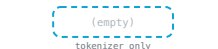


spaCy

An overview of spaCy, the industrial-strength NLP library that turns raw text into tokens, entities, and parses at scale.

Load a pipeline, get a Doc

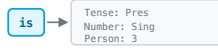
A trained pipeline maps text to an annotated Doc.

PURPOSE	COMMAND	RESULT
Load a trained pipeline.	<code>nlp = spacy.load("en_core_web_sm")</code>	
Process text into a Doc.	<code>doc = nlp("Apple is looking at buying U.K. ...")</code>	
A Doc is a sequence of tokens.	<code>len(doc), doc[0], doc[1:3]</code>	
Split into sentences.	<code>list(doc.sents)</code>	
Get back the original text.	<code>doc.text, token.text_with_ws</code>	
Start blank (no model).	<code>nlp = spacy.blank("en")</code>	

`spacy.load(name)` builds the nlp pipeline; `nlp(text)` returns a Doc you index like a sequence (`doc[i]` Token, `doc[i:j]` Span)

Tokens & Linguistic Features

Each Token carries POS, lemma, dependency, and morphology.

PURPOSE	COMMAND	RESULT
Part-of-speech tag.	<code>token.pos_, token.tag_</code>	
Lemma (dictionary form).	<code>token.lemma_</code>	
Token shape & flags.	<code>token.is_stop, token.is_alpha token.like_num</code>	
Morphological features.	<code>token.morph.to_dict()</code>	
Explain a tag in words.	<code>spacy.explain("VBG")</code>	
Iterate tokens & their heads.	<code>[(t.text, t.pos_, t.head.text) for t in doc]</code>	

String annotations end in an underscore (`pos_`, `lemma_`, `dep_`); the bare name returns an integer hash, not text

Named Entities (doc.ents)

Real-world spans the model labels (ORG, GPE, MONEY, ...).

PURPOSE	COMMAND	RESULT
List the entities.	<code>doc.ents</code>	
Read one entity's fields.	<code>ent.text, ent.label_, ent.start_char</code>	start_char=0 end_char=5
Decode a label.	<code>spacy.explain("GPE")</code>	
Per-token entity tags (BIO).	<code>token.ent_iob_, token.ent_type_</code>	B begins, I inside, O outside
Visualize entities (HTML).	<code>displacy.render(doc, style="ent")</code>	colored highlight marks

doc.ents is a tuple of Span objects; each ent has .text, .label_, and char offsets; per-token the same info is BIO tags

The Dependency Parse

Every token has one head; arcs are typed syntactic relations.

PURPOSE	COMMAND	RESULT
Find the sentence root.	<code>[t for t in doc if t.dep_ == "ROOT"]</code>	no head
Read a token's relation.	<code>token.dep_, token.head</code>	child -> head
Walk a token's children.	<code>list(token.children)</code>	
Extract base noun phrases.	<code>doc.noun_chunks</code>	chunks
Decode a relation label.	<code>spacy.explain("nsubj")</code>	
Visualize the parse (SVG).	<code>displacy.render(doc, style="dep")</code>	arc diagram

The parse is a tree: each token points to one head via a typed dep_ relation; the lone self-headed token is the ROOT

Spans & the Matcher

Slice the Doc into spans; find token patterns with the Matcher.

PURPOSE	COMMAND	RESULT
Slice tokens into a Span.	<code>span = doc[2:4]</code>	
Build a span from char offsets.	<code>doc.char_span(0, 5, label="ORG")</code>	
Token-attribute patterns.	<code>m = Matcher(nlp.vocab)</code> <code>m.add("LOOK", [{"LOWER": "looking"}, ...])</code>	
Run the matcher.	<code>for mid, start, end in m(doc):</code>	
Match exact phrases fast.	<code>pm = PhraseMatcher(nlp.vocab)</code> <code>pm.add("CO", [nlp.make_doc("Apple")])</code>	
Store named span groups.	<code>doc.spans["sc"] = [doc[0:1], doc[5:6]]</code>	

a Span is a slice of a Doc; Matcher finds token-attribute patterns, PhraseMatcher matches exact phrases fast

Similarity & Word Vectors

Compare meaning with vectors. Use md/lg, not sm, for real vectors.

PURPOSE	COMMAND	RESULT
Read a token's vector.	<code>token.vector</code> , <code>token.vector.shape</code>	
Does it have real vectors?	<code>token.has_vector</code> , <code>token.is_oov</code>	
Compare two docs or tokens.	<code>doc1.similarity(doc2)</code> <code>token1.similarity(token2)</code>	
Load a vectors pipeline.	<code>nlp = spacy.load("en_core_web_md")</code>	
Small models warn, not error.	<code>en_core_web_sm</code> tensors, no vectors	
Manage expectations.	<code># similarity reflects training data</code>	

`.vector` and `.similarity()` need real word vectors: load `en_core_web_md` or `_lg`, not `_sm` (tensors only)

Pipeline components

nlp is a sequence of named components you can inspect and edit.

PURPOSE	COMMAND	RESULT						
List the components.	<code>nlp.pipe_names</code>							
Add a built-in component.	<code>ruler = nlp.add_pipe("entity_ruler", before="ner")</code>							
Add rule-based entity patterns.	<code>ruler.add_patterns([{"label": "ORG", "pattern": "spacy"}])</code>							
Run without some components.	<code>with nlp.select_pipes(disable=["ner"]): nlp(text)</code>							
Register a custom component.	<code>@Language.component("my_comp") nlp.add_pipe("my_comp", last=True)</code>							
Inspect what each needs / sets.	<code>nlp.analyze_pipes(pretty=True)</code>	<table border="1"> <thead> <tr> <th>component</th><th>assigns</th><th>requires</th></tr> </thead> <tbody> <tr> <td>tagger</td><td>token.pos</td><td>tok2vec</td></tr> </tbody> </table>	component	assigns	requires	tagger	token.pos	tok2vec
component	assigns	requires						
tagger	token.pos	tok2vec						

nlp is an ordered list of named components: nlp.pipe_names lists them, add_pipe inserts, select_pipes disables, analyze_pipes explains the order

Process at Scale with nlp.pipe

Stream many texts efficiently in batches.

PURPOSE	COMMAND	RESULT
Stream many texts.	<code>docs = nlp.pipe(texts)</code>	
Tune the batch size.	<code>nlp.pipe(texts, batch_size=50)</code>	
Run components in parallel.	<code>nlp.pipe(texts, n_process=4)</code>	
Keep metadata with each text.	<code>nlp.pipe(data, as_tuples=True)</code>	
Skip unused components.	<code>nlp.pipe(texts, disable=["parser", "ner"])</code>	
Process, then collect results.	<code>for doc in nlp.pipe(texts): doc.ents</code>	

nlp.pipe(texts) batches internally for a big speedup; tune batch_size, n_process, as_tuples, and disable instead of looping nlp(text)