

Claude Certified Architect — Foundations Certification

مطالعہ گائیڈ (سرکاری امتحانی گائیڈ کی بنیاد پر)

تعارف

سرٹیفیکیشن اس بات کی تصدیق کرتی ہے کہ ایک ماہر حقیقی Claude Certified Architect — Foundations Claude Code، Claude Agent SDK، Claude API، اور Model Context Protocol (MCP) — یعنی بنیادی علم کا جائزہ لیتا ہے — کے ساتھ پروڈکشن ایپلیکیشنز بنائی جاتی ہیں۔ Claude و بنیادی ٹیکنالوجیز جن سے

agentic systems امتحانی سوالات حقیقت پسندانہ صنعتی منظرناموں پر مبنی ہوتے ہیں: کسٹمر سپورٹ کے لیے، multi-agent research pipelines، ڈیزائن کرنا، Claude Code کو CI/CD میں ضم کرنا، developer productivity، structured data بنانا، اور غیر ساختہ دستاویزات سے tools نکالنا۔

هدف امیدوار

کرتا ہے آپ ship کے ساتھ پروڈکشن ایپلیکیشنز ڈیزائن اور Claude کے جو solution architect مثالی امیدوار ایک کے پاس کم از کم 6 ماہ کا عملی تجربہ ہونا چاہیے:

- Claude Agent SDK** — multi-agent orchestration، subagents کو delegate کرنا، tool integration، lifecycle hooks
- Claude Code** — CLAUDE.md، MCP servers، Agent Skills، planning mode
- Model Context Protocol (MCP)** — backend integration کے لیے tools اور resources
- Prompt engineering** — JSON schemas، few-shot examples، data extraction templates
- Context windows** — لمبی دستاویزات کے ساتھ کام، multi-agent context passing
- CI/CD pipelines** — automated code review، test generation
- Escalation and reliability** — error handling، human-in-the-loop

امتحانی فارمیٹ

پیرامیٹر	قدر
سوال کی قسم	Multiple choice (میں سے 1 درست 4)
اسکورنگ	100–1000 scale، passing score 720
Guessing penalty	نہیں (سوال کا جواب دیں!)
Scenarios	8 4 (randomly selected) میں سے

امتحانى مواد: 5 ڈومينز

وزن	ڈومین
27%	1. Agent architecture and orchestration
18%	2. Tool design and MCP integration
20%	3. Claude Code configuration and workflows
20%	4. Prompt engineering and structured output
15%	5. Context management and reliability

امتحان منظر نام □

Scenario 1: Customer Support Agent

account اور، returns، billing disputes استعمال کرتے ہوئے Claude Agent SDK بناتے ہیں جو آپ ایک agent MCP tools (`get_customer`، `lookup_order`، `process_refund`، `escalate_to_human`) سنبھالتا ہے، مناسب first-contact resolution + % استعمال کرتا ہے 80 دفعہ (`escalation`، `escalate_to_human`) ساتھ۔

Scenario 2: Claude Code Code Generation

code generation, refactoring, debugging, documentation اور custom slash commands کے ساتھ، آپ کو اسے development میں استعمال کرتے ہوئے Claude Code کے ساتھ integrate کرنا چاہیے۔

Scenario 3: Multi-Agent Research System

تیار کرنا، جانچنا، reports کو مکمل citations سسٹم کو report generation اور synthesis tasks کو specialized subagents اور coordinator agent سونپنا: web research, document analysis,

Scenario 4: Developer Productivity Tools

□ agent engineers کو unfamiliar codebases explore کرنے، boilerplate code اور routine tasks automate □□ میں مدد دیتا □□ Built-in tools (Read, Write, Bash, Grep, Glob) MCP servers اور استعمال □□ کے □□ □□

Scenario 5: Claude Code for Continuous Integration

pull اور automated code reviews، test generation میں ضم کریں تاکہ CI/CD pipeline کو Claude Code request feedback کم سے کم false positives ڈیزائن ہونے کا نثر کے Prompts حاصل ہوں۔

Scenario 6: Structured Data Extraction

کرتا ہے، validate کے ذریعے JSON schemas کو output، سسٹم غیر ساختہ دستاویزات سے معلومات نکالتا ہے اور کو درست طور پر سنبھالنا چاہیے۔ edge cases پر برقرار رکھتا ہے اسے high accuracy اور

Scenario 7: Conversational AI Architecture Patterns

multi-turn conversational systems میں جن میں context window management، turns متعدد، memory strategies، tool design کے لیے execution محفوظ، کا برقرار رکھنا instructions میں اور مہم یا متضاد، user inputs کو سنبھالنا شامل ہے۔

Scenario 8: Agentic AI Tools (ہماری مدد کریں — مواد موجود نہیں)

میں شامل نہیں ہے۔ guide کے لیے دیے گئے لیکن ابھی تک اس candidates کی اطلاع امتحان دینے والے scenario اس کریں share میں GitHub Issues کے سوالات دیکھیں، تو ان میں scenario ہے اگر آپ نے اصل امتحان میں اس شامل کر سکیں۔ آپ کا تعاون۔ اس شخص کی مدد کرے گا جو امتحان کی تیاری کر coverage تاکہ ہم مکمل رہے۔

سرکاری دستاویزات

وسیلہ	URL
Claude API — Messages	https://platform.claude.com/docs/en/api/messages
Claude API — Tool Use	https://platform.claude.com/docs/en/build-with-claude/tool-use
Claude API — Message Batches	https://platform.claude.com/docs/en/build-with-claude/message-batches
Claude Agent SDK — Overview	https://platform.claude.com/docs/en/agent-sdk/overview
Claude Agent SDK — Hooks	https://platform.claude.com/docs/en/agent-sdk/hooks
Claude Agent SDK — Subagents	https://platform.claude.com/docs/en/agent-sdk/subagents
Claude Agent SDK — Sessions	https://platform.claude.com/docs/en/agent-sdk/sessions
Model Context Protocol (MCP)	https://modelcontextprotocol.io/
MCP — Tools	https://modelcontextprotocol.io/docs/concepts/tools
MCP — Resources	https://modelcontextprotocol.io/docs/concepts/resources
MCP — Servers	https://modelcontextprotocol.io/docs/concepts/servers
Claude Code — Documentation	https://code.claude.com/docs/en/overview
Claude Code — CLAUDE.md and Memory	https://code.claude.com/docs/en/memory
Claude Code — Skills (incl. slash commands)	https://code.claude.com/docs/en/skills

Claude Code — Hooks	https://code.claude.com/docs/en/hooks
Claude Code — Sub-agents	https://code.claude.com/docs/en/sub-agents
Claude Code — MCP Integration	https://code.claude.com/docs/en/mcp
Claude Code — GitHub Actions CI/CD	https://code.claude.com/docs/en/github-actions
Claude Code — GitLab CI/CD	https://code.claude.com/docs/en/gitlab-ci-cd
Claude Code — Headless (non-interactive mode)	https://code.claude.com/docs/en/headless
Prompt Engineering Guide	https://platform.claude.com/docs/en/build-with-claude/prompt-engineering/overview
Extended Thinking	https://platform.claude.com/docs/en/build-with-claude/extended-thinking
Anthropic Cookbook (code examples)	https://github.com/anthropics/anthropic-cookbook

نظریاتی بنیادیں: حصہ ۱

یہ حصہ تمام نظریاتی احاطہ کرتا ہے جس کی آپ کو امتحان میں کامیابی کے لیے ضرورت ہے۔ مواد کو کے مطابق — اس سے آپ کے domains کے مطابق ترتیب دیا گیا ہے، نہ کہ امتحانی concepts ٹیکنالوجیز اور موضوع کی زیادہ گہری سمجھ پیدا کر سکتے ہیں۔

کی بنیادیں Claude API — Model Interaction: باب 1

دستاویزات: [Messages API](#) | [Prompt Engineering](#)

1.1 API Request Structure

Claude API کی request کی ر Claude Messages API کی پیروی کرتا ہے request-response model ایک شامل ہوتا ہے:

```
{
  "model": "claude-sonnet-4-6",
  "max_tokens": 1024,
  "system": "You are a helpful assistant.",
  "messages": [
    {"role": "user", "content": "Hi!"},
    {"role": "assistant", "content": "Hello!"},
    {"role": "user", "content": "How are you?"}
  ],
  "tools": [...],
  "tool_choice": {"type": "auto"}
}
```

۱۰:۰۰ م فیلڈز

- `model` — model selection (`claude-opus-4-6` , `claude-sonnet-4-6` , `claude-haiku-4-5`)
- `max_tokens` — response کی تعداد میں زیادہ سے زیادہ tokens
- `system` — system prompt (model کی تعریف کرنا)
- `messages` — conversation history (history تسلسل برقرار رکھنے کا لیجئے آپ کو مکمل)
()
- `tools` — definitions کی tools دستیاب
- `tool_choice` — tool selection strategy

1.2 Message Roles

استعمال کرتا □□ instructional role اور ایک conversational roles دو array messages

- `user` — user messages, tool results (ایک `user` role message اندر `tool_result` کی صورت میں بھیجے جاتے ہیں، ایک الگ content block)
- `assistant` — model responses (history میں شامل کیے جاتے ہیں history)، tool use requests (tool_use content blocks)
- `system` — top-level `system` field سے سیٹ کیا جا سکتا ہے (turn پر `system` inline میں {"role": "system", ...} کی صورت میں) یا (سے لاگو ہوتا ہے turn پر) کے ذریعے سیٹ کیا جا سکتا ہے (turn پر `system` inline میں {"role": "system", ...} کی صورت میں) اس نقطہ سے آگے) شامل کیا جا سکتا ہے inline کی صورت میں {"role": "system", ...} میں messages (کے تابع — نیچے دیکھیں placement rules لاگو ہوتا ہے، مخصوص

کے طور پر user-role message کے طور پر نہیں بھیجا جاتا۔ یہ ایک message والے role "tool" کو Tool results کے content block tool_result میں ایک content پر بھیج دیا جاتا ہے۔ یہیں جس کے شامل ہوتا ہے۔

```
{
  "role": "user",
  "content": [
    {
      "type": "tool_result",
      "tool_use_id": "toolu_01...",
      "content": "...
    }
  ]
}
```

top-level کے طور پر ظاہر ہو سکتا ہے، نہ صرف role میں بھی ایک array messages براہ راست system کے بغیر invalidate prefix کو cached field کے top-level کے ذریعے اس کا مقصد system parameter placement rules شامل کرنا ہے اس کے مخصوص instructions کے درمیان conversation:

- پر ختم ہونے والے server tool use یا (ہوں blocks tool_result بشمول وہ جس میں) user turn یا ایک assistant turn کے فوراً بعد آنا چاہیے۔
- ہونا چاہیے element کا آخری array سد پہلے آنا چاہیے یا assistant turn کسی کسی
- ملتا error کے درمیان نہیں آ سکتا — ایسا کرنے پر 400 tool_result اور اس کے block tool_use کسی کسی
- پہلے والوں پر اور بعد (کے درمیان شامل کیے گئے conversation بشمول) system messages بعد میں آنے والے پر فوقیت رکھتے ہیں۔ field system کے top-level کے لیے turns کے

Model requests بھیجی جا رہی ہیں conversation history میں آپ کو مکمل API request انٹیلیجنس: م۔ ہر آزاد ہوتا ہے call محفوظ نہیں رکھتا — ہر state درمیان

1.3 Response میں stop_reason Field

کیوں روکی generation کے model شامل ہوتا ہے، جو بتاتا ہے کہ stop_reason میں Claude API response

Value	Description	Action
"end_turn"	مکمل کر لیا response کے Model	کو دکھائیں user نتیجہ
"tool_use"	کرنا چاہتا ہے Model tool call	Tool execute اور کریں result دیں
"max_tokens"	Token limit ختم ہو گئی	Response truncated ہے ممکن ہے limit ہے
"stop_sequence"	Stop sequence مل گئی	کریں handle کے مطابق application logic اپنی

Agentic systems میں "end_turn" اور "tool_use" — یہی agent loop سب سے اہم ہیں —

1.4 System Prompt

متعین کرنا ہے behavioral rules اور context کے جو instruction ایک خاص System prompt

- میں دیا جاتا ہے field system کا حصہ نہیں ہوتا؛ الگ array messages

- user messages پر فوقیت رکھتا ہے
- میں لاگو رہتا ہے conversation ہوتا ہے اور پوری load ایک بار
- role, constraints, اور output format کے لیے استعمال ہوتا ہے

پیدا کر سکتی ہے مثال کے طور پر tool associations کو wordings کی system prompt میں لیا جائے گا۔
کو ضرورت سے زیادہ get_customer کو model کی تصدیق کریں ”جیسی کہ ادیت customer پر،“ ہمیشہ
استعمال کرنے پر مجبور کر سکتی ہے، حتیٰ کہ جب یہ ضروری نہ ہو۔

1.5 Context Window

Context window text (tokens) جس کی کل مقدار اس میں process ایک وقت میں model کی کل مقدار اس میں کر سکتا شامل:

- System prompt
- مکمل message history
- Tool definitions
- Tool results

مسائل context-window

1. **Lost-in-the-middle effect:** models input سب سے زیادہ اعتماد کو زیادہ سے زیادہ پر موجود معلومات پر مبنی ہے۔ اگرچہ یہ ایک مشکل ہے، مگر درمیان کی تفصیلات process کر سکتے ہیں۔ حل: اہم معلومات شروع یا آخر کے قریب miss کرتے ہیں، مگر درمیان کی تفصیلات رکھیں۔
2. **Tool results** واپس کر کے 40+ fields tool شامل کرنا۔ اگر output میں tool call context کا جمع ہونا: ہر tool کا بڑا حصہ ضائع ہو جاتا ہے۔ context میں صرف 5 اہم ہوتے ہیں، تو
3. **Progressive summarization:** history compress کرتے وقت numeric values, percentages, اور dates اکثر گم ہو کر غیر واضح ہو جاتے ہیں ("تقریباً"، "کچھ"، "چند")۔

tool_use اور Tools: باب 2

Tool Use: دستاویزات

2.1 tool_use ☐ کیا

برای Model کو اجازت دینا کہ external functions کو Claude جو mechanism ایک tool_use result کو اور execute اس code بنانا؛ آپ کا structured tool call request میں چلانا — وہ code راست واپس کرتا

2.2 Tool Definition

کے: کیا جاتا ہے define کے ذریعے JSON schema کی tool

Tool description ك: انتقائي ام يلو

1. **Description tool selection mechanism** □□□ LLM tools کی بنیاد پر descriptions کا انتخاب کرتا □□□ Minimal descriptions (“Customer information retrieve □□ کرتا □□”) □□□ اُن مواقع پر غلطیوں کا □□□ کرتا □□□ overlap □□□ ایک دوسرے □□□ tools باعث بنتی □□□ ہیں جب

2. Description میں شامل کریں

- کیا کرتا ہے اور کیا واپس کرتا ہے Tool
- Input formats اور example values
- Edge cases اور constraints
- tool similar alternatives میں کب استعمال ہو

3. `analyze_document` اور `analyze_content` سے بچیں اگر **overlapping descriptions** ایک جیسے یا انہیں خلط ملط کر دے گا model تقریباً ایک جیسی ہوں، تو descriptions

4. **Built-in tools** □ **MCP tools:** agents similar functionality □ built-in tools (Read, Grep) کو MCP tools — مضبوط بنائیں MCP tool descriptions پر ترجیح دے سکتے ہیں □ اس سے بچنے کا لہذا concrete advantages, unique data, □ built-in tools نمایاں کریں جو context یا وہ فراہم نہیں کر سکتے □

2.3 tool_choice Parameter

کیسے منتخب کرتا model tools کنٹرول کرتا tool_choice

Value	Behavior	کب استعمال کریں
<code>{ "type": "auto" }</code>	Model کو tool call خود طے کرنا یا text دینا	default زیادہ تر حالات میں
<code>{ "type": "any" }</code>	Model کو tool call کو لازماً کوئی	guaranteed جب آپ کو structured output چاہیے
<code>{ "type": "tool", "name": "extract_metadata" }</code>	Model کو tool call کو لازماً ایک مخصوص	forced first step / execution order جب آپ کو چاہیے

2. **Nullable fields:** استعمال کریں جو ممکنہ طور پر `"type": ["string", "null"]` ایسی معلومات کے لیے `null` Model موجود نہ ہوں۔
3. **Enums with "other":** predefined categories کے لیے `"other"` شامل کریں `"other"` + ایک detail string
4. **Enum "unclear":** ایماندار model category اُن صورتوں کے لیے جب `"unclear"` سے بہتر category غلط

