

Single Day Event

Time to Engineer Your Context

Xiaochang (Jack) Liu
Johnson & Johnson



Disclaimer

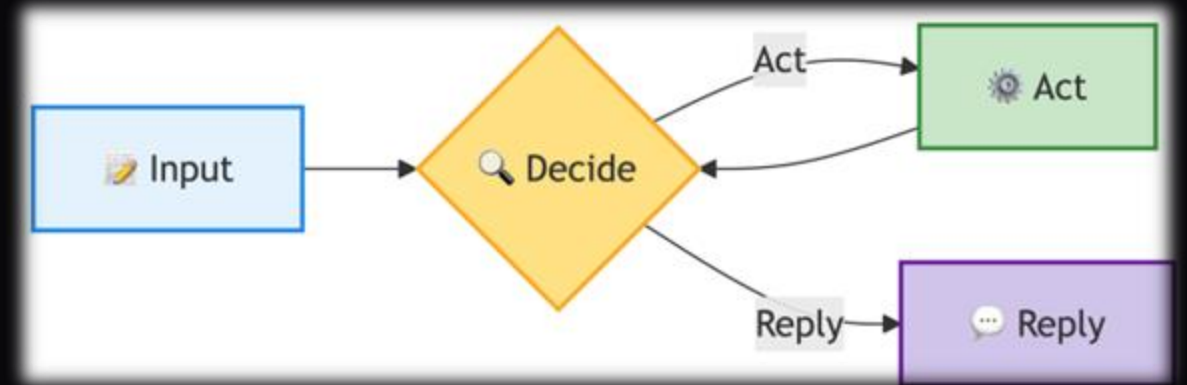
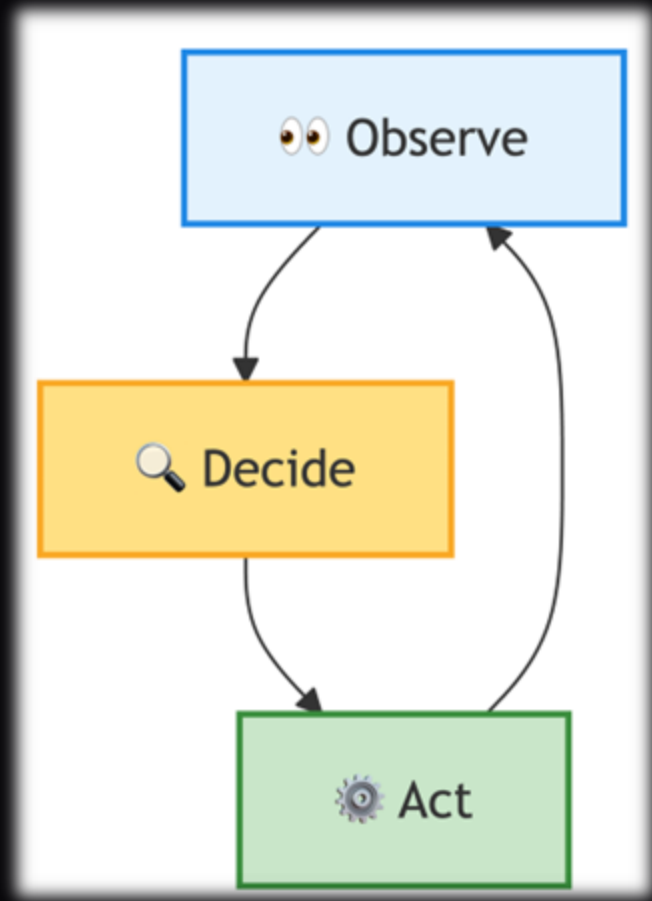
This presentation is for informational purposes only and does not represent professional guidance or advice. Any views and opinions expressed during this presentation are those of the presenters and do not necessarily reflect the views or policies of Johnson & Johnson. This presentation may not be reproduced, modified, or copied, in whole or in part, without the express written permission of Johnson & Johnson.



Time to Engineer Your Context

Xiaochang (Jack) Liu
Johnson & Johnson

LLM AGENT



Simplified Reasoning and Acting (ReAct) Agent Flow

Agent is a Loop!

