

Lesson 2.1 – Introduction to the ChatGPT Block

Duration

50–55 minutes

Learning Objectives

1. **Identify** each parameter of the ChatGPT Request block (prompt, result variable, temperature, length, mode, session).
 2. **Experiment** with setting different parameter values (especially temperature, length, and session) to see how they affect ChatGPT's responses.
 3. **Recognize** that ChatGPT's output can be influenced by the prompt (brief mention of prompt quality).
-

Preparation Steps

Resource Materials

- **Tutorial Reference:** [“The ChatGPT Block”](#) from CreatiCode's documentation.
 - **Parameter Cheat Sheet:** A concise handout or slide listing each parameter (prompt, result, mode, length, temperature, session), with brief definitions. See the appendix of this document.
-

Lesson Outline

1. Introduction to ChatGPT & the Request Block (5 minutes)

1. What is ChatGPT?

- ChatGPT is a text-generating AI from OpenAI that can “chat” or generate text based on a prompt.
- ChatGPT works by predicting the next word repeatedly based on the vast amount of training data it has seen. It does not truly understand the meaning of the sentences. For example, the word that follows “United States of” is most likely “America”.
- Mention that it can produce mistakes (hallucinations) and that “prompts” (our request) can significantly affect its output.

2. The ChatGPT Block Overview

- Show the block in the “AI” category.
 - Emphasize that **all** parameters in this block are critical to controlling ChatGPT’s behavior.
-

2. Parameter Deep Dive (25 minutes)

2.1 – Request (Prompt)

- **Definition:** The text or question we give to ChatGPT.
- **Brief Tip:** More specific prompts generally yield better answers, but *we’ll only touch on this briefly* today.

2.2 – Result (Variable)

- **Definition:** The variable where ChatGPT’s response is stored.
- **Demo:** Turn on the variable monitor to watch the response appear.
- **Teacher Demo:** input the prompt “hello” and then click this block to run it, and display the value of the response variable on the stage. Then change the prompt to “write a poem” and run this block again.
- **Student Practice:** try some new prompts like questions or writing some text.

2.3 – Mode (Waiting vs. Streaming)

- **Waiting**
 - The current script **pauses** at this block until it receives the entire response from ChatGPT and stores it in the result variable. This waiting behavior is similar to the “broadcast [MESSAGE] and wait” block.
 - Useful when we’re okay waiting for a complete answer before continuing the script.

- **Streaming**
 - The block **doesn't pause**, so other blocks below this block will continue to run immediately before any response from ChatGPT is received;
 - ChatGPT's response is *streamed* into the variable in chunks.
 - Initially, the variable will be empty. Over time, it will fill up. You can see a partial response appear and get longer and longer. The response will end with an extra check mark emoji ✅, which can be used to check if the entire response has been received.
 - Good for longer responses or "live" chat experiences.
- **Teacher Demo**: Create or select a variable called "result."
 - Run this block with "mode = waiting." Use a short prompt, e.g., "Write a 10-line poem about cats." Make sure the "result" variable is displayed on the stage.
 - Observe that the script pauses, then "result" suddenly contains the whole poem.
 - Try again with mode = streaming and observe that the result will stream in.
- **Student Practice**
 - Encourage students to replicate the result with different prompts.

2.4 – Length (Maximum Tokens) (**obsolete**)

- **Definition**: A cap on how long the response can be (a "token" ~ word/part-of-word).
- **Tip**: If set too low, answers might be cut off. Also, it is a common misunderstanding that this will control the length of the response. It is **ONLY** a limit, and the length of the response is actually determined by the request and other parameters like temperature. For example, if the request says "be concise" or "respond within 50 words", then the response tends to be shorter.

2.5 – Temperature (**obsolete**)

- **Definition**: A number between 0 and 1 that controls how "creative" or "random" the AI can get.
 - **0**: Generate the most deterministic, predictable responses.
 - **1**: Generate the most creative, random responses.

2.6 – Session Type (New Chat vs. Continue)

- **Definition**: Whether ChatGPT will treat the current request as completely new or continue a prior conversation. If "continue", previous chat messages are sent to ChatGPT along with the current prompt, so it can use them as the "context".

- **Note:** When the green flag button is clicked, the first time this block is used, a new session will start with no prior conversation, even if you select “continue”. This helps ensure a new run will not be affected by previous runs.
 - **Note:** There is a token limit on the length of prior chat messages that can be handled by ChatGPT. If the conversation reaches this limit, earlier messages will be discarded.
 - **Teacher Demo**
 - Show how “New Chat” ignores previous context: ask, “What is 2+2?” (ChatGPT responds 4). Then ask “Is that an even number?” in “new” session—ChatGPT might be confused.
 - Switch to “Continue” and ask, “Is that an even number?”—now it remembers context from the prior question.
 - **Student Practice**
 - Students do a short 2-step Q&A. “What’s your favorite color?” followed by “Why do you like that color?”
 - First with “new chat” each time (ChatGPT forgets prior answers).
 - Then, with “continue” each time (ChatGPT builds a mini conversation).
-

3. Basics of Prompting (15 minutes)

- **Teacher Overview**
 - Remind students that **prompt design** is essential for good responses.
 - **Teacher Demonstration**
 - Write an accurate and detailed prompt
 - Assign a role to ChatGPT
 - **Student Practice**
 - Compare results from a generic prompt and a much more detailed prompt
 - E.g. “recommend 3 birthday gifts” vs “recommend 3 birthday gifts for a 10-year boy who lives in Spain and loves soccer.”
 - Compare results from a generic prompt and a prompt with role assignment
 - E.g. “Explain solar power to me” vs “You are a 5th grade teacher and I’m your student. Explain solar power to me.”
-

4. Wrap-Up & Q&A (5 minutes)

- **Recap**
 - Highlight each parameter: **Prompt, Result Variable, Mode (waiting vs. streaming), Length, Temperature, Session.**
- **Questions**
 - Invite students to share any surprising results they got or ask clarification about how the parameters work.

