

Tidying Data in the Tidyverse

Megan Lewis

2025-05-29

Introduction to session

- What is “untidy” data?
- What is “tidy” data?
- An intro to the **tidyverse**
- A couple **tidyverse** packages

The problem with data

- Real world data often messy
- Often not collected in tidy formats
- Can slow down analysis and increase frustration

Messy, untidy data

“Happy families are all alike; every unhappy family is unhappy in its own way.” - Leo Tolstoy

“Tidy datasets are all alike, but every messy dataset is messy in its own way.” - Hadley Wickham

According to an [article](#) in Forbes:

“Data preparation accounts for about 80% of the work of data scientists”

Common problems with messy data sets

- Column headers are values but should be variable names
- Single column with multiple variables
- Variables entered in both rows and columns
- Multiple types of observational units stored in the same table
- A single observational unit is stored in multiple tables

Examples from: [Hadley Wickham \(2014\)](#)

Column headers are values but should be variables

religion	<\$10k	\$10–20k	\$20–30k	\$30–40k	\$40–50k	\$50–75k
Agnostic	27	34	60	81	76	137
Atheist	12	27	37	52	35	70
Buddhist	27	21	30	34	33	58
Catholic	418	617	732	670	638	1116
Don't know/refused	15	14	15	11	10	35
Evangelical Prot	575	869	1064	982	881	1486
Hindu	1	9	7	9	11	34
Historically Black Prot	228	244	236	238	197	223
Jehovah's Witness	20	27	24	24	21	30
Jewish	19	19	25	25	30	95

Single column with multiple variables

country	year	column	cases
AD	2000	m014	0
AD	2000	m1524	0
AD	2000	m2534	1
AD	2000	m3544	0
AD	2000	m4554	0
AD	2000	m5564	0
AD	2000	m65	0
AE	2000	m014	2
AE	2000	m1524	4
AE	2000	m2534	4
AE	2000	m3544	6
AE	2000	m4554	5
AE	2000	m5564	12
AE	2000	m65	10
AE	2000	f014	3

Variables entered in both rows and columns

id	year	month	element	d1	d2	d3	d4	d5	d6	d7	d8
MX17004	2010	1	tmax	—	—	—	—	—	—	—	—
MX17004	2010	1	tmin	—	—	—	—	—	—	—	—
MX17004	2010	2	tmax	—	27.3	24.1	—	—	—	—	—
MX17004	2010	2	tmin	—	14.4	14.4	—	—	—	—	—
MX17004	2010	3	tmax	—	—	—	—	32.1	—	—	—
MX17004	2010	3	tmin	—	—	—	—	14.2	—	—	—
MX17004	2010	4	tmax	—	—	—	—	—	—	—	—
MX17004	2010	4	tmin	—	—	—	—	—	—	—	—
MX17004	2010	5	tmax	—	—	—	—	—	—	—	—
MX17004	2010	5	tmin	—	—	—	—	—	—	—	—

Multiple types of observational units stored in the same table

year	artist	track	time	date.entered	wk1	wk2	wk3
2000	2 Pac	Baby Don't Cry	4:22	2000-02-26	87	82	72
2000	2Ge+her	The Hardest Part Of ...	3:15	2000-09-02	91	87	92
2000	3 Doors Down	Kryptonite	3:53	2000-04-08	81	70	68
2000	98~0	Give Me Just One Nig...	3:24	2000-08-19	51	39	34
2000	A*Teens	Dancing Queen	3:44	2000-07-08	97	97	96
2000	Aaliyah	I Don't Wanna	4:15	2000-01-29	84	62	51
2000	Aaliyah	Try Again	4:03	2000-03-18	59	53	38
2000	Adams, Yolanda	Open My Heart	5:30	2000-08-26	76	76	74

