



Schema Beats Prompt

Designing Trustworthy AI APIs

Presented at ISCon Banja Luka 2026

Edvin Teskeredzic | Senior AI Software Engineer

MARCH 2026



What we'll talk about

1. What “trustworthy” means for an API
2. FastAPI + Pydantic: Contracts as the Default
3. Breaking stuff with AI
4. Enforcing structure on LLMs
5. Designing usable AI APIs
6. How this affects DevEx
7. Make AI APIs boring again



About me

- Senior AI Software Engineer @ Infobip





About me

- Senior AI Software Engineer @ Infobip
- Published some research papers 📖





About me

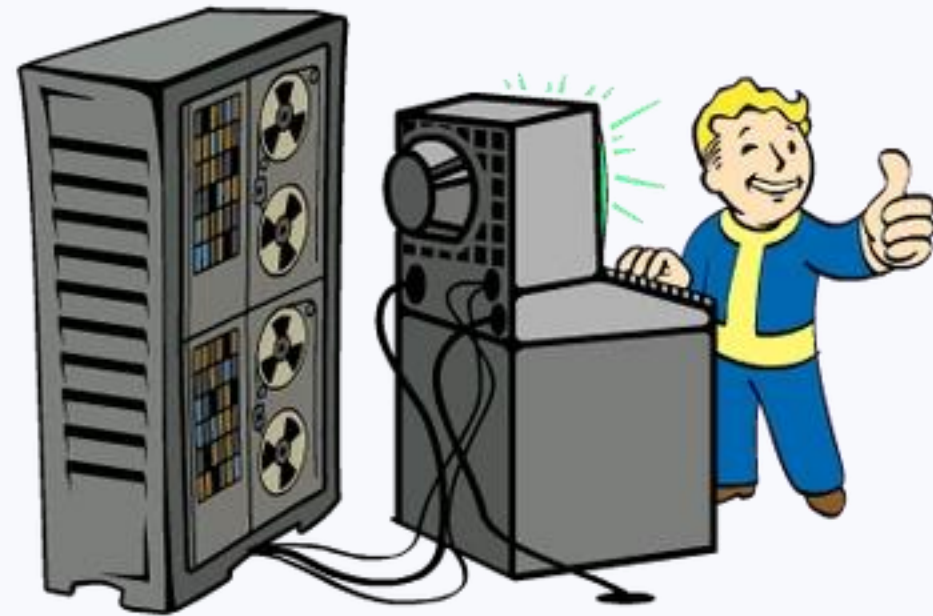
- Senior AI Software Engineer @ Infobip
- Published some research papers 📖
- Some of my (nerdy) hobbies:
 - Pub quizzes
 - Sci-Fi and Fantasy
 - Gaming





What I do at work

- AI and ML (obviously)
- ... but also...
- Developer Experience (DevEx)
- Systems Design and ML infra
- Tooling, observability, traceability
- API design



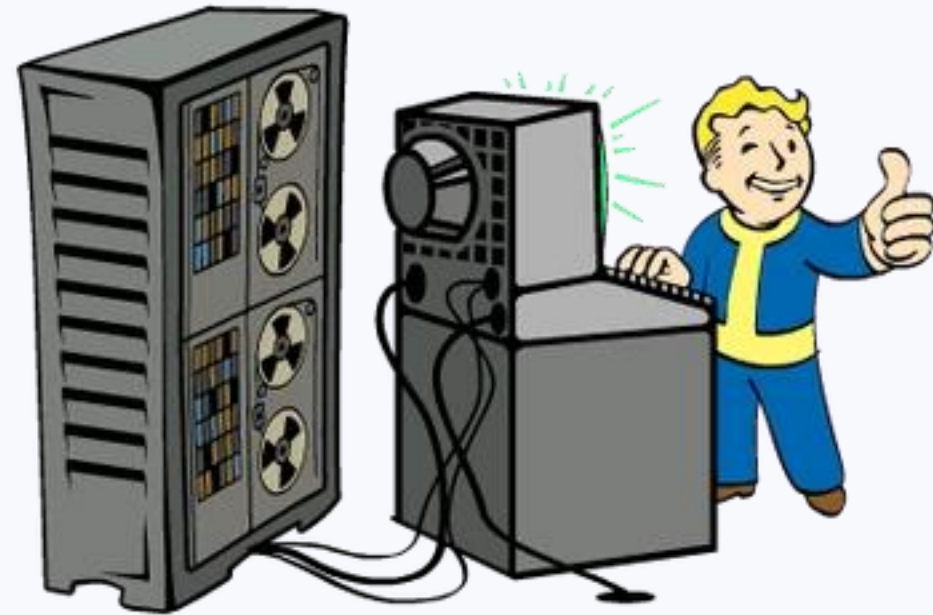


What I do at work

- AI and ML (obviously)
- ... but also...
- Developer Experience (DevEx)
- Systems Design and ML infra
- Tooling, observability, traceability
- **API design**



We'll be talking
about this today!





Describe an API in three words

- **Deterministic** (regarding their contract)
- **Explicit**
- **Predictable**



Describe an API in three words

- **Deterministic** (regarding their contract)
- **Explicit**
- **Predictable**

Now, imagine we let a little kid play with our API responses before sending them back to the user...



Describe an API in three words

- **Deterministic** (regarding their contract)
- **Explicit**
- **Predictable**

Now, imagine we let a little kid play with our API responses before sending them back to the user...

...this happens when we introduce AI (LLMs) to our API



AI APIs...

- Are **probabilistic**
- Have **best-effort outputs**
- Fail **creatively**



AI APIs...