WHAT IS CONTEXT

“Context = Everything that helps the LLM understand *what you mean* and *what matters*.”



Context for LLM Agent

CONTEXT WINDOW

User: "What's THIS FILE about?"

LLM: "I need to read this file: /path/to/file"

LLM: "use tool: read_file('/path/to/file')"

Tool: "File Content"

User: "Modify the file to implement ..."

LLM: "I need to modify this file: /path/to/file"

LLM: "use tool: write_file('/path/to/file', 'content')"

Tool: "Success"

User: "How about this data..."

LLM: "I need to get data from this API"

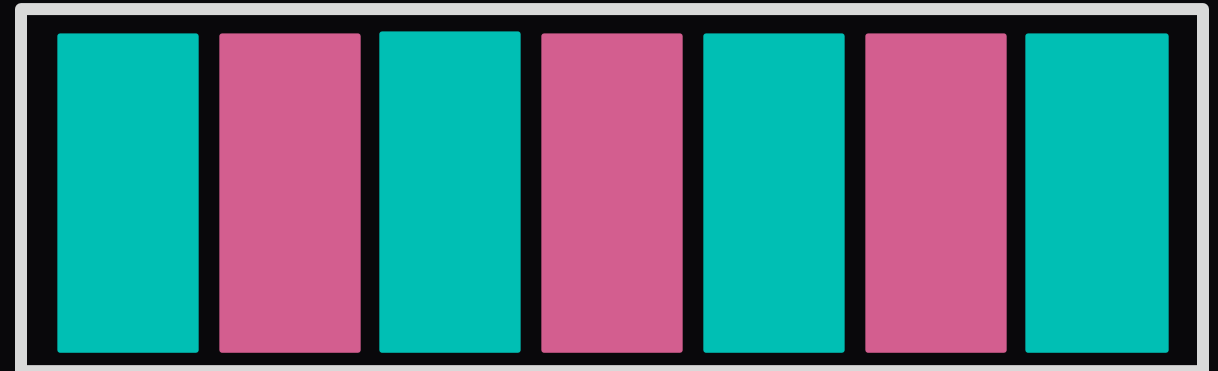
LLM: "use tool: get_data('API_ENDPOINT', 'request')"

Tool: "Data Content"

LLM: "Here's the answer..."

User: "What if..."

