

FAIR DATA PRINCIPLES

F

Findable

Consistent names · logical folders · searchable metadata · persistent IDs (DOI, GEO, Ensembl)

A

Accessible

Documented paths · readable formats · open without special software · clear metadata

I

Interoperable

Open formats (`.csv` , `.tsv`) · standard gene IDs · numeric data as numbers, not PDFs

R

Reusable

README · samplesheet · tidy structure · descriptive headers · documented pipeline + parameters

Gene IDs: Use Ensembl IDs (`ENSG00000141510`) as primary keys — never gene symbols. Symbols change across annotation versions and cause silent mismatches.

METADATA — WHAT TO CREATE

DOCUMENT	CONTENTS	LOCATION
<code>samplesheet.csv</code>	One row per file: filename, condition, treatment, replicate, collection date	<code>data/</code>
<code>design.tsv</code>	One row per sample: condition + treatment variables for analysis code	<code>analysis/input/</code>
<code>README.txt</code>	Project overview, directory map, file naming convention, tool versions	Project root

SAMPLESHEET STRUCTURE (TIDY DATA)

FILENAME	LIGHT	GENOTYPE	MEDIA	TIMEPOINT	DATE
<code>LD_phyA_off_t04_2020-08-12.norm.txt</code>	LD	phyA	off	t04	2020-08-12
<code>LD_phyA_on_t04_2020-07-14.norm.txt</code>	LD	phyA	on	t04	2020-07-14
<code>SD_phyB_off_t04_2020-08-13.norm.txt</code>	SD	phyB	off	t04	2020-08-13

... one row per file ...

One variable per column · one observation per row · one value per cell

✓ README MUST-HAVES

- | | |
|---|---|
| ☐ Project title and one-sentence description | ☐ File naming convention — every field defined |
| ☐ Organism (species · strain/cell line) | ☐ Samplesheet location and what columns mean |
| ☐ Reference genome build + annotation version | ☐ Pipeline name + version + key parameters |
| ☐ Date created and point of contact (name + email) | ☐ Statistical cutoffs and normalization method |
| ☐ Directory map — what each folder contains | ☐ Known issues or data caveats |

FILE NAMING RULES

Pattern: `{light}_{genotype}_{media}_{timepoint}_{date}.`
`{datatype}.txt`

RULE	✓ DO THIS	✗ NOT THIS
No spaces	<code>Mock_DMS0_rep1.fastq.gz</code>	<code>Mock DMS0 rep1.fastq.gz</code>
No special characters	<code>results_v2.csv</code>	<code>results (v2)!.csv</code>
Zero-pad	<code>s01, s02, s12</code>	<code>s1, s2, s12</code>
ISO 8601 date	<code>2024-07-16</code>	<code>07-16-24</code> · <code>JUL-16</code>
Consistent case	<code>LD_phyA</code> · <code>LD_phyB</code>	<code>LD_phyA</code> · <code>ld_phyb</code>
Content-based	<code>PCA_all_samples.png</code>	<code>fig_3_a.png</code>
Variables first	<code>LD_phyA_on_t04_2024-07.norm.txt</code>	<code>2024-07-14_s1_LD.txt</code>

PROJECT DIRECTORY TEMPLATE

```
project_name/
├── data/                ← raw data + metadata (READ-ONLY)
│   ├── raw/            ← unprocessed files, never modify
│   ├── norm/           ← normalized / processed outputs
│   ├── samplesheet.csv
│   └── *.fastq.gz       ← reference files
├── analysis/           ← code, pipeline inputs & outputs
│   ├── script.Rmd
│   ├── input/
│   └── output/
├── doc/                ← protocols, notes, documentation
├── results/            ← final figures and tables for paper
└── README.txt          ← who/what/when/where/how + map
```

- One project per directory
- Separate raw from processed — treat raw as read-only
- No `OLD/` , `misc/` , or version number suffixes
- Name files by content, not manuscript position

▶ ONE THING TO FIX THIS WEEK

No README?

Write one for your most-used project. Start with: a description, where is the raw data, what does each folder contain.

Results in .xlsx?

Export your key analysis result to `.csv`

No samplesheet?

Create a `.csv` with one row per file: sample name · condition · replicate · file name. Move collection dates out of filenames and into this table.

Gene symbols as IDs?

Add an `ensembl_id` column (`ENSG...`)

Inconsistent file names?

Write down the convention you wish you had used — pick one and document it in your README. Apply it to all new files going forward.

No folder structure?

Start by separating datatypes into `raw/` and `results/`.