- Are **probabilistic**
- Have **best-effort outputs**
- Fail **creatively**

When a traditional API fails, it's explicit.
When an LLM fails, it **confidently makes something up – implicit failure!**



A simple example

```
1  {  
2    "summary": "User is eligible",  
3    "confidence": 0.82  
4  }  
5
```

An LLM could do something like:

- "confidence": "high"
- Missing **summary** field
- Additional fields we did not ask for
- Invalid JSON...



bro please
respond in valid json format
without errors and
make super sure the
syntax is extra correct
i'm begging you... and
please, pretty please,
don't make up answers
my career depends on it bro



**PROMPT
ENGINEER**



So, what's the issue?

- Is this a model problem?
 - Nope.
- Oh, so a prompt problem then?
 - No, try again.



So, what's the issue?

- Is this a model problem?
 - Nope.
- Oh, so a prompt problem then?
 - No, try again.

**This is an API
design problem!**



TAKEAWAY

Prompting helps the model
understand intent.

But **schemas** are what make
APIs trustworthy!



What we'll show today

Show how **the Python ecosystem** lets us enforce structure even *when the thing generating the data is non-deterministic!*

Remember: The goal isn't smarter prompts.

The goal is making AI APIs boring again!





**What “trustworthy”
means for an API**



What does it mean to trust an API?

A trustworthy API does a few **boring** things extremely well:

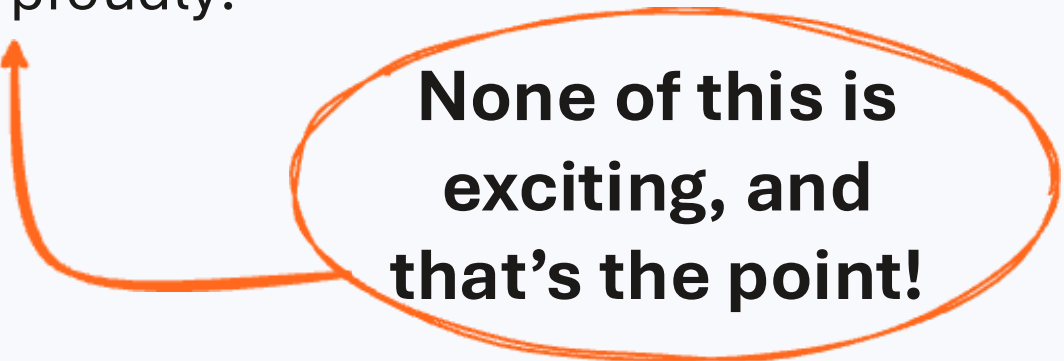
1. It rejects invalid input
2. It never returns invalid output (we rather have it fail!)
3. Stable + explicit contract
4. Fail loudly and proudly!



What does it mean to trust an API?

A trustworthy API does a few **boring** things extremely well:

1. It rejects invalid input
2. It never returns invalid output (we rather have it fail!)
3. Stable + explicit contract
4. Fail loudly and proudly!



**None of this is
exciting, and
that's the point!**



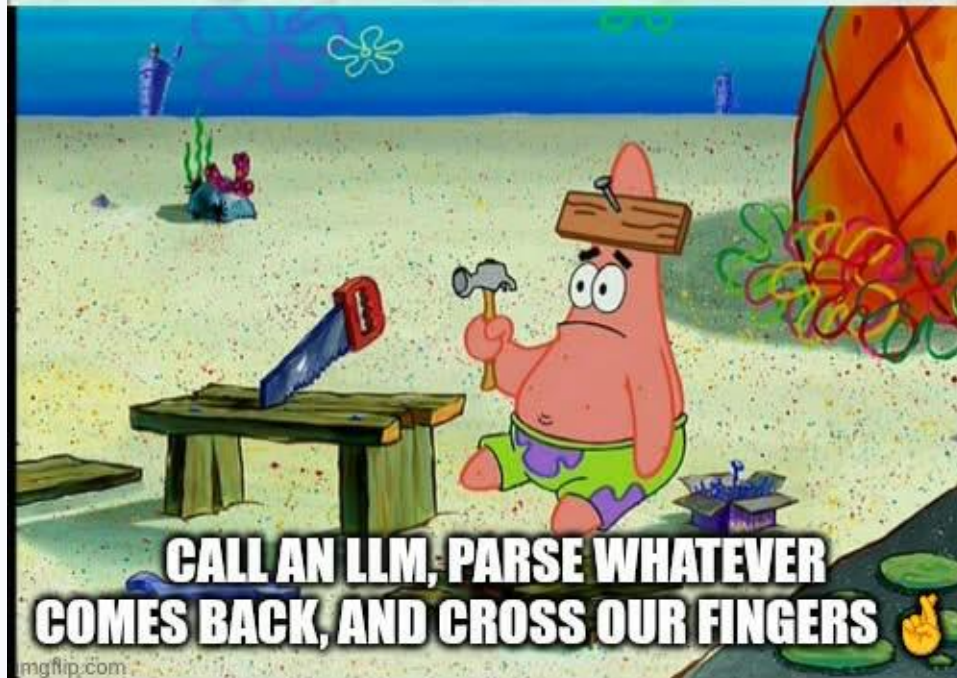
We know how to be boring

1. Request schemas
2. Response schemas
3. Validation at the boundary
4. Clear error codes

**But something
strange happens
when we add LLMs
to the mix...**



VALIDATE INPUT & DOCUMENT RESPONSE SHAPE



**CALL AN LLM, PARSE WHATEVER
COMES BACK, AND CROSS OUR FINGERS**

**We treat an
unpredictable
component as trusted!**



THIS IS ABOUT CONTRACT ENFORCEMENT

If your API contract stops at the LLM boundary, you don't really have a contract.



