

Named identifiers

Allowed: letters, digits, period (.) and underscore (_).
Must start with letter or period. If starts with period, cannot be followed by digit.
Cannot use keywords such as TRUE, IF, etc.

```
variable <- "value"  
typeof(variable) -> "logical"
```

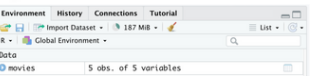
Useful Facts

As you create (i.e., assign) objects in an R session, the objects will appear in the **Environment**.
Use ? to look for help in R.
To install a package, do `install.packages("name of the package")`

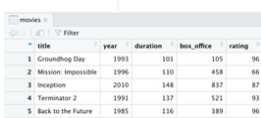
Dataframes

part of dplyr, in turn **tidyverse** package.
a table with **columns representing variables** and **rows representing observations**.
(think of a 2x2 matrix array. **All vectors in data frame must have the same number of elems.**)
Can also use **data.frame()**, a part of R's standard library

```
movies <- data.frame(  
  title = c("Groundhog Day", "Mission: Impossible", "Inception", "Terminator 2", "Back to the Future"),  
  year = c(1993, 1996, 2010, 1991, 1985),  
  duration = c(101, 110, 148, 137, 116),  
  box_office = c(105, 458, 837, 521, 389),  
  rating = c(96, 66, 87, 93, 96)  
)
```



```
movies <- data.frame(  
  title = c("Groundhog Day", "Mission: Impossible", "Inception", "Terminator 2", "Back to the Future"),  
  year = c(1993, 1996, 2010, 1991, 1985),  
  duration = c(101, 110, 148, 137, 116),  
  box_office = c(105, 458, 837, 521, 389),  
  rating = c(96, 66, 87, 93, 96)  
)
```

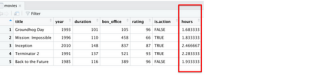


get column values: `name_of_data_frame$name_of_column` `> movies$rating`
[1] 96 66 87 93 96

Suppose you are asked to create a vector whose elements are TRUE if each of the corresponding movies is longer than 2 hours, and FALSE otherwise
`> movies$duration > 120`
[1] FALSE FALSE TRUE TRUE FALSE

create new column: `name_of_dataframe$new_col <- movies$duration/60`

```
moviesHours <- movies$duration/60
```



NULL VALUES

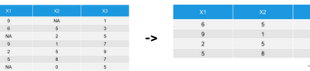
By default, if there is a missing value, the mean will evaluate to NA
`> mean(x)`
[1] NA

`na.rm()` is a logical value indicating whether NA values should be ignored. default is FALSE
`> mean(x, na.rm = TRUE)`
[1] 2.5

`is.na(x)` indicates which elements are missing.
`> is.na(x)`
[1] FALSE FALSE FALSE FALSE TRUE

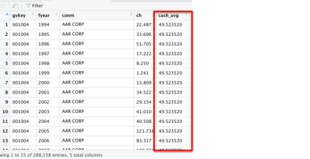
`sum(is.na(x))` adds up all the elements that evaluate to 1, thus giving the **number of missing elements**.is.na() also works with data frames.
`> sum(is.na(x))`
[1] 1

`data_omit <- na.omit(data)` any rows with NA are deleted



`group_by()`
Creates groups of observations based on one or more variables

Ex. to get average cash for each firm, enter:
`companies_with_avg_ch <- companies %>% group_by(gvkey) %>% mutate(cash_avg = mean(ch, na.rm = TRUE))`



`group_by()` WITH TWO VARIABLES

- To do this, enter into the Console:
`filter(!is.na(naichsh), !is.na(fyear)) %>% group_by(naichsh, fyear) %>% summarise(
 min_at = min(at, na.rm = TRUE),
 avg_at = mean(at, na.rm = TRUE),
 max_at = max(at, na.rm = TRUE))`

It is always a good practice to remove observations with NA values in the columns that will be used as group_by variables, before doing a group_by.

`group_by()` WITH TWO VARIABLES

- The first few rows of at_summary will look like the following:



Operations & Data Structures

AND: &
OR: |
NOT: ! Missing Values: NA

Vectors
must be of same data type.
`city <- c("Vancouver", "Victoria", "Delta", "Surrey")`

`length(x)`
`> city <- c("Vancouver", "Victoria", "Delta", "Surrey")
> length(city)
[1] 4`

Subsetting
R IS ONE INDEXED.
Use - to remove elements, : to return range. can assign subsets to new vectors.