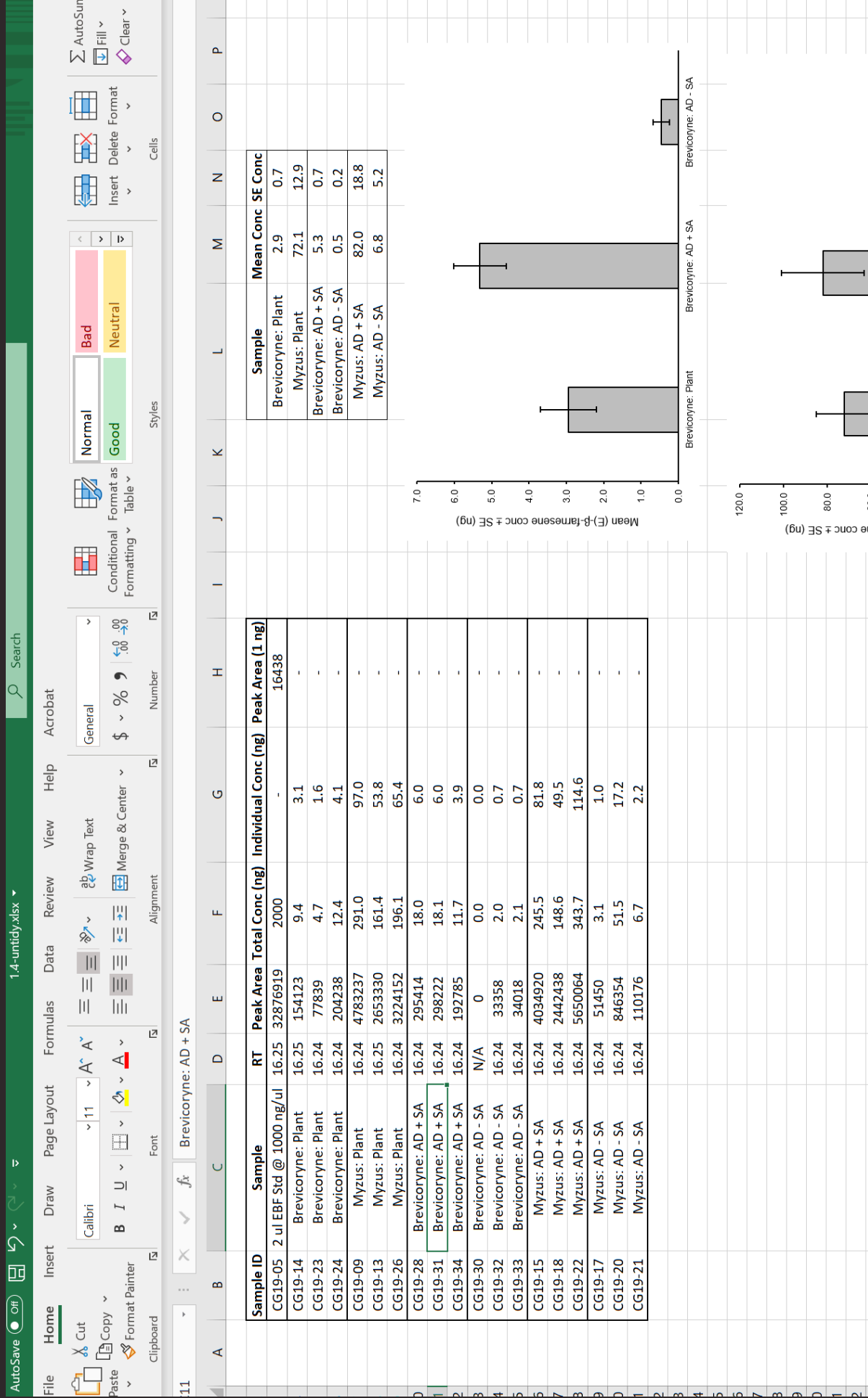
One type in multiple tables

- Values about a single type of observational unit
- In several tables, split by another variable (e.g., year, person, location etc.)

Other issues - Excel formatting...

- Merged cells
- Explanatory text (should be elsewhere)
- Unnecessary table *decoration*
- Summary stats/info

What can untidy data look like?



