

Regex, the 20 pieces that cover 99%.

D8A.ACADEMY/RESOURCES/REGEX-FOR-DATA

ANCHORS & CLASSES

<code>^start · end\$</code>	edges of the string
<code>\d \w \s</code>	digit word space
<code>.</code>	any single char
<code>[A-Za-z]</code>	build your own class

QUANTIFIERS

<code>+ · * · ?</code>	<code>1+ · 0+ repeats</code>
<code>?</code>	optional
<code>{2,4}</code>	between 2 and 4
<code>.*</code>	← greedy. eats everything.

GROUPS

<code>(abc)</code>	capture it
<code>(?:abc)</code>	group, no capture
<code>a b</code>	or
<code>\1</code>	reuse capture 1

PATTERNS TO STEAL

<code>\d{4}-\d{2}-\d{2}</code>	ISO date
<code>[\w.]+@[\w.]+\. \w+</code>	email-ish
<code>\d+[.,]\d+</code>	decimal number
<code>\s{2,}</code>	extra spaces

IN PANDAS

<code>s.str.contains(r'...')</code>	filter rows
<code>s.str.extract(r'(...)')</code>	pull the group
<code>s.str.replace(r'...', '')</code>	clean

IN SQL

<code>x ~ 'regex'</code>	postgres match
<code>REGEXP_LIKE(x, r)</code>	most other dialects
<code>REGEXP_REPLACE(x, r, '')</code>	clean in place

WORTH MEMORIZING

01 Columns out of chaos

One messy log column in, three clean columns out.

```
# "order #4521 shipped 2026-08-12 (US)"
s = df["raw"]
df["order_id"] = s.str.extract(r"#(\d+)")
df["date"] = s.str.extract(
    r"(\d{4}-\d{2}-\d{2})")
df["country"] = s.str.extract(
    r"\((\w{2})\)")
```

→ extract returns whatever the (group) captures. One pattern per concept beats one giant pattern.

02 Trap: greedy by default

The star eats as much as it can, then walks back. Usually too far.

PATTERN	ON THIS INPUT	IT MATCHES
<code>(.)</code>	<code>a</code> <code>b</code>	<code>a</code> <code>b</code> : too much
<code>(.*?)</code>	<code>a</code> <code>b</code>	<code>a</code> : lazy stops early
<code>^\d+\$</code>	<code>abc123</code>	nothing: anchors save you

→ Add `?` after `*` or `+` to make it lazy. And anchor with `^` `$` whenever you mean the whole string.