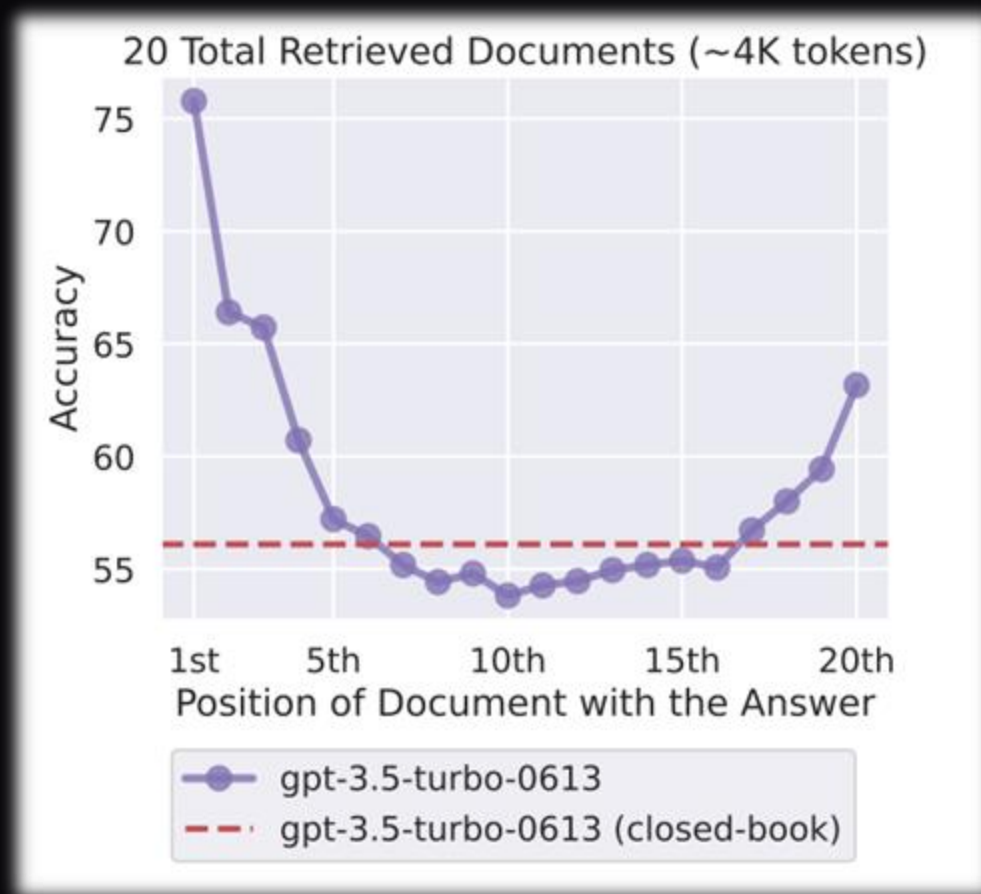
Context Window



70%! 80%! 90%! 100%!

100%!

LOST IN THE MIDDLE



Nelson F. Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. 2024. Lost in the Middle: How Language Models Use Long Contexts. Transactions of the Association for Computational Linguistics, 12:157–173.

CONTEXT ENGINEERING

“End users shouldn’t engineer prompts - systems should.”

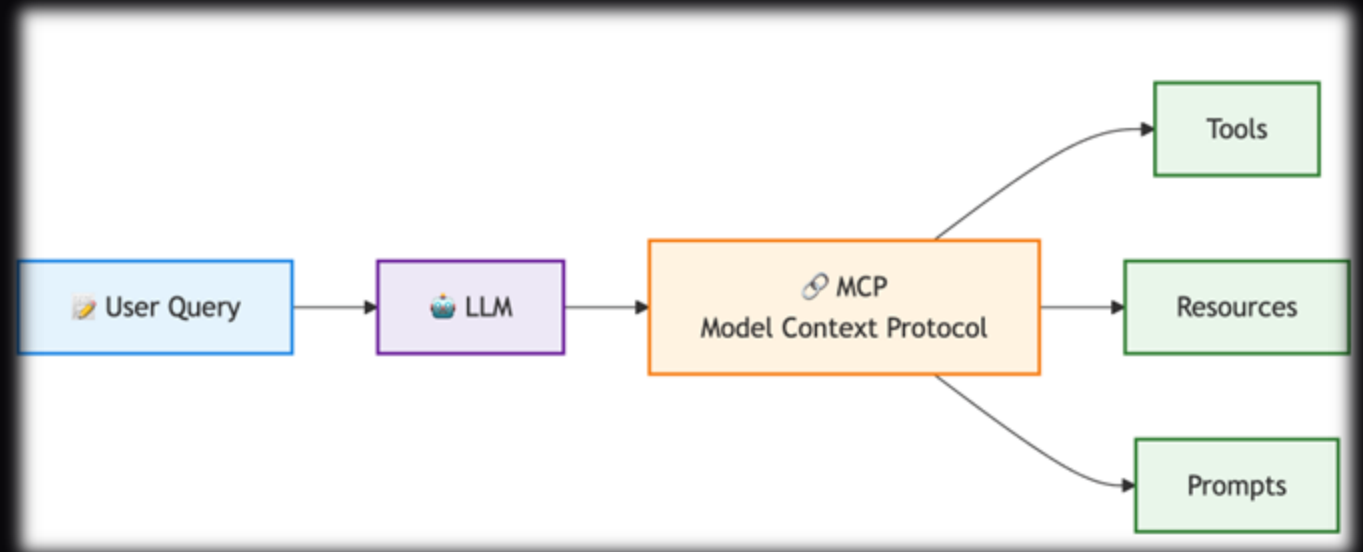
“Dynamic Context Injection and Reason Patterns for LLM (Agent).”



Dynamic Context Enabled

MODEL CONTEXT PROTOCOL

“An open standard that makes it easy for developers to connect LLMs with tools, resources, and prompts—enabling dynamic context injection without custom integrations.”



