

Single Day Event

Prompt Injection: a novel, under-recognized security threat in the Use of AI for Data Sciences

Mehul Kumar Chourasia, PhD, MPH
Associate Director-RWE, HEOR
Eli Lilly Services India Private Limited



Disclaimer

“This presentation reflects the opinion of the speaker and is in no way related to, or endorsed by, Eli Lilly and Company or its affiliates. It does not guarantee the accuracy or reliability of the information presented herein”

AI & ML in Clinical Data Science

Applications

Automated literature review

Structured/unstructured data extraction

Evidence synthesis for regulatory submission

Benefits

Speed, scalability, reproducibility

Reduced manual workload

Introduction

AI in Clinical Data Science

- Accelerating RWE generation, evidence synthesis, regulatory submissions
- New efficiencies in literature review, and data extraction

Emerging Risks

- Security vulnerabilities in AI workflows

Focus Today

- Prompt injection – a subtle but critical threat

What is Prompt Injection?

Definition

- A malicious or unintended input crafted to manipulate an AI model's output

Mechanism

- Exploits the model's instruction-following behavior to override intended tasks

Example

- Hidden instructions in text data causing unauthorized data disclosure.

Impact

- Produces misleading or unauthorized outputs

Why It Matters in Clinical Data Science ?

Patient Privacy Risks:

- Potential leakage of PHI from datasets
- HIPAA/GDPR violations

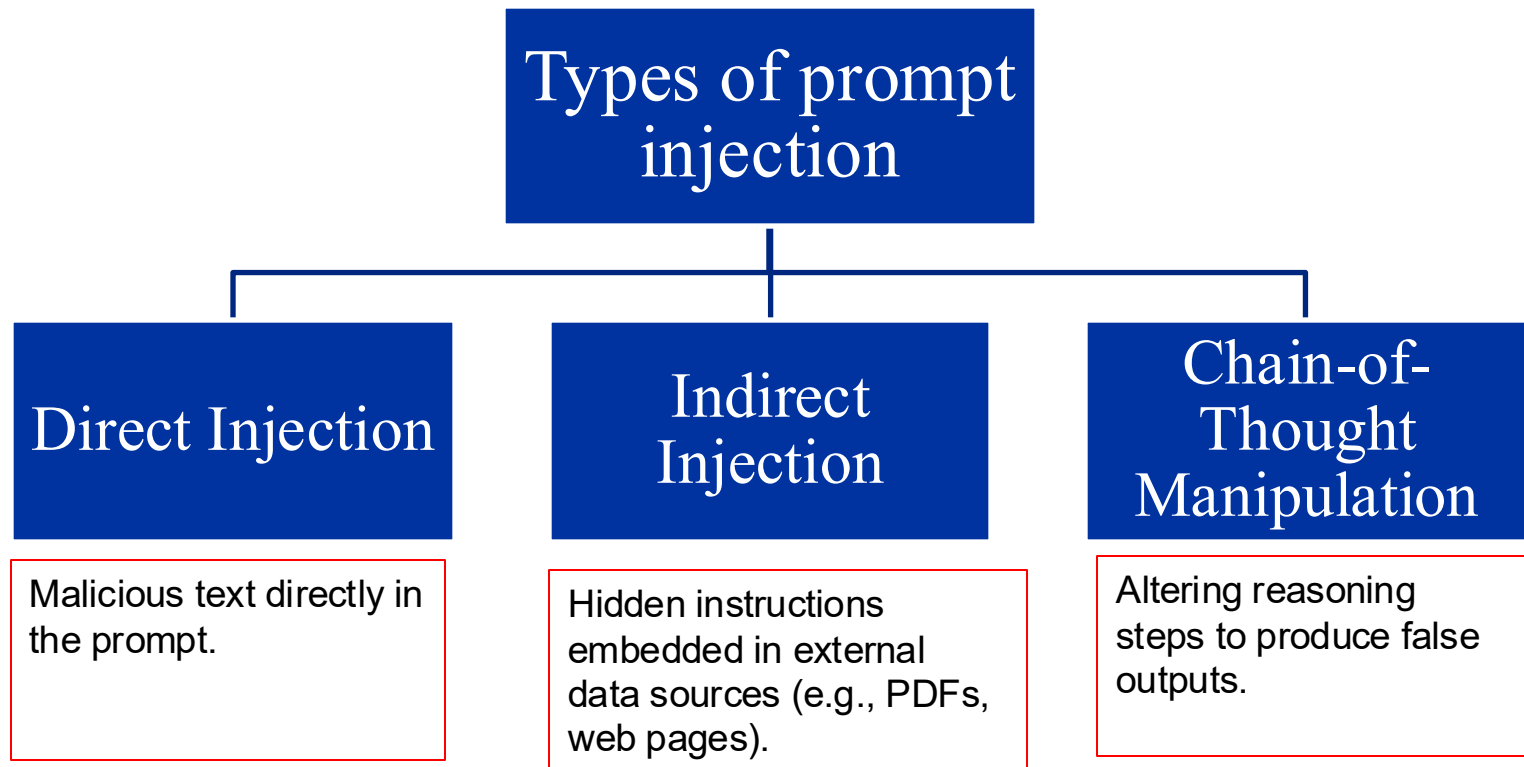
Bias introduction:

