

DASK FOR PARALLEL COMPUTING CHEAT SHEET



Install Dask with conda	<code>conda install dask</code>
Install Dask with pip	<code>pip install dask[complete]</code>
DASK COLLECTIONS	EASY TO USE BIG DATA COLLECTIONS
DASK DATAFRAMES	PARALLEL PANDAS DATAFRAMES FOR LARGE DATA
Import	<code>import dask.dataframe as dd</code>
Read CSV data	<code>df = dd.read_csv('my-data.*.csv')</code>
Read Parquet data	<code>df = dd.read_parquet('my-data.parquet')</code>
Filter and manipulate data with Pandas syntax	<code>df['z'] = df.x + df.y</code>
Standard groupby aggregations, joins, etc.	<code>result = df.groupby(df.z).y.mean()</code>
Compute result as a Pandas dataframe	<code>out = result.compute()</code>
Or store to CSV, Parquet, or other formats	<code>result.to_parquet('my-output.parquet')</code>
EXAMPLE	<code>df = dd.read_csv('filenames.*.csv') df.groupby(df.timestamp.day) \ .value.mean().compute()</code>
DASK ARRAYS PARALLEL NUMPY ARRAYS FOR LARGE DATA	Import
Create from any array-like object	<code>import dask.array as da</code>
Including HFD5, NetCDF, or other on-disk formats.	<code>import h5py dataset = h5py.File('my-data.hdf5')['/group/dataset'] x = da.from_array(dataset, chunks=(1000, 1000))</code>
Alternatively generate an array from a random distribution.	
Perform operations with NumPy syntax	<code>da.random.uniform(shape=(1e4, 1e4), chunks=(100, 100))</code>
Compute result as a NumPy array	<code>y = x.dot(x.T - 1) - x.mean(axis=0)</code>
Or store to HDF5, NetCDF or other on-disk format	<code>result = y.compute()</code>
EXAMPLE	<code>out = f.create_dataset(...) x.store(out) with h5py.File('my-data.hdf5') as f: x = da.from_array(f['/path'], chunks=(1000, 1000)) x -= x.mean(axis=0) out = f.create_dataset(...) x.store(out)</code>
DASK BAGS	PARALLEL LISTS FOR UNSTRUCTURED DATA
Import	<code>import dask.bag as db</code>
Create Dask Bag from a sequence	<code>b = db.from_sequence(seq, npartitions)</code>
Or read from text formats	<code>b = db.read_text('my-data.*.json')</code>
Map and filter results	<code>import json records = b.map(json.loads) .filter(lambda d: d["name"] == "Alice")</code>
Compute aggregations like mean, count, sum	<code>records.pluck('key-name').mean().compute()</code>
Or store results back to text formats	<code>records.to_textfiles('output.*.json')</code>
EXAMPLE	<code>db.read_text('s3://bucket/my-data.*.json') .map(json.loads) .filter(lambda d: d["name"] == "Alice") .to_textfiles('s3://bucket/output.*.json')</code>

DASK COLLECTIONS (CONTINUED)	
ADVANCED	
Read from distributed file systems or cloud storage	<code>df = dd.read_parquet('s3://bucket/myfile.parquet')</code>
Prepend prefixes like <code>hdfs://</code> , <code>s3://</code> , or <code>gcs://</code> to paths	<code>b = db.read_text('hdfs:///path/to/my-data.*.json')</code>
Persist lazy computations in memory	<code>df = df.persist()</code>
Compute multiple outputs at once	<code>dask.compute(x.min(), x.max())</code>
CUSTOM COMPUTATIONS FOR CUSTOM CODE AND COMPLEX ALGORITHMS	
DASK DELAYED LAZY PARALLELISM FOR CUSTOM CODE	
Import	<code>import dask</code>
Wrap custom functions with the <code>@dask.delayed</code> annotation	<code>@dask.delayed</code> <code>def load(filename):</code> <code>...</code>
Delayed functions operate lazily, producing a task graph rather than executing immediately	<code>@dask.delayed</code> <code>def process(data):</code> <code>...</code>
Passing delayed results to other delayed functions creates dependencies between tasks	<code>load = dask.delayed(load)</code> <code>process = dask.delayed(process)</code>
Call functions in normal code	
Compute results to execute in parallel	<code>data = [load(fn) for fn in filenames]</code> <code>results = [process(d) for d in data]</code>
CONCURRENT FUTURES ASYNCHRONOUS REAL-TIME PARALLELISM	
Import	
Start local Dask Client	<code>from dask.distributed import Client</code>
Submit individual task asynchronously	<code>client = Client()</code>
Block and gather individual result	<code>future = client.submit(func, *args, **kwargs)</code>
Process results as they arrive	<code>result = future.result()</code>
EXAMPLE	<code>for future in as_completed(futures):</code> <code>...</code>
	<code>L = [client.submit(read, fn) for fn in filenames]</code> <code>L = [client.submit(process, future) for future in L]</code> <code>future = client.submit(sum, L)</code> <code>result = future.result()</code>
SET UP CLUSTER HOW TO LAUNCH ON A CLUSTER	
MANUALLY	
Start scheduler on one machine	<code>\$ dask-scheduler</code> Scheduler started at <code>SCHEDULER_ADDRESS:8786</code>
Start workers on other machines	<code>host1\$ dask-worker SCHEDULER_ADDRESS:8786</code>
Provide address of the running scheduler	<code>host2\$ dask-worker SCHEDULER_ADDRESS:8786</code>
Start Client from Python process	<code>from dask.distributed import Client</code> <code>client = Client('SCHEDULER_ADDRESS:8786')</code>
ON A SINGLE MACHINE	
Call <code>Client()</code> with no arguments for easy setup on a single host	<code>client = Client()</code>
CLOUD DEPLOYMENT	
See dask-kubernetes project for Google Cloud	
<code>pip install dask-kubernetes</code>	
See dask-ec2 project for Amazon EC2	
<code>pip install dask-ec2</code>	
MORE RESOURCES	
User Documentation dask.pydata.org	
Technical documentation for distributed scheduler distributed.readthedocs.org	
Report a bug github.com/dask/dask/issues	