MCP: A standard bridge for injecting tools, resources, and prompts into LLM workflows.

TOOL ENGINEERING

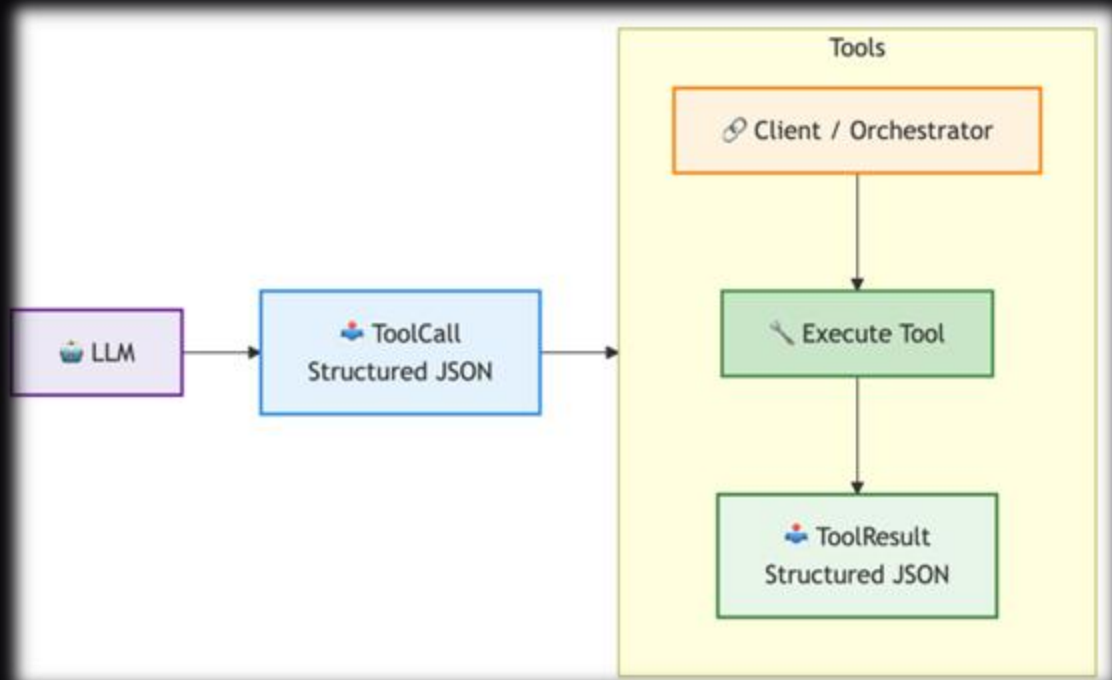
“Tool = Enablement + Empowerment”



- **Get real-time data**
- **Do complex tasks**
- **Improve accuracy**
- **Stay secure & compliant**
- **Control cost & speed**
- **Enable auditability**

- LLMs can't access fresh info alone.
- APIs, databases, code execution.
- Reduce hallucinations with trusted sources.
- Validate inputs, protect PII.
- Use caching, retries, rate limits.
- Log and trace every external call.

TOOL ENGINEERING



Tool Execution Process

```
"type": "tool_call",  
"name": "get_weather",  
"arguments": { "city": "Shanghai", "units": "metric" },  
"request_id": "uuid-1234"
```

