

sentence-transformers

sentence-transformers turns sentences into embeddings for semantic search, covered here through similarity, retrieval, and reranking for RAG.

Load a Model

One line pulls a pretrained embedding model from the Hub.

PURPOSE	COMMAND	RESULT
Load a small, fast default model.	<pre>model = SentenceTransformer("sentence-transformers/all-MiniLM-L6-v2")</pre>	
Check the embedding dimension.	<pre>model.get_sentence_embedding_dimension()</pre>	
See the input length limit.	<pre>model.max_seq_length</pre>	
Pick the device (CPU vs GPU).	<pre>SentenceTransformer("all-mpnet-base-v2", device="cuda")</pre>	
Pick a stronger model (slower).	<pre>SentenceTransformer("all-mpnet-base-v2")</pre>	

`SentenceTransformer("name")` loads a Hub model; all-MiniLM-L6-v2 (384 dims) is the fast default, all-mpnet-base-v2 (768) is stronger

Encode Text to Vectors

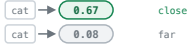
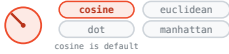
`model.encode` turns strings into dense embedding vectors.

PURPOSE	COMMAND	RESULT
Encode a list of sentences.	<pre>emb = model.encode(["a cat...", "a dog..."])</pre>	
Get a torch tensor instead.	<pre>model.encode(texts, convert_to_tensor=True)</pre>	
Normalize to unit length.	<pre>model.encode(texts, normalize_embeddings=True)</pre>	
Encode a query vs a document.	<pre>model.encode_query(q) model.encode_document(doc)</pre>	
Show progress on long jobs.	<pre>model.encode(texts, show_progress_bar=True)</pre>	
Encode a single string too.	<pre>model.encode("just one sentence")</pre>	

`model.encode(texts)` returns an (n, dim) float32 ndarray; `convert_to_tensor=True` for torch, `normalize_embeddings=True` for cosine = dot

Cosine Similarity

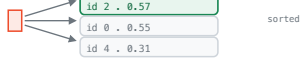
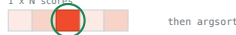
Compare embedding vectors; cosine in [-1, 1], higher means closer.

PURPOSE	COMMAND	RESULT
Score every pair (matrix).	<pre>sim = model.similarity(emb, emb)</pre>	
Read one pair's similarity.	<pre>float(sim[0, 1])</pre>	
Pairwise (row i vs row i only).	<pre>model.similarity_pairwise(a, b)</pre>	
Cosine without a model handle.	<pre>cos_sim(emb[0], emb[1])</pre>	
Switch the similarity metric.	<pre>SentenceTransformer(name, similarity_fn_name="dot")</pre>	

`model.similarity(a, b)` gives the full pairwise cosine matrix; `similarity_pairwise` compares rows one-to-one; cosine is the default metric

Semantic Search

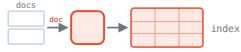
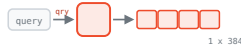
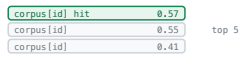
Rank a whole corpus against a query by vector similarity.

PURPOSE	COMMAND	RESULT
Embed the corpus once.	<pre>corpus_emb = model.encode(corpus, convert_to_tensor=True)</pre>	
Embed the query.	<pre>q_emb = model.encode(query, convert_to_tensor=True)</pre>	
Top-k search in one call.	<pre>hits = semantic_search(q_emb, corpus_emb, top_k=3)</pre>	
Score the matrix by hand.	<pre>cos_sim(q_emb, corpus_emb)</pre>	
Use dot product for speed.	<pre>semantic_search(q_emb, corpus_emb, score_function=dot_score)</pre>	

embed the corpus once, embed each query, then `semantic_search` returns the top-k {corpus_id, score} hits sorted best-first

Build & Query an Index

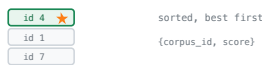
Encode the corpus once, reuse it for every query.

PURPOSE	COMMAND	RESULT
Build the index (asymmetric).	<pre>corpus_emb = model.encode_document(corpus, convert_to_tensor=True)</pre>	
Save the index to disk.	<pre>torch.save(corpus_emb, "corpus.pt")</pre>	
Load it back later.	<pre>corpus_emb = torch.load("corpus.pt")</pre>	
Embed the query side.	<pre>q_emb = model.encode_query("how to reset password")</pre>	
Query the index, read hits.	<pre>for h in semantic_search(q_emb, corpus_emb, top_k=5)[0]:</pre>	
Keep adding documents.	<pre>corpus_emb = torch.cat([corpus_emb, model.encode_document(new)])</pre>	

encode_document builds the index once; torch.save / torch.load persist it; encode_query each incoming query and semantic_search the saved tensor

Rerank with a CrossEncoder

A cross-encoder reads (query, doc) together for sharper scores.

PURPOSE	COMMAND	RESULT
Load a reranker model.	<pre>ce = CrossEncoder("cross-encoder/ms-marco-MiniLM-L6-v2")</pre>	
Score query-document pairs.	<pre>scores = ce.predict([(query, d) for d in docs])</pre>	
Rank candidates directly.	<pre>ranks = ce.rank(query, docs, top_k=3)</pre>	
Get 0..1 probabilities.	<pre>ce.predict(pairs, activation_fn=Sigmoid())</pre>	
Bi-encoder vs cross-encoder.	<pre># two vectors then compare vs one joined pass</pre>	

a CrossEncoder reads (query, doc) jointly for a sharper score; too slow for a whole corpus, so run it only on the top-k from retrieval

Batch Encode at Scale

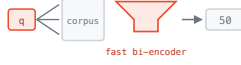
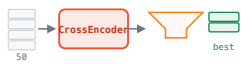
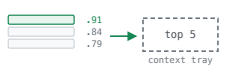
Tune batch size, precision, and parallelism for big corpora.

PURPOSE	COMMAND	RESULT
Encode in tuned batches.	<code>model.encode(texts, batch_size=256, show_progress_bar=True)</code>	
Half precision for GPU speed.	<code>model.half()</code> then <code>model.encode(texts)</code>	
Quantize embeddings (smaller).	<code>model.encode(texts, precision="int8")</code>	
Truncate dims (Matryoshka).	<code>model.encode(texts, truncate_dim=128)</code>	
Parallel multi-process pool.	<code>pool = model.start_multi_process_pool()</code> <code>model.encode(texts, pool=pool)</code>	
Quantize plus normalize on disk.	<code>model.encode(texts, precision="int8", normalize_embeddings=True)</code>	

raise batch_size and show_progress_bar; half()/int8/truncate_dim shrink and speed it up; pass pool= for CPU parallelism

The retrieve-and-rerank (RAG) flow

Bi-encoder retrieves top-k, cross-encoder reranks, then you answer.

PURPOSE	COMMAND	RESULT
1. Retrieve top-k fast.	<code>hits = semantic_search(model.encode_query(q), corpus_emb, top_k=50)</code>	
2. Rerank the candidates.	<code>ranks = ce.rank(q, [corpus[h['corpus_id']] for h in hits[0]])</code>	
3. Take the best passages.	<code>top = [corpus[r['corpus_id']] for r in ranks[:5]]</code>	
4. Build the prompt context.	<code>context = "\n\n".join(top)</code>	
Why two stages.	<code># recall (millions) then precision</code>	
Build the two-stage retriever.	<code>bi = SentenceTransformer(name)</code> <code>ce = CrossEncoder(reranker_name)</code>	

retrieve-and-rerank: a cheap bi-encoder finds many candidates (recall), a cross-encoder reranks the few (precision), the top passages become LLM context