Let's be boring

The question becomes: **How do we enforce the same boring guarantees (validation, schemas, explicit failures) when the data is generated by an AI?**



Let's be boring

The question becomes: **How do we enforce the same boring guarantees (validation, schemas, explicit failures) when the data is generated by an AI?**



And this is where



gets interesting!



FastAPI + Pydantic: Contracts as the Default





Refresher for non-Python people

- **FastAPI** – High-performance web framework
- **Pydantic** – Data validation with Python models





Python is great for API design

In **FastAPI**, you do not describe an API in three places:

1. You don't define a request schema
2. You don't define validation logic
3. You don't define documentation

You define one thing – a Pydantic model.



Pydantic model example

```
1 class UserReq(BaseModel):
2     """Create user request."""
3     email: EmailStr = Field(..., description="Valid email")
4     name: str = Field(..., min_length=2, description="2+ chars")
5
6     @field_validator("name")
7     @classmethod
8     def strip_name(cls, v: str) -> str:
9         v = v.strip()
10        if not v: raise ValueError("name required")
11        return v
```



Pydantic model example

```
1 class UserReq(BaseModel):  
2     """Create user request."""  
3     email: EmailStr = Field(..., description="Valid email")  
4     name: str = Field(..., min_length=2, description="2+ chars")  
5  
6     @field_validator("name")  
7     @classmethod  
8     def strip_name(cls, v: str) -> str:  
9         v = v.strip()  
10        if not v: raise ValueError("name required")  
11        return v
```

Includes:

- Request schema ●



Pydantic model example

```
1 class UserReq(BaseModel):  
2     """Create user request."""  
3     email: EmailStr = Field(..., description="Valid email")  
4     name: str = Field(..., min_length=2, description="2+ chars")  
5  
6     @field_validator("name")  
7     @classmethod  
8     def strip_name(cls, v: str) -> str:  
9         v = v.strip()  
10        if not v: raise ValueError("name required")  
11        return v
```

Includes:

- Request schema ●
- Documentation ●



Pydantic model example

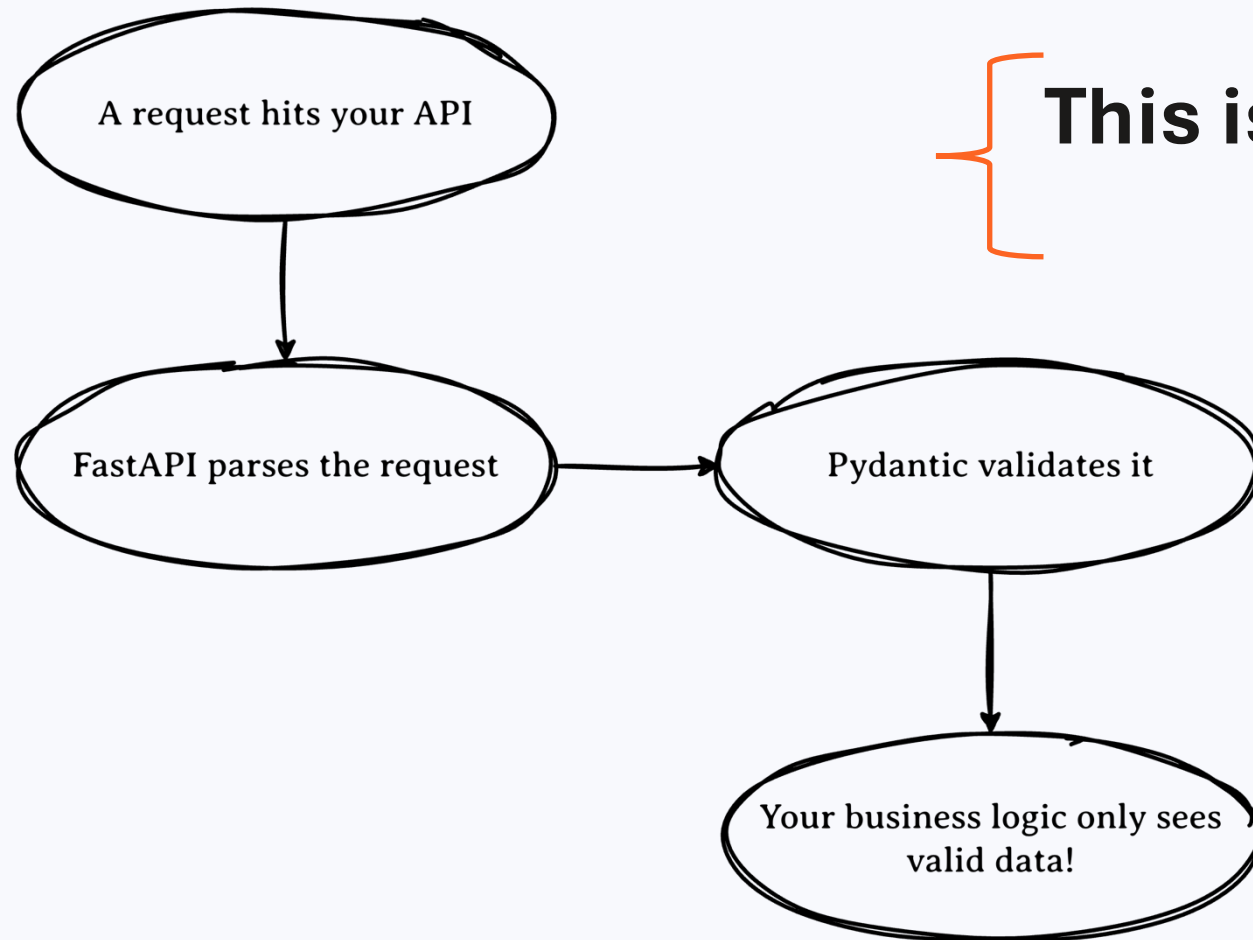
```
1 class UserReq(BaseModel):  
2     """Create user request."""  
3     email: EmailStr = Field(..., description="Valid email")  
4     name: str = Field(..., min_length=2, description="2+ chars")  
5  
6     @field_validator("name")  
7     @classmethod  
8     def strip_name(cls, v: str) -> str:  
9         v = v.strip()  
10        if not v: raise ValueError("name required")  
11        return v
```

Includes:

- Request schema ●
- Documentation ●
- Validation ●



This is how we do



This is not optional, this is the default



But wait, there's more!