Example Tool Call

TOOL ENGINEERING

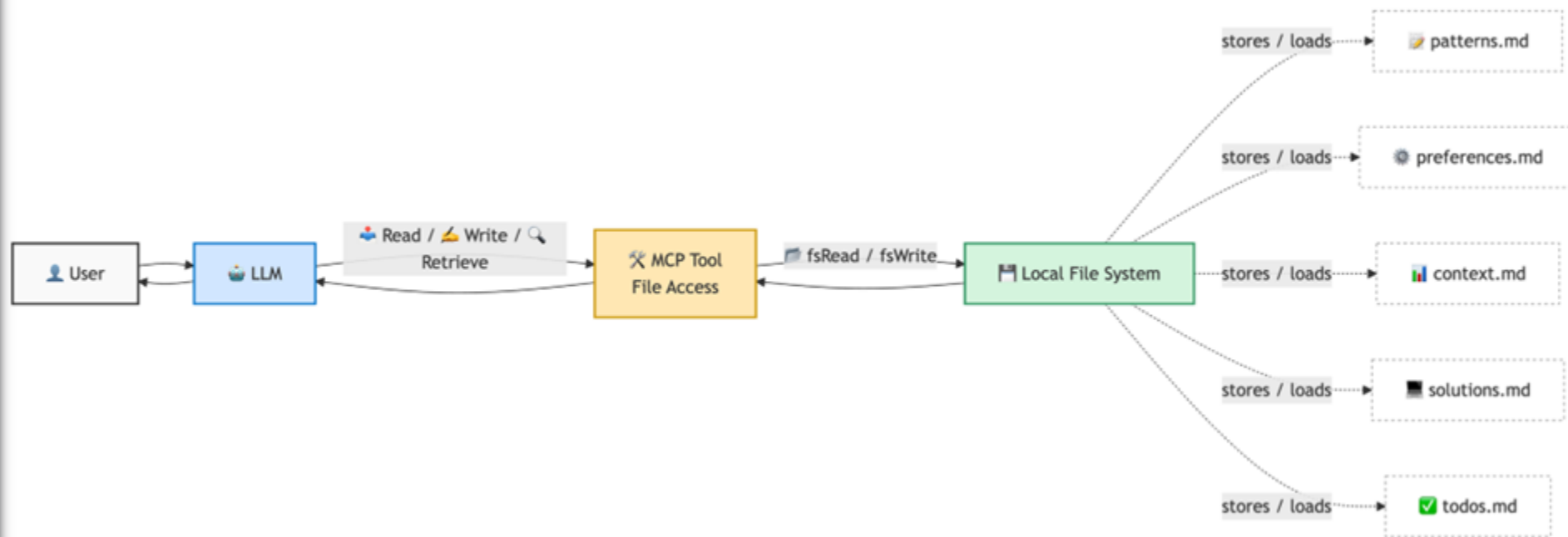
- 1. Clear Purpose:** Each tool should do one thing well
- 2. Structured I/O:** Define inputs and outputs explicitly for predictable behavior.
- 3. Composable:** Tools should integrate easily with other tools and workflows.
- 4. Stateless or Minimal State:** Reduce hidden dependencies; make behavior transparent.
- 5. Safe & Secure:** Validate inputs, enforce policies, and prevent misuse.

Principles of Tool Design

```
"tools": {  
  {  
    "name": "get_weather",  
    "description": "Fetches current weather conditions for a given location.",  
    "input_schema": {  
      "type": "object",  
      "required": ["location"],  
      "properties": {  
        "location": { "type": "string", "description": "City or region name." },  
        "units": { "type": "string", "enum": ["metric", "imperial"], "default":  
          "metric" }  
      }  
    },  
    "output_schema": {  
      "type": "object",  
      "properties": {  
        "temperature": { "type": "number", "description": "Current temperature." },  
        "condition": { "type": "string", "description": "Weather description." },  
        "humidity": { "type": "number", "description": "Humidity percentage." }  
      }  
    }  
  }  
}
```