Extensions & Differentiation

For Advanced Students

- **Complex Prompts:** Let them try multi-sentence prompts with more detail to see how the response changes.

For Students Needing Extra Support

- **Guided Parameters:** Provide a short list of recommended parameter values (e.g., “Try temperature=1, length=50, mode=streaming”).
- **Pair Work:** Work in pairs so one student controls parameters while the other observes and records changes.

Notes for Teachers

- The “system request” block is mentioned in the tutorial, but it will be introduced in later lessons.

Assessment

Wayground format:

https://wayground.com/admin/assessment/6914e4879c733752b2ef3637?source=lesson_share

Questions (1 point each)

Q1. What is the purpose of the **Result** parameter in the ChatGPT Request block?

- A) It sets how creative the response is
- B) It selects which variable stores ChatGPT’s reply
- C) It chooses between new chat and continue
- D) It sets the maximum number of tokens

Q2. In **streaming** mode, which statement is correct?

- A) The script pauses until the full response is ready
- B) The result variable fills in gradually while the script keeps running
- C) The response arrives all at once into the variable
- D) Streaming is only for very short answers

Q3. What does **Length (maximum tokens)** do?

- A) It exactly forces the response to be that long
- B) It's just a suggestion; the model can ignore it
- C) It's a hard cap; too small can truncate the answer
- D) It makes the model more creative

Q4. Briefly explain the difference between **New Chat** and **Continue** session types, and give one reason to use each.

Q5. You set **Length = 20** and see "TOKEN LIMIT REACHED." Name **one** change you can make to get a complete answer next time.

Answers and Rubrics (1 point each)

- Q1: B — It selects the variable that stores ChatGPT's reply text. (1 pt)
- Q2: B — The result streams into the variable while the script continues. (1 pt)
- Q3: C — Length is a hard upper bound; too low can cut off output. (1 pt)
- Q4 (1 pt)
 - 1.0: States that New Chat starts without prior context; Continue includes previous messages as context. Gives a sensible reason for each (e.g., New Chat for fresh, unrelated requests; Continue to keep a conversation thread).
 - 0.5: Mentions only one correctly or vaguely references "memory" without a clear reason.

- *0.0: Incorrect or off-topic.*
- **Q5 (1 pt)**
 - *1.0: Provides one valid fix, such as increase Length (more tokens) or ask for a shorter response in the prompt (e.g., “answer in 50 words”).*
 - *0.5: Generic fix without specifics (e.g., “change parameters”).*
 - *0.0: Incorrect or off-topic.*

Appendix - ChatGPT Block Parameter Cheat Sheet

1. Prompt

- **Definition:**
The text or question you send to ChatGPT to generate a response.
 - **Example:**
"Tell me a fun fact about space."
-

2. Result Variable

- **Definition:**
The variable that stores ChatGPT's response. You can monitor this variable to see the AI's answer.
 - **Example:**
Variable Name: `result`
-

3. Mode

- **Definition:**
Determines how ChatGPT sends back its response.
 - **Options:**
 - **Waiting:**
 - **Behavior:** The program pauses until the full response is received.
 - **Use When:** You want the complete answer before moving on.
 - **Streaming:**
 - **Behavior:** The response is sent piece-by-piece as it's generated.
 - **Use When:** You want to see the answer build up in real time.
 - **Visual Cue:**
 - **Waiting Mode:** Response appears all at once.
 - **Streaming Mode:** Response appears gradually.
-

4. Length (Maximum Tokens)

- **Definition:**
Sets the maximum size of ChatGPT's response.
 - **Details:**
 - **Token:** Approximately a word or part of a word.
 - **Function:** It does not directly control the length of the response but only serves as an upper limit. To force a shorter answer, explicitly say so in the prompt.
 - **Recommendation:**
 - Start with a moderate length (e.g., 500 tokens) and adjust based on your needs.
-

5. Temperature

- **Definition:**
Controls the creativity and randomness of ChatGPT's responses.
 - **Range:**
 - 0 to 1
 - **Settings:**
 - 0:
 - **Effect:** Responses are very predictable and focused.
 - **Use When:** You need clear and accurate answers.
 - 1:
 - **Effect:** Highly creative and varied responses.
 - **Use When:** You want imaginative or diverse answers.
-

6. Session

- **Definition:**
Decides whether ChatGPT remembers previous conversations.
- **Options:**
 - **New Chat:**
 - **Behavior:** Starts a fresh conversation without any memory of past interactions.
 - **Use When:** You want an independent response.
 - **Continue:**
 - **Behavior:** Remembers the context from previous messages in the session.
 - **Use When:** You want a coherent, ongoing conversation.
- **Note:**
 - **Session Start:**
 - When the green flag button is clicked, a new session will start with no prior conversation, even if you select "continue"
 - **Token Limit:**

- **Problem:** Conversations with many messages can hit a limit, causing ChatGPT to forget earlier messages.