- Docs are generated automatically – the same model becomes your OpenAPI schema
- Clients don't have to guess
- SDKs can be generated safely

**If it's not in the schema,
it doesn't exist**



You thought it was over? Think again!

- Most people think validation stops at input – it does not.
- FastAPI + Pydantic can validate **responses**
- Outgoing data is checked against the declared schema
- Extra fields are dropped, missing fields raise errors

**Your API cannot lie.
Not even accidentally.**



IT JUST WORKS...

...if our code produces the
data.

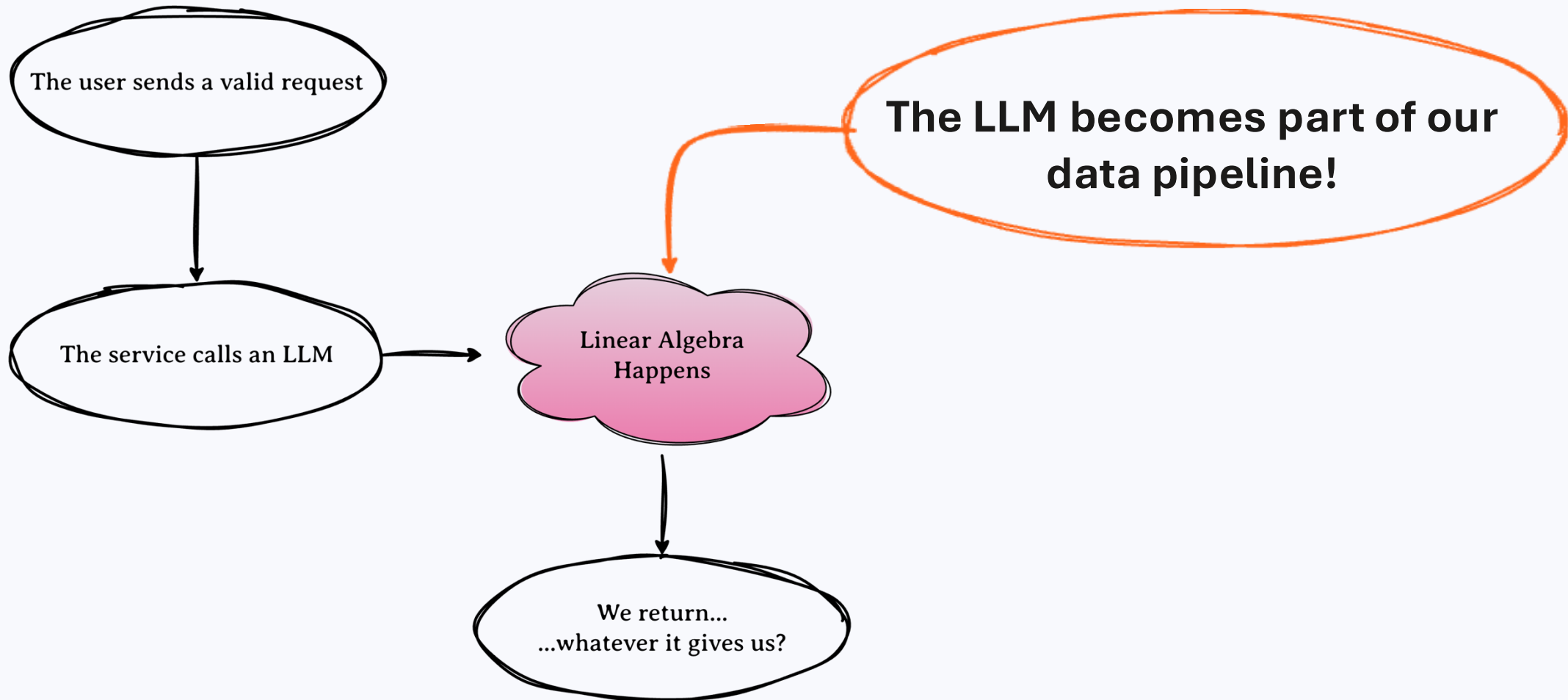
But what happens when the
data comes from an LLM?



Breaking stuff with AI



This is how we do (with AI)





This is how we do (with AI)

**The moment the LLM
generates the response,
we lose all our
guarantees**



This is how we do (with AI)

- JSON that ***almost*** parses
- Missing required fields
- Extra fields we did not ask for
- Numbers represented as strings
- ...

In short: **We have broken the API contract**



Can we fix it?

“We’ll just improve the prompt!”



Can we fix it?

“We’ll just improve the prompt!”





Can we fix it?

“We’ll just improve the prompt!”



- Prompts reduce the probability of failure, **but they do not eliminate it!**
- They are not contracts!
- They do not fail loudly!