Tool Description

MEMORY

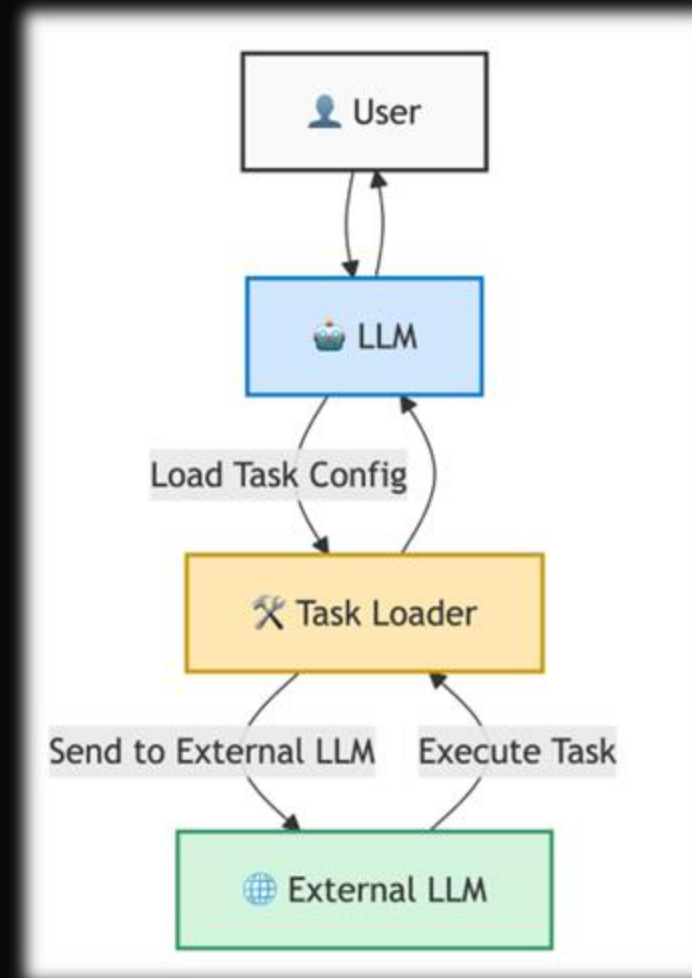


“Memory via local file system ensures persistent, secure context without external dependencies.”

TASK OFFLOADING

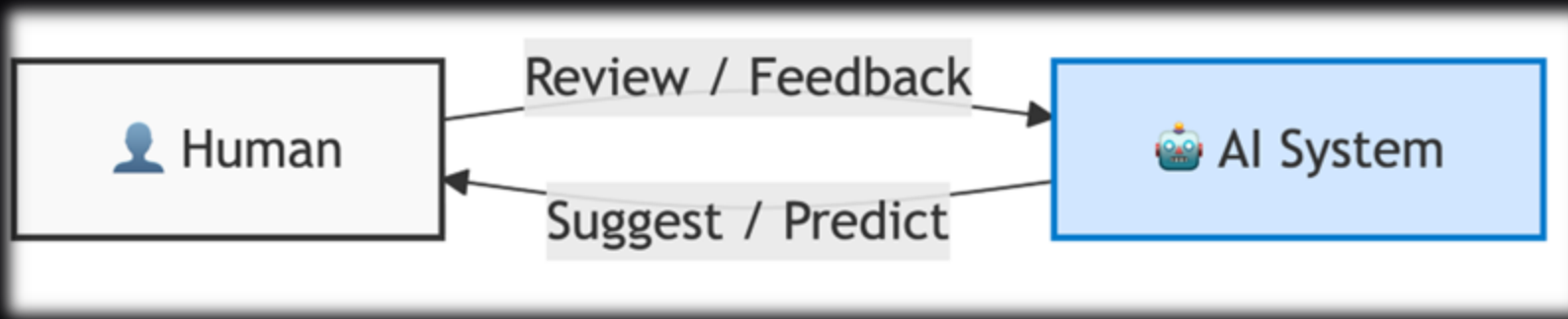
“**Task Offloading** means preloading instructions, context, or workflows into the LLM so it can execute a specific task efficiently without re-prompting every detail.”

- **Reduces prompt complexity** → No need to repeat long instructions.
- **Keeps task-specific memory** → Maintains relevant context for multi-step operations.
- **Improves consistency** → Ensures the LLM follows the same rules across steps.