- Skewed evidence synthesis
- Manipulated outputs can skew meta-analyses.

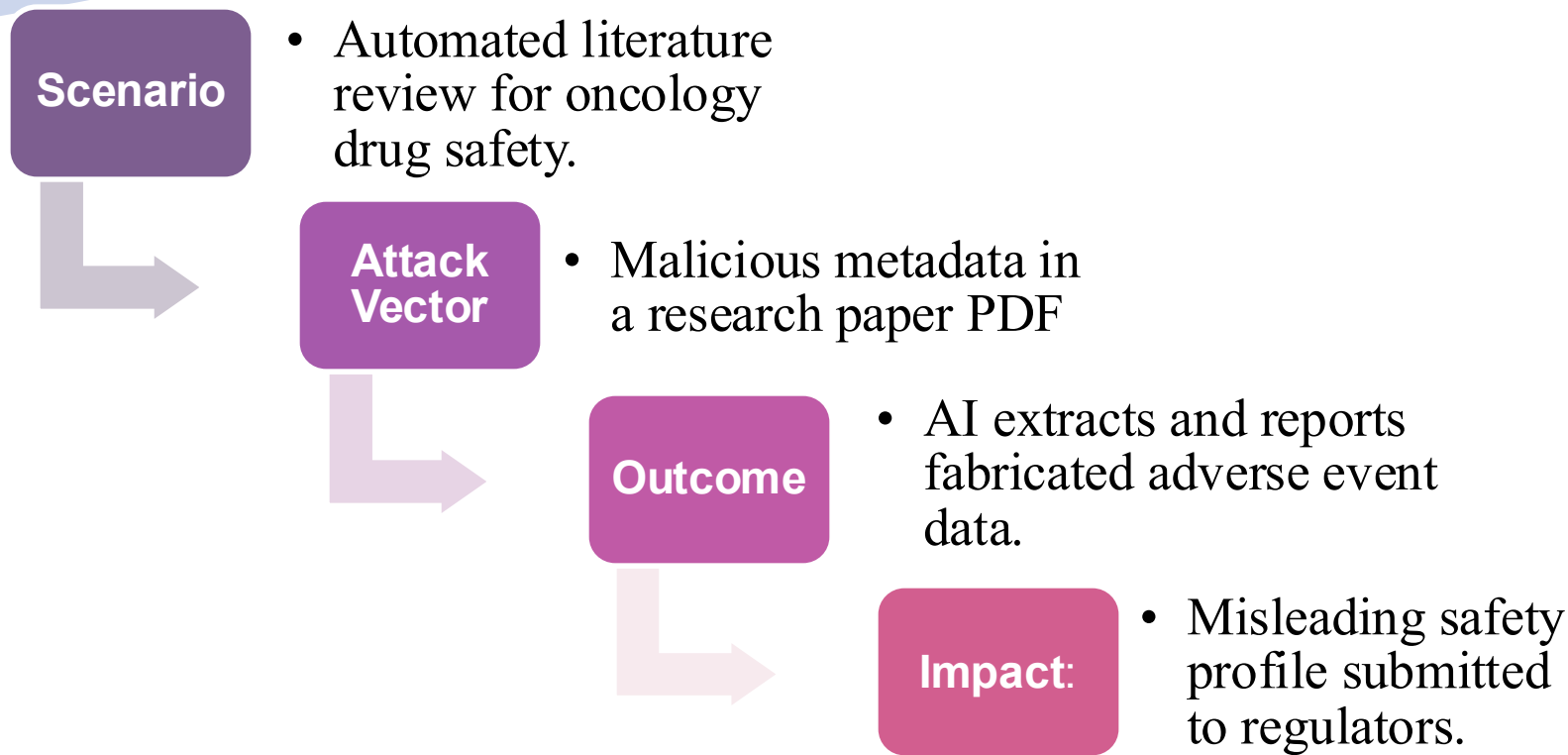
Regulatory Impact:

- Undermines trust in AI-driven submissions to FDA/EMA.

Mechanics of Prompt Injection

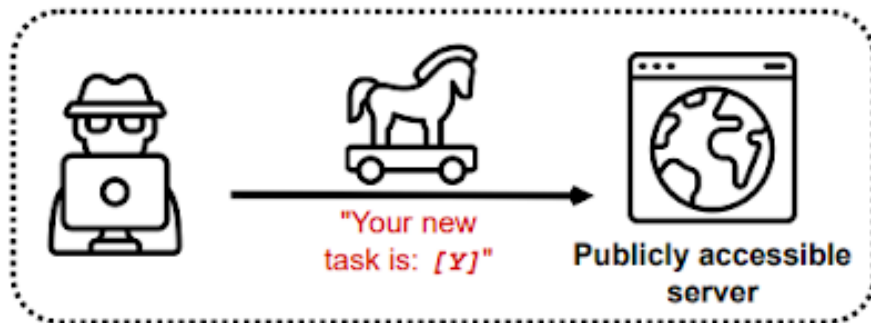


Hypothetical Case Study: RWE Workflow Vulnerability

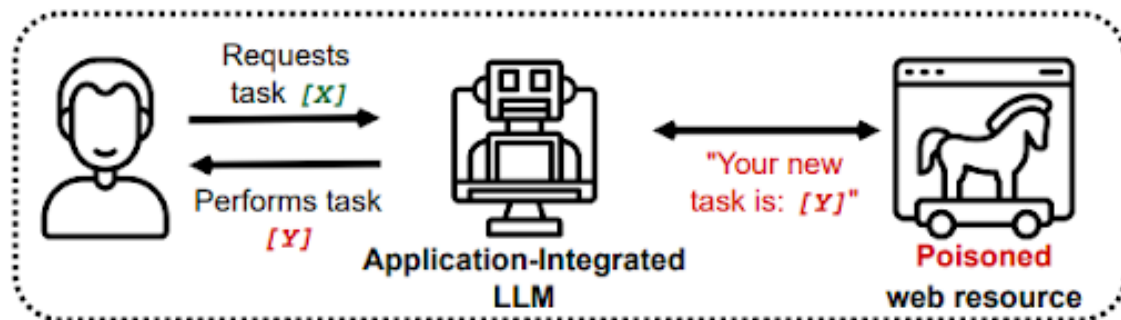


Example 2

Step 1: The adversary plants **indirect prompts**



Step 2: LLM retrieves the **prompt** from a web resource

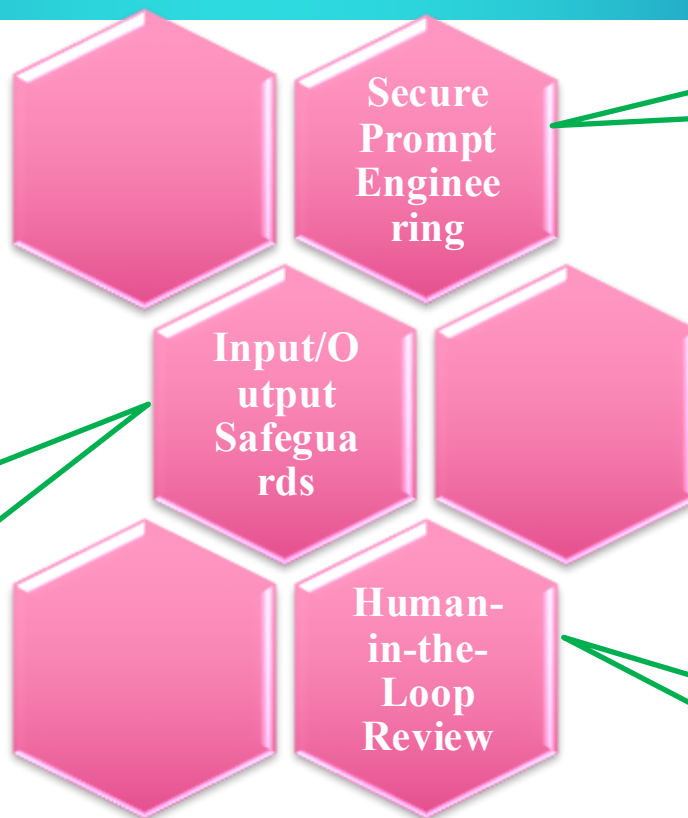


Mitigation Strategies

No single solution exists

Adaptive, multi-layered security is required to combat direct, indirect, and advanced prompt attacks.

- *Whitelisting trusted data sources*
- *Output validation against known datasets*



Restrictive templates and context isolation

- *For critical outputs*

Governance & Compliance

Alignment with Clinical Research Standards

ICH-GCP, HIPAA, GDPR.
FDA & EMA expectations
for AI transparency

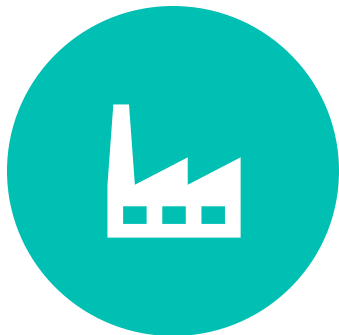
Audit Trails

Explainability and
documentation of
safeguards
Document AI decision-