**What if we treated LLM
output the same way we treat
user input?**



**What if an AI response had to
earn the right to be correct
and leave our API?**



Enforcing Structure on LLMs



(a.k.a. *“Finally, the code!”*)





Do you wanna build an AI endpoint?

```
1 @app.post("/analyze")
2 async def analyze_text(text: str):
3     response = await call_llm(
4         f"Analyze this text and return sentiment and confidence: {text}"
5     )
6     return response
```



Do you wanna build an AI endpoint?

```
1 @app.post("/analyze")
2 async def analyze_text(text: str):
3     response = await call_llm(
4         f"Analyze this text and return sentiment and confidence: {text}"
5     )
6     return response
```

What is the response type here?



Do you wanna build an AI endpoint?

```
1 @app.post("/analyze")
2 async def analyze_text(text: str):
3     response = await call_llm(
4         f"Analyze this text and return sentiment and confidence: {text}"
5     )
6     return response
```

What is the response type here?

We don't know – and neither does our API!



Let's make it a bit better



```
1 prompt = ""  
2 Return JSON with:  
3 - sentiment: positive | neutral | negative  
4 - confidence: number between 0 and 1  
5 ""
```

This helps. A lot, actually.



TAKEAWAY

We need to treat the LLM like
any untrusted external system



Let's turn guidance into a contract

```
1  
2 from pydantic import BaseModel, Field  
3 from typing import Literal  
4  
5 class SentimentResult(BaseModel):  
6     sentiment: Literal["positive", "neutral", "negative"]  
7     confidence: float = Field(ge=0.0, le=1.0)
```

First, we'll define an output contract.



Enter PydanticAI

Pydantic AI is a Python framework for quickly and reliably building production-grade Generative AI applications and workflows



Enter PydanticAI

Features include:

1. Model-agnostic
2. Observability & Traceability out of the box
3. Full type-safety
4. Evaluation and performance monitoring
5. Support for MCP, A2A, and UI event streams
6. Human-in-the-Loop
7. Graph support for complex flows
8. And much more!



PydanticAI Agent with structured output

```
1 from pydantic_ai import Agent
2
3 agent = Agent(
4     model="openai:gpt-4o-mini",
5     result_type=SentimentResult,
6     system_prompt=(
7         "Analyze the sentiment of the text and return a structured result."
8     ),
9 )
```



THE CORE IDEA:

The LLM does not decide what is valid. The schema does.



Enter PydanticAI



```
1 result = await agent.run(text)
```

The result is not a dict. It is not a raw JSON. It is a **validated Pydantic model!**



```
1 result.sentiment  
2 result.confidence
```



BUT WHAT IF THE MODEL MISBEHAVES?

That's fine – because failure is
now visible!



Putting it all together



```
1 @app.post("/analyze", response_model=SentimentResult)
2 async def analyze_text(text: str):
3     result = await agent.run(text)
4     return result
```

What we get:

- Validated request
- Validated AI output
- Validated API response
- Accurate OpenAPI schema



TAKEAWAY

The prompt helps the model.
The schema protects the
system.



Designing usable AI APIs



Designing usable AI APIs

- So now our AI output is validated. Great.
- But that still doesn't automatically make our API good.
- Validation solves correctness. Not usability. Not stability. Not long-term contracts.

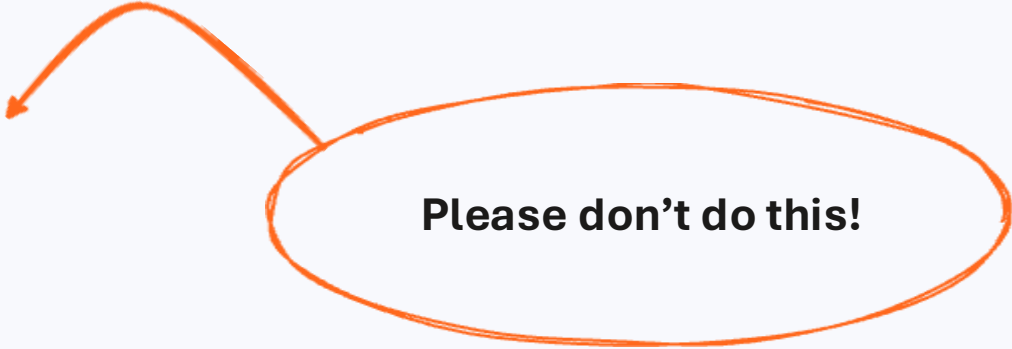
**An AI API still must
behave like any normal
API**



Pattern 1: Never expose raw LLM output



```
1 @app.post("/analyze")
2 async def analyze(text: str):
3     result = await agent.run(text)
4     return result
```



Please don't do this!



Instead, separate agent output from API output



```
1 class SentimentInternal(BaseModel):
2     sentiment: Literal["positive", "neutral", "negative"]
3     confidence: float
4     reasoning: str
5
6
7 class SentimentResponse(BaseModel):
8     sentiment: Literal["positive", "neutral", "negative"]
9     confidence: float
```



Instead, separate agent output from API output – then transform!

```
1 result = await agent.run(text)
2
3 return SentimentResponse(
4     sentiment=result.sentiment,
5     confidence=result.confidence
6 )
```



