

Anthropic Claude SDK

The Anthropic Python SDK at a glance, from a single Claude message to streaming, tool use, images, and prompt caching.

Client & First Message

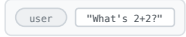
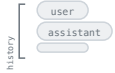
Construct the client, send one turn, read the text back.

PURPOSE	COMMAND	RESULT
Construct the client (reads env key).	<pre>client = anthropic.Anthropic()</pre>	
Send the simplest message.	<pre>msg = client.messages.create(model=..., max_tokens=1024, messages=[{...}])</pre>	
Read the first text block.	<pre>msg.content[0].text</pre>	
Walk all text blocks safely.	<pre>for b in msg.content: if b.type == "text": print(b.text)</pre>	
See why it stopped.	<pre>msg.stop_reason</pre>	
Check token usage.	<pre>msg.usage.input_tokens, .output_tokens</pre>	

`client.messages.create(...)` returns a `Message` whose `.content` is a list of typed blocks; read `.stop_reason` and `.usage` for why and what it cost

Messages & Roles

The conversation is a list of role-tagged turns you resend every call.

PURPOSE	COMMAND	RESULT
One user turn (shorthand).	<pre>messages = [{"role": "user", "content": "..."}]</pre>	
Multi-turn history (alternating).	<pre>messages = [{"role": "user", ...}, {"role": "assistant", ...}, ...]</pre>	
Append the model's reply back.	<pre>messages.append({"role": "assistant", ...})</pre>	
Add the next user turn.	<pre>messages.append({"role": "user", ...})</pre>	
Content can be a block list.	<pre>"content": [{"type": "text", "text": "Hi"}]</pre>	
First turn must be user.	<pre>messages[0]["role"] == "user"</pre>	

The API is stateless: a conversation is a list of `{role, content}` turns, always starting with user, that you resend in full every call

System Prompt & Parameters

Steer behavior with system, bound output with max_tokens, deepen reasoning with thinking.

PURPOSE	COMMAND	RESULT
Set a system prompt.	<pre>create(..., system="You are terse.")</pre>	
Bound the output length.	<pre>create(..., max_tokens=1024)</pre>	
Turn on adaptive thinking.	<pre>thinking={"type": "adaptive"}</pre>	
Tune effort vs cost.	<pre>output_config={"effort": "high"}</pre>	
Show summarized reasoning.	<pre>thinking={"type": "adaptive", "display": "summarized"}</pre>	
Pick the model.	<pre>model="claude-opus-4-8"</pre>	

system steers, max_tokens is a required hard cap; thinking={"type": "adaptive"} deepens reasoning (hidden unless display="summarized")

Streaming

Stream tokens as they arrive; collect the final Message at the end.

PURPOSE	COMMAND	RESULT
Open a streaming context.	<pre>with client.messages.stream(model=..., max_tokens=1024, messages=msgs) as stream:</pre>	
Print text as it arrives.	<pre>for text in stream.text_stream: print(text, end="", flush=True)</pre>	
Get the complete Message at the end.	<pre>final = stream.get_final_message()</pre>	
Handle event types by hand.	<pre>for event in stream: if event.type == "content_block_delta": ...</pre>	
Why stream at all.	<pre># long outputs / high max_tokens</pre>	
Stream with thinking visible. thinking display="summarized"	<pre>thinking={"type": "adaptive", "display": "summarized"}; stream.text_stream</pre>	

with client.messages.stream(...) as stream: iterate stream.text_stream for tokens, then get_final_message() for the whole Message

Tool Use

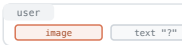
Define tools, get a tool_use block, run it, send a tool_result back.

PURPOSE	COMMAND	RESULT
Define a tool (JSON schema).	<pre>tools = [{"name": "get_weather", "input_schema": {"city": str}}]</pre>	
Offer the tools on a call.	<pre>messages.create(..., tools=tools)</pre>	
Model asks to call a tool.	<pre>msg.stop_reason == "tool_use"</pre>	
Run it and build the result.	<pre>{"type": "tool_result", "tool_use_id": block.id, "content": ...}</pre>	
Send the result, loop until done.	<pre>messages += [assistant, tool_result] create(...) again until end_turn</pre>	
Let the SDK run the loop.	<pre>client.beta.messages.tool_runner(...)</pre>	

tools= sends JSON schemas; stop_reason "tool_use" carries a tool_use block; reply with a matching tool_result and loop to end_turn

Multimodal (Images)

Put image blocks in the content list, base64 or URL, alongside text.

PURPOSE	COMMAND	RESULT
Encode a local image.	<pre>data = base64.standard_b64encode(open("cat.png", "rb").read()).decode()</pre>	
Send a base64 image block.	<pre>{"type": "image", "source": {"type": "base64", "media_type": "image/png", "data": data}}</pre>	
Send an image by URL.	<pre>{"type": "image", "source": {"type": "url", "url": "https://.../cat.png"}}</pre>	
Mix image and a question.	<pre>content=[image_block, {"type": "text", ...}]</pre>	
Read the description back.	<pre>msg.content[0].text</pre>	
Supported formats note.	<pre># jpeg, png, gif, webp</pre>	

An image is just another content block in a user turn: base64 or url, sat beside your text block, answered as ordinary text

Prompt Caching

Mark a stable prefix with `cache_control` to reuse it cheaply across calls.

PURPOSE	COMMAND	RESULT
Cache a big system prefix.	<pre>system=[{"type":"text","text":BIG_DOC,"cache_control":{"type":"ephemeral"}}]</pre>	
Let the SDK auto-place it.	<pre>client.messages.create(..., cache_control={"type":"ephemeral"})</pre>	
First call writes the cache.	<pre>msg.usage.cache_creation_input_tokens</pre>	
Later calls read the cache.	<pre>msg.usage.cache_read_input_tokens</pre>	
Pick a longer TTL.	<pre>"cache_control":{"type":"ephemeral","ttl":"1h"}</pre>	
Cache is a strict prefix match.	<pre># any byte change before the mark = miss</pre>	

`cache_control` marks a stable prefix: the first call writes it (~1.25x), later identical-prefix calls read it (~0.1x); any byte change before the mark misses

Token Counting & Models

Count before you send; pick the model by cost, context, and capability.

PURPOSE	COMMAND	RESULT
Count tokens before sending.	<pre>client.messages.count_tokens(messages=...)</pre>	
Read the count.	<pre>count.input_tokens</pre>	
List available models.	<pre>client.models.list()</pre>	
Inspect one model's limits.	<pre>client.models.retrieve(id).max_input_tokens</pre>	
Choose by tier.	<pre>model=opus vs sonnet vs haiku</pre>	
Avoid deprecated spellings.	<pre>budget_tokens= · temperature= thinking={"type": "adaptive"}</pre>	

`count_tokens` sizes a prompt cheaply; `models.list` / `retrieve` read live limits; pick `opus` / `sonnet` / `haiku` by capability, context, and cost