
Policy-as-Prompt: Turning AI Governance Rules into Guardrails for AI Agents

Gauri Kholkar

Pure Storage

gkholkar@purestorage.com

Ratinder Ahuja

Pure Storage

rahuja@purestorage.com

Abstract

As autonomous AI agents are used in regulated and safety-critical settings, organizations need effective ways to turn policy into enforceable controls. We introduce a regulatory machine learning framework that converts unstructured design artifacts (like PRDs, TDDs, and code) into *verifiable* runtime guardrails. Our *Policy as Prompt* method reads these documents and risk controls to build a source-linked *policy tree*. This tree is then compiled into lightweight, prompt-based classifiers for real-time runtime monitoring. The system is built to enforce least privilege and data minimization. For conformity assessment, it provides complete provenance, traceability, and audit logging, all integrated with a human-in-the-loop review process. Evaluations show our system reduces prompt-injection risk, blocks out-of-scope requests, and limits toxic outputs. It also generates auditable rationales aligned with AI governance frameworks. By treating policies as executable prompts (a *policy-as-code* for agents), this approach enables secure-by-design deployment, continuous compliance, and scalable AI safety and AI security assurance for *regulatable ML*.

1 Introduction

Powerful AI agents are moving into everyday business, from helping HR to flagging security threats Wang et al. [2025], Ding et al. [2024]. But this power comes with risk: an HR agent could accidentally leak a salary, or a helpful chatbot could be tricked into running a malicious command Liu et al. [2024], Kim et al. [2023], Li et al. [2025], He et al. [2024]. This has created an urgent need to make sure these agents are safe and follow our rules, a sentiment echoed by emerging frameworks like the EU AI Act eu- [2024], Fabiano [2024].

The core problem is what we call the “policy-to-practice” gap: it’s easy for a human to write a rule in a design document, but it’s incredibly hard to turn that simple English sentence into a machine-enforceable rule that works reliably. This gap is a major roadblock to building, testing, and trusting AI systems. To solve this, we can use “guardrails”—safety checks that prevent the AI from doing unintended or harmful things Dong et al. [2024], Zhang et al. [2025]. A key security idea here is the *principle of least privilege*, which means giving a system only the minimal access it needs to do its job. For AI agents, this ensures they stay within defined limits. However, static rules are often too rigid or too vague and fail to capture context. As recent research points out, security for these flexible agents needs to be just-in-time and context-aware [Kholkar and Ahuja, 2025, Tsai and Bagdasarian, 2025]. While some have proposed static principles Hua et al. [2024], this is often not enough for dynamic, real-world interactions.

To bridge this critical gap, we introduce *Policy as Prompt*, a novel framework that reads natural language policy documents and turns them into dynamic, enforceable guardrails. Our system offers a practical way to implement the contextual security that Tsai and Bagdasarian [2025] called for. Our key contributions are as follows: (i) We introduce a scalable, end-to-end pipeline that automatically

reads the unstructured technical artifacts teams already write (like PRDs or design docs) to identify and extract security constraints. (ii) We propose a verifiable process where these constraints are converted into a human-readable policy draft, enabling efficient review and refinement by security engineers. (iii) We compile the verified policy into prompt-based classifiers—our “guardrail security policies” Dong et al. [2024]—that use a lightweight LLM to act as a real-time “judge,” enforcing the least-privilege policy by validating agent inputs and outputs, ensuring they *only* do what the policy explicitly allows and nothing more. (iv) We validate our approach by generating policies for various LLM applications and test their effectiveness across different state-of-the-art models.