2.5 Syntax Semantic Errors بمقابلہ

Error type	Example	Mitigation
Syntax	Invalid JSON, field type غلط	JSON schema کے ساتھ <code>tool_use</code> (ختم کر دیتا)
Semantic	Totals match نہیں کرتے، value field غلط، hallucination	Validation checks, feedback کے ساتھ retry, self-correction

بنانا Claude Agent SDK — Agentic Systems: باب 3

دستاویزات: [Agent SDK](#) | [Hooks](#) | [Subagents](#) | [Sessions](#)

3.1 Agentic Loop کیا ہے

Agentic loop کی ایک actions صرف جواب نہیں دیتا — وہ Model کے pattern کا بنیادی task execution خودکار sequence انجام دیتا ہے:

1. Claude کو request کے ساتھ tools بھیجیں
2. Response وصول کریں
3. stop_reason چیک کریں:
 - "tool_use" -> tool execute کریں، result history میں append کریں step 1 پر واپس جائیں
 - "end_turn" -> task مکمل، result user کو دکھائیں
4. کریں repeat مکمل ہونے تک

خود منتخب tool کی بنیاد پر اگلا tool results اور پچھلے context اپنے Claude **model-driven approach** کے لیے: `max_iterations=5` (مثلاً) استعمال کرنا `arbitrary iteration limit` کے طور پر `Primary stop condition` کے لیے `hard-coded decision trees` کرتا ہے۔

Anti-patterns (ان سے بچیں):

- Completion detect کے لیے assistant text parse کرنا ("Task completed")
- Primary stop condition کے طور پر `arbitrary iteration limit` (مثلاً `max_iterations=5`)
- سمجھنا completion signal دیا گیا یا نہیں، اس کے textual content نہ assistant کے چیک کرنا

درست طریقہ: `completion signal` کا واحد قابل اعتماد `stop_reason == "end_turn"`

3.2 AgentDefinition Configuration

AgentDefinition Claude Agent SDK میں agent configuration object ہے:

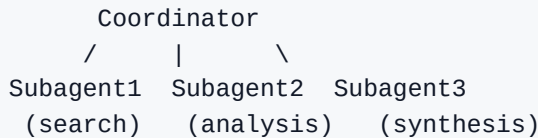
```
agent = AgentDefinition(
    name="customer_support",
    description="Handles customer requests for returns and order issues",
    system_prompt="You are a customer support agent...",
    allowed_tools=["get_customer", "lookup_order", "process_refund",
"escalate_to_human"],
    # Coordinator کے لیے:
    # allowed_tools=["Task", "get_customer", ...]
)
```

parameters:

- name / description — agent کی شناخت اور وضاحت
- system_prompt — system prompt کی ہدایات والا
- allowed_tools — (least privilege اصول) کی فہرست tools اجازت یافتہ

3.3 Hub-and-Spoke: Coordinator اور Subagents

Multi-agent architecture میں بنائی جاتی ہے hub-and-spoke topology عموماً



Coordinator کی ذمہ داریاں:

- Task کو subtasks میں تقسیم کرنا
- dynamic selection (درکار کے subagents میں سے طے کرنا کہ کون سا)
- کو سونپنا subagents کام
- کرنے والے نتائج جمع اور validate
- Errors اور retries handle کرنا
- تک پہنچانا user نتائج

الگ الگ context subagents کا: اصول

- Subagents conversation history کو coordinator سے خود بخود inherit نہیں کرتے
- میں واضح طور پر دینا ہوتا ہے subagent prompt context تمام ضروری
- Subagents memory share کے درمیان نہیں کرتے
- observability اور error control کے ذریعے گزرتی ہے communication coordinator تمام

3.4 Subagents **Task Tool** بنانے کے لیے

Subagents **Task** tool کے ذریعے spawn کیے جاتے ہیں:

```
# Coordinator کے allowedTools میں "Task" ہونا چاہیے
coordinator_agent = AgentDefinition(
    allowed_tools=["Task", "get_customer"]
)
```

Explicit context passing لازمی ہے:

```
# Bad: subagent کے پاس context نہیں ہے
Task: "Analyze the document"

# Good: مکمل context prompt میں
Task: "Analyze the following document.
Document: [full document text]
Prior search results: [web search results]
Output format requirements: [schema]"
```

Parallel spawning: subagents parallel — کر سکتا ہے Task s call میں متعدد response ایک coordinator کے چلنے کے لیے:

```
# Coordinator کے ایک response میں شامل ہے:
Task 1: "Search for articles about X"
Task 2: "Analyze document Y"
Task 3: "Search for articles about Z"
# تینوں ایک ساتھ چلتے ہیں
```

3.5 Agent SDK میں Hooks

Hooks agent lifecycle کی اجازت دیتے ہیں transformation اور interception پر points کے مخصوص

PostToolUse tool result کو model پر لے کر intercept تک پہنچانے کے لیے:

```
# کریں formats normalize سے تاریخ کے MCP tools مثال: مختلف
@hook("PostToolUse")
def normalize_dates(tool_result):
    # Unix timestamp -> ISO 8601
    # "Mar 5, 2025" -> "2025-03-05"
    return normalized_result
```

Outgoing-call interception hook policy کو block کرنے کے لیے:

```
# کریں refunds block مثال: 500$ سے اوپر
@hook("PreToolUse")
def enforce_refund_limit(tool_call):
    if tool_call.name == "process_refund" and tool_call.args.amount > 500:
        return redirect_to_escalation(tool_call)
```

Hooks اور prompt instructions میں فرق

Attribute	Hooks	Prompt instructions
Guarantee	Deterministic (100%)	Probabilistic (>90%, 100% نہیں)
کب استعمال کریں	Critical business rules, financial operations, compliance	General preferences, recommendations, formatting
مثال	Refunds > \$500 block کریں	"Try to solve before escalating"

استعمال کریں hooks، نہیں prompts — مالی، قانونی، یا حفاظتی نتائج کو failure قاعدہ: جب

باب 4: Model Context Protocol (MCP)

[Servers](#) | [Resources](#) | [Tools](#) | [MCP](#): دستاویزات

4.1 MCP کیا ہے

MCP سے جوڑنا Claude کو external systems جو open protocol ایک (MCP) Model Context Protocol

کرتا ہے resource types define تین بنیادی:

- Tools** — functions (CRUD operations, API calls, command execution) agent actions جنہیں لیں
- Resources** — context agent جسے data کے لیے (documentation, database schemas, content catalogs) پڑھ سکتا ہے
- Prompts** — predefined prompt templates کے لیے tasks عام

4.2 MCP Servers

فرام کرتا ہے جب tools/resources کرتا ہے اور MCP protocol implement کے جو process ایک MCP server

کرتا ہے: MCP server سے connect

- جو جائے ہیں discover خود بخود تمام tools
- ایک ساتھ دستیاب ہوتے ہیں tools کے connected servers تمام
- انہیں کیسے استعمال کرے گا model طے کرتی ہیں کے Tool descriptions

4.3 MCP Servers کو Configure کرنا

Project configuration (`.mcp.json`) — team usage کے لیے:

```
{
  "mcpServers": {
    "github": {
      "command": "npx",
      "args": ["-y", "@modelcontextprotocol/server-github"],
      "env": {
        "GITHUB_TOKEN": "${GITHUB_TOKEN}"
      }
    },
    "jira": {
      "command": "npx",
      "args": ["-y", "mcp-server-jira"],
      "env": {
        "JIRA_TOKEN": "${JIRA_TOKEN}"
      }
    }
  }
}
```

اہم نکات:

- میں `.mcp.json` project root میں محفوظ ہوتا ہے اور version control میں ہوتا ہے
- Environment variables (`${GITHUB_TOKEN}`) secrets کے لیے استعمال ہوتے ہیں
- commit خود tokens — کے لیے استعمال ہوتے ہیں
- contributors کے لیے دستیاب ہوتا ہے
- تمام project کے

User configuration (`~/.claude.json`) — personal/experimental servers کے لیے:

- User کے home directory میں محفوظ ہوتا ہے
- Version control کے ذریعے share ہوتا ہے
- Personal experiments اور testing کے لیے

Servers کا انتخاب:

- Standard integrations (Jira, GitHub, Slack) موجود ہیں
- community MCP servers کے لیے پبلک سرورس موجود ہیں
- کو ترجیح دینا
- کے لیے بنائیں
- unique, team-specific workflows صرف custom servers اپنی

4.4 MCP میں isError Flag

کو بتانا کہ agent استعمال ہوتا ہے یا نہیں `isError: true` میں response آتا ہے تو error میں MCP tool جب call fail ہوئی ہے

Structured error (اچھا):

```
{
  "isError": true,
  "content": {
    "errorCategory": "transient",
    "isRetryable": true,
    "message": "The service is temporarily unavailable. Timeout while calling the orders API.",
    "attempted_query": "order_id=12345",
    "partial_results": null
  }
}
```

Generic error (anti-pattern):

```
{
  "isError": true,
  "content": "Operation failed"
}
```

escalate بدلیں، یا query، کریں retry کو فیصلہ سازی کے لیے کوئی معلومات نہیں دیتا — کیا generic error agent کریں؟

4.5 MCP Resources

کیا: بغیر طلب کر سکتا actions حاصل کرنے کے لیے agent context میں جو ایسا Resources

- Content catalogs (hierarchical navigation، کی فہرست project tasks مثلاً تمام)
- Database schemas (data structure کے لیے سمجھنے کے لیے)
- Documentation (API references، internal guides)
- Issue/task summaries

Resource کا فائدہ: agent کے لیے exploratory tool calls کی ضرورت نہیں پڑتی کے لیے سمجھنے کے لیے agent کا فائدہ: Resource کا فائدہ: agent کے لیے exploratory tool calls کی ضرورت نہیں پڑتی کے لیے سمجھنے کے لیے

5 اور Claude Code — Configuration Workflows

دستاویزات: [Claude Code](#) | [Memory / CLAUDE.md](#) | [Skills](#) | [MCP](#) | [Hooks](#) | [Sub-agents](#) | [GitHub Actions](#) | [Headless](#)

5.1 CLAUDE.md Hierarchy

hierarchy استعمال کرتا ہے اس کی تین سطحی Claude Code میں جو instruction file(s) و CLAUDE.md ہے:

1. User-level: ~/.claude/CLAUDE.md
 - ہر لاگو user صرف اسی
 - نسخہ ہوتا ہے کہ share ذریعہ VCS
 - working style اور preferences ذاتی
2. Project-level: .claude/CLAUDE.md یا root CLAUDE.md
 - ہر لاگو project contributors تمام
 - ہوتا ہے کہ manage ذریعہ VCS
 - Coding standards, testing standards, architectural decisions
3. Directory-level: subdirectories میں CLAUDE.md
 - کہ ساتھ کام ہو تو لاگو ہوتا ہے کہ files کی directory جب اسی
 - conventions کہ اُس حصہ کہ مخصوص codebase

~/.claude/CLAUDE.md میں ملتی کیونکہ و project instructions کو team member عام غلطی: ایک نئے میں رکھی گئی تھیں (user-level) ~/.claude/CLAUDE.md کہ بجائے (project-level)

5.2 @path Syntax (File Imports)

و جاتی configuration modular کر سکتا ہے، جس سے refer کہ ذریعہ @path کو files بیرونی CLAUDE.md ہے:

Project CLAUDE.md

میں بیان کیہ گئے ہیں @./standards/coding-style.md
Coding standards
میں ہیں @./standards/testing-requirements.md
Test requirements
میں ہیں @package.json dependencies میں اور @README.md
Project overview

کہ: قواعد @path

- (نہیں space کوئی) آئے file path کہ فوراً بعد @
- Relative اور absolute paths دونوں supported ہیں
- Relative paths کہ مطابق file اُس resolve میں جس میں import ہوتا ہے
- Maximum import nesting depth 5 ہے

شامل ہوتا ہے۔ ہیں standards میں صرف متعلقہ package کم ہوتی ہے اور duplication اس سے

5.3 .claude/rules/ Directory

کرنے کہ organize کہ حساب سے topic کو rules کا متبادل ہے، جو monolithic CLAUDE.md .claude/rules/ لیے استعمال ہوتا ہے:

```
.claude/rules/
testing.md          -- testing conventions
api-conventions.md  -- API conventions
deployment.md       -- deployment rules
react-patterns.md   -- React patterns
```

YAML frontmatter کے ساتھ `paths` conditional loading:

```
---
paths: ["src/api/**/*.ts"]
---
```

استعمال کریں `explicit error handling` کے ساتھ `async/await` کہ لیں `API files` واپس کریں `endpoint standard response wrapper` سے

```
---
paths: ["**/*.test.tsx", "**/*.test.ts"]
---
```

استعمال ہونہ چاہئیں `describe/it blocks` میں `Tests` نہیں `hardcoding` استعمال کریں `Data factories` استعمال کریں `Database mock` — نہ کریں `test database`

یہ کیسے کام کرتا ہے:

- `paths` pattern سے `match` کرتا ہے جو `file edit` ایسا `Claude Code` ہوتا ہے جب `load` صرف اس وقت `Rule` ہو
- نہیں `rules load` ہوتا ہے — غیر متعلقہ `tokens` اور `context` اس سے
- میں `codebase files` کرنے دیتا ہے، چاہے `conventions apply` کے لحاظ سے `file type` آپ کو `Glob patterns` (کے لیے بہت مفید `tests`) کے ہیں بھی

کب استعمال کریں `directory-level CLAUDE.md` بمقابلہ `.claude/rules/` والے `paths`:

- `.claude/rules/` with `paths` — جب `conventions` کئی `directories` میں بکھری `files` (tests, migrations) پر لاگو ہوں
- `Directory-level CLAUDE.md` — جب `conventions` کسی `directory` اور `کے` میں `و` اور `وابستہ` `و` `ضروری` `نہ`

5.4 Custom Slash Commands اور Skills

`skills` کو (`.claude/commands/`) `custom commands` میں `Claude Code version` نوٹ: موجود بناتے `commands` `/name` `formats` کر دیا گیا ہے دونوں `unify` کے ساتھ (`.claude/skills/`) `supported` اب بھی `format` کا ذکر کرتی ہے — `.claude/commands/` میں امتحانی گائیڈ

یہ `invoke` کے ذریعے `/name` میں جو `Slash commands reusable prompt templates`:

`.claude/commands/` `format (legacy, supported)`:

```
.claude/commands/
review.md      -- /review -- standard code review
test-gen.md    -- /test-gen -- test generation
```

.claude/skills/ **format (current):**

```
.claude/skills/
review/SKILL.md -- /review -- frontmatter configuration
test-gen/SKILL.md
```

Project commands (**.claude/commands/** یا **.claude/skills/**):

- کرنے والے سب کے لیے دستیاب ہوتے ہیں repo clone میں محفوظ ہوتے ہیں اور VCS
- یقینی بناتے ہیں consistent workflows ٹیم میں

User commands (**~/.claude/commands/** یا **~/.claude/skills/**):

- ان میں ہوتے ہیں share کے ذریعے VCS جو commands ذاتی
- کے لیے workflows انفرادی

5.5 Skills — .claude/skills/

Skills advanced commands میں جو SKILL.md frontmatter کے ذریعے configure میں ہوتے ہیں:

```
---
context: fork
allowed-tools: ["Read", "Grep", "Glob"]
argument-hint: "Path to the directory to analyze"
---
```

کا تجزیہ کریں۔ code structure میں directory دی گئی
تیار کریں۔ report پر ایک architectural patterns اور Dependencies

Frontmatter parameters:

Parameter	Description
<code>context:</code> <code>fork</code>	نہیں کرتا pollute کو verbose output main session میں چلانا یا isolated subagent کو
<code>allowed-</code> <code>tools</code>	نہیں کرتا skill files delete مثلاً — security) دستیاب ہوں گے tools یا محدود کرتا ہے کون سے (سکے گی اگر اجازت نہ ہو
<code>argument-</code> <code>hint</code>	مانگنے کا اشارہ argument ہو تو invoke کے بغیر skill parameters جب

Skill **CLAUDE.md** کے استعمال کریں **بمقابلہ**:

- **Skill** — task مخصوص کے لیے on-demand invocation (review, analysis, generation)
- **CLAUDE.md** — general standards اور conventions ہمیشہ لوڈ ہونے والے

Personal skills (~/ .claude/skills/):

- متاثر نہ ہونے teammates مختلف ناموں سے بنائیں تاکہ personal variants اپنی

5.6 Planning Mode بمقابلہ Direct Execution

Planning mode:

- نہیں کرتا changes کرتا ؛ plan اور investigate صرف Model
- استعمال کرتا Read، Grep، Glob کرنے کے لیے Codebase explore
- کرتا approve user تیار کرتا جسے implementation plan ایک
- safe exploration کے side effects بغیر

Planning mode کب استعمال کریں:

- files (درجنوں) changes بڑے
- Multiple plausible approaches (microservices: boundaries کیسے define ہوں؟)
- Architectural decisions (کیا framework کون سا؟ structure؟)
- Unfamiliar codebase (change پر عمل سمجھنا ضروری ہے)
- Library migrations 45+ files کو متاثر کریں

Direct execution کب استعمال کریں:

- Single-file fixes ساتھ stack trace واضح
- شامل کرنا validation check ایک
- changes اچھی طرح سمجھ گئی، غیر مبہم

Combined approach:

- Investigation اور design کے لیے planning mode
- User plan approve کر
- Approved plan implement کے لیے direct execution

Explore subagent — codebase explore کے لیے specialized subagent:

- Verbose output کو main context سے isolate کرتا
- واپس کرتا summary صرف
- Multi-phase tasks میں context-window exhaustion روکتا

5.7 /compact Command

/compact context compress کے لیے built-in command ہے:

- میں جگہ خالی کرتا context window کو summarize کو history چھپی
- Long investigation sessions میں verbose tool output context سے ہٹا دیتا جب
- Risk: exact numeric values، dates، اور specific details summarization میں گم ہو سکتا ہے

5.8 /memory Command

```
/memory sessions درمیان memory manage کرنا کے لیے built-in command:
```

- محفوظ کیے جا سکیں context اور notes، preferences، لیکے کھولنے، تاکہ file editing (CLAUDE.md)
- بوتی load پر خود بخود startup کے پار برقرار رہتی اور Information sessions
- محفوظ context current work اور frequently used commands، user preferences، Project conventions، کر کے لیکے مفید
- دوبارے سمجھانے کا متبادل instructions میں وہی session پر

5.9 Claude Code CLI for CI/CD

-p (or **--print**) flag:

```
claude -p "Analyze this pull request for security issues"
```

- Non-interactive mode: prompt process ☐ کرتا ، stdout پر print اور ☐ کرتا exit ☐ جاتا ☐
- User input کا انتظار نہ ☐ میں کرتا
- CI/CD pipelines میں Claude ☐ کا واحد درست طریقہ ☐

CI ☐ ☐ structured output:

```
claude -p "Review this PR" --output-format json --json-schema
'{"type":"object",...}'
```

- `--output-format json` — output کو JSON میں دیتا ہے
- `--json-schema` — output کو schema مطابق validate کرتا ہے
- `Result` کو automatically inline PR comments میں دے سکتا ہے

Session context isolation: میں کم مؤثر review کیا و، اکثر code generate جس نہ Claude session و ی اور اپنہ فیصلوں کو reasoning context اپنی (model) وتی کر۔ کا امکان کم وتا challenge ساتھ رکھتا و اور اپنہ فیصلوں کو reasoning context اپنی (model) وتی Review و independent instance استعمال کریں،

Duplicate comments □ **بچھا** □ **review** □ **وقت** □ **review results context** □ **کرتا** □ **بعد دوبار** □ **commits** □ **سدا** □ **نہا** □ **Claude** □ **میں شامل کریں اور** □ **unresolved issues report** □ **صرف نہا** □ **یا** □ **کرنے** □ **کی** □ **دایت دیں** □

5.10 fork_session اور Session Management

`--resume <session-name>` named session کو resume کرتا ہے

```
claude --resume investigation-auth-bug
```

- conversation □□ جاری رکھتا □□ ساتھ بچھلی saved context

- Long investigations پر پھیلی ہوئی ان کے لیے مفید sessions جو متعدد
- Risk: tool results stale ہوں تو files کے بعد session اگر پچھلی

fork_session shared context سے independent branch بناتا ہے:

```
Codebase investigation
|
fork_session
/      \
Approach A:  Approach B:
Redux       Context API
```

- کرتے ہیں context inherit تک forks branch point دونوں
- کرتے ہیں diverge اس کے بعد وہ آزادانہ طور پر
- Approaches compare یا strategies test کرنے کے لیے مفید

Resume نئی session کے بجائے نئی:

- Tool results stale ہوں (files بدل چکی ہوں)
- ہو گیا وہ context degrade کافی وقت گزر چکا ہو اور
- کرنے کے بجائے یہ بہتر ہے کہ "م نہ جو پایا اس کا مختصر خلاصہ یہ ہے": resume کے ساتھ tool data پرانے کیا جائے restart "... سے"

6 Prompt Engineering — Advanced Techniques

دستاویزات: [Prompt Engineering](#) | [Anthropic Cookbook](#)

6.1 Few-shot Prompting

Few-shot prompting میں 2–4 prompt دکھانے کے لیے expected behavior وہ طریقہ ہے جس میں input/output examples شامل کیے جاتے ہیں۔

Few-shot text descriptions سے زیادہ مؤثر کیوں ہے:

- "ہوں" کو کئی طریقوں سے سمجھا جا سکتا ہے precise جیسے "زیادہ" instruction میں
- دکھاتا ہے decision logic اور expected format غیر میں طور پر Example
- (دہراتا نہیں examples صرف) کرتا ہے generalize پر cases کو نئے Model pattern

Few-shot examples کی اقسام اور استعمال:

1. Ambiguous scenarios کے لیے examples:

Request: "My order is broken"

Action: Call get_customer -> lookup_order -> check status.

Rationale: "broken" damaged item میں درکار order details ہو سکتا ہے؛ آپ کو

Request: "Get me a manager"

Action: escalate_to_human call کریں فوراً

Rationale: Customer human نہ کریں واضح طور پر

2. Output formatting کے لیے examples:

Finding example:

```
{
  "location": "src/auth/login.ts:42",
  "issue": "SQL injection in the username parameter",
  "severity": "critical",
  "suggested_fix": "Use a parameterized query"
}
```

3. Acceptable problematic code کے لیے examples:

```
// Acceptable (flag نہ کریں):
const items = data.filter(x => x.active);

// Problem (flag کریں):
const items = data.filter(x => x.active == true); // Use strict equality ===
```

4. Different document formats سے extraction کے examples:

Document with inline citations:

"As shown in the study (Smith, 2023), the rate is 42%."

-> {"value": "42%", "source": "Smith, 2023", "type": "inline_citation"}

Document with bibliography references:

"The rate is 42%. [1]"

-> {"value": "42%", "source": "reference_1", "type": "bibliography"}

5. Informal measurements کے لیے examples:

Text: "about two handfuls of rice"

-> {"amount": "~100g", "original_text": "two handfuls", "precision": "approximate"}

Text: "a pinch of salt"

-> {"amount": "~1g", "original_text": "a pinch", "precision": "approximate"}

Few-shot اور non-standard measurement units میں مؤثر جو purely rule-based instructions کے لیے بہت متنوع ہوتی ہیں

Prompts میں format normalization rules: structured output کے لیے strict JSON schemas استعمال شامل کریں normalization rules میں prompt، کریں

Normalization:

- Dates: ہمیشہ ISO 8601 (YYYY-MM-DD); "yesterday" -> absolute date compute کریں
- Currency: numeric amount + currency code; "five bucks" -> {"amount": 5, "currency": "USD"}
- Percentages: decimal fraction; "half" -> 0.5

وہ values inconsistent ہیں مگر JSON syntactically کم ہیں۔ ان semantic errors اس سے

6.2 Explicit Criteria بمقابلہ Vague Instructions

Bad (vague):

چیک کریں accuracy کی Code comments
کریں۔ high-confidence findings report رہیں - صرف Conservative

Good (explicit criteria):

کریں جب flag صرف اُس صورت میں problematic کو Comment:
1. Comment کا actual code behavior سے CONTRADICT ہو کرتا
2. Comment کا حوالہ دہ variable یا function کسی غیر موجود
3. TODO/FIXME comment ایسے bug کی طرف اشارہ کرے جو کسی

Flag نہ کریں:
- کہ لحاظ سے پرانہ ہوں style جو صرف comments ایسے
- wording inaccuracies معمولی
- Missing comments (الگ category یہ الگ)

Severity criteria examples کے ساتھ define کریں:

CRITICAL: User کے لیے runtime failure
Example: Payment process کرتے وقت NullPointerException

HIGH: Security vulnerability
Example: SQL injection, XSS, missing authorization checks

MEDIUM: Logic bug بغیر فوری impact کے
Example: Wrong sorting, off-by-one error

LOW: Code quality
Example: Duplication, کہ لیے suboptimal algorithm

6.3 Prompt Chaining

میں توڑتا ہے sequence کی focused steps کو ایک complex task Prompt chaining

Step 1: auth.ts analyze کریں (صرف local issues)

-> Output: auth.ts میں issues کی list

Step 2: database.ts analyze کریں (صرف local issues)

-> Output: database.ts میں issues کی list

Step 3: Integration pass (cross-file dependencies)

-> Output: module boundaries پر issues

یہ کیوں اہم ہے:

- **Attention dilution** — جب files ایک ساتھ دیکھتا ہے تو کچھ files ایک ساتھ دیکھتا ہے، model سے بچاتا ہے — جب bugs miss کیے جاتے ہیں
- **Shallow commentary** — shallow commentary کر سکتا ہے اور باقی پر
- **Consistent analysis quality** — consistent analysis quality کر سکتا ہے
- **Cross-file interactions** — cross-file interactions الگ الگ کر سکتا ہے

کب استعمال کریں **dynamic decomposition** بمقابلہ **prompt chaining**:

- **Prompt chaining** — predictable, repeatable tasks (code review, file migrations)
- **Dynamic decomposition** — open-ended investigations ان subtasks execution کے دوران واضح ہوتے ہیں

6.4 “Interview” Pattern

Solution implement کرنے کے لیے Claude clarifying questions پوچھتا ہے:

Claude: "API کے لیے caching implement کرنے سے پہلے چند سوالات"

1. event-based یا TTL آپ پسند کرتے ہیں؟ cache invalidation strategy میں سے کون سی
2. stale data ہو تو کیا cache unavailable اگر قابل قبول ہے؟
3. Caching per-user یا global؟
4. Cache کے لیے expected data volume کیا ہے؟

یہ کب مفید ہے:

- Unfamiliar domain (fintech, healthcare, legal systems)
- Implications tasks (cache strategies, failure modes) واضح
- Viable approaches context میں بہترین انتخاب

6.5 Validation اور Retry-with-Feedback

extracted data validation fail کر کے:

Step 1: Document سے data extract کریں

Step 2: Validate کریں (Pydantic, JSON Schema, business rules)

Step 3: retry کریں کہ ساتھ context — error اگر:

- Original document
- Previous (incorrect) extraction
- Specific error: "Field 'total' = 150, but sum(line_items) = 145. Re-check values."

کب مؤثر ہوگا **Retry**:

- Format errors (date غلط format میں)
- Structural errors (field جگہ غلط)
- Arithmetic inconsistencies (model check دوبارہ کر سکتا ہے)

Retry گا کہ مدد نہ میں کرے گا:

- Information source document میں موجود ہے نہ میں
- Required context external (data اور کسی document میں کیا گیا)

Pydantic validation tool: Pydantic Python library جو schema-based data validation کے لیے استعمال ہوتی ہے۔ امتحان کے لیے اہم نکات یہ ہیں:

- **Structural validation:** types, requiredness, enum constraints JSON کے بعد میں چیک کیے جاتے ہیں۔
- **Semantic validation:** custom validators business logic enforce کرتے ہیں (sum of items equals total! start_date < end_date)
- **Validate-retry loops:** Pydantic validation failure پر error message بنا کر Claude کو error context میں prompt ساتھ دوبارہ کریں
- **JSON Schema generation:** Pydantic models سے JSON Schema generate کر سکتے ہیں۔ `tool_use` فراہم کرتا ہے single source of truth میں، جو ایک

6.6 Self-correction

Internal contradictions detect کرنے کا pattern:

```
{
  "stated_total": "$150.00",
  "calculated_total": "$145.00",
  "conflict_detected": true,
  "line_items": [
    {"name": "Widget A", "price": 75.00},
    {"name": "Widget B", "price": 70.00}
  ]
}
```

Model stated value اور computed value میں تو مختلف ہیں — اگر وہ `conflict_detected` آپ دونوں نکالتا ہے —

7: Message Batches API

دستاویزات: [Message Batches](#)

7.1 Overview

Message Batches API کو asynchronous processing کے لیے requests کے batches submit دینا ہے:

Attribute	Value
Savings	50% مقابلہ میں synchronous calls
Processing window	24 hours (latency SLA guarantee) زیادہ سے زیادہ
Multi-turn tool calling	Supported (request = ایک response = ایک)
Correlation	request اور response کو جوڑنے کے لیے custom_id field

7.2 Batch API Synchronous API بمقابلہ کب استعمال کریں

Task	API	Why
Pre-merge PR check	Synchronous	Developer 24 hours قابل قبول نہیں انتظار کر رہا ہے
Overnight tech-debt report	Batch	Result 50% savings صبح تک چاہیے
Weekly security audit	Batch	50% savings فوری نہیں
Interactive code review	Synchronous	درکار response فوری
10,000 documents process کرنا	Batch	Bulk processing! savings میں

7.3 custom_id کا استعمال

```
{
  "custom_id": "doc-invoice-2024-001",
  "params": {
    "model": "claude-sonnet-4-6",
    "max_tokens": 1024,
    "messages": [{"role": "user", "content": "Extract data from: ..."}]
  }
}
```

آپ کو یہ کرنے دیتا ہے custom_id:

- Result کو original document link سے
- Failure کو submit دوبارہ failed documents کی صورت میں صرف
- Successful documents کو process دوبارہ

7.4 Batch میں Failures کو Handle کرنا

1. 100 documents کا batch submit کریں
2. 95 succeed 5 fail (context limit exceeded) ہوں
3. Failures کو custom_id سے identify کریں
4. Strategy میں تقسیم کریں chunks کو long documents (مثلاً) بدلیں

کریں submit کو دوبارہ failed documents صرف اُن 5.5

7.5 SLA Planning

تک لگ سکتے ہیں hours کو 24 Batch API چاڑھ اور result میں hours اگر آپ کو 30

- Submission window: $30 - 24 = 6$ hours
- کریں submit پڑھ hours کم از کم 24 deadline کو Batches
- میں تقسیم کریں 4-hour windows لیں 4 Frequent submissions

8 Task Decomposition Strategies

8.1 Fixed Pipelines (Prompt Chaining)

ہوتا ہے defined پڑھ step ر:

```
Document -> Metadata extraction -> Data extraction -> Validation -> Enrichment -> Final output
```

کب استعمال کریں:

- Task structure predictable (reviews ایک ہی template follow کریں)
- پڑھ معلوم ہوں steps تمام
- چاڑھ stability اور reproducibility آپ کو

8.2 Dynamic Adaptive Decomposition

ہوتا ہے ہیں generate کی بنیاد پر Subtasks intermediate results

- "Legacy codebase لیں tests add کریں"
- > Glob, Grep کریں structure map پڑھ
- کہ ساتھ partial coverage، 2 tests بغیر modules ملا: 3 ->
- > Prioritize: payments module شروع کریں (high risk)
- دریافت ہوئی dependency پر external API: کام کہ دوران ->
- > Adapt: tests پڑھ mock add کریں کہ لیں external API لکھنے سے پڑھ

کب استعمال کریں:

- Open-ended investigative tasks
- شروع میں معلوم نہ ہو scope جب مکمل
- پر منحصر ہو results کہ step پچھلا step جب ر

8.3 Multi-pass Code Review

کے لیے pull requests والی 10+ files

```
Pass 1 (per-file): auth.ts analyze کریں -> local issues کی list
Pass 1 (per-file): database.ts analyze کریں -> local issues کی list
Pass 1 (per-file): routes.ts analyze کریں -> local issues کی list
...
Pass 2 (integration): files درمیان کہ relationships analyze کریں
-> Cross-file issues: inconsistent types, circular dependencies
```

کیوں خراب pass پر ایک سی 14 files:

- Attention dilution: shallow کچھ کی، deep analysis کی files کچھ
- Inconsistent comments: approve دوسری میں، pattern flag میں ایک file
- Missed bugs: cognitive overload واضح سی و سی errors کی وجہ سے

9 Escalation اور Human-in-the-Loop: باب 9

9.1 کریں Escalate کی طرف Human کب

Escalation triggers (rules واضح):

Situation	Action
Customer "get me a manager" واضح طور پر کہہ	کریں؛ حل کرنے کی کوشش نہ کریں escalate فوراً
Policy request نہ کرتی سی و cover کو	Escalate policy جب competitor price matching (مثلاً) کریں (سی و)
Agent پیش رفت نہ کر سکا	کریں escalate کے بعد attempts مناسب تعداد میں
Threshold سی و financial operation اوپر	Escalate enforce، prompt کے ذریعہ hook ترجیحاً) کریں (نہیں)
Customer multiple matches تلاش کرتے وقت	مانگیں؛ انداز نہ لگائیں identifiers مزید

نہیں سی و reliable trigger جو:

Unreliable method	Why it fails
Sentiment analysis	نہیں کرتا correlate سی و Customer کا mood case complexity
Model self-rated confidence (1–10)	کمزور سی و calibration غلط سی و ہو سکتا سی و؛ Model confidently
Automatic classifier	موجود نہ سی و؛ training data ممکن؛ Overengineering

9.2 Escalation Patterns

Immediate escalation:

Customer: "I want to speak to a manager"
Agent: [کرتا == escalate_to_human call فوراً]
NOT: "I can help with your issue, let me..."