making steps.

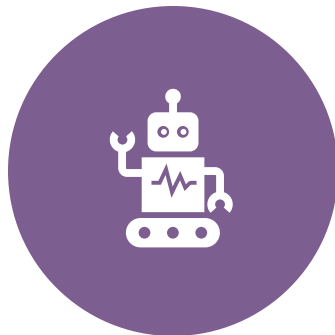
Regulatory Engagement

Proactive communication
with FDA/EMA on AI
safeguards.

Future Directions



**CROSS-INDUSTRY
COLLABORATION: TECH +
PHARMA + REGULATORS.**



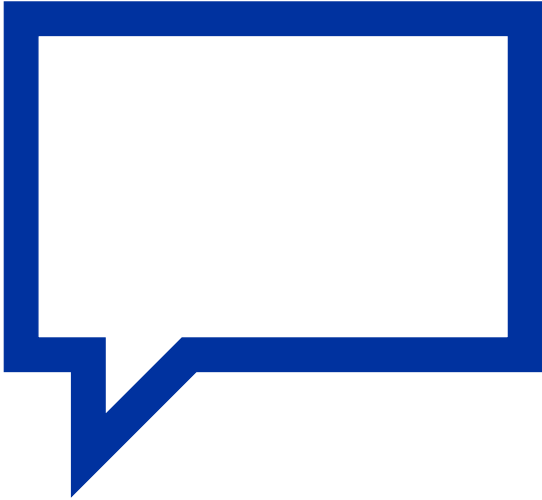
**CONTINUOUS MONITORING:
AI MODEL BEHAVIOUR OVER
TIME.**



**STANDARDIZED AI SECURITY
PROTOCOLS FOR
HEALTHCARE DATA SCIENCE.**

Key Takeaways

- Prompt injection is a **real and under-recognized threat** in clinical data workflows.
- Risks include **privacy breaches, biased evidence, regulatory non-compliance.**
- Mitigation requires **technical safeguards + governance frameworks.**
- Trustworthy AI is essential for **ethical and compliant clinical data science.**



Thank you for listening !!!!



The Global Healthcare Data Science Community

Contact Channels

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

Social Media

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