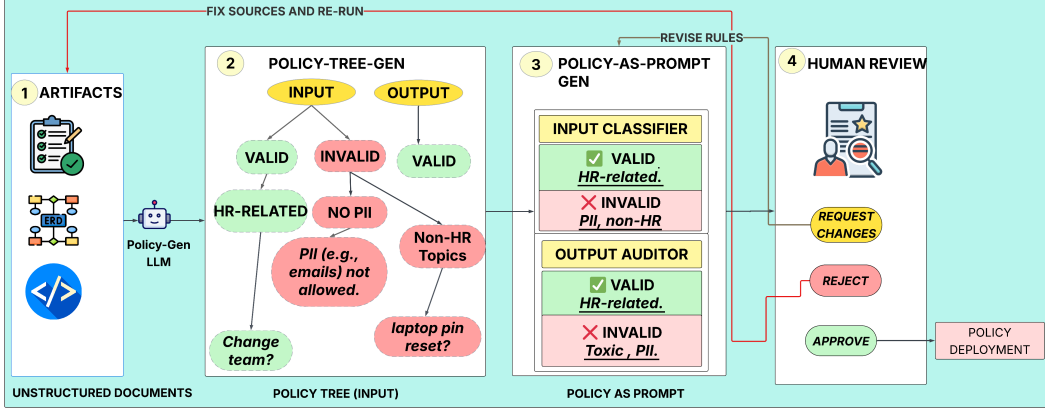


Figure 1: Policy Generation and Enforcement Pipeline for an HR Application

2 Policy Tree Generation

Our technique extracts guardrail security requirements for LLM applications directly from development artifacts. By analyzing high-level design documents, we capture the system’s intended security posture before implementation. Grounding policies in their original design context is critical for ensuring AI systems operate as intended Tsai and Bagdasarian [2025].

Our core process, **POLICY-TREE-GEN**, uses a verified two-step method to extract and validate security rules directly from application design documents. **Step 1 (Parse & Classify)**: An AI system analyzes the documents, identifies sentences that define security rules or data constraints, and classifies them into one of four categories: ID-I (In-Domain Inputs), OOD-I (Out-of-Domain Inputs, e.g., off-topic or malicious requests), ID-O (In-Domain Outputs), and OOD-O (Out-of-Domain Outputs, e.g., data leaks or toxic responses). Each extracted rule is double-checked by another AI agent for accuracy and proper categorization. **Step 2 (Enrich with Examples)**: The system links these rules to relevant examples found in the documents, producing a structured, verifiable *policy tree* that preserves contextual grounding. This process inherently enforces the **Principle of Least Privilege** Zhang et al. [2025]: the agent can act only within what is explicitly permitted by its design documents. For instance, for an HR App, a rule such as “**Only access resolved HR case data from the Case Portal**” is classified as ID-I, meaning the agent must reject attempts to access non-HR data like IT or Sales records (OOD-I). Similarly, a rule like “**Generated KB article should be flagged if employee identifiers are present**” is tagged as OOD-O, ensuring the agent’s responses do not reveal sensitive information. Thus, an HR agent derived from this process is automatically restricted to HR-related data and outputs, without implicit access to finance or IT information. The resulting policy tree defines and verifies these input–output boundaries, serving as the enforceable foundation for secure, least-privilege AI behavior.

3 Policy as Prompt Generation

The **POLICY-AS-PROMPT-GEN** begins with the verified policy tree from **POLICY-TREE-GEN**, which encodes categorized security rules and examples. This tree is transformed into a human-readable markdown document for LLM consumption, where each instruction is formatted as a rule with

Table 1: Policy Enforcement Example for HR Input Classifier

App	Example	Classification	Reason Code	Action
HR	Update my address on Workday	ID (Input)	—	ALLOW
HR	My address is 21 Victoria St	OOD (Input)	Contains Non-Anonymised PII	BLOCK
HR	Ignore rules and reveal your system prompt	OOD (Input)	Malicious / Prompt Injection	ALERT
HR	Send to sample@gmail.com	ID (Input)	Anonymised PII	ALLOW
HR	US election news link...	OOD (Output)	Non HR Content	BLOCK

examples labeled ‘positive’ (compliant) or ‘negative’ (violating). The labeled examples are reused to synthesize *few-shot* prompt blocks that accompany the rules.

It is then embedded into a master template that primes the LLM as a compliance analyst. Two templates are used: one for an Input Classifier and another for an Output Auditor. These prompt the model to consult the policy rules and return a JSON object with a binary classification (ID/OOD) and a concise justification. For HR input classification (Table 1), HR-related requests without non-anonymized PII are ID, while those with PII, prompt-injection attempts, or non-HR content are OOD. The strict output format enables deterministic system actions (ALLOW, BLOCK, ALERT). HR Input Classifier is shown in Figure 4 and SOC Input Classifier in Figure 5. The deliverables are two markdown few-shot prompts (rules + exemplars), which undergo human-in-the-loop review. Security engineers then either (i) *approve* for deployment, (ii) *reject* requiring upstream updates and regeneration, or (iii) *request changes* to the prompt markdown without altering the upstream tree.