[illegible]

Why we don't want messy data

- Difficult to analyse
- Easy to make mistakes
- Not reproducible

Messy data might not be perfect, but we can work with it...



But what does “tidy” data look like?

country	year	cases	population
Afghanistan	1999	745	19987071
Afghanistan	2000	1666	20595360
Brazil	1999	37737	172006362
Brazil	2000	80488	174504898
China	1999	212258	1272915272
China	2000	213766	1280428583

Tidy data

One column: One variable

country	year	cases	population
Afghanistan	1999	745	19987071
Afghanistan	2000	1666	20595360
Brazil	1999	37737	172006362
Brazil	2000	80438	174504898
China	1999	212258	1272915272
China	2000	213766	1280428583

Variables

One row: One *observation*

country	year	cases	population
Afghanistan	1999	745	1000707
Afghanistan	2000	1000	20050000
Bhutan	1999	87707	1720000
Bhutan	2000	90100	1715010
China	1999	212259	1272015
China	2000	210700	1200420

Observations

One cell: One *value*

country	year	cases	population
Afghanistan	1999	745	19987071
Afghanistan	2000	1666	20595360
Brazil	1999	37737	172006362
Brazil	2000	80488	174504898
China	1999	212258	1272915272
China	2000	218756	1280428583

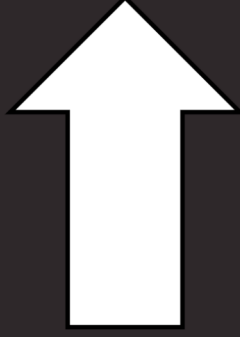
Values

Tidy data visual example

ID	BP1	BP2
A	100	120
B	140	115
C	120	125

Tidy data visual example

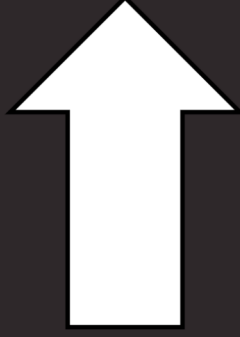
ID	BP1	BP2
A	100	120
B	140	115
C	120	125



ID	measurement	value
A	BP1	100
A	BP2	120
B	BP1	140
B	BP2	115
C	BP1	120
C	BP2	125

Tidy data visual example

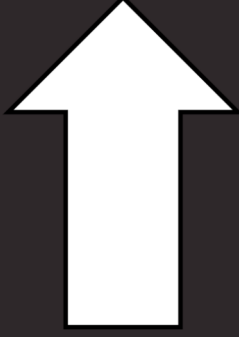
ID	BP1	BP2
A	100	120
B	140	115
C	120	125



ID	measurement	value
A	BP1	100
A	BP2	120
B	BP1	140
B	BP2	115
C	BP1	120
C	BP2	125

Tidy data visual example

ID	BP1	BP2
A	100	120
B	140	115
C	120	125



ID	measurement	value
A	BP1	100
A	BP2	120
B	BP1	140
B	BP2	115
C	BP1	120
C	BP2	125

Why we like ‘tidy’ data

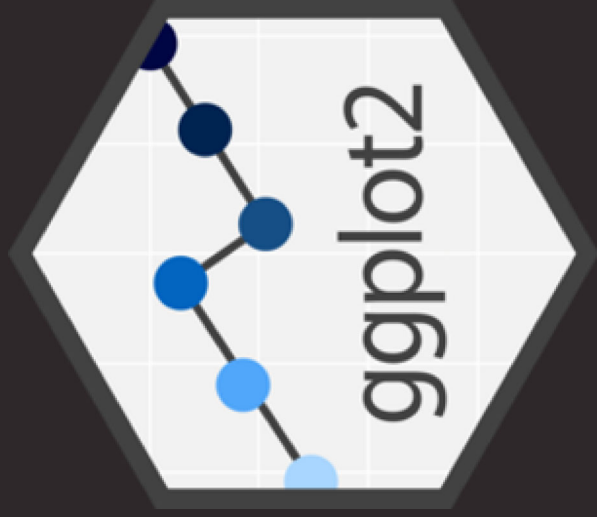
- Easier to wrangle, plot and model
- Gives us a standard structure that works beautifully with the **tidyverse**

Dramatic entrance from the Tidyverse...