Resolve escalation: کی کوشش کے بعد

Customer: "My refrigerator broke two days after purchase"
Agent: [کرتا == warranty replacement offer، چیک کرتا == order]
escalate -> مطمئن نہ ہو customer اگر

Nuanced escalation (acknowledge → resolve → reiteration پر escalate):

Customer: "This is outrageous, I'm very unhappy with the quality!"
Agent: [frustration acknowledge] "سمجھتا ہوں frustration میں آپ کی"
[resolution offer] "پیش کر سکتا ہوں replacement یا refund میں"
Customer: "No, I want to talk to someone!"
Agent: [customer دوبارہ insist کرہ -> immediate escalation]

کریں escalate دیں، اور صرف اسی وقت concrete solution کریں، پھر emotion acknowledge بنیادی اصول: پل۔
manager (ی۔) نہ کریں escalate کے پل۔ اطلاع پر dissatisfaction کی خواہش دہرائیں customer human جب
(مانگنے کے برابر نہیں ہیں)

Policy gap کے لیے escalation:

Customer: "Competitor X has this item 30% cheaper—give me a discount"
Policy: کرتی == price adjustments cover پر site صرف اپنی
Agent: [escalate – policy competitor price matching کو cover نہیں کرتی]

9.3 Structured Handoff Protocols

کو دینی چاہیے structured summary human agent پر Escalation:

```
{
  "customer_id": "CUST-12345",
  "customer_name": "Ivan Petrov",
  "issue_summary": "Refund request for a damaged item",
  "order_id": "ORD-67890",
  "root_cause": "Item arrived damaged; photos attached",
  "actions_taken": [
    "Verified customer via get_customer",
    "Confirmed order via lookup_order",
    "Offered a standard replacement – customer insists on a refund"
  ],
  "refund_amount": "$89.99",
  "recommended_action": "Approve a full refund",
  "escalation_reason": "Customer requested to speak with a manager"
}
```

دیکھنا summary تک رسائی نہیں ہوتی — صرف یہ full conversation transcript کو Human operator

ہونی چاہیے self-contained اس لیے یہ مکمل اور

9.4 Confidence Calibration اور Human Oversight

Data extraction systems □J □K:

1. **Field-level confidence scores:** model ر extracted field ل confidence score ن
2. **Calibration:** labeled validation sets استعمال کر thresholds tune کریں
3. **Routing:**
 - High confidence + stable accuracy -> automated processing
 - Low confidence ل ambiguous sources -> human review

Stratified random sampling:

- High-confidence extractions □□ sample audit □□ لیجی بھی باقاعدگی سے
- Aggregate 97% accuracy □□ errors □□ لیجی 40 document type کسی مخصوص
- Accuracy □□ analyze □□ حساب سے field اور document type □□ ہیں overall کو صرف

10 Error Handling میں Multi-agent Systems: پایب

10.1 Error Categories

Category	Examples	Retryable	Agent action
Transient	Timeout, 503, network failure	Yes	Exponential backoff retry ساتھ
Validation	Invalid input format, missing required field	No (input fix کریں)	Request modify کریں اور retry کریں

Business	Policy violation, threshold exceeded	No	پیش alternative کو سمجھائیں؛ User کریں
Permission	Access denied	No	Escalate کریں

10.2 Error-handling Anti-patterns

Anti-pattern	Problem	Correct approach
Generic status “search unavailable”	Coordinator decide نہیں کر سکتا کہ recover کیسے کرنا ہے	Error type, query, partial results, alternatives واپس کریں
Silent suppression (empty result = success)	Coordinator match سمجھتا ہے کہ failure ملا، حالانکہ	“no results” اور “search failure” میں فرق کریں
workflow پر پورا ایک failure abort کرنا	ضائع ہو جاتا ہے partial results تمام	Partial results کے ساتھ جاری رکھیں؛ gaps annotate کریں
Subagent میں infinite retries	Latency اور resources کا ضیاع	Local recovery (1–2 retries)، پھر coordinator تک propagate

10.3 Structured Subagent Error

```
{
  "status": "partial_failure",
  "failure_type": "timeout",
  "attempted_query": "AI impact on music industry 2024",
  "partial_results": [
    { "title": "AI Music Generation Report", "url": "...", "relevance": 0.8 }
  ],
  "alternative_approaches": [
    "Try a narrower query: 'AI music composition tools'",
    "Use an alternative data source"
  ],
  "coverage_impact": "Not covered: AI impact on music production"
}
```

کو فیصلہ کرنے کے لیے ضروری معلومات دیتا ہے coordinator

- Modified query کے ساتھ retry کریں؟
- Partial results استعمال کریں؟
- delegate کو subagent کسی اور
- کریں؟ gap annotate کے بغیر جاری رکھیں اور section اس

10.4 Final Synthesis میں Coverage Annotations

Report: AI Impact on Creative Industries

Visual Art (FULL COVERAGE)

[research results]

Music (PARTIAL COVERAGE – search agent timeout)

[partial results]

⚠ Note: اس section کی coverage search agent timeout کی وجہ سے محدود ہے۔

Literature (FULL COVERAGE)

[research results]

11 Context میں Production Systems: باب 11 Management

11.1 Facts کو Extract میں Separate Block کریں

Conversation history (جو summarization کے دوران degrade ہو جاتی ہے) پر انحصار کرنے کے بجائے، key facts کو structured block میں extract کریں:

```
=== CASE FACTS (کیا جائے update آئے fact جب بھی نیا) ===  
Customer ID: CUST-12345  
Order ID: ORD-67890  
Order Date: 2025-01-15  
Order Amount: $89.99  
Issue: Damaged item on delivery  
Customer Request: Full refund  
Status: Pending manager approval  
===
```

کیوں نہ ہو؟ summarized کتنی ہی history میں شامل کریں، چاہے prompt پر block یں

11.2 Tool Results کو Trimming کرنا

کے لیے صرف 5 درکار ہیں current task واپس کرتا ہے مگر 40+ fields lookup_order اگر

```
# PostToolUse hook: رکھیں صرف relevant fields
@hook("PostToolUse", tool="lookup_order")
def trim_order_fields(result):
    return {
        "order_id": result["order_id"],
        "status": result["status"],
        "total": result["total"],
        "items": result["items"],
        "return_eligible": result["return_eligible"]
    }
```

کم noise ہوتی ہے اور context اس سے

11.3 Position-aware Input

Lost-in-the-middle effect رکھیں critical information کو ذہن میں رکھتے ہوئے:

```
[KEY FINDINGS – اوپر]
3 critical vulnerabilities ملیں...

[DETAILED RESULTS – درمیان]
=== File auth.ts ===
...
=== File database.ts ===
...

[ACTION ITEMS – آخر میں]
Priority: merge auth.ts vulnerabilities fix کریں.
```

11.4 Scratchpad Files

Long investigations میں agent key findings scratchpad file لکھ سکتا ہے:

```
# investigation-scratchpad.md
## Key findings
- PaymentProcessor in src/payments/processor.ts inherits from BaseProcessor
- refund() موتی ہے call سے places تین OrderController, AdminPanel, CronJob
- External PaymentGateway API limit: 100 req/min
- Migration #47 نہ refund_reason (NOT NULL) add 01-12-2024 کیا
```

scratchpad دوبارہ چلائے کہ بجائے agent discovery (میں session یا نئی) ہو جائے context degrade کر سکتا ہے

11.5 Context Subagents کو Delegate کرنا

Main agent: "payments module کی dependencies investigate کریں"

-> Subagent (Explore): 15 files read کرتا، imports trace کرتا

-> Returns: "Payments depends on AuthService, OrderModel, and the external PaymentGateway API"

Main agent: context 15 files میں رکھتا ہے کہ line

Separate context layer: Multi-agent systems میں subagent context budget — میں کام کرتا ہے کہ
کے طور پر separate context layer ایک Coordinator ملتی ہے کہ information کے لیے ضروری task اسے صرف اپنے
کرتا context allocate کرتا ہے، اور global state store، کرتا ہے کہ subagent outputs aggregate کرتا ہے کہ: وہ
ہے کہ بھر دیتا ہے کہ جو information کو ایسی agent window نہیں ہوتا، جہاں ایک "context leakage" ہے کہ اس سے
دوسروں کے لیے غیر متعلق ہے کہ

Subagents کے constrained context budgets:

- Minimal context: specific task + ضروری data بھیجیں
- Subagent کو raw data dumps کے بجائے structured results کی ہدایت دیں
- اور کم distractions کا مطلب کم tools کریں — کم limit کو toolset کے subagent سے allowedTools context cost

11.6 Structured State Persistence (crash recovery کے لیے)

کرتا ہے کہ export پر ایک location ایک معلوم state اپنی agent

```
// agent-state/web-search-agent.json
{
  "status": "completed",
  "queries_executed": ["AI music 2024", "AI music composition"],
  "results_count": 12,
  "key_findings": [...],
  "coverage": ["music composition", "music production"],
  "gaps": ["music distribution", "music licensing"]
}
```

Coordinator resume پر ایک manifest load کرتا ہے کہ:

```
// agent-state/manifest.json
{
  "web-search": "completed",
  "doc-analysis": "in_progress",
  "synthesis": "not_started"
}
```

کو محفوظ رکھنا Provenance: باب 12

12.1 Attribution Loss Problem

گم ہو سکتا ہے link "claim → source" کی جانچ میں، تو results summarize سے sources جب متعدد

Bad: "The AI music market is estimated at \$3.2B." (No source, no year.)

Good:

```
{
  "claim": "The AI music market is estimated at $3.2B.",
  "source_url": "https://example.com/report",
  "source_name": "Global AI Music Report 2024",
  "publication_date": "2024-06-15",
  "confidence": 0.9
}
```

12.2 Conflicting Data کو Handle کرنا

دیں values مختلف sources جب دو

```
{
  "claim": "Share of AI-generated music on streaming platforms",
  "values": [
    {
      "value": "12%",
      "source": "Spotify Annual Report 2024",
      "date": "2024-03",
      "methodology": "Automated classification"
    },
    {
      "value": "8%",
      "source": "Music Industry Association Survey",
      "date": "2024-07",
      "methodology": "Survey of 500 labels"
    }
  ],
  "conflict_detected": true,
  "possible_explanation": "Methodology اور time period کا فرق"
}
```

کو coordinator کے ساتھ محفوظ رکھیں اور attribution کو اندھا دھند منتخب نہ کریں۔ دونوں کو value کسی ایک فیصلہ کریں دیں۔

12.3 Correct Interpretation کے لیے Dates شامل کریں

سمجھ لیا جا سکتا ہے contradiction کو temporal differences کے بغیر Dates

Bad: "Source A says 10%, source B says 15%. Contradiction."

Good: "Source A (2023) says 10%, source B (2024) says 15%. Likely 5+ % growth."

12.4 Content Type مطابق Render کریں

میں نڈ ٹھونسیں format ر چیز کو ایک ی:

- Financial data -> tables
- News اور analysis -> prose
- Technical findings -> structured lists
- Time series -> chronological ordering

13 Claude Code Built-in Tools

13.1 Tool Selection Reference

Task	Tool	Example
تلاش کرنا files سڈ pattern نام	Glob	<code>**/*.test.tsx</code> , <code>src/components/**/*.ts</code>
File contents اندر search کرنا	Grep	Function name, error message, import
File مکمل پڑھنا	Read	Analysis کرنا file load کر لیا
بنانا file نیا	Write	Scratch سڈ file create کرنا
precise edit میں file موجود	Edit	Unique text match ذریعہ specific snippet replace کرنا
چلانا Shell command	Bash	git, npm, tests, build

13.2 Incremental Investigation Strategy

نڈ پڑھیں سمجھ بوجھ کو رفتہ رفتہ بنائیں files ایک ساتھ تمام:

- Grep: entry points تلاش کریں (function definition, export)
- Read: پڑھیں files ملنے والی
- Grep: usages تلاش کریں (import, calls)
- Read: consumer files پڑھیں
- کریں repeat نہ مل جائے picture جب تک مکمل

13.3 Fallback: Edit بجائے Read + Write

و جائے fail کی وجہ سڈ Edit non-unique text match جب

1. Read — لوڈ کریں file content مکمل
2. Content کو programmatically modify کریں
3. Write — updated version لکھ دیں

امتحانی ڈومین نوٹس: II حصہ

Orchestration اور Agent Architecture: ڈومین 1 (27%)

1.1 Autonomous Task Execution اور Agentic Loops Design

Key knowledge:

- Agent loop lifecycle: Claude request بھیجیں، `stop_reason` ("tool_use" بمقابلہ "end_turn") چیک واپس کریں results کو iteration کے لیے اگلی tools execute کریں،
- Tool results conversation history میں append ہیں تاکہ model کو action اگلا model کو
- Model-driven decision making (Claude tool اگلا) (چنتا) بمقابلہ hard-coded decision trees

Key skills:

- Flow control: جب `stop_reason = "tool_use"` تو loop جاری رکھیں اور `"end_turn"` پر رک جائیں
- Iterations درمیان context کو tool results کو append کرنا
- Anti-patterns بچنا: completion کے لیے assistant text parse کرنا، primary stopping mechanism پر استعمال کرنا arbitrary iteration limits

1.2 Multi-agent Systems (Coordinator–Subagent) Orchestrate

Key knowledge:

- Hub-and-spoke architecture: coordinator تمام inter-agent communication، error handling، routing اور مالک ہوتا ہے
- Subagents isolated context میں کام کرتے ہیں — coordinator کی history inherit کرتے ہیں
- Coordinator responsibilities: task decomposition، delegation، result aggregation، dynamic selection of subagents
- Coordinator کا خطرہ overly narrow decomposition

Key skills:

- کم duplication کے درمیان تقسیم کریں تاکہ subagents کو Research coverage

- Iterative refinement loops implement کریں (coordinator synthesis evaluate اور tasks re-route کریں)
- Observability تمام communication coordinator کے ذریعے route کے لیے تمام

1.3 Subagent Calls, Context Passing, اور Spawning Configure کرنا

Key knowledge:

- Task tool subagents spawn؛ coordinator کے allowedTools میں "Task" ہونا چاہیے
- Subagent context prompt میں explicitly؛ ہونا چاہیے
- AgentDefinition configuration: descriptions, system prompts, tool constraints
- Alternatives explore کے لیے fork_session کے ذریعے session management

Key skills:

- Previous agents کے full outputs subagent prompt میں شامل کریں
- Context pass کے وقت data اور metadata کو separate کے لیے structured formats کریں
- Coordinator turn میں متعدد Task calls کے ذریعے parallel subagents spawn کریں
- Coordinator prompts goals اور quality criteria، step-by-step instructions، طور پر لکھیں

1.4 Enforcement اور Handoff Patterns ساتھ Multi-step Workflows Implement کرنا

Key knowledge:

- Workflow ordering کے لیے programmatic enforcement (hooks, preconditions) اور prompt guidance کے درمیان فرق
- deterministic guarantees کے ساتھ operations (مثلاً مالی) درکار ہوں، prompts identity verification کے لیے، کافی نہیں ہوتے
- Escalation کے دوران structured handoff protocols (customer ID, reason, recommended action)

Key skills:

- مکمل نہ ہونے والے steps کریں جب تک کہ downstream calls کو block نہ کر لیں (مثلاً get_customer کے verified ID تک واپس کرنے کے لیے process_refund block کریں)
- Multi-aspect customer requests کے items کو الگ میں تقسیم کریں
- Human کو escalate کے وقت structured summaries produce کریں