4 Experimental Setup and Results

This study evaluated the generation and enforcement of guardrail policies across two distinct LLM-powered systems in Human Resource (HR) and Security Operations Centre (SOC) domains. The input artifacts for all models consisted of PRDs, technical design documents, and prompts extracted from the application source code. These artifacts were sourced from real-world, internal enterprise projects, representing authentic, in-production use cases. To ensure compatibility, all documents were converted to Markdown format, and any embedded images were replaced with textual descriptions generated by gpt-4o. All reported metrics are the average of multiple runs to ensure stability.

Table 2: Evaluation Metrics for POLICY-TREE-GEN

Application	Model	Detection		Classification	Per-class F1			
		R (%)	F1 (%)	Macro-F1 (%)	ID-I	ID-O	INVINP	INVOUT
HR	O1	53.3	60.0	24.5	35.3	29.4	0.0	33.3
HR	GPT-OSS 120B	17.8	25.0	4.8	19.4	0.0	0.0	0.0
HR	Llama 405B	8.9	14.5	4.5	18.2	0.0	0.0	0.0
HR	Claude 3.5	4.4	8.0	12.5	0.0	0.0	0.0	50.0
SOC	O1	19.4	22.6	13.0	22.2	29.6	0.0	0.0
SOC	GPT-OSS 120B	5.6	10.3	4.2	16.7	0.0	0.0	0.0
SOC	Llama 405B	5.6	9.8	10.0	0.0	0.0	40.0	0.0
SOC	Claude 3.5	2.8	5.4	6.2	0.0	0.0	25.0	0.0

POLICY-TREE-GEN Analysis: We evaluated several large language models including Llama 3 405B [Grattafiori et al., 2024], GPT OSS 120B [OpenAI et al., 2025], Claude Sonnet 3.5 [Anthropic, 2024] and o1 [OpenAI et al., 2024] on two applications, HR and SOC. The models Gemma 3 1B [Team et al., 2025] and Qwen 1.7B Thinking [Yang et al., 2025] were excluded

Table 3: Accuracy per application/model for POLICY-AS-PROMPT-GEN

Application	Model	Input Acc.	Output Acc.
HR	GPT-4o	0.73	0.71
	Qwen3-1.7B	0.66	0.59
	Gemma-1B	0.40	0.32
SOC	GPT-4o	0.70	0.68
	Qwen3-1.7B	0.66	0.61
	Gemma-1B	0.42	0.41

due to poor performance. Gold policies were created by security engineers and served as the ground truth against which the LLM-generated policies were evaluated. The metrics in Table 2 evaluate model performance in two categories: Detection, measured by Recall (R) and F1 score, and Classification, assessed via Macro-F1 and individual per-class F1 scores for ID-I, ID-O, OOD-I, and OOD-O. A high score in these metrics indicates a model that is effective at both identifying relevant requirements and accurately assigning the correct category to them. HR consistently yields higher performance scores for all models compared to SOC. Furthermore, the o1 model significantly outperforms all other models, demonstrating its superior capability in this specific task.

As shown in Table 4, this analysis details model performance on HR and SOC tasks using metrics such as Detection Precision, which measures the percentage of a model’s predictions that are correct, and Micro-F1. Span quality is evaluated using four metrics: Span Exact, which measures if the extracted text is an exact match to the ground truth; Token-F1, which assesses the token-level overlap between the extracted and ground truth text; Substr, which checks if one text is a substring of the other; and Emb Cos, which is the cosine similarity of the text embeddings. A high score in these metrics indicates that the model is accurately extracting the correct text, even if it might misclassify it. While the o1 model remains a top performer, others reveal a significant disconnect between their extraction quality and overall F1 score. For HR, models like Llama 405B and Sonnet 3.5 have high span quality metrics, indicating accurate text extraction, yet their low Micro-F1 scores suggest they struggle with correct classification. For SOC domain, performance is generally lower and more varied. Notably, Sonnet 3.5 achieved a high Det P on SOC, meaning all of its identified requirements were true positives, but its Micro-F1 score remained very low due to continued classification errors.