- Collection of R packages for data science
- Built around tidy data principles
- Consistent syntax & philosophy



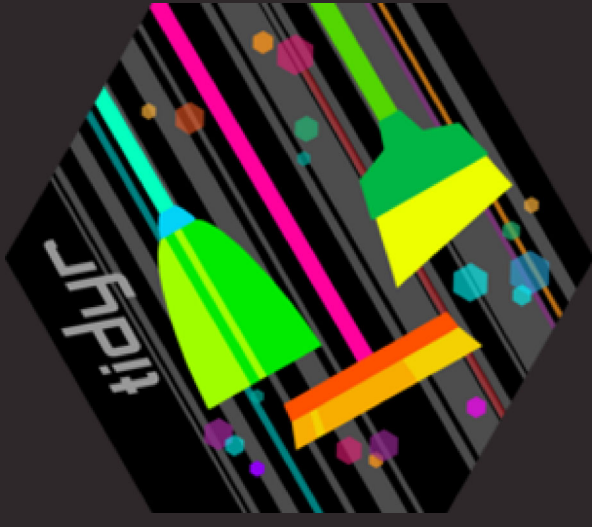
What's in the Tidyverse?



Visualize your
data



Manipulate your
data



Tidy your data

What's in the Tidyverse?



Read
rectangular
data



Work with
functions and
vectors



Re-imagining
the data frame

What's in the Tidyverse?



Work with
strings

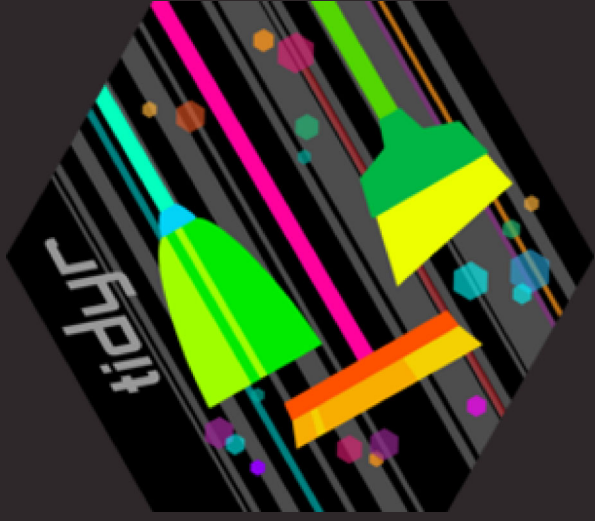


Work with
factors



Work with dates
and times

Today's focus



Tidying



Wrangling

A couple quirks...

The tidyverse can have a couple weird but helpful features...

- Tibbles
- Pipes

Tibbles

- Special type of data frame used in the tidyverse
- Designed for large datasets
 - Only show the first few rows and columns that fit on one screen
 - can use `View()` to see data in interactive, scrollable and filterable view
 - or `glimpse()` to see all variables in console

Pipes

- Two forms:
 - Native pipe `| >`
 - `magrittr` package pipe `%>%`
 - (*package loaded as part of tidyverse*)
- Keyboard short cut
 - `ctrl + shift + m`
 - `cmd + shift + m`
- Need to change R settings to use native pipe with shortcut

Why pipes?