1.5 Tool Calls Intercept اور Data Normalize کرنے کے لیے Agent SDK Hooks

Key knowledge:

- Hook patterns (مثلاً PostToolUse) tool results کو model کے consume کے لیے intercept کرنے کے لیے

- Hooks outgoing calls intercept `compliance rules enforce` کرتے ہیں (مثلاً threshold) اوپر refunds سے block کرنا)
- Hooks **deterministic guarantees** دیتی ہیں، جبکہ **probabilistic compliance** prompt instructions دیتے ہیں

Key skills:

- `PostToolUse` hooks سے data formats normalize کریں (Unix timestamps, ISO 8601, numeric status codes)
- Policy-violating actions block کریں اور interception hooks سے redirect کریں
- منتخب کریں hooks کے بجائے prompts چاہیں اور تو **guaranteed compliance** کو business rules جب

1.6 Complex Workflows کے لیے Task Decomposition Strategies

Key knowledge:

- **Fixed pipelines** (prompt chaining) **dynamic adaptive decomposition** کی بنیاد پر intermediate results بمقابلہ
- Prompt chaining: sequential steps (integration pass، پھر analyze کو الگ file ر)
- Adaptive investigation plans جو discovered information کی بنیاد پر subtasks generate کرتے ہیں

Key skills:

- Predictable multi-aspect reviews کے لیے prompt chaining؛ open-ended investigations استعمال کریں؛ **dynamic decomposition**
- Large code reviews کو per-file analysis + separate cross-file integration pass میں تقسیم کریں
- Open-ended tasks کو degrade کریں؛ prioritized plan پھر structure map کریں؛

1.7 Session State، Resuming، اور Forking

Key knowledge:

- Named sessions continue کریں `--resume <session-name>`
- Shared context سے independent investigation branches بنانے کے لیے `fork_session`
- Resume کی اہمیت file changes کو agent کرتے وقت
- Structured summary کے ساتھ session، stale results کے ساتھ resuming سے زیادہ reliable سے زیادہ

Key skills:

- Named investigation sessions continue کریں `--resume`
- Approaches parallel compare کریں `fork_session`
- Resuming (context ابھی current) سے شروع کریں session اور نئی (results stale) درمیان انتخاب کریں

5.10 fork_session اور Session Management

--resume <session-name> named session کو resume کرنا:

```
claude --resume investigation-auth-bug
```

- جاری رکھتا fork_session conversation کے ساتھ پچھلی saved context
- Long investigations sessions کے لیے مفید جو متعدد sessions
- Risk: tool results stale ہوں بدل چکی ہوں تو files کے بعد session اگر پچھلی

fork_session shared context سے independent branch بنانا:

```
Codebase investigation
  |
  fork_session
  /      \
Approach A: Approach B:
Redux      Context API
```

- کرتا fork_session context inherit تک forks branch point دونوں
- کرتا fork_session diverge اس کے بعد وہ آزادانہ طور پر
- کرتا fork_session مفید strategies test کرنے یا Approaches compare

Resume session کے بجائے نئی Resume:

- Tool results stale ہوں (files بدل چکی ہوں)
- ہو گیا fork_session context degrade کافی وقت گزر چکا ہو اور
- کرنے کے بجائے یہ بہتر ہے کہ "م نہ جو پایا اس کا مختصر خلاصہ یہ ہے: resume کے ساتھ tool data پرانے کیا جائے restart "... سے

2 Tool Design اور MCP Integration (18%)

2.1 Clear Descriptions کے ساتھ Tool Interfaces Design کرنا

Key knowledge:

- Tool descriptions minimal descriptions کرتا ہے؛ LLM tools select میں جن سے mechanism و بنیادی Tool descriptions
- Input formats, example queries, edge cases, اور applicability boundaries کی اہمیت
- Ambiguous یا overlapping descriptions misrouting میں سبب بنتی
- System prompt wording tools کے associations کے ساتھ غیر ارادی بنا سکتی

Key skills:

- کریں distinguish سے واضح طور پر similar alternatives کو tool لکھیں جو ر descriptions ایسی

- Functional overlap (مثلاً) `analyze_content` -> `extract_web_results` tools rename کریں ختم کرنے کے لیے
- واضح ہوں input/output contracts میں تقسیم کریں جن کے specialized tools کو General-purpose tools

2.2 MCP Tools کے لیے Structured Error Responses Implement کرنا

Key knowledge:

- MCP tool responses میں `isError` flag
- **Transient errors** (timeouts), **validation errors** (bad input), **business errors** (policy violations), اور **access/permission errors** کے درمیان فرق
- Generic errors ("Operation failed") میں recovery decisions صحیح روک دیتے ہیں
- Retryable اور non-retryable errors میں فرق

Key skills:

- `errorCategory` (transient/validation/permission), `isRetryable` اور human-readable message جیسی structured metadata واپس کریں
- Business-rule violations کے ساتھ user-facing explanations کے لیے واضح `retryable: false` استعمال کریں
- Transient failures کے لیے subagents میں local recovery صرف وہ errors propagate کریں؛ صرف وہ حل نہ کر سکیں
- Access failures (retry decision) اور valid empty results (no matches) میں فرق کریں

2.3 Agents میں Tools Allocation اور `tool_choice` Configure کرنا

Key knowledge:

- کم کرتی tool selection reliability (مثلاً 18 کے بجائے 4-5) tools بہت زیادہ
- Specialization کے لیے agents رکھنے والے tools سے باہر
- Scoped tool access: cross-role utilities + role-relevant tools صرف
- `tool_choice: "auto"`, `"any"` اور forced tool selection (`{"type": "tool", "name": "..."}`)

Key skills:

- تک محدود رکھیں relevant tools صرف subagent کے
- General tools کو constrained alternatives سے بدلیں (مثلاً `fetch_url` -> `load_document`)
- یقینی tool call کے بجائے text answer استعمال کریں تاکہ `tool_choice: "any"`
- Specific tool force execution order assured کریں تاکہ

2.4 MCP Servers اور Claude Code میں Integrate کرنا

Key knowledge:

- MCP server scope: user کے لیے experiments بمقابلہ (`.mcp.json`) project ٹیموں کے لیے (`~/claude.json`)
- Secrets management کے لیے `.mcp.json` میں environment variable substitution (مثلاً `${GITHUB_TOKEN}`)
- available ہوتے ہیں اور ایک ساتھ discover پر tools connection کے connected MCP servers تمام ہیں
- MCP resources بطور “content catalogs” (task summaries, database schemas) exploratory tool calls کم کرتے ہیں

Key skills:

- Shared MCP servers کو project `.mcp.json` میں env-var-based tokens ساتھ configure کے
- Personal/experimental servers کو `~/claude.json` میں رکھیں
- Standard integrations کے لیے community MCP servers کو custom servers پر ترجیح دیں

2.5 Built-in Tools (Read, Write, Edit, Bash, Grep, Glob) منتخب اور Apply کرنا

Key knowledge:

- **Grep**: file contents اندر search (function names, error messages, imports)
- **Glob**: name/extension patterns سے files تلاش کرنا
- **Read/Write**: full-file operations؛ **Edit**: unique text matches کے ذریعے precise changes
- **Read + Write** پر واپس جائیں تو fail کی وجہ سے Edit non-unique matches

Key skills:

- Content search کے لیے Grep اور pattern-based discovery کے لیے Glob استعمال کریں
 - Read کرنے کے لیے flows trace پھر Grep کے لیے entry points: سمجھ بوجھ رفتہ رفتہ بنائیں
 - Wrapper modules کے ذریعے function usage trace کریں
-

3 اور Claude Code Configuration ڈومین Workflows (20%)

3.1 Hierarchy, Scope, اور Modular Organization کے ساتھ CLAUDE.md Configure کرنا

Key knowledge:

- CLAUDE.md hierarchy: user (`~/claude/CLAUDE.md`), project (`./claude/CLAUDE.md`) یا root (`CLAUDE.md`), اور directory-level (subdirectories میں CLAUDE.md)
- User-level settings میں share کو VCS پر لاگو کو میں اور user صرف ایک
- External files refer (مثلاً `@./standards/coding-style.md`) syntax `@path` کو لے کر CLAUDE.md modular بنے
- Monolithic CLAUDE.md کو topic-focused rule files کو لے کر `./claude/rules/` directory بجائے

Key skills:

- user-اس لیے نئی میں ملتی کیونکہ وہ instructions کو team member (نئی) کریں Hierarchy issues diagnose (میں نئی میں تھیں project-level)
- (مثلاً) @path کریں کہ لیے standards selectively include میں CLAUDE.md کے package ر استعمال کریں (@./standards/testing.md)
- files (testing.md, api-conventions.md, deployment.md) .claude/rules/ کو متعدد CLAUDE.md بڑے میں تقسیم کریں

3.2 Custom Slash Commands اور Skills بنانا اور Configure کرنا

Key knowledge:

- `~/.claude/commands/` بمقابلہ `(shared) (VCS ذریعہ)` **Project commands** میں `~/.claude/commands/` **user commands** میں
- `~/.claude/skills/` میں SKILL.md frontmatter: `context: fork` , `allowed-tools` , `argument-hint`
- `context: fork` skill کو isolated subagent context `context: fork` سے جدا کرتا ہے تاکہ main session pollute نہ ہو
- Personal skill variants `~/.claude/skills/` میں `context: fork` سے جدا کر سکتے ہیں

Key skills:

- Project slash commands `.claude/commands/` پوری team میں رکھیں تاکہ پوری
- Verbose output `context: fork` استعمال کریں `skills isolate` والے `کرنے کے لیے`
- Skill `allowed-tools` محدود کرنے کے لیے `tools` استعمال کر کے استعمال کریں
- Required parameters `argument-hint` مانگنے کے لیے استعمال کریں

- (independent) کب sequentially اور (interdependent) میں کب دیں message ایک issues تمام

Key skills:

- Transformation requirements 3–2 لیجئے کہ concrete input/output examples فراہم کریں
- Implementation expected behavior, edge cases, اور performance requirements کے ساتھ test sets بنائیں
- Design aspects (cache invalidation, failure modes) کے لیے interview pattern سامنے لائیں
- Edge cases کے لیے sample inputs اور expected outputs کے ساتھ concrete test cases دیں

3.6 Claude Code کو CI/CD Pipelines میں Integrate کرنا

Key knowledge:

- Automated pipelines میں non-interactive mode `-p` (یا `--print`) flag
- CI میں structured output `--output-format json` اور `--json-schema`
- CLAUDE.md project context (testing standards, review criteria) فراہم کرتا ہے
- Session context isolation:** code generate والی same session review میں کم مؤثر ہوتی ہے

Key skills:

- Interactive input پر hang نہ لیں۔ CI کو Claude Code میں -p کے ساتھ چلائیں
- Structured results (مثلاً inline PR comments) کے لیے --output-format json + --json-schema استعمال کریں
- new/unfixed issues (صرف) شامل کریں review results کرتے وقت پچھلے run کے بعد دوبارہ commits نہ رپورٹ کریں
- Test generation quality کے لیے available fixtures اور testing standards میں CLAUDE.md بیکتر کریں۔ کے لیے document
- Duplication میں شامل کریں تاکہ context کو existing test files کرتے وقت tests generate نہ کرے اور style consistent رہے

4 ڈومین: Prompt Engineering اور Structured Output (20%)

4.1 Accuracy ☐ Prompts ☐ ساتھ Explicit Criteria ☐ تر کرنے کے لیے ☐ Design کرنا

Key knowledge:

- Explicit criteria vague instructions (مثلاً) “flag comments only when they contradict code” (مقابلہ) “check comment accuracy”
- “be more conservative” generic guidance concrete categorical criteria کم مؤثر

- درست high false-positive rates میں categories پر اثر انداز ہوتے ہیں: کچھ False positives developer trust کم کر دیتی ہیں trust پر بھی categories

Key skills:

- Review criteria define کرنا (minor style) ignore اور کیا (bugs, security) کرنا report کریں: کیا
- High false-positive rates والی categories طور پر disable عارضی
- Explicit severity criteria define کرنا ساتھ code examples کے لیے level پر

4.2 Output Consistency Few-shot Prompting بہتر کرنے کے لیے استعمال کرنا

Key knowledge:

- Few-shot examples consistently formatted, actionable output کا سب سے مؤثر method پیدا کرنے کا سب سے
- Few-shot cases (tool selection, test coverage gaps) کو دکھا سکتا ہے
- Few-shot model میں defaults میں مدد دیتا ہے، صرف generalize پر patterns کو نئے Few-shot model
- Few-shot extraction tasks میں hallucinations کو کم کر سکتا ہے

Key skills:

- Ambiguous scenarios کے لیے targeted examples کے ساتھ 4–2 rationale
- Output format (location, issue, severity, suggested fix) شامل کریں examples دکھانے والے
- میں فرق دکھائیں real issues اور acceptable code patterns دیں جو examples ایسے
- دیں examples کے extraction سے درست documents والے structures مختلف

4.3 JSON Schemas کے ساتھ tool_use اور Structured Output کرنا Enforce

Key knowledge:

- JSON syntax errors کی ضمانت دینے اور tool_use with JSON Schemas schema-conformant output طریقہ کے سب سے reliable کرنے کا سب سے
- کرنی ہوتی tool call میں اسے لازماً "any" واپس کر سکتا ہے؛ model text میں tool_choice: "auto" منتخب کرتی ہے forced selection specific tool
- Strict JSON Schemas syntax errors سے روکتیں semantic errors ختم کرتی ہیں مگر (totals؛ میں ملے؛ values غلط fields غلط)
- Schema design: required vs optional fields؛ enums کے ساتھ "other" + detail string extensibility کے لیے

Key skills:

- JSON Schemas کے ساتھ extraction tools define کریں اور tool_use results سے data parse کریں
- استعمال کریں tool_choice: "any" یقینی بنانے کے لیے structured output کے ساتھ تو schemas متعدد
- Specific tool call force کریں: tool_choice: {"type": "tool", "name": "extract_metadata"}
- Source information میں fields کے لیے values fabricate نہ ہونے کی صورت میں optional/nullable رکھیں

- Extensible categorization `enum` values + detail fields استعمال کریں اور `"unclear"` اور `"other"` کے لیے

4.4 Extraction Quality ☐ لی ☐ Validation, Retries, اور Feedback Loops Implement کرنا

Key knowledge:

- Retry-with-error-feedback: concrete validation errors میں retry prompt کو لیں correction guide شامل کریں
- فوائد:
 - retries میں موجود نئی و تو information source اگر
 - Feedback loop design: track کریں (`detected_pattern`) pattern کو trigger کرنے والا
 - Semantic errors (`totals reconcile`) (جو `tool_use` address سے `tool_use`) (نہیں `totals reconcile`)

Key skills:

- دیں follow-up prompts کے ساتھ specific validation errors اور، incorrect extraction، Original document،
- تو اس کی شناخت کریں (میں کے اور external document صرف required info کے فوائد کے و retry جب
- False positives analyze کے لیے detected_pattern fields میں findings کرنے کے لیے
- Discrepancies detect کے لیے calculated_total اور stated_total دونوں extract کے کر کے self-correction design کریں

4.5 Efficient Batch Processing Strategies Design کرنا

Key knowledge:

- Message Batches API: 50% savings, 24-hour processing window, latency SLA guarantee
- Batch processing non-blocking tasks (overnight reports, audits) اور blocking tasks (pre-merge checks) کے لیے مناسب
- Batch API single request میں multi-turn tool calling support کرتا ہے
- custom_id fields batch request/response correlate کرتے ہیں

Key skills:

- استعمال Batch API کے لیے overnight/weekly workloads اور synchronous API کے لیے blocking checks کریں
- SLA needs کے مطابق batch submission cadence plan کریں (24-hour guarantee مثلاً 30) (4-hour processing کے ساتھ 4-hour windows)
- Failed documents (identified by `custom_id`) کو دوبارہ submit کرنے کے لیے failures handle کریں
- Large-scale processing کے لیے sample کے ساتھ iterate کریں

4.6 Multi-instance اور Multi-pass Review Architectures Design کرنا

Key knowledge:

- Self-review limitations: model کا challenge ساتھ رکھنا اور اپنے فیصلوں کو reasoning context اپنی model امکان کم ہوتا ہے
- Independent review instances (generation context) subtle issues میں بہتر ہوتی ہیں (کم بغیر generation context)
- Multi-pass review: per-file local analysis + cross-file integration pass attention dilution سے بچانا ہے

