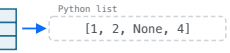


PyArrow

The essentials of PyArrow, Python's gateway to the Arrow columnar format, moving data between pandas, Polars, and Parquet at zero-copy speed.

Arrays & ChunkedArray

The column is the atom; Arrow arrays are typed, contiguous, and null-aware.

PURPOSE	COMMAND	RESULT
Build a typed column from a list.	<code>arr = pa.array([1, 2, None, 4])</code>	
Pin the type explicitly.	<code>pa.array([1, 2, 3], type=pa.int16())</code>	
Stitch many arrays into one column.	<code>ca = pa.chunked_array([[1, 2], [3, 4]])</code>	
Collapse chunks into one buffer.	<code>ca.combine_chunks()</code>	
Slice without copying.	<code>arr.slice(1, 2)</code>	
Read values back to Python.	<code>arr.to_pylist()</code>	

A column is a contiguous typed buffer plus a null bitmap; chunked_array grows it across buffers and slice is a zero-copy view

Table & Schema

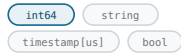
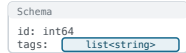
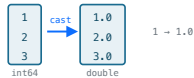
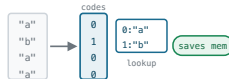
A Table is named columns plus a Schema; columns stay columnar.

PURPOSE	COMMAND	RESULT
Build a Table from a dict.	<code>t = pa.table({"id": [...], "name": [...]}))</code>	
Inspect the schema.	<code>t.schema</code>	
Grab one column.	<code>t.column("id")</code>	
Add a derived column.	<code>t.append_column("flag", pa.array([...]))</code>	
Project a subset of columns.	<code>t.select(["name"])</code>	
Dump back to Python / batches.	<code>t.to_pydict(), t.to_batches()</code>	

pa.table({...}) builds named ChunkedArray columns + a Schema; select / append_column / column stay columnar and cheap

Types & Nullability

Every column has an explicit type and a per-value null bit.

PURPOSE	COMMAND	RESULT
Scalar and temporal types.	<code>pa.int64(), pa.string(), pa.timestamp("us")</code>	
Build a Schema with a list type.	<code>pa.schema([("id", pa.int64()), ("tags", pa.list_(pa.string()))])</code>	
Make a field non-nullable.	<code>pa.field("id", pa.int64(), nullable=False)</code>	
Cast a column to another type.	<code>arr.cast(pa.float64())</code>	
Dictionary-encode (categoricals).	<code>arr.dictionary_encode()</code>	
Nested struct columns.	<code>pa.array([{"a": 1, "b": "z"}])</code>	

Declare type + nullability per field in a Schema; cast changes type, dictionary_encode is Arrow's categorical, struct/list nest columns

Parquet read & write

Parquet is the columnar file format; read whole or just the columns you need.

PURPOSE	COMMAND	RESULT
Write a Table to Parquet.	<code>pq.write_table(t, "data.parquet")</code>	
Read a Parquet file back.	<code>pq.read_table("data.parquet")</code>	
Read only some columns.	<code>pq.read_table("data.parquet", columns=["id"])</code>	
Compress on write.	<code>pq.write_table(t, "data.parquet", compression="zstd")</code>	
Peek metadata without loading.	<code>pq.read_metadata("data.parquet").num_rows</code>	
Write a partitioned dataset.	<code>pq.write_to_dataset(t, "out/", partition_cols=["year"])</code>	

pq.write_table / read_table round-trip a Table; pass columns=[...] to prune, compression= to shrink, read_metadata to peek without loading

Dataset over many files

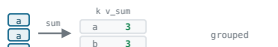
One logical table over a folder; push filters and columns down to the files.

PURPOSE	COMMAND	RESULT
Open a folder as one dataset.	<code>dataset = ds.dataset("out/", format="parquet")</code>	
Read partition keys from paths.	<code>ds.dataset("out/", partitioning="hive")</code>	
Materialize to a Table.	<code>dataset.to_table()</code>	
Push a filter down to files.	<code>dataset.to_table(filter=ds.field("year") == 2024)</code>	
Read only some columns.	<code>dataset.to_table(columns=["v"])</code>	
Stream in batches.	<code>dataset.scanner().head(2)</code>	

`ds.dataset()` is lazy; `to_table(filter=..., columns=...)` pushes predicates and projection down to skip whole partitions and unread bytes

Compute Kernels

`pyarrow.compute` runs vectorized C++ kernels over columns.

PURPOSE	COMMAND	RESULT
Aggregate a whole column.	<code>pc.sum(col), pc.mean(col)</code>	
Element-wise arithmetic.	<code>pc.multiply(col, 2)</code>	
Build a boolean mask.	<code>mask = pc.greater(col, 2)</code>	
Filter rows by a mask.	<code>t.filter(pc.equal(t["k"], "a"))</code>	
Group and aggregate.	<code>t.group_by("k").aggregate([("v", "sum")])</code>	
Sort the Table.	<code>t.sort_by([("v", "descending")])</code>	

`pc.*` are vectorized C++ kernels over columns; Table verbs `filter`, `group_by(...).aggregate(...)`, and `sort_by` run on those same kernels

Bridge to pandas, Polars, numpy

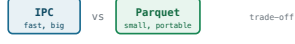
Share the same buffers; the conversion is often zero-copy.

PURPOSE	COMMAND	RESULT
Arrow Table to a pandas DataFrame.	<code>df = t.to_pandas()</code>	
Keep Arrow-backed pandas dtypes.	<code>t.to_pandas(types_mapper=pd.ArrowDtype)</code>	
pandas DataFrame back to Arrow.	<code>pa.Table.from_pandas(df)</code>	
Hand off to / from Polars.	<code>pl.from_arrow(t), pl_df.to_arrow()</code>	
Arrow array to a numpy array.	<code>arr.to_numpy(zero_copy_only=False)</code>	
numpy array to an Arrow array.	<code>pa.array(np.array([1.0, 2.0]))</code>	

Arrow is the shared layout under pandas and Polars, so handoffs reuse the same buffers; nulls into numpy force a copy

IPC & Feather

Write the in-memory layout straight to a file or stream, no re-encoding.

PURPOSE	COMMAND	RESULT
Write a Feather (IPC) file.	<code>feather.write_feather(t, "t.arrow")</code>	
Read a Feather file back.	<code>feather.read_table("t.arrow")</code>	
Write a length-prefixed IPC stream.	<code>with pa.ipc.new_stream(sink, t.schema) as w: w.write_table(t)</code>	
Read an IPC stream.	<code>with pa.ipc.open_stream(buf) as r: r.read_all()</code>	
Random-access IPC file format.	<code>pa.ipc.new_file(...), pa.ipc.open_file(buf)</code>	
When to pick which.	<code>IPC / Feather vs Parquet</code>	

Feather / IPC writes the in-memory layout with almost no encode step: fastest to persist or move a Table; reach for Parquet when you want small, portable files