Pattern 2: Make OpenAPI honest

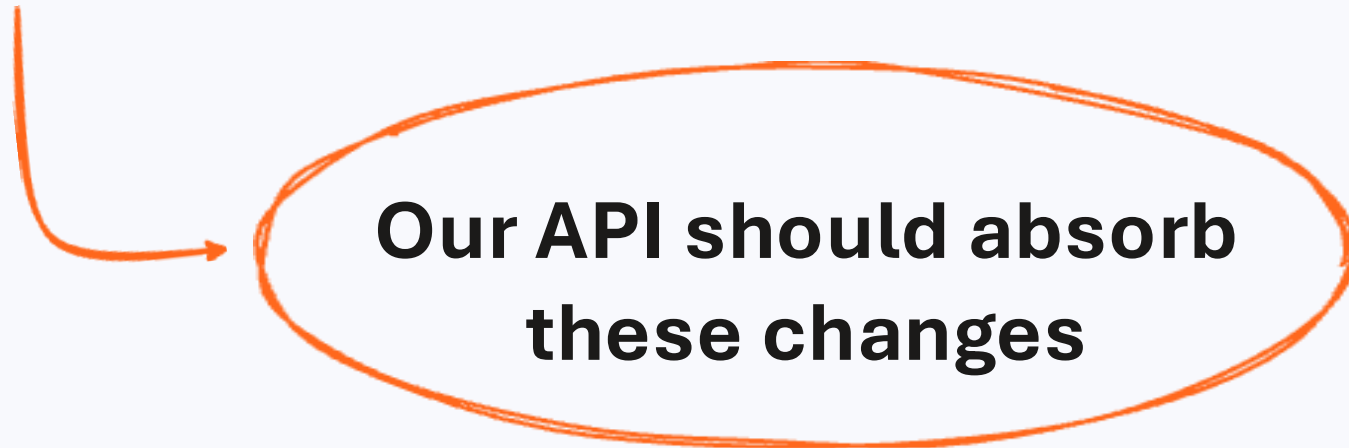
```
1 @app.post("/analyze", response_model=SentimentResponse)
```

**Your OpenAPI schema
now represents exactly
what clients receive**



Pattern 3: Design for evolution

- Prompts change
- Models change
- Tools change





Pattern 3: Design for evolution

Techniques:

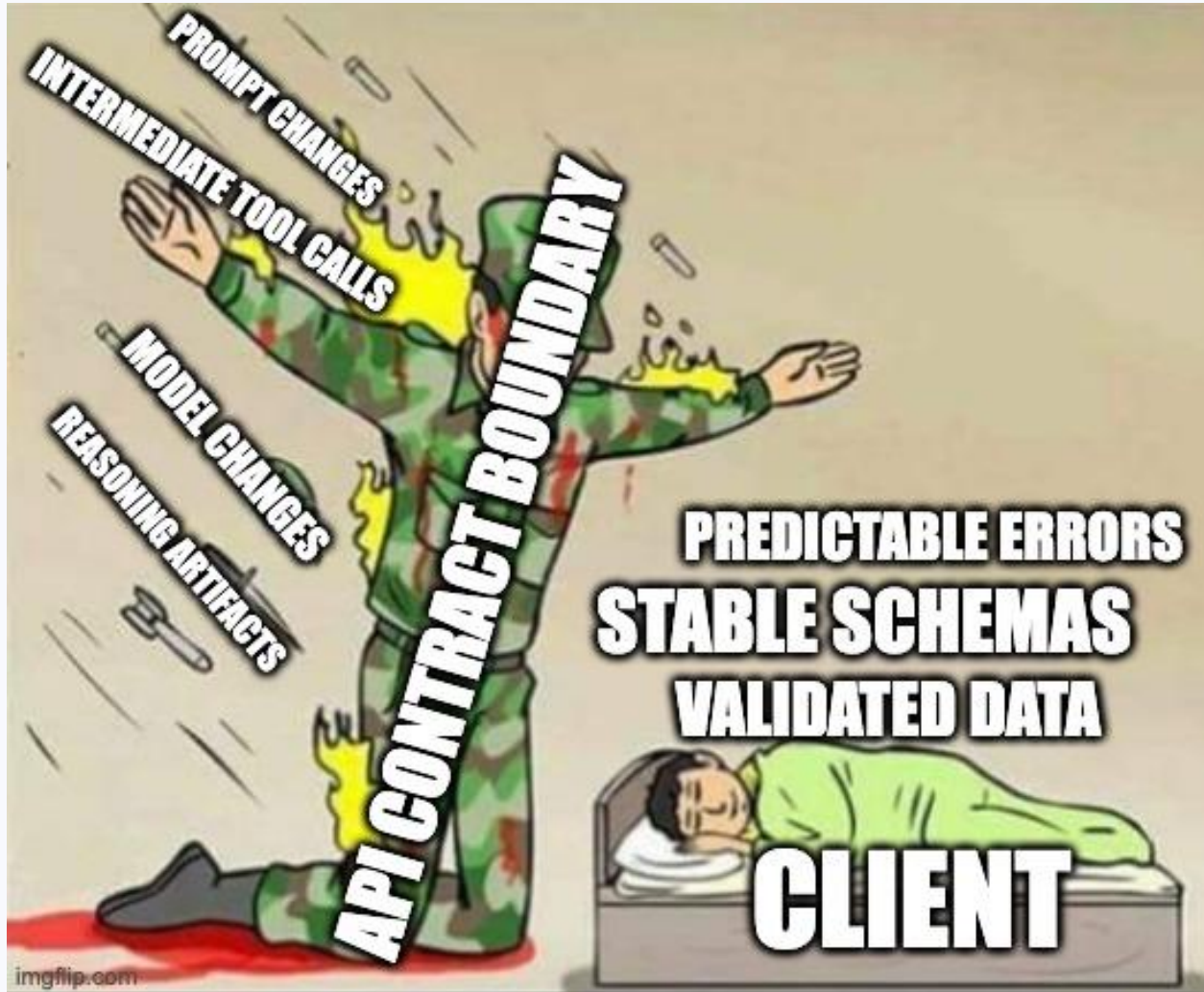
- Stable response models
- Optional fields for expansion
- Versioned endpoints when needed

```
1 class SentimentResponse(BaseModel):  
2     sentiment: str  
3     confidence: float  
4     explanation: Optional[str] = None
```

**The AI can get smarter without
breaking clients**



The bigger design principle





The bigger design principle

Once schemas define :
everything

- Validation
- Documentation
- AI output
- API responses

**This is where the Python
ecosystem becomes a
developer's best friend**

We get a powerful feedback loop!