`city[2]` returns "Victoria"
`city[1:3]` returns a vector "Vancouver" "Victoria" "Delta"
`city[c(1,2,4)]` returns specified elements "Vancouver", "Victoria", "Surrey"
`city[-3]` returns "Vancouver" "Victoria" "Surrey"

Can update the value(s) of the subset of a vector.
Can also add new elements to an existing vector.

If you perform an operation with vectors, the operation will be applied to each element of the vector.
`> city_temp + 2
[1] 27 22 29 32`

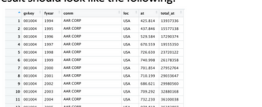
Function
`max(x)` the maximum value of x
`min(x)` the minimum value of x
`range(x)` the minimum and maximum values of x
`length(x)` the number of elements in x
`sum(x)` the sum of the elements in x
`mean(x)` the average value of the elements in x
`sd(x)` the standard deviation of x
`var(x)` the variance of x
`sqrt(x)` the square root of x
`round(value_to_round, decimal_places_to_round)` ex. `round(x/y, 3)`

sum() is another aggregate function that can be used within summarise() or mutate(). Returns the total values of the vector passed as an argument.

Ex. Calculate the sum of total assets of all firms by country of headquarters (loc) and fiscal year (fyear). Call this new variable tot_at. Retain all the observations in the original data.

Can use the sum() function with mutate().

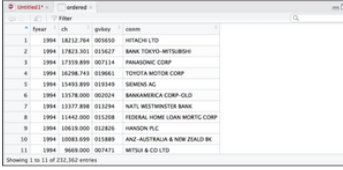
```
tot_at <- companies %>%  
  filter(!is.na(fyear), !is.na(loc)) %>%  
  group_by(fyear, loc) %>%  
  mutate(total_at = sum(at, na.rm = TRUE)) %>%  
  select(gvkey, fyear, comm, loc, at, total_at)
```

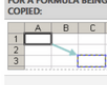


All firms located in the same loc and the same fyear are grouped together as a subset, and then R returns the total of at for each subset – the country/year combination.

arrange() FUNCTION
orders the rows by variable(s) specified, in ascending order.

Ex. Arrange observations in companies in ascending order of fyear and descending order of ch within each fyear.
`ordered <- companies %>%
 filter(!is.na(fyear), !is.na(ch)) %>%
 arrange(fyear, desc(ch)) %>%
 select(fyear, ch, gvkey, comm)`



FOR A FORMULA BEING COPIED:	IF THE REFERENCE IS:	IT CHANGES TO:
	\$A\$1 (absolute column and absolute row)	\$A\$1 (the reference is absolute)
	A\$1 (relative column and absolute row)	C\$1 (the reference is mixed)
	\$A1 (absolute column and relative row)	\$A3 (the reference is mixed)
	A1 (relative column and relative row)	C3 (the reference is relative)

pipe operator
When the pipe operator (i.e., %>%) is used, the dplyr functions (like select() and filter()) take the data from the previous step as their input.

`companies %>% filter(loc=="CAN")` is identical to `filter(companies, loc=="CAN")` examples:

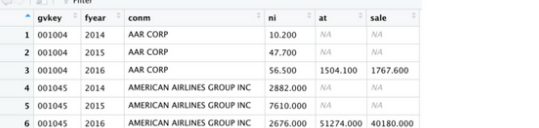
inner_join(dataset_x, dataset_y) FUNCTION
returns all rows from x where there are matching values in y, and all columns from x and y.

`merged1 <- inner_join(example1, example2)`
To specify this match more clearly (which is good practice), you can do the following:
`merged1 <- inner_join(example1, example2, by = c("gvkey", "fyear"))`

left_join(dataset_, dataset_y) FUNCTION
returns all rows left_join(x, y) returns all rows from x, and all columns from x and y

`merged2 <- left_join(example1, example2)`

R will return Joining, by = c("gvkey", "fyear")



Observations in the "left" dataset with no matching observations in the other dataset will be retained in the resultant dataset. Observations in the "right" dataset which do not have matching datasets, will not be returned.

Non One To One Matching

When you merge two datasets whose key variables uniquely identify each observation within the dataset, this kind of merge is called a one-to-one merge.

