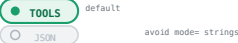


instructor

A look at instructor, which hands back clean, typed Pydantic objects from an LLM instead of raw text, with automatic retries.

Patch a Client

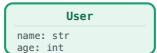
Wrap any OpenAI or Anthropic client; the same instructor surface comes out.

PURPOSE	COMMAND	RESULT
Build + patch from a model string (recommended).	<pre>client = instructor.from_provider("openai/gpt-5-nano")</pre>	 one call: builds + patches
Same call, an Anthropic model.	<pre>client = instructor.from_provider("anthropic/claude-4-5-haiku-latest")</pre>	 agnostic
Patch a client you already configured (OpenAI).	<pre>from openai import OpenAI client = instructor.from_openai(OpenAI())</pre>	 use when you hold a configured client
Patch an Anthropic client object.	<pre>from anthropic import Anthropic client = instructor.from_anthropic(Anthropic())</pre>	 accent ring
Async client, same shape.	<pre>client = instructor.from_provider("openai/gpt-5-nano", async_client=True)</pre>	
Pick the structured-output mode.	<pre>instructor.from_provider(..., mode=instructor.Mode.TOOLS)</pre>	 default avoid mode= strings

from_provider("provider/model") builds + patches in one call; from_openai / from_anthropic wrap a client you already hold

Define a response_model

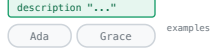
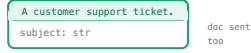
A plain Pydantic model is the schema you want back.

PURPOSE	COMMAND	RESULT
Declare the shape you want.	<pre>class User(BaseModel): name: str age: int</pre>	
Pass it as the response_model.	<pre>client.create(response_model=User, messages=[...])</pre>	
The return value is a real instance.	<pre>user.name is str, user.age is int</pre>	
Make a field optional.	<pre>nickname: str None = None</pre>	
Constrain a value.	<pre>age: int = Field(ge=0, le=130)</pre>	
Restrict to a fixed set.	<pre>role: Literal["admin", "user", "guest"]</pre>	

A response_model is plain Pydantic: typed fields, optional with | None, ranges with Field(ge, le), fixed choices with Literal[...]

Field Descriptions Guide the Model

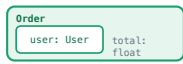
Field(description=...) is prompt engineering inside the schema.

PURPOSE	COMMAND	RESULT
Describe what a field means.	<pre>full_name: str = Field(description="first and last name")</pre>	
Steer the formatting.	<pre>date: str = Field(description="ISO 8601, e.g. 2026-06-18")</pre>	
Constrain and describe together.	<pre>score: int = Field(ge=0, le=100, description="confidence 0-100")</pre>	
Steer with examples.	<pre>Field(description="...", examples=["Ada", "Grace"])</pre>	
The error message guides the re-ask.	<pre>assert "@" in v, "must be an email" # inside a @field_validator</pre>	
Whole-model context.	<pre>"""A customer support ticket."""</pre>	

Field(description=...), examples=, constraints, and the model docstring are all sent to steer output, shape it in the schema not the prompt

Nested + list models

Compose models; ask for many at once.

PURPOSE	COMMAND	RESULT
Nest one model inside another.	<pre>class Order(BaseModel): user: User total: float</pre>	
A list of sub-objects.	<pre>items: list[Item]</pre>	
Extract a list as the result.	<pre>client.create(response_model=list[User], ...)</pre>	
Optional nested block.	<pre>address: Address None = None</pre>	
Deeply typed list value.	<pre>tags: list[Literal["new", "vip"]]</pre>	
Self-documenting nested schema.	<pre>city: str = Field(description=...)</pre>	

Schemas are plain Pydantic, so they compose: nest models, use list[Model] for many objects, and type every element.

Extract One Typed Object

Ragged model text in, a validated object out.

PURPOSE	COMMAND	RESULT
Extract from a sentence.	<pre>user = client.create(response_model=User, messages=[{"role":"user", ...}])</pre>	
Use the OpenAI-shaped call too.	<pre>client.chat.completions.create(response_model=User, messages=[...])</pre>	
Add a system instruction.	<pre>messages=[{"role":"system", "content": "Extract the user."}, {"role":"user", ...}]</pre>	
Anthropic needs max_tokens.	<pre>client.create(response_model=User, max_tokens=1024, messages=[...])</pre>	
Bad input raises an error.	<pre>model returnsage: "old" # raisespydantic.ValidationError</pre>	

create() and chat.completions.create() both accept response_model=; you get back a validated User, not a string to parse

Automatic retries

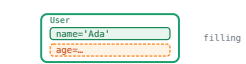
A failed validation is re-asked with the error, not crashed.

PURPOSE	COMMAND	RESULT
Validation failure triggers a retry.	<pre>client.create(response_model=User, max_retries=3, messages=[...])</pre>	
A custom rule that can fail.	<pre>@field_validator("name") def caps(cls, v): assert v.istitle()</pre>	
The model sees the error message.	<pre>"name must be Title Case" -> model</pre>	
Fine-grained retry control (tenacity).	<pre>from tenacity import Retrying, stop_after_attempt max_retries=Retrying(stop=stop_after_attempt(3))</pre>	
Give up after the budget.	<pre>retries exhausted -> InstructorRetryException</pre>	
Omit it: one retry by default.	<pre>client.create(response_model=User, messages=[...])</pre>	

max_retries feeds the validation error back to the model and re-asks; a tenacity Retrying(...) gives full control, InstructorRetryException raises when the budget runs out

Stream Partial Objects

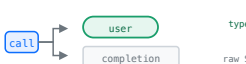
Watch fields fill in live instead of waiting for the whole object.

PURPOSE	COMMAND	RESULT
Stream one object as it fills.	<pre>stream = client.create_partial(response_model=User, stream=True, ...)</pre>	
Consume the partial stream.	<pre>for partial in stream: print(partial)</pre>	
Stream many objects (iterable).	<pre>for user in client.create_iterable(response_model=User, messages=[...]): ...</pre>	
Async streaming.	<pre>async for partial in stream: ...</pre>	
Literal fields need the mixin.	<pre>class User(BaseModel, PartialLiteralMixin):</pre>	
Final frame is a normal object.	<pre># last yielded value is a complete User</pre>	

`create_partial(stream=True)` yields one object filling in; `create_iterable` streams many; the last value is a complete, validated object

with_completion (raw access)

Get the typed object AND the raw response together.

PURPOSE	COMMAND	RESULT
Get object + raw completion.	<pre>user, completion = client.create_with_completion(response_model=User, messages=[...])</pre>	
Read token usage.	<pre>completion.usage</pre>	
Reach the raw model id / metadata.	<pre>completion.model, completion.id</pre>	
Inspect what the model sent.	<pre>completion.choices[0].message</pre>	
Hook into every call (logging).	<pre>client.on("completion:kwargs", log_fn) client.on("completion:error", err_fn)</pre>	
The object is unchanged.	<pre>user == create(...) result</pre>	

`create_with_completion` returns (object, completion); read `completion.usage` / `.model` / `.id` for the raw SDK response, hook calls with `client.on(...)`