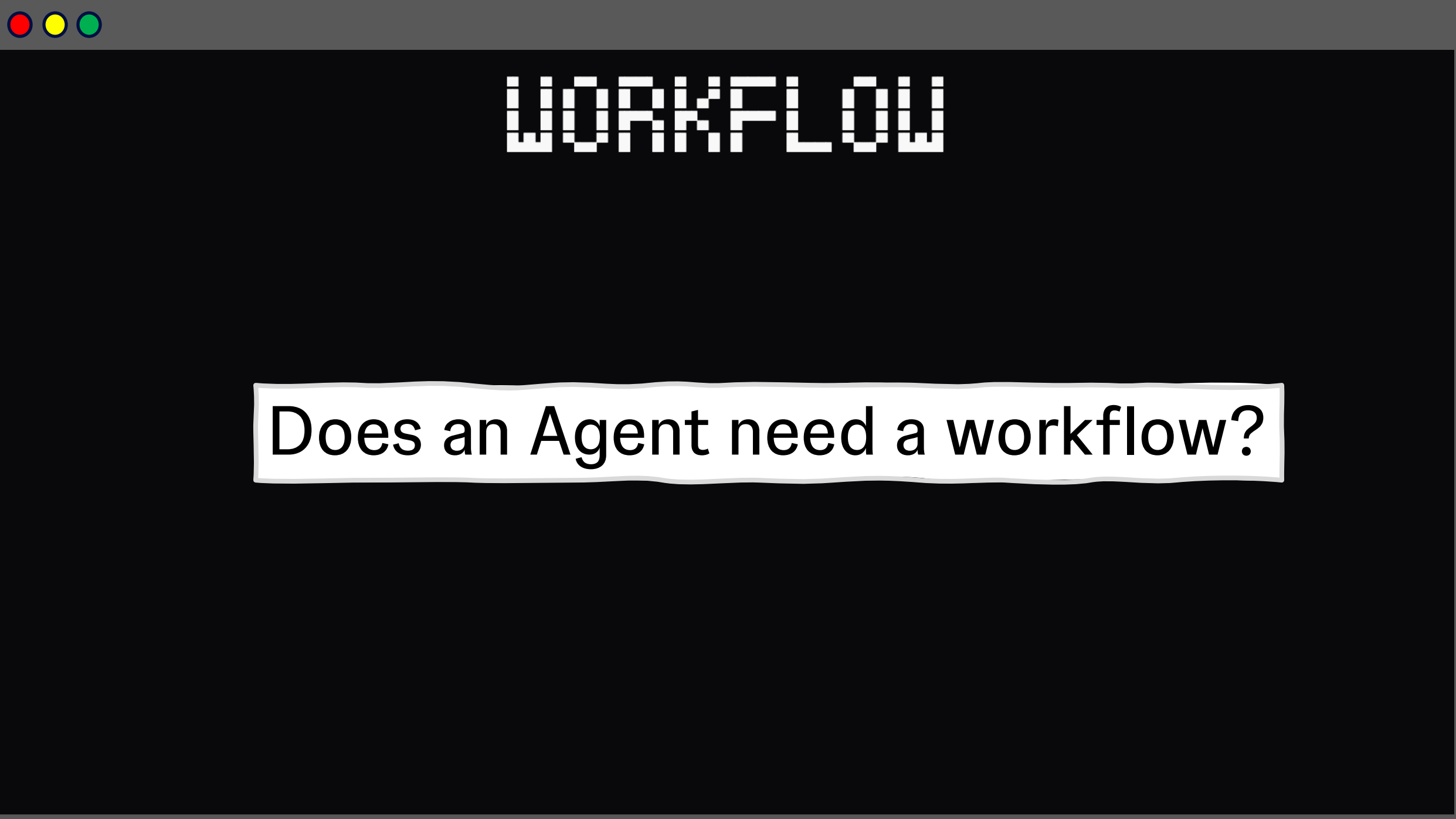


External LLM Task Offloading

HUMAN IN THE LOOP



“Human in the Loop (HITL) means keeping a person involved in the AI workflow to **review, guide, or approve decisions**, ensuring safety, accuracy, and accountability.”

A terminal window with a dark background and a light gray title bar at the top. The title bar contains three colored circles (red, yellow, green) on the left. The word "WORKFLOW" is displayed in a white, pixelated font at the top center. Below it, a white rectangular box with a slightly wavy top edge contains the text "Does an Agent need a workflow?".

WORKFLOW

Does an Agent need a workflow?

WORKFLOW

Is it a workflow?





**The Global Healthcare
Data Science Community**

Contact Channels

📞 UK +44 1843 609600
✉ office@phuse.global
🌐 phuse.global

Social Media

🐦 [@phusetwitta](https://twitter.com/phusetwitta)
📘 [/phusebook](https://facebook.com/phusebook)
📺 [/phusetube](https://youtube.com/phusetube)
🌐 [/company/phuse](https://linkedin.com/company/phuse)

THANK YOU