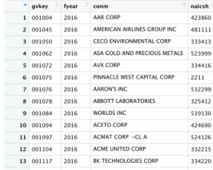
- However, the attribute(s) (i.e., column(s)) used as a basis of matching the observations in two data sets do not have to be key variables for both datasets.

In such situations, an observation in one dataset can match with multiple observations in the other datasets through non-one-to-one matching.

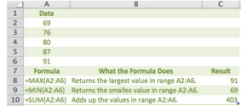
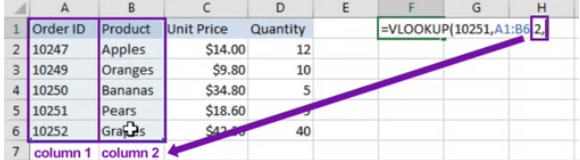
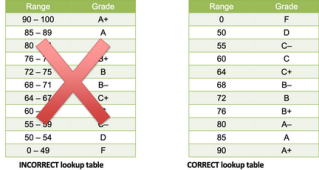
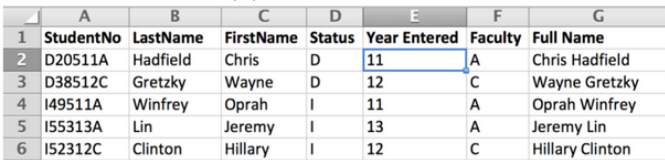
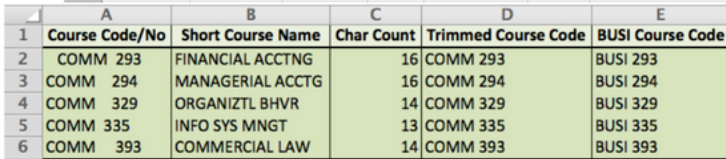
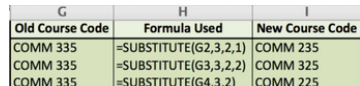
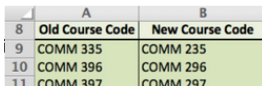
However, the variable or variable(s) to perform the merge should uniquely identify at least one of the datasets.

You want to merge the two datasets, so that each observation will have the following variables: gvkey, fyear, naichsh, and NAICS_description. This will attempt to add the industry description to every observation in the Example3 dataset.

To do this, enter into the Console:
`merged3 <- left_join(example3, NAICS_2_6_digit_codes, by = c("naichsh" = "NAICS"))`



NAICS	NAICS Description
1	11 Agriculture, Forestry, Fishing and Hunting
2	111 Crop Production
3	1111 Oilseed and Grain Farming
4	11111 Soybean Farming
5	111110 Soybean Farming
6	11112 Oilseed (except Soybean) Farming
7	111120 Oilseed (except Soybean) Farming
8	11113 Dry Pea and Bean Farming
9	111130 Dry Pea and Bean Farming
10	11114 Wheat Farming
11	111140 Wheat Farming