- Takes whats on the left and passes it along to the function on it's right
- Reduces need for nesting or use of intermediate objects
- Native `|>` vs `magrittr %>%`
 - generally behave identically
 - `%>%` was introduced in 2014 in R
 - `|>` included in R in 2021 but now part of base R

Quick syntax with and without pipes

No pipe

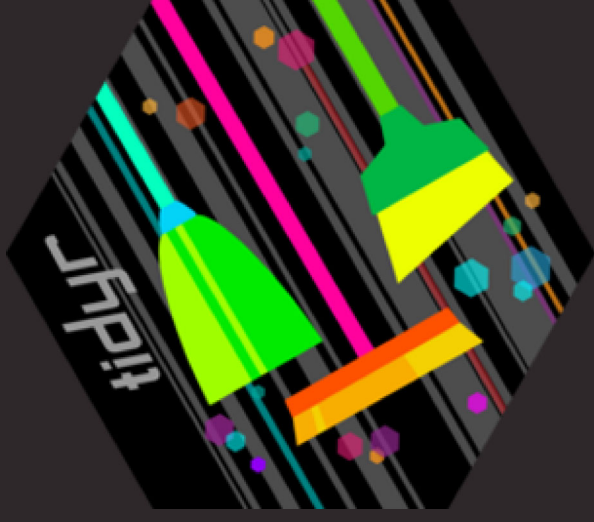
- `function(data, arguments)`
- Data is first argument
- Function wraps around data

With Pipe

- `data | > verb(arguments)`
- Data flows into function
- Easier to read and chain multiple steps

The **tidyr** package

- Help tidy messy data
- Focus on reshaping and reorganizing
- Built around tidy data principles



tidyr functions

- `pivot_longer()`: Reshapes wide data into long format
- `pivot_wider()`: Turns long data into wide format
- `separate()`: Splits one column into multiple based on a delimiter
- `unite()`: Combines multiple columns into one
- `drop_na()` / `replace_na()`: Remove or fill in missing values in a tidy-friendly way

pivot_longer()

tidyr

```
1 # Pivot week columns into long format
2 billboard_long <- billboard_sm |>
3   pivot_longer(
4     cols = starts_with("wk"),
5     names_to = "week",
6     values_to = "rank"
7   )
```

Base R

```
1 # Reshape wide week columns to long
2 billboard_long_base <- reshape(
3   billboard_sm,
4   varying = grep("^wk", names(billboard_sm), value = TRUE),
5   v.names = "rank",
6   timevar = "week",
7   times = grep("^wk", names(billboard_sm), value = TRUE),
8   direction = "long"
9 )
```

separate()

tidyr

```
1 # Split date.entered into components
2 billboard_sep_date <- billboard_long |>
3   separate(date.entered, into = c("year", "month", "day"), sep = "-")
```

Base R

```
1 # Use strsplit to separate date into components
2 date_parts <- do.call(rbind, strsplit(as.character(billboard_long_base$date
3 billboard_long_base$year <- date_parts[, 1]
4 billboard_long_base$month <- date_parts[, 2]
5 billboard_long_base$day <- date_parts[, 3])
```

Demo ft. glittery dung data

Does the addition of glitter to dung affect brood ball production in dung beetles?

Ecological
Entomology



Royal
Entomological
Society

METHODS | Open Access |

A simple technique to assess resource use in dung beetle breeding studies

Megan J. Lewis , Jacob D. Berson, Raphael K. Didham, Theodore A. Evans

First published: 11 September 2023 | <https://doi.org/10.1111/een.13277>

YOU CANNOT POLISH DUNG, BUT YOU CAN ROLL IT IN GLITTER

Ecological
Entomology



631

Demo ft. glittery dung data

- Independent Variables
 - Dung beetle species (*Onthophagus taurus* & *Euoniticellus fulvus*)
 - Four glitter colours + no glitter control
- Response variable
 - Number of brood balls
- Individual dung beetles in individual tubes

Demo!