Key skills:

- Changes review استعمال second independent Claude instance کے بغیر generation context کے لیے کریں
- Multi-file reviews cross-file dataflow میں تقسیم کریں تاکہ per-file passes + integration passes analysis ہو سکے
- Self-rated confidence route طریقہ سے calibrated reviews استعمال کر کے verification passes کے ساتھ کریں

5 Context Management اور Reliability (15%) ڈومین 5

5.1 Critical Information Conversation Context محفوظ رکھنا اور Manage کرنا

Key knowledge:

- Progressive summarization کے خطرات: numeric values, percentages, اور dates vague summaries میں بدل جاتے ہیں
- Lost-in-the-middle effect: models long inputs کے شروع اور آخر کو زیادہ reliable سے کر کے process طریقہ سے findings miss ہیں، لیکن درمیان کے
- Tool outputs relevance میں context میں disproportionately جمع ہو سکتے ہیں (40+ fields جب 5 fields کے مقابلہ میں درکار ہوں)
- Subsequent API requests کی اہمیت full conversation history میں

Key skills:

- Transactional facts کو summarized history سے باہر ایک persistent “case facts” block میں نکالیں
- Verbose tool outputs کو relevant fields تک trim کریں
- Key findings کو aggregated data کے شروع سے explicit section headings کے ساتھ رکھیں
- Subagents کے structured outputs میں metadata (dates, sources) شامل کروائیں

5.2 Effective Escalation Patterns Design اور Ambiguity Resolve کرنا

Key knowledge:

- Suitable escalation triggers: human کی explicit request, policy gaps/exceptions, پیش رفت نہ ہو پانا
- Immediate escalation (explicit request) بمقابلہ attempt-to-resolve (agent scope اندر)
- Sentiment analysis اور model confidence self-ratings case complexity کی unreliable proxies
- Multiple customer matches کی heuristic guessing کی بجائے additional identifiers چائیں

Key skills:

- System prompt میں explicit escalation criteria few-shot examples کے ساتھ دیں
- Human کی explicit request کو بغیر اضافی investigation فوراً execute کے
- Escalate نہ تو silent یا ambiguous کے لیے کسی policy جب
- Tool results میں multiple matches کے تو additional identifiers مانگیں

5.3 Multi-agent Systems میں Error Propagation Strategies Implement کرنا

Key knowledge:

- Structured error context (failure type, query, partial results, alternatives) coordinator کو smarter recovery دیتا ہے
- Access failures (timeouts retry decision میں) اور valid empty results (no matches) میں فرق کریں
- Generic error statuses (“search unavailable”) coordinator کے valuable context چھپا دیتے ہیں
- Silent suppression کے anti-patterns کرنا دونوں workflow abort پر پورا failure یا ایک

Key skills:

- Structured error context میں failure type, partial results, possible alternatives کی واپس کریں
- Access failures کو valid empty results الگ کریں
- Transient failures کے لیے subagents میں local recovery صرف non-recoverable errors partial results کے ساتھ propagate کریں
- Synthesis میں coverage annotate کی اچھی طرح supported اور کے ان gaps ہیں

5.4 Large Codebases Investigate وقت Context Efficiently Manage کرنا

Key knowledge:

- Long sessions میں context degradation: model unstable answers اور مخصوص classes دینے لگتا ہے
- Scratchpad files key findings کو context boundaries میں پار محفوظ رکھتے ہیں
- Subagents کو delegate کے verbose discovery output isolate کے جاتی ہیں

- Structured state persistence crash recovery □□ ممکن بناتی

Key skills:

- Specific questions میں رکھیں high-level coordination main agent spawn کریں جبکہ subagents spawn کر کے لیے
- Key findings استعمال کریں scratchpad files کرنے کے لیے refer محفوظ رکھنے اور بعد میں
- phase اگلا subagents spawn کرنے کے لیے summarize key findings کرنے کے لیے
- Long investigations /compact کم کرنے کے لیے context usage کے دوران

5.5 Human Oversight اور Confidence Calibration کے ساتھ Workflows Design کرنا

Key knowledge:

- Aggregate metrics (97% overall accuracy) specific document types یا fields سے چھپا سکتا ہے
- Stratified random sampling high-confidence extractions error rates measure کرتا ہے
- Field-level confidence calibration labeled validation sets کے ذریعے
- Automation کے ذریعے document type اور field segment کے حساب سے accuracy validate کریں

Key skills:

- New error patterns detect کریں stratified random sampling implement کریں
- Stable performance validate کریں accuracy analyze کریں حساب شد field اور document type کریں
- Field-level confidence scores output کریں review thresholds calibrate کریں ذریعہ labeled data اور
- Low-confidence یا ambiguous-source extractions human review کریں route کی طرف

5.6 Multi-source Synthesis میں Provenance Preserve کرنا اور Uncertainty Handle کرنا

Key knowledge:

- محفوظ نہ رہیں mappings “claim → source” کہو جاتی ہیں اگر attribution کا دوران Summarization
- برقرار رہتی چاہئیں structured mappings کا دوران Aggregation
- conflict annotate کا ساتھ attribution منتخب کرنے کا بجائے value ایک arbitrarily کو conflicting statistics کریں handle کا
- publication/collection dates سمجھنے سے بچنے کا لیے contradiction کو Temporal differences شامل کریں

Key skills:

- Subagents کے “claim → source” mappings (URL, document name, quotes) output کروائیں
- Reports کو stable findings اور disputed ones کے structure میں الگ کریں
- Conflicting values کے annotations اور ساتھ محفوظ رکھیں
- Correct temporal interpretation کے publication dates کے لیے شامل کریں

- Content type مطابق render کریں: financial data tables میں، news prose میں، technical findings میں structured lists میں

Claude Certified Architect — Foundations Certification

مطالعہ گائیڈ (سرکاری امتحانی گائیڈ کی بنیاد پر)

تعارف

Claude Certified Architect — Foundations سرٹیفیکیشن اس بات کی تصدیق کرتی ہے کہ ایک ماہر Claude-based Claude Code، Claude Agent SDK، Claude API، اور Model Context Protocol (MCP) بنیادی ٹیکنالوجیز جن کو جانچتا ہے — یعنی وہ بنیادی علم کو جانچتا ہے۔ ساتھ پروڈکشن ایپلیکیشنز تیار کی جاتی ہیں۔

بنانا agentic systems امتحان کے سوالات حقیقی دنیا کے منظرناموں پر مبنی ہوتے ہیں: کسٹمر سپورٹ کے لیے، multi-agent research pipelines میں شامل کرنا CI/CD کو Claude Code، ڈیزائن کرنا developer productivity tools نکالنا structured data بنانا، اور غیر ساختہ دستاویزات سے tools

هدف امیدوار

کرتا ہے آپ ship کے ساتھ پروڈکشن ایپلیکیشنز ڈیزائن اور Claude کے جو solution architect مثالی امیدوار ایک کے پاس کم از کم 6 ماہ کا عملی تجربہ ہونا چاہیے:

- Claude Agent SDK — multi-agent orchestration، subagents کو delegate کرنا، tool integration، lifecycle hooks
- Claude Code — CLAUDE.md، MCP servers، Agent Skills، planning mode
- Model Context Protocol (MCP) — backend integration کے لیے tools اور resources
- Prompt engineering — JSON schemas، few-shot examples، data extraction templates
- Context windows — لمبی دستاویزات، multi-agent context passing
- CI/CD pipelines — automated code review، test generation
- Escalation and reliability — error handling، human-in-the-loop

امتحانی فارمیٹ

پیرامیٹر	قدر
سوال کی قسم	Multiple choice (1 درست 4)

اسکورنگ	100–1000 scale, passing score 720
Guessing penalty	نہیں (ہر سوال کا جواب دیں!)
Scenarios	8 4 میں سے (randomly selected)

امتحانی مواد: 5 ڈومینز

ڈومین	وزن
1. Agent architecture and orchestration	27%
2. Tool design and MCP integration	18%
3. Claude Code configuration and workflows	20%
4. Prompt engineering and structured output	20%
5. Context management and reliability	15%

امتحانی منظر نامہ

Scenario 1: Customer Support Agent

ایک agent جو آپ کے Claude Agent SDK کے ساتھ returns, billing disputes, اور account issues کے ساتھ agent MCP tools (get_customer, lookup_order, process_refund, escalate_to_human) کے ساتھ 80% first-contact resolution کے ساتھ 80% استعمال کرتا ہے۔

Scenario 2: Claude Code کے ساتھ Code Generation

code generation, refactoring, debugging, documentation کے ساتھ Claude Code کے ساتھ development کے ساتھ custom slash commands اور CLAUDE.md configuration کے ساتھ integrate اور planning mode کے ساتھ

Scenario 3: Multi-Agent Research System

coordinator agent specialized subagents کے ساتھ tasks: web research, document analysis, synthesis, اور report generation کے ساتھ citations کے ساتھ

Scenario 4: Developer Productivity Tools

agent engineers کے ساتھ unfamiliar codebases explore اور boilerplate code کے ساتھ automate Built-in tools (Read, Write, Bash, Grep, Glob) اور MCP servers کے ساتھ

Scenario 5: Claude Code for Continuous Integration

Claude Code کو CI/CD pipeline میں ضم کریں تاکہ automated code reviews، test generation، pull request feedback اور pull request feedback کم سے کم false positives بننے چاہئیں کہ Prompts حاصل ہو request feedback

Scenario 6: Structured Data Extraction

کرتا ہے، validate کے ذریعے JSON schemas کو output، سسٹم غیر ساختہ دستاویزات سے معلومات نکالتا ہے اور درست طور پر سنبھالنے چاہئیں کہ edge cases برقرار رکھتا ہے اس high accuracy اور

Scenario 7: Conversational AI Architecture Patterns

multi-turn conversational systems میں جن میں context window management، turns متعدد، memory strategies، execution محفوظ، tool design کے لیے instructions میں اور مہم یا متضاد، user inputs کو سنبھالنا شامل ہے

Scenario 8: Agentic AI Tools (امداد موجود نہیں — ہماری مدد کریں)

میں شامل نہیں ہے، guide نہ دی لیکن ابھی تک اس candidates کی اطلاع امتحان دینے والے scenario اس میں share میں GitHub Issues کے سوالات دیکھیں، تو انہیں scenario اگر آپ نہ اصل امتحان میں اس شامل کر سکیں آپ کا تعاون اس شخص کی مدد کرے گا جو امتحان کی تیاری کر coverage تاکہ ہم مکمل رہے

نظریاتی بنیادیں: حصہ 1

یہ حصہ و بنیادی نظریہ بیان کرتا ہے جس کی آپ کو امتحان میں کامیابی کے لیے ضرورت ہے مواد کو topic کے مطابق — اس سے domains کے مطابق ترتیب دیا گیا ہے، نہ کہ امتحانی concepts اور technologies کی گہری سمجھ بھنی ہے

کی بنیادیں Claude API — Model Interaction: باب 1

دستاویزات: Messages API | Prompt Engineering

1.1 API Request Structure

Claude API میں عموماً Claude Messages API request پر کام کرتا ہے request-response model ایک شامل ہوتا ہے:

```
{
  "model": "claude-sonnet-4-6",
  "max_tokens": 1024,
  "system": "You are a helpful assistant.",
  "messages": [
    {"role": "user", "content": "Hi!"},
    {"role": "assistant", "content": "Hello!"},
    {"role": "user", "content": "How are you?"}
  ],
  "tools": [...],
  "tool_choice": {"type": "auto"}
}
```

۱۰:۰۰ م فیلڈز

- `model` — model selection (`claude-opus-4-6` , `claude-sonnet-4-6` , `claude-haiku-4-5`)
- `max_tokens` — response `توکنز میں زیادہ سے زیادہ`
- `system` — system prompt (model کی تعریف)
- `messages` — conversation history (**full history** `پوری تاریخچه`)
- `tools` — available tools کی definitions
- `tool_choice` — tool selection strategy

1.2 Message Roles

استعمال کرتا □□ instructional role اور ایک conversational roles دو array messages

- `user` — user messages, tool results (ایک `user` role message کے اندر `tool_result` کی صورت میں بھیجے جاتے ہیں، ایک الگ content block)
- `assistant` — model responses (بھیجے وقت شامل ہونے والے history میں tool use requests (`tool_use` content blocks))
- `system` — top-level `system` field سے سیٹ کیا جا سکتا ہے (turn کے ذریعے سیٹ کیا جا سکتا ہے) یا (سے لاگو ہوتا ہے) اس نقطے سے آگے) شامل کیا جا سکتا ہے inline کی صورت میں `{"role": "system", ...}` میں (تبعی — نیچے دیکھیں placement rules لاگو ہوتا ہے، مخصوص)

کے طور پر user-role message کے طور پر نہیں بھیجا جاتا۔ ایک message والے role "tool" کو Tool results کے content block tool_result میں ایک content پر بھیجے جاتا ہے۔ میں جس کے شامل ہوتا ہے۔

```
{
  "role": "user",
  "content": [
    {
      "type": "tool_result",
      "tool_use_id": "toolu_01...",
      "content": "..."
    }
  ]
}
```

top-level کا طور پر ظاہر ہو سکتا ہے، نہ صرف role میں بھی ایک array messages براہ راست system کیے بغیر invalidate کو cached prefix کا system field top-level کے ذریعے اس کا مقصد parameter system placement rules شامل کرنا اس کے مخصوص instructions درمیان conversation:

- پر ختم ہونے والے server tool use یا (tools_result blocks بشمول وہ جس میں) user turn یا ایک assistant turn فوراً بعد آنا چاہیے۔
- ہونا چاہیے element کا آخری array سے پہلے آنا چاہیے یا assistant turn کسی
- ملتا error کے درمیان نہیں آ سکتا — ایسا کرنے پر 400 tool_result اور اس کے block tool_use کسی
- پہلے والوں پر اور بعد (کے درمیان شامل کیے گئے conversation بشمول) system messages بعد میں آنے والے پر فوقیت رکھتے ہیں۔ system field top-level کے لیے turns

محفوظ state کے درمیان Model requests بھیجیں conversation history میں مکمل API request ام بات: ہر آزاد ہوتی ہے call نہیں رکھتا؛ ہر

1.3 stop_reason

Claude API response میں stop_reason کے جو بتاتا ہے کہ:

Value	Description	Action
"end_turn"	نہ جواب مکمل کر لیا	کو نتیجہ دکھائیں user
"tool_use"	Model tool call چاہتا ہے	واپس دیں result چلائیں اور tool
"max_tokens"	Token limit ختم ہو گئی	response truncated؛ limit؛
"stop_sequence"	Stop sequence مل گئی	کریں handle کے مطابق application logic

کرتے loop control ہیں — یہی signals سب سے اہم "end_turn" اور "tool_use" میں Agentic systems ہیں۔

1.4 System Prompt

System prompt متعین کرنا ہے behavioral rules اور context جو instruction ایک خاص

- messages array الگ system field کا حصہ نہیں ہوتا؛ الگ
- user messages پر فوقیت رکھتا ہے
- ہو کر پوری گفتگو میں لاگو رہتا ہے load ایک بار
- کی وضاحت کے لیے استعمال ہوتا ہے output format، اور role، constraints،

کو غیر ارادی طور پر متاثر کر سکتی ہے مثال tool selection wording کی system prompt: امتحان کے لیے اہم ضرورت سے زیادہ get_customer کو model کی تصدیق کریں جیسی ہدایت customer کے طور پر "میش استعمال کرنے پر لا جا سکتی ہے

1.5 Context Window

Context window اس میں شامل process ایک وقت میں model جس (tokens) و کل متن کر سکتا ہے:

- System prompt
- message history مکمل
- Tool definitions
- Tool results

اہم مسائل:

1. **Lost-in-the-middle effect:** درمیان کی input لمبے، مؤثر ہوتی ہیں، آغاز اور اختتام کی معلومات زیادہ مؤثر ہوتی ہیں، حل: اہم معلومات شروع یا آخر میں رکھیں۔ details miss
2. **Tool results** واپس tool 40+ fields شامل کرتی ہیں اگر output میں tool call context کا جمع ہونا: ہر tool results ضائع ہونا۔ context مگر صرف 5 اہم ہوں تو
3. **Progressive summarization:** history compress پر numeric values، percentages، اور dates اکثر vague ہوتی ہیں۔

2 tool_use اور Tools: باب 2

Tool Use: دستاویزات

2.1 tool_use کیا ہے

Model کو Claude external functions call جس کے ذریعے mechanism و tool_use واپس کرتا result کر کے execute اسے code بنانا ہے؛ آپ کا structured tool call بنانا چلتا ہے — و code درست ہے