AND THE SURPRISING THING IS:

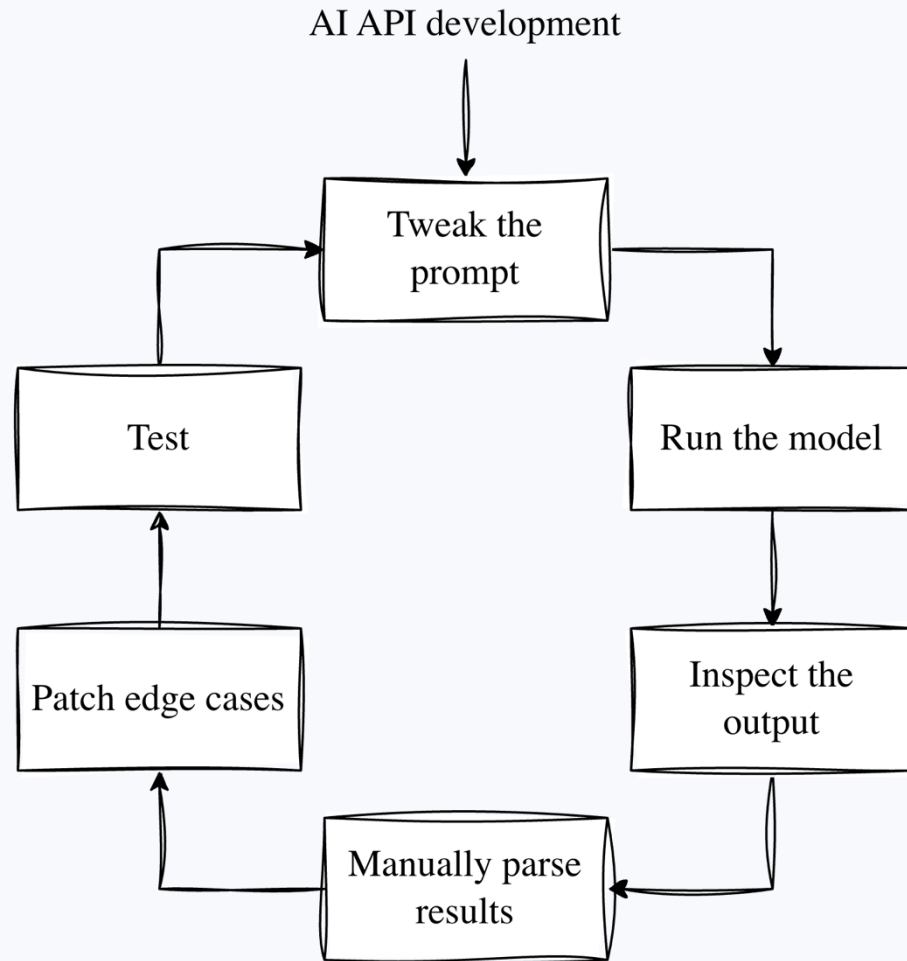
This does not slow developers
down; **it makes them faster!**



How this affects DevEx



AI development normally feels chaotic



This leads to:

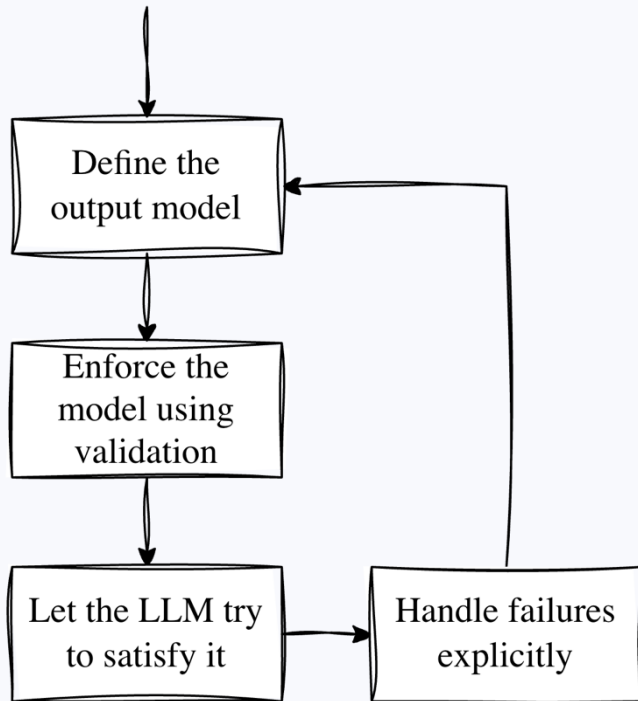
- Fragile parsing logic
- Implicit assumptions
- Constant regressions
- Hard-to-debug failures

... it does not feel like software engineering at all!



Let's reframe using what we've learned

Schema-driven AI
API development



Now our development loop is:

- Deterministic
- Testable
- Inspectable

... instead of chasing weird outputs, you define what “correct” means!



The hidden productivity benefits

- **Autocomplete**
- **Safer refactors**
 - The IDE becomes our ally
- **Faster collaboration**
 - Schema as a shared language
- **Reliable documentation**
 - OpenAPI is always accurate – no doc drift!



```
1 result["sentiment"]  
2 # vs  
3 result.sentiment
```



THE PSYCHOLOGICAL SHIFT

“How do I make the **model**
behave?”