Introducing **dp1yr**

- Data manipulation in the tidyverse
- Consistent, readable grammar
 - First argument is always a data frame
 - Other arguments describe which column to operate on
 - Output always a new data frame
- Works with tibbles and pipes

Core **dp1**yr verbs

- Each verb does one thing
- Solving complex problems require combining multiple verbs with pipes
- Organised into four groups based on what they operate on:
 - rows
 - columns
 - groups
 - tables

Row verbs

These verbs work across or modify rows of your dataset:

- `filter()` – keep only rows that match a condition
- `arrange()` – reorder rows
- `distinct()` – keep only unique rows

Filtering

dplyr

```
1 # Filter to include only humans
2 starwars |>
3   filter(species == "Human")
```

Base R

```
1 # Filter using indexing
2 starwars[starwars$species == "Human", ]
```

Arranging

dplyr

```
1 # Arrange by height in ascending order
2 starwars |>
3   arrange(height)
```

Base R

```
1 # Use order() to sort by height
2 starwars[order(starwars$height), ]
```

Distinct

dplyr

```
1 # Get unique species values in a data frame
2 starwars |>
3   distinct(species)
```

Base R

```
1 # To get a data frame equivalent to dplyr
2 data.frame(species = unique(starwars$species))
```

Column verbs

These verbs help you work with columns of data:

- `mutate()` – create or transform columns
- `select()` – keep/drop columns
- `rename()` – rename columns
- `relocate()` – reorder columns

Mutate

dplyr

```
1 # Create a new column with height in meters
2 starwars |>
3 mutate(height_m = height / 100)
```

Base R ::: {fragment}

```
1 # Same result using base R column creation
2 starwars$height_m <- starwars$height / 100
```

...

Select

dplyr

```
1 # Select specific columns
2 starwars |>
3   select(name, species, height)
```

Base R

```
1 # Use column indexing
2 starwars[c("name", "species", "height")]
```

Rename

dplyr

```
1 # Rename the 'name' column
2 starwars |>
3   rename(character_name = name)
```

Base R

```
1 # Rename column manually
2 names(starwars)[names(starwars) == "name"] <- "character_name"
```

Relocate

dplyr

```
1 # Move 'species' before 'name'
2 starwars |>
3   relocate(species, .before = name)
```

Base R

```
1 # Reorder columns manually
2 cols <- names(starwars)
3 new_order <- c("species", setdiff(cols, "species"))
4 starwars[new_order]
```

Group verbs

These verbs are for grouped operations, often used in summarising:

- `group_by()` – group dataset by one or more variables
- `summarise()` – calculate summary statistics per group
- `slice_*`() – select specific rows within groups (e.g., `slice_max()`, `slice_head()`)

Group by & Summarise

dplyr

```
1 # Group by species and calculate average height
2 starwars |>
3   group_by(species) |>
4   summarise(mean_height = mean(height, na.rm = TRUE))
```

Base R

```
1 # Use aggregate to get same result
2 aggregate(height ~ species, data = starwars, FUN = mean, na.rm = TRUE)
```

Slice

dplyr

```
1 # Take first 3 rows
2 starwars |>
3   slice(1:3)
```

Base R

```
1 # Base R equivalent
2 starwars[1:3, ]
```

Piping it all together...