2.2 Tool Definition

تو define کے ذریعے tool JSON schema ہے:

```
{
  "name": "get_customer",
  "description": "Finds a customer by email or ID. Returns the customer profile, including name, email, order history, and account status. Use this tool BEFORE lookup_order to verify the customer's identity. Accepts an email (format: user@domain.com) or a numeric customer_id.",
  "input_schema": {
    "type": "object",
    "properties": {
      "email": {"type": "string", "description": "Customer email"},
      "customer_id": {"type": "integer", "description": "Numeric customer ID"}
    },
    "required": []
  }
}
```

Tool description نکات:

- Minimal descriptions کا سبب بنتی selection غلط Mechanism کا بنیادی tool selection کی Description ہیں
- edge cases، کیا input format، کیا کرتا، کیا واپس کرتا tool: میں یہ شامل کریں Description، اور یہ کب استعمال ہو
- Overlapping descriptions سے بچیں
- Built-in tools اور MCP tools overlap کو تو واضح بنائیں MCP tool

2.3 tool_choice

Value	Behavior	کب استعمال کریں
<code>{ "type": "auto" }</code>	خود طے کرتا Model	default عام
<code>{ "type": "any" }</code>	کرنی tool call کو لازماً Model کوگی	structured output جب چاہیے
<code>{ "type": "tool", "name": "extract_metadata" }</code>	کوگی call لازماً tool خاص	forced first step / order کی لیں

2.4 JSON Schemas for Structured Output

reliable لیں کا سبب structured output سے Claude استعمال کرنا `tool_use` کی سائنہ JSON schemas کی semantic correctness نافذ کرتا، مگر required structure کم کرتا اور syntax errors طریقہ سے ضمانت نہیں دیتا

Design اصول:

- رکھیں جو ہمیشہ ملنے چاہئیں fields required صرف وہ
- استعمال کریں nullable fields کی لیں missing data ممکنہ
- enum values جیسے "unclear" اور "other" information شامل کریں تاکہ

2.5 Syntax اور Semantic Errors

Error type	Example	Mitigation
Syntax	Invalid JSON، field غلط	JSON schema کے ساتھ tool_use
Semantic	Totals match نہیں کرتے، value field غلط	Validation checks، feedback retry

3 بنانا Claude Agent SDK — Agentic Systems

دستاویزات: [Agent SDK](#) | [Hooks](#) | [Subagents](#) | [Sessions](#)

3.1 Agentic Loop

Agentic loop میں Claude کے actions کی sequence چلاتا ہے، جو اب اس کے دیئے ہوئے tools کو بھیجتا ہے:

- Claude کو tools کے ساتھ request بھیجتا ہے
- Response لیتا ہے
- stop_reason چیک کرتا ہے:
 - "tool_use" → tool چلائے، result history کو دوبارہ request میں ڈالیں، پھر دوبارہ
 - "end_turn" → task مکمل، user کو دینے کے نتیجے میں
- مکمل ہونے تک دہرائیں

Text parsing یا arbitrary iteration دیکھیں کہ `stop_reason == "end_turn"` کے لیے completion signal: درست limit قابلِ اعتماد نہیں ہے

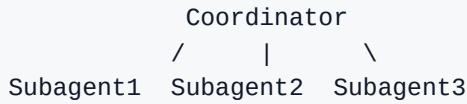
3.2 AgentDefinition

```
agent = AgentDefinition(  
    name="customer_support",  
    description="Handles customer requests for returns and order issues",  
    system_prompt="You are a customer support agent...",  
    allowed_tools=["get_customer", "lookup_order", "process_refund",  
    "escalate_to_human"],  
)
```

اصول اپنائیں: Least privilege کے allowed_tools اور name، description، system_prompt کے لیے مہیا کیے گئے ہیں

3.3 Coordinator اور Subagents

Multi-agent systems میں hub-and-spoke topology اکثر استعمال کرتی ہے:



Coordinator task کو توڑتا ہے، subagents کے نتائج جوڑتا ہے، کام سونپتا ہے، منتخب کرتا ہے، errors handle کرتا ہے، جواب پلانچانا ہے final تک user کو، اور

نہیں کرتے، inherit خود بخود parent history الگ ہوتا ہے؛ وہ context کا subagents کا ماحول:

3.4 Task Tool

Subagents Task tool کے ذریعے spawn کیے جاتے ہیں۔ Coordinator کے allowed_tools میں "Task" ہونا چاہیے۔

context passing:

- Task: "Analyze the document" کافی نہیں
- Task: "Analyze this document. Full text: ... Output schema: ..." درست ہے

Coordinator ایک response میں multiple Task s لے سکتا ہے، لہذا parallel subagents بھیج سکتا ہے،

3.5 Hooks

Hooks lifecycle کے مخصوص points پر work کیے جاتے ہیں۔

PostToolUse tool result کو model کے لیے normalize تک پلانچانا، سہل ہونا:

```

@hook("PostToolUse")
def normalize_dates(tool_result):
    return normalized_result

```

PreToolUse policy violation block کیے جاتے ہیں:

```

@hook("PreToolUse")
def enforce_refund_limit(tool_call):
    if tool_call.name == "process_refund" and tool_call.args.amount > 500:
        return redirect_to_escalation(tool_call)

```

Critical business rules ہوتے ہیں۔ prompt instructions probabilistic ہوتے ہیں؛ hooks deterministic: یاد رکھیں

باب 4: Model Context Protocol (MCP)

4.1 MCP کیا ہے

MCP resource سے جوڑنا اس کے تین بنیادی Claude کو external systems جو open protocol ایک types ہیں:

1. **Tools** — actions کے functions
2. **Resources** — context data (documentation, schemas, catalogs)
3. **Prompts** — predefined task templates

4.2 MCP Servers

MCP server ایک process جو tools/resources expose کرتا ہے Connect پر tools discover ہوتا ہے۔ مطابقت استعمال ہوتا ہے۔ descriptions ہیں اور

4.3 Configuration

Project config (`.mcp.json`) کے لیے:

```
{
  "mcpServers": {
    "github": {
      "command": "npx",
      "args": ["-y", "@modelcontextprotocol/server-github"],
      "env": {"GITHUB_TOKEN": "${GITHUB_TOKEN}"}
    }
  }
}
```

User config سے آتی ہے۔ secrets environment variables میں ہوتی ہے؛ Project config version control (`~/.claude.json`) personal/experimental servers کے لیے ہوتی ہے۔

4.4 isError

MCP tool errors میں `isError: true` Generic errors کے بجائے structured error کو واپس کریں:

```
{
  "isError": true,
  "content": {
    "errorCategory": "transient",
    "isRetryable": true,
    "message": "The service is temporarily unavailable."
  }
}
```

4.5 Resources

Resources: schemas, docs, content catalogs, summaries, exploratory tool calls, agent action, data, and agent. کم کرتے ہیں۔

5 اور Claude Code — Configuration Workflows: باب 5

دستاویزات: [Claude Code](#) | [Memory / CLAUDE.md](#) | [Skills](#) | [MCP](#) | [Hooks](#) | [Sub-agents](#) | [GitHub Actions](#) | [Headless](#)

5.1 CLAUDE.md Hierarchy

CLAUDE.md تین سطحوں پر ہو سکتا ہے:

- User-level:** `~/.claude/CLAUDE.md` — پر لاگو user صرف اسی
- Project-level:** `.claude/CLAUDE.md` یا `root CLAUDE.md` — team پوری
- Directory-level:** subdirectories میں CLAUDE.md — folder conventions مخصوص

5.2 @path Syntax

CLAUDE.md میں @path external files refer تاکہ جا سکتے ہیں تاکہ config modular کیے جا سکتے ہیں:

Coding standards @./standards/coding-style.md میں ہیں
Test rules @./standards/testing-requirements.md میں ہیں

5.3 .claude/rules/

.claude/rules/ rules کو topic لحاظ سے organize کرنا اور paths کے ذریعے conditional loading کرتا ہے۔ دیتا ہے:

```
---  
paths: ["src/api/**/*"]  
---
```

استعمال کریں explicit error handling اور async/await کہ لیم API files

5.4 Commands اور Skills

`.claude/commands/` project commands کے لیے، جبکہ `~/claude/commands/` personal commands کے لیے۔
ملتی ہیں۔ skills کے ساتھ frontmatter میں `.claude/skills/` کے لیے۔

5.5 Skills اور context: fork

context: fork skill کو isolated context میں چلانا، allowed-tools سے tool access اور محدود کریں،
argument-hint سے required input مانگیں۔

5.6 Planning Mode

Planning mode کے لیے، بہترین unfamiliar codebases اور multiple approaches، تبدیلیوں architectural بڑی
کے لیے، بہتر changes کے واضح Direct execution اور واضح۔

5.7 /compact اور /memory

/compact context compress کرتا ہے، اور /memory notes اور preferences میں مدد دیتا ہے۔

5.8 Claude Code CLI for CI/CD

-p / --print non-interactive mode کے لیے استعمال ہوتا ہے:

```
claude -p "Analyze this pull request for security issues"
```

Structured output کے لیے --output-format json اور --json-schema استعمال کریں۔

باب 6: Prompt Engineering — Advanced Techniques

دستاویزات: [Prompt Engineering](#) | [Anthropic Cookbook](#)

6.1 Few-shot Prompting

Few-shot prompting میں 2–4 examples کے expected behavior کے لیے دکھایا جاتا ہے۔
vague instructions کے لیے example ambiguity زیادہ مؤثر ہے کیونکہ کم کرتا ہے۔

6.2 Explicit Criteria

decide severity کرنا، اور ignore کرنا، کیا flag دیں: کیا clear criteria بجائے vague instructions کرنا

6.3 Prompt Chaining

Complex task کو focused steps میں توڑیں: local analysis، integration pass، attention dilution کم کرتا

6.4 Interview Pattern

domain سامنے لاتا، خاص طور پر جب design ambiguity پوچھنا clarifying questions سے Claude unfamiliar

6.5 Validation and Retry

کریں retry کے ساتھ، specific validation error اور، original document، incorrect output، Extraction fail کو

6.6 Self-correction

کریں discrepancy handle کو تو conflict دونوں نکال؛ Model stated value اور calculated value

7: Message Batches API

دستاویزات: [Message Batches](#)

7.1 Overview

Message Batches API asynchronous processing 50% savings، up to 24 hours window، اور، multi-turn tool calling supported، custom_id کے requests اور responses link میں

7.2 کب استعمال کریں

Pre-merge checks اور interactive tasks کے لیے synchronous API، overnight reports اور bulk jobs کے استعمال کریں Batch API کے لیے

7.3 custom_id

custom_id result کو original document سے جوڑنا، failures identify اور صرف failed items کو دوبارہ submit کرنے میں مدد دیتا ہے

8 Task Decomposition Strategies

8.1 Fixed Pipelines

Predictable tasks کے لیے fixed pipeline جیسے Document → Extraction → Validation → Final output

8.2 Dynamic Decomposition

Open-ended tasks میں subtasks intermediate results کی بنیاد پر بنتے ہیں

8.3 Multi-pass Code Review

Large PRs کے لیے per-file analysis کے بعد cross-file integration pass تاکہ attention dilution کم ہو

9 Escalation اور Human-in-the-Loop

9.1 Escalate کب کریں

Explicit user request، policy gaps، progress نہ ہونا، threshold سے اوپر financial action، یا ambiguous matching میں sentiment analysis اور self-rated confidence unreliable کے لیے escalate کی صورت میں

9.2 Escalation Patterns

human کو hand off یا explicit insistence یا policy gap کے لیے کی کوشش کریں، پھر اگر resolve نہ ہو، مسئلہ کریں

9.3 Structured Handoff

Escalation پر human شامل کریں تاکہ customer ID، issue summary، actions taken، اور recommended action کے مکمل context

9.4 Confidence Calibration

Field-level confidence scores, calibration on labeled data, اور stratified sampling استعمال کریں

10 Error Handling میں Multi-agent Systems: باب 10

10.1 Error Categories

Transient errors retryable ہیں، validation errors input fix ہیں مانگتی، business errors policy-driven ہوتی ہیں کی جاتی ہیں۔ permission errors generally escalate ہیں، اور

10.2 Anti-patterns

Generic “search unavailable” response, silent suppression, workflow abort یا پورا کرنا نقصان دہ۔ Structured errors واپس کریں

10.3 Structured Subagent Error

Structured error میں failure type, attempted query, partial results, alternative approaches, اور coverage شامل کریں impact

10.4 Final Synthesis میں Coverage

Report میں کون سی sections fully covered ہیں، کون سی partially covered ہیں، اور کون سی gaps باقی ہیں

11 Context میں Production Systems: باب 11

11.1 Facts Separate Block میں

Transactional facts کو persistent case facts block میں رکھیں تاکہ summarization باوجود critical data نہ ہو

11.2 Tool Results Trimming

PostToolUse hook رکھیں relevant fields سے صرف

11.3 Position-aware Input

کم □□ lost-in-the-middle effect آخر میں رکھیں تاکہ □□ action items شروع میں اور Important findings

11.4 Scratchpad Files

محفوظ □□□□ context میں long investigations میں لکھیں تاکہ □□ scratchpad findings verbose

11.5 Subagents کو Delegate کرنا

رکھیں □□ summary میں صرف main context میں دیں اور exploration subagents

11.6 Structured State Persistence

ممکن □□ crash recovery کر □□ تاکہ □□ export میں JSON file یا state manifest اپنی agent □□

12 کو محفوظ رکھنا Provenance: باب 12

12.1 Attribution Loss

رکھیں □□ confidence اور ، publication date ، source URL ک □□ ساتھ Claim

12.2 Conflicting Data

کریں □□ conflict flag ک □□ ساتھ محفوظ رکھیں اور attribution دیں تو دونوں کو values مختلف sources اگر

12.3 Dates

دیکھیں □□ publication dates سمجھنے □□ سد □□ □□ contradiction کو Temporal differences

12.4 Content Type □□ Render

time series میں ، اور technical findings structured lists ، news prose میں ، میں financial data tables chronological order میں دکھائیں □□

باب 13: Claude Code Built-in Tools

13.1 Tool Selection Reference

Task	Tool	Example
تلاش کرنا file pattern	Glob	<code>**/*.test.tsx</code>
file contents میں search	Grep	function name, error message
پڑھنا file	Read	analysis
بنانا file نیا	Write	scratch
precise edit	Edit	unique text match
shell command	Bash	git, npm, tests

13.2 Incremental Investigation

دیکھیں consumer files، پھر usages، پھر entry points، پڑھیں؛ پھر files ایک ساتھ سب

13.3 Read + Write Fallback

کریں Write، پھر programmatically کریں، تبدیلی Read، تو Edit fail اگر

امتحانی ڈومین نوٹس: II حصہ

1 ڈومین: Agent Architecture اور Orchestration (27%)

Agentic loops, coordinator–subagent architecture, Task، spawning, hooks, escalation, اور multi-step workflows پر فوکس کریں

2 ڈومین: Tool Design اور MCP Integration (18%)

Tool descriptions, tool_choice, JSON schema output, structured errors, MCP servers, resources, اور built-in tools، فرق کو سمجھیں

3 ڈومین: Claude Code Configuration اور Workflows (20%)

CLAUDE.md hierarchy, `.claude/rules/`, commands, skills, planning mode, `context: fork`, اور CI/CD integration ☐ م ☐ یں ☐ ا

4 ڈومین: Prompt Engineering اور Structured Output (20%)

Few-shot prompting, explicit criteria, prompt chaining, validation/retry loops, JSON schema enforcement, batch-friendly prompting

5 ڈومین: Context Management اور Reliability (15%)

Summarization risks, lost-in-the-middle, trimming tool output, confidence calibration, escalation, اور
provenance preservation □ یاد رکھیں

Scenario 7: Conversational AI Architecture Patterns

سوال 61

کرنے کے لیے preview سے پلاٹ اثرات کا execution میں tool remove_team_member صورتحال: آپ کو براہ راست agent سے معلوم ہوتا ہے کہ Production monitoring parameter dry_run: boolean ہے۔ اگر dry_run=false ہے تو preview step کو bypass کر کے call کے ساتھ، deletion کیا ہو اور user کو preview سے پلاٹ ایک

سوال 62

فولٹا fail وقت 12% tool `search_catalog` سڈ معلوم ہوتا ہے آپ کا Production monitoring صورتحال میں جو کبھی query syntax errors پر کامیاب ہوتے ہیں، اور 4 retry میں جو network timeouts: 8 ہوتے ہیں، جس سڈ ہڈ فائدہ return یکساں طریقہ سڈ error types کامیاب نہیں ہوتے فی الحال دونوں retries ہوتی ہیں