“What does a **correct answer**
look like?”



With the right tools, it's easy



**In many ecosystems,
achieving this requires a
lot of glue code.**

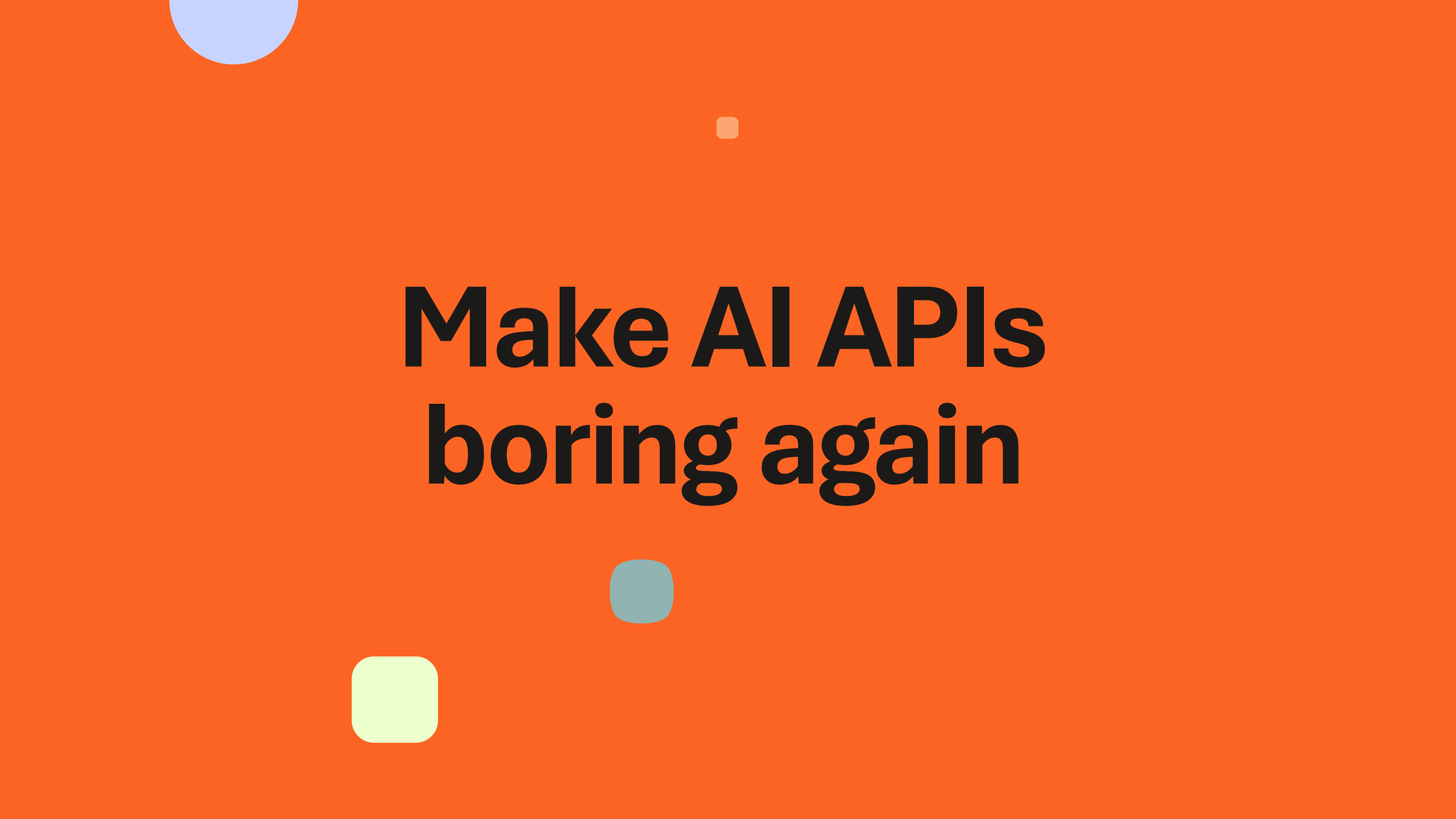
**In Python, it mostly falls
out of the tools**





TAKEAWAY

Schemas do not slow down AI
development – **they make it
safe to move fast!**



**Make AI APIs
boring again**



Let's recap

- Traditional APIs are ***deterministic***
- AI systems are ***probabilistic***
 - LLMs sometimes produce output that is ***almost right***



Let's recap

- Traditional APIs are ***deterministic***
- AI systems are ***probabilistic***
 - LLMs sometimes produce output that is ***almost right***



**This is where
most problems
begin**



Let's recap – the core mistake

- When building AI APIs, we often focus on the model
 - *Better prompts*
 - *Better context*
 - *Better tools*



Let's recap – the core mistake

- When building AI APIs, we often focus on the model
 - *Better prompts*
 - *Better context*
 - *Better tools*



But reliability does
not come from
prompts – **reliability
comes from
contracts**



Let's recap – the grand idea

The AI can be creative internally – the API should never be!

So, please:

- Validate your inputs
- Enforce schemas on outputs
- Separate internal models from public APIs
- Treat LLMs as untrusted dependencies





FINAL TAKEAWAY

Prompt engineering helps the
model understand the task.

Schema engineering helps the
system stay trustworthy!



**In practice, schema
beats prompt – every
time.**



Thank you!

Edvin Teskeredzic

Senior AI Software Engineer

Edvin.Teskeredzic@infobip.com

www.infobip.com