dplyr

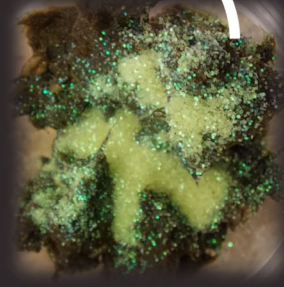
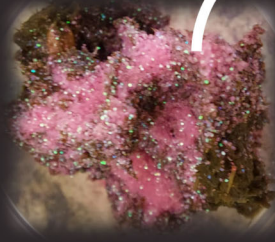
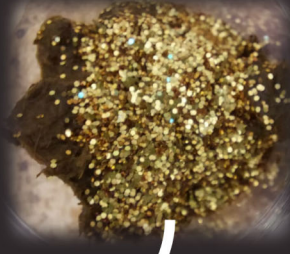
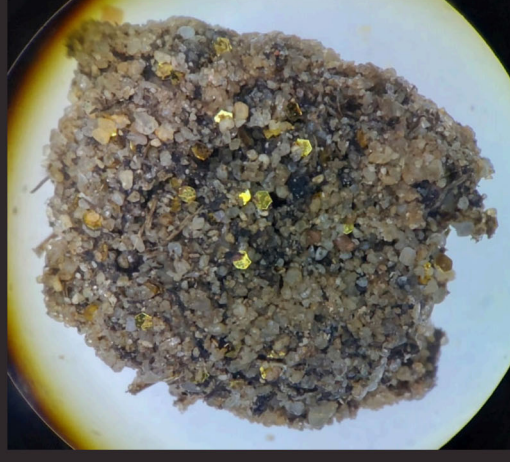
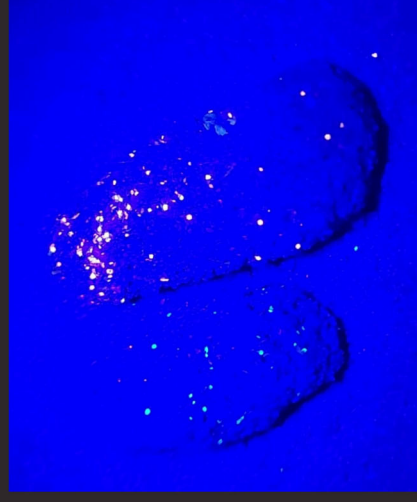
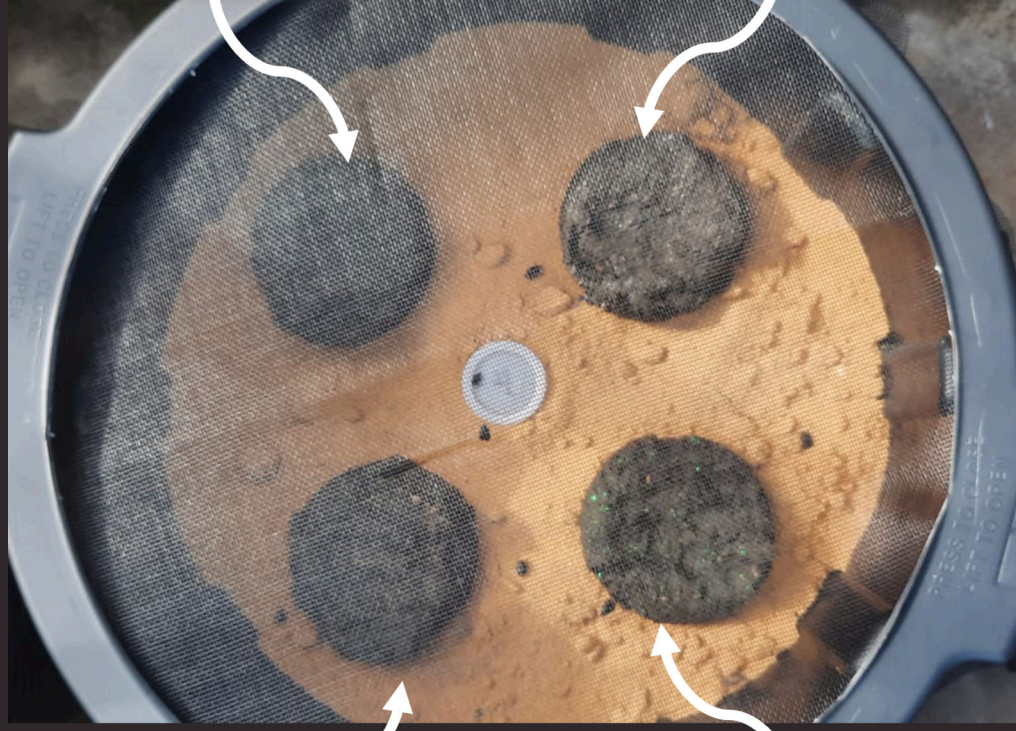
```
1 starwars |>
2   filter(!is.na(mass)) |>           # Remove rows where mass is NA
3   group_by(species) |>             # Group data by species
4   summarise(avg_mass = mean(mass)) |> # Mean mass for each species
5   arrange(desc(avg_mass))         # Sort results by mean mass, descending
```

Base R

```
1 # Filter out NAs
2 filtered <- starwars[!is.na(starwars$mass), ]
3
4 # Mmean mass by species
5 avg_mass <- aggregate(mass ~ species, data = filtered, FUN = mean)
6
7 # Sort descending
8 avg_mass <- avg_mass[order(-avg_mass$mass), ]
```

Live demo - ft. more glittery dung

Will dung beetles show a dung preference based on glitter colour when offered multiple options?



