 90, 141))
block_and_cluster_ra(blocks = blocks,
 clusters = clusters)
```

## Multi Arm Trials

Set the number of arms with `num_arms` or with `conditions`.

```
complete_ra(N = 100, num_arms = 3)
complete_ra(N = 100, conditions = c("control",
 "placebo", "treatment"))
```

The `*_each` arguments in `randomizr` functions specify design parameters for each arm separately.

```
complete_ra(N = 100, m_each = c(20, 30, 50))
complete_ra(N = 100,
 prob_each = c(0.2, 0.3, 0.5))
```

If the design is the **same** for all blocks, use `prob_each`:

```
blocks <- rep(c("A", "B", "C"),
 c(50, 100, 200))
block_ra(blocks = blocks,
 prob_each = c(.1, .1, .8))
```

If the design is **different** in different blocks, use `block_m_each` or `block_prob_each`:

```
block_m_each <- rbind(c(10, 20, 20),
 c(30, 50, 20),
 c(50, 75, 75))
block_ra(blocks = blocks,
 block_m_each = block_m_each)

block_prob_each <- rbind(c(.1, .1, .8),
 c(.2, .2, .6),
 c(.3, .3, .4))
block_ra(blocks = blocks,
 block_prob_each = block_prob_each)
```

If `conditions` is numeric, the output will be **numeric**.

If `conditions` is not numeric, the output will be a **factor** with levels in the order provided to conditions.

```
complete_ra(N = 100, conditions = -2:2)
complete_ra(N = 100, conditions = c("A", "B"))
```

## Declaration

Learn about assignment procedures by “declaring” them with `declare_ra()`

```
declaration <-
 declare_ra(N = 100, m_each = c(30, 30, 40))
```

`declaration` # print design information

Conduct a random assignment:

```
conduct_ra(declaration)
```

Obtain observed condition probabilities (useful for inverse probability weighting if probabilities of assignment are not constant)

```
Z <- conduct_ra(declaration)
obtain_condition_probabilities(declaration, Z)
```

## Sampling

All assignment functions have sampling analogues: Sampling is identical to a two arm trial where the treatment group is sampled.

| Assignment                          | Sampling                             |
|-------------------------------------|--------------------------------------|
| <code>simple_ra()</code>            | <code>simple_rs()</code>             |
| <code>complete_ra()</code>          | <code>complete_rs()</code>           |
| <code>block_ra()</code>             | <code>strata_rs()</code>             |
| <code>cluster_ra()</code>           | <code>cluster_rs()</code>            |
| <code>block_and_cluster_ra()</code> | <code>strata_and_cluster_rs()</code> |
| <code>declare_ra()</code>           | <code>declare_rs()</code>            |
| <code>conduct_ra()</code>           | <code>draw_rs()</code>               |

## Stata

A Stata version of `randomizr` is available, with the same arguments but different syntax:

```
ssc install randomizr
set obs 100
complete_ra, m(50)
```

`randomizr` is part of the `DeclareDesign` suite of packages for designing, implementing, and analyzing social science research designs.