کریں گے؟ modify کیسے error handling کی tool آپ

- A) System prompt میں few-shot examples شامل کریں جو network errors demonstrate کرنے میں فرق کیسے کریں syntax errors
- B) کریں uniformly exponential backoff retry logic apply پر errors تمام
- C) Tool syntax errors کریں؛ automatic retry with backoff implement لیں network timeouts اندر Tool parameter validation details [درست] سڈ ساتھ فوری واپس کریں
- D) سڈ ساتھ واپس کریں error type details اور boolean flag `retryable` کو errors تمام

کو tool — abstraction boundary کرنا صحیح retries handle لیں transient errors پر Tool level: کیوں C پر agent کر سکتا ہے بغیر deterministic retry logic implement کے بارے میں یقینی علم اور و error type Uniform backoff (B) syntax errors کر (A) follow prompt instructions کر یا (D) interpret flag انحصار کے و errors پر وقت ضائع کرتا ہے جو کبھی کامیاب نہیں ہوں گے

سوال 63

بہت کم risk tolerance اور "میری user، میں turns کی گفتگو کے کئی Investment strategy صورتحال کرنا چاہتا ہے؟ invest کرنا چاہتا ہوں" اب و پوچھتا ہے: "مجھے کس چیز میں returns maximize پھر" میں اپنا

سڈ priority کی اصل recommendation user سب سڈ بہتر طریقہ سڈ یقینی بنانا کے approach کونسا ہے؟

- A) سڈ پوچھیں کے کونسی چیز زیادہ اہم ہے [درست] user تضاد کو سامنے لائیں اور
- B) فراہم کریں recommendations کے لیں الگ الگ scenarios دونوں
- C) کے ساتھ آگے بڑھیں preference سب سڈ حال میں بیان کی گئی
- D) تجویز کریں balanced portfolio تضاد کو حل کیے بغیر

مانگنا کی واحد clarification برا راست متضاد ہوں، تو تضاد سامنے لانا اور preferences کی user کیوں: جب A low risk کرنا اور Returns maximize کے اصل ارادہ سڈ ہے آہنگ و recommendation user طریقہ کے بنیادی طور پر ناقابل مطابقت ادا ہیں جن کے لیں انسانی فیصلہ درکار tolerance

سوال 64

jazz کے "مجھے User کرتے ہیں playlist preferences refine میں conversation turns صورتحال پسند ہیں؟ genres پوچھتا ہے "آپ کو کون سا Claude، بعد messages پسند ہے" کے دو

سب سڈ زیادہ ممکنہ وجہ کیا ہے؟

- A) Claude درکار vector database connection برقرار رکھنے کے لیں conversation memory کو

- B) Model کی context window exceed ہو گئی
- C) Claude API کو session_id parameter درکار
- D) شامل نہیں کر رہی [درست] messages array messages میں پچھلے application آپ کی

D کیوں: Claude پاس server-side memory کے request stateless ہیں — API call کے پاس Claude turns کے پاس پچھلے conversation history میں مکمل messages array کا کوئی علم نہیں؛ دو Claude architecture (C) session_id اور (A) Vector databases میں ہیں؛ ناممکن (B) context window overflow کے لیے exchange کے messages

سوال 65

History تک پہنچ گئی 78,000 tokens conversation، بعد cooking session صورتحال: 40 منٹ کے شامل ہیں آپ کو اہم معلومات discussion اور عمومی allergies, recipe scaling, clarified cooking terms، کم کرنے کے لیے tokens محفوظ رکھتے ہوئے

کدام درمیان بہترین توازن رکھتا ہے؟ token reduction اور approach preservation کونسا

- A) خلاصہ کریں conversation history پوری
- B) رکھیں tokens صرف سب سے حالیہ 20,000
- C) خلاصہ کریں، اور discussion نکالیں، عمومی structured data (allergies, quantities, preferences) م رکھیں [درست] verbatim کو exchanges حالیہ
- D) کریں retrieve کے ذریعے متعلقہ حصہ semantic search کریں اور conversation externally store مکمل

C کیوں: Allergies معلومات محفوظ رکھتی ہے valuable پر سب سے زیادہ cost سب سے کم Hybrid approach summarization) ہیں extract میں compact structured block ایک critical facts جیسے recipe quantities اور exchanges ہوتی ہیں، اور حالیہ discussion summarize عمومی، (سب سے بچاتے ہوئے precision loss کے دوران conversational coherence کے لیے verbatim میں Options A اور B dietary information کے لیے کوئی ضرورت نہیں؛ cooking session ایک D ہیں؛

سوال 66

ka preferences اور topics پہلے کے assistant میں conversations رپورٹ کرتے ہیں کے طویل Users صورتحال رکھتی ہے message pairs صرف آخری 25 implementation حساب کھو دیتا ہے آپ کی موجود

کیا ہے؟ solution سب سے مؤثر

- A) Hybrid approach: messages verbatim [درست] خلاصہ کرتے ہوئے حالیہ messages پرانے
- B) vector similarity search پر conversation history مکمل
- C) Window 50 کو message pairs تک بڑھائیں
- D) آگے جوڑیں running summary خلاصہ کریں اور dropped messages میں turn

A کیوں: محفوظ رکھتی ہے exact recent context کے دونوں پہلوؤں کو حل کرتی ہے Hybrid approach (conversational coherence کے لیے اہم) preferences جبکہ پہلے کی (کے لیے اہم) miss کو context م (B) Vector search صرف اسی مسئلہ کو ملتوی کرتا ہے (C) بڑھانا Window رکھتی ہے مشابہ نہیں ہو سکتا semantically سب سے current query کر سکتی ہے جو

سوال 67

بڑھ جاتی latency سے زیادہ ٹوکن ہیں تو 50 turns conversations رپورٹ کرتے ہیں کہ جب Users: صورتحال بھی costs ہیں اور

بنیادی وجہ کیا ہے؟

- A) شامل ہوتی ہیں [درست] conversation history کے ساتھ پوری API request
- B) Model gradually responses generate کرتا ہے
- C) History کے ساتھ database operations میں ہو جاتا ہے بڑھنے کے ساتھ
- D) Model internal user profile کی ضرورت والا processing زیادہ بناتا ہے

A میں پوری array messages کو request — کے stateless مکمل طور پر API کی Claude: کیوں A زیادہ request بڑھتی ہیں، کے conversations شامل کرنی ہوتی ہیں جیسے جیسے conversation history کے درمیان کوئی Model calls دونوں بڑھانا cost اور directly processing latency لے جاتی ہے، جو tokens internal state رکھتا ہے (D غلط ہے)

سوال 68

تک بڑھ گئی ہیں جب conversation history 85,000 tokens کے بعد weekly sessions: صورتحال: تین مہینوں کے discussions پچھلی assistant کے بارے میں کیا نتیجہ اخذ کیا تھا؟، تو theme کے isolation پوچھتا ہے "م نہ user کرنے کے بجائے عام جوابات دیتا ہے reference

کیا ہے؟ approach سب سے مؤثر

- A) Rolling window truncation
- B) Key conclusions capture کے ساتھ progressive summarization
- C) Relevant exchanges retrieve کے لیے [درست] semantic embeddings
- D) Discussion conclusions mark کے لیے structured XML tags شامل کریں

C scale کو discussion کے جو تین مہینوں کی approach واحد semantic search پر Conversation history Rolling window کر سکتی ہے specific relevant exchanges surface پر demand کر سکتی ہے جبکہ truncation (A) history کے Progressive summarization (B) discussions کا بیشتر حصہ ضائع کر دے گی history (A) truncation پوچھتا ہے ہیں users کہو دیتی ہے جن کے بارے میں specific conclusions، کرتی ہے compress میں abstractions

سوال 69

کرتا ہے، system prompt guidelines follow میں 10-15 Claude، کے دوران QA testing: صورتحال کے اندر token limits ابھی Conversation آ جانی ہیں deviation میں responses لیکن بعد کے

کیا ہے؟ solution بہترین

- A) Behavioral guidelines کے user message کو پلے میں منتقل کریں
- B) 20 turns کے بعد نئی conversation شروع کریں
- C) Conversation breakpoints پر guidelines reinforce والے user-role messages insert [درست] کریں
- D) Non-compliant responses کو regenerate کے لیے post-response validation کرنے کے استعمال کریں

history ، کا برا راست مقابلہ کرتا ، Behavioral reminders کا periodic injection instruction drift ، کیوں C user کو پلائی Guidelines کر کے constraints re-establish regular intervals کو ساتھ accumulate ، Post-response validation (D) corrective کم کرتا authority ان کا (A) میں منتقل کرنا message ، نہ ، preventive اور significant latency شامل کرتی

سوال 70

adaptation rules اور teaching methodology کو جو 2,800-token system prompt میں AI tutor 2,800 صورتحال: آپ کے کو نظر انداز کرنا شروع کر دیتا ، assistant proficiency levels ، بعد turns کرنا 12 define

کیا ؟ fix سب سے مؤثر

- A) 5–4 turns میں inject reminders کریں
- B) Verbose rules کو proficiency-level adaptation demonstrate والے few-shot examples سے replace کریں [درست]
- C) Critical rules کو system prompt میں رکھیں
- D) Responses evaluate کر کے difficulty level match اور اگر regenerate نہ ہو تو

abstract rules کیونکہ vulnerable کے لیے 2,800-token system prompt drift Declarative rules کیوں B سے concrete few-shot examples کو Verbose rules کا تقاضا کرتے ہیں ، reasoning سے model میں turn clear کرنے کے لیے match کو model ، کریں correct proficiency-level adaptation demonstrate کرنا جو replace ہوتا ہے ، زیادہ abstract instructions میں turns دیتا ہے — یہ کئی behavioral patterns

سوال 71

clarifying کرنا، اور reasoning explain برقرار رکھنا، اپنی enthusiastic tone کو assistant صورتحال: آپ کے کوونی چاہئیں؟ define ان behavioral guidelines پوچھنے ہیں یہ questions

کوونی چاہئیں؟ define ان behavioral guidelines یہ

- A) user message سے پہلے prepend کریں
- B) System prompt [درست] میں
- C) assistant message سے پہلے
- D) Environment variables میں

persistent behavioral کوونی والی apply میں conversation خاص طور پر پوری System prompt کیوں B redundant (A) سے پہلے شامل کرنا user message کے لیے ڈیزائن کیا گیا ہے ، اور guidelines constraints deviate سے statements اپنے پہلے model کیونکہ (C) assistant message ، overhead پر اثر نہیں ڈالتے behavior کے model (D) Environment variables کر سکتا ہے

سوال 72

جیسے "یقیناً!" اور "میں مدد کرنے میں خوش ہوں!" رپورٹ کرتے Users repetitive response openings صورتحال ہیں

کیا ؟ approach سب سے مؤثر

- A) Direct response opening ☐ ساتھ partial assistant message append ☐ [درست] کریں
- B) Temperature setting ☐ کم کریں
- C) Greetings ☐ لیڈ ☐ ٹائٹل ☐ responses کو post-process کریں ☐
- D) System prompt ☐ شامل کریں ☐ instructions ☐ سہ ☐ بچہ ☐ کی phrases میں ان

A greeting پر generation level کرنا prefill جواب کی شروعات سے direct response کو Assistant کیوں: patterns روکتا — model prefill سے continue نہ بچائے opening phrases generate کرتا Post-processing (C) میں reliable مدد کر سکتی ہیں لیکن کم (D) System prompt instructions fragile ایک (B) randomness control کرتا ، specific phrase patterns میں۔

سوال 75

turns rules follow استعمال کرتا ہے ابتدائی assistant contractor-persona system prompt صورتحال: آپ کا Conversation length 2,500 tokens صرف دیتا ہے assistant generic advice تک 7 turn کرتا ہے، لیکن

سب سے زیادہ ممکنہ وجہ کیا ہے؟

- A) System prompts initial behavior establish کرتے ہیں صرف
- B) Turns accumulate کمزور ہوتی ہے attention کی model ہونے کے ساتھ
- C) Accumulated assistant responses system prompt اثر کو dilute [درست] کرتے ہیں
- D) System prompt صرف ایک بار بھیجا جاتا ہے

system prompt جمع ہونے میں assistant responses conversation history کیوں جیسے جیسے C behavioral constraints کو reflect کرنے والے assistant-generated content بڑھنے والے proportion کا text کرنے والے relative بجائے اپنے پیمائش پر بھی token lengths Model کم ہوتا جاتا ہے relative system prompt instructions compound drift کرتا ہے، کرتا ہے follow کو زیادہ patterns کے outputs

سوال 76

کئی Assistant "میں مدد کر سکتا ہے؟" میں جیسے "کیا آپ requests میں Users صورتحال abandonment % جس سے 40، (کیا ہے؟ deadline؟ کیا مدد؟ report کون سی) کرتا ہے respond سوالات پوچھ کر ہوتا ہے

کیا ہے؟ solution بہترین

- A) Reasonable assumptions میں انہیں explicitly کریں، اور adjustments [درست] بیان کریں
- B) Respond smaller model کرنے سے پہلے ambiguity classify کریں کے ساتھ
- C) Assumptions predefined interpretations بیان کیے بغیر استعمال کریں
- D) Assistant turn کو رکنے تک محدود کریں clarifying question میں ایک turn کو

رکھتے ہیں in control اور informed کو user کے ساتھ آگے بڑھنا Reasonable stated assumptions کیوں A کرتی ہے confuse کو users (C) Silent predefined interpretations کو مکمل طور پر ختم کرتا ہے back-and-forth کی back-and-forth turns بھی (D) One-question limit، match ان کے ارادے سے response میں جب infrastructure اور latency مسئلہ حل کیے بغیر (B) core UX Smaller classification model ضرورت ہے شامل کرتا ہے complexity

عملی مشقیں

Exercise 1: Multi-tool Agent with Escalation Logic

Goal: tool integration، structured errors، اور escalation کے ساتھ agent loop بنائیں

Exercise 2: Configuring Claude Code for Team Development

Goal: CLAUDE.md, custom commands, path-specific rules, اور MCP servers configure کریں

Exercise 3: Structured Data Extraction Pipeline

Goal: JSON schema, validation/retry loops, اور Message Batches API کے ساتھ extraction pipeline بنائیں

Exercise 4: Designing and Debugging a Multi-agent Research Pipeline

Goal: subagent orchestration, explicit context passing, error propagation, اور source tracking implement کریں

Appendix: Technologies and Concepts

Technology	Key aspects
Claude Agent SDK	agent loops, <code>stop_reason</code> , hooks, <code>Task</code> , <code>allowedTools</code>
MCP	servers, tools, resources, <code>isError</code> , <code>.mcp.json</code>
Claude Code	CLAUDE.md, <code>.claude/rules/</code> , <code>.claude/commands/</code> , <code>.claude/skills/</code> , planning mode
Claude Code CLI	<code>-p</code> , <code>--output-format json</code> , <code>--json-schema</code>
Claude API	<code>tool_use</code> , <code>tool_choice</code> , <code>stop_reason</code> , <code>max_tokens</code>
Message Batches API	50% savings, 24-hour window, <code>custom_id</code>
JSON Schema	required/optional, nullable, enum values
Few-shot prompting	ambiguous scenarios کے لیے examples
Prompt chaining	sequential decomposition
Context window	token budget, lost-in-the-middle
Confidence calibration	field-level scoring, stratified sampling

Out-of-Scope Topics

یہ موضوعات امتحان میں شامل نہیں ہیں:

- Claude models کی fine-tuning
- Authentication, billing, یا account management
- implementation کی تفصیلی programming languages مخصوص
- MCP servers کی hosting/infrastructure
- Claude کی internal architecture یا model weights
- Constitutional AI, RLHF, یا safety training
- Embeddings یا vector databases کی implementation details
- Computer use / browser automation
- Vision / image analysis
- Streaming API یا SSE
- Rate limiting اور detailed API cost calculations
- OAuth اور API key rotation details
- Cloud-provider-specific configs
- Performance benchmarks
- Prompt caching implementation details
- Token counting algorithms

Preparation Recommendations

1. بنائیں agent کا ساتھ ایک Claude Agent SDK
2. کریں configure پر real project کو Claude Code
3. کریں test اور MCP tools design
4. بنائیں Data extraction pipeline
5. کریں Prompt engineering practice
6. سیکھیں Context management patterns
7. سمجھیں human-in-the-loop workflows اور Escalation
8. حل کریں Practice exam