Live demo - ft. more glittery dung

- Subset of data from larger experiment (5 out of 100 bins)
- Four dung pats per bin
 - Each dung pat with different colour glitter
- Six females, six males added per bin
- Seven days to produce brood balls

Demo!

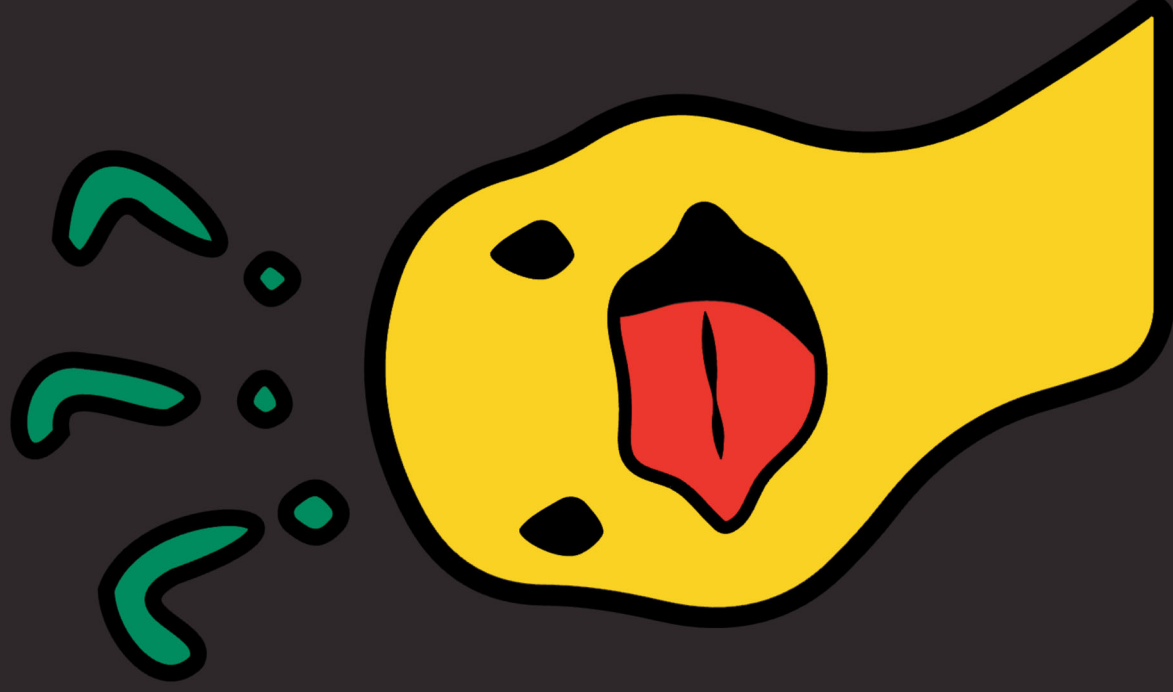


Wrap up

- Messy data is normal – but it can be tamed!
- Tidy data: Smoother analysis, faster, and more reproducible
- Tidyverse: a consistent, powerful toolkit data prep
- Pipes: build readable, step-by-step workflows
- Small, simple verbs — like `filter()`, `mutate()`, `group_by()` — combine to solve complex problems

You can't polish a turd... but you can roll it in glitter and pipe it through `dplyr()`

Questions



Resources

- [Hadley Wickham \(2014\)](#)
- [R for Data Science \(2e\)](#)
- [Tidyverse documentation](#)
- [R Studio cheat sheets](#)
- [Tidy Tuesday Project](#)