What	Syntax	Description	Example	• Some common logical operators
IF	=IF(logical_test, [value_if_true], [value_if_false])	Logical Operator	=IF(B2>D2, "Good", "Bad")	= Equal to
MAX	=MAX(number1, [number2], ...)	Values can be numbers (e.g. 10), a cell reference that contains numbers (e.g. A8), a cell range/array that contains numbers (e.g. A8:A10), or a combination of the three.		<> Not equal to
MIN	=MIN(number1, [number2], ...)			< Less than
SUM	=SUM(number1, [number2], ...)			> Greater than
AND	=AND(logical1, [logical2], ...)	Used to evaluate if all logicals are TRUE or FALSE		<= Less than or equal to
OR	=OR(logical1, [logical2], ...)	Used to evaluate if at least one of the logicals are TRUE or FALSE		>= Greater than or equal to
COUNTIFS	=COUNTIFS(criteria_range1, criteria1, ...)	Count values in a specific range that meet certain criteria. At least 1 criterion is needed, but up to 127 criteria can be tested. Criteria range refers to range of cells you want to test for	Ex. How many NY airports had at least 10 different airlines and at least 1 million passenger departures in 2014? 1) airlines >=10 2) passenger departures >=1000000	
		Criteria can be in the form of a number, cell reference, or text. Do NOT use formulas/functions for criteria..		
			=COUNTIFS(D3:D7,">=10",C3:C7,">=1000000")	
			returns amount that fulfills the criteria.	
SUMIFS	=SUMIFS(sum_range, criteria_range1, criteria1, ...)	Adds values in a specified range or multiple ranges that meet specified criteria. At least 1 criterion is needed, but up to 127 criteria can be tested. Same as COUNTIFS, but sums it up afterward.	Ex. What are the total passenger departures at NY airports that had at least 10 different airlines and at least 1 million passenger departures in 2014? Our sum_range is passenger departures (C3:C7).	
	sum_range : the cells you want to sum criteria_range : the values you're checking criteria1 : what do you want it to equal to?			
	=VLOOKUP(lookup_value, table_array, col_index_num, [range_lookup])			
VLOOKUP	=VLOOKUP(VALUE, table, index_number, approx_match)	The VLOOKUP function looks up a value in the first column of a range of cells, and then returns a value from any cell on the same row of the range.		
	value The value to look for in the first column of the table.			
	table : Two or more columns of data that's sorted in ascending order			
	index_number : The column number in the table from which the matching value must be returned. The first number is one.			
	approximate_match : Optional. FALSE = EXACT MATCH, TRUE = APPROXIMATE MATCH (if omitted, TRUE is DEFAULT)	For example, you can use VLOOKUP to assign a grade of B+ for students who receive a grade from 76 to 79.		
		If the VLOOKUP is used to find an exact match, the rows can be in any order.	Approximate Matches • If you want to look up a value in a range (e.g. 76 to 79), you MUST arrange the lookup table so that the data are sorted from the lowest to the highest. • You MUST only include the lowest value in the range (in this case, 76). It is called the breakpoint . • Do NOT use the complete range in Column 1 (i.e. NOT 76 – 79).	
LEFT	=LEFT(text,num_chars)	LEFT function returns the first character(s) in a string of text based on the specified number of characters.	Faculty is =RIGHT(A2,1) Status is =LEFT(A2,1) Year Entered is =MID(A2,5,2) FULL NAME is =CONCATENATE(C2," ",B2.) OR =C2&" "&B2	
RIGHT	=RIGHT(text,num_chars)			
		return the last character(s) in a string of text based on the specified number of characters.		
MID	=MID(text,start_num,num_chars)	MID function returns character(s) in a string of text, starting at a specified position and based on the specified number of characters.		
	start_num : the (one indexed, inclusive) number to start at. num_chars : the number of characters to include.			
CONCATENATE	=CONCATENATE(text1,[text2],...) =C2&" "&B2 (alternate)	CONCATENATE function join two or more strings of text into one. Up to 255 strings of text can be joined.		
		Instead of using the CONCATENATE function, we could also use the ampersand (&) operator to join multiple strings of text into one.		
LEN	=LEN(text)	Returns the number of characters in a text string. Spaces are counted, too.		
TRIM	=TRIM(text)	TRIM function removes all spaces from text strings, EXCEPT for single spaces between words.		
SUBSTITUTE	=SUBSTITUTE(text,old_text,new_text,[instance_num]) Instance_num : optional. If you use instance number (1, 2, 3, etc.), Excel will only substitute the new_text for the old_text in that specific instance.	Substitutes a new text string for an old text string.	Trimmed Course Code is =TRIM(A2) Char Count is =LEN(A2) BUSI Course Code is =SUBSTITUTE(D2,"COMM","BUSI").	New Course Code is =REPLACE(A9,6,1,2)
REPLACE	=REPLACE(old_text,start_num,num_chars,new_text)	Ex. Dean wants to change COMM 335 to COMM 235. If you decide to use SUBSTITUTE function to change the course number, instance_number is required because you only want to change the number 3 the first time Excel finds the number 3. The instance number is, therefore, 1.		
FIND	=FIND(find_text,within_text,[start_num]) Find_text : refers to the text you want to find. If find_text is not in the within_text, the FIND function will return the error #VALUE! Start_num : is optional. Start_num refers to the character position at which FIND function will start looking for the text.	The FIND function does NOT allow wildcard characters. The FIND function is case sensitive .		
SEARCH	=SEARCH(find_text,within_text,[start_num])	The SEARCH function allows wildcard characters and is NOT case sensitive .		
Wildcard characters				
• ? is used to find any single character. One ? is equal to one character.				
• * is used to find any number of characters.				
• ~ is used to find an actual ?, *, or ~ inside a text string.				