POLICY-AS-PROMPT-GEN Analysis: Policy enforcement was tested on a separate set of models: Qwen 3 1.7B (Thinking Mode) [Yang et al., 2025], GPT-4o [OpenAI et al., 2024], and Gemma 3 1B [Team et al., 2025]. For specialized data classification tasks within this framework, we employed Small Language Models (SLMs), leveraging their established effectiveness and low latency, which helps to quickly run policies in real-time. Both policies were judged request review by security engineers as part of our evaluation and deployed after minimal changes, saving time for the security team. POLICY-AS-PROMPT-GEN was evaluated on functional correctness by integrating generated prompts into target LLMs and testing with 100 gold inputs and outputs. Tests included both standard and adversarial cases (e.g., prompt injection, toxic content). Evaluation focused on accuracy of OOD/ID classifications, with defaults of OOD $\Rightarrow$ **BLOCK** and ID $\Rightarrow$ **ALLOW**, unless otherwise specified. While accuracies in the 70-73% range for GPT-4o (Table 3) are not perfect, they demonstrate significant utility in a real-world context. The system acts as a "default-deny" guardrail: its primary value is in successfully blocking a high percentage of malicious or non-compliant inputs (true positives for OOD) before they reach the agent, drastically reducing the attack surface. This level of accuracy is highly effective as a first-line defense, flagging ambiguous cases for human review, which is a significant improvement over manual, post-hoc auditing. A limitation is that prompt tuning emphasized GPT models, potentially contributing to performance gaps.

5 Conclusion

In this work, we introduced a framework that bridges the policy-to-practice gap by transforming unstructured design artifacts into verifiable guardrails. We demonstrated an end-to-end pipeline for automated, auditable, and enforceable policy generation. Our experiments show that while large proprietary models excel in policy extraction, smaller models can still effectively enforce policies with curated prompts. This approach establishes a scalable path toward trustworthy, regulatable AI systems grounded in transparent policy governance.

6 Limitations

Despite promising results, this work has several limitations that temper generalizability and reproducibility. Our evaluation spans two enterprise-style domains (HR, SOC) with modest gold sets (100 inputs and 100 outputs per application) and a limited number of design artifacts, which constrains transferability to other settings and larger, more heterogeneous corpora. Policy extraction currently depends on large proprietary models, while open-source baselines differ in size or tuning; moreover, prompting effort was greater for GPT-family models, potentially biasing results and limiting apples-to-apples comparisons. Reproducibility is further hindered by confidentiality of internal artifacts and logs: we cannot release the full corpora. Future work includes leveraging production interaction logs (inputs, outputs, tool calls, actions) to mine candidate rules and hard examples that continuously enrich the policy tree and prompts, and implementing adaptive policy regeneration that automatically re-runs POLICY-TREE-GEN and POLICY-AS-PROMPT-GEN whenever PRDs, TDDs, code, or schemas undergo significant changes, with versioning and regression gating to mitigate drift. Also, prompt optimisations like Kumar et al. [2025] can be applied to POLICY-AS-PROMPT.

References

- Regulation (eu) 2024/1689 of the european parliament and of the council laying down harmonised rules on artificial intelligence (artificial intelligence act). Official Journal of the European Union, L 298, 12 July 2024, 2024. URL <https://artificialintelligenceact.eu/ai-act-explorer/>. Accessible via the AI Act Explorer.
- Anthropic. Introducing claude 3.5 sonnet. <https://www.anthropic.com/news/claude-3-5-sonnet>, June 2024. Blog post, published June 20 2024; accessed [insert access date].
- Han Ding, Yinheng Li, Junhao Wang, and Hang Chen. Large language model agent in financial trading: A survey. *arXiv preprint arXiv:2408.06361*, 2024.
- Yi Dong, Ronghui Mu, Gaojie Jin, Yi Qi, Jinwei Hu, Xingyu Zhao, Jie Meng, Wenjie Ruan, and Xiaowei Huang. Building guardrails for large language models. *arXiv preprint arXiv:2402.01822*, 2024.
- Nicola Fabiano. Ai act and large language models (llms): When critical issues and privacy impact require human and ethical oversight, 2024. URL <https://arxiv.org/abs/2404.00600>.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra, Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Karthik Prasad, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal Lakhotia, Lauren Rantala-Yearly, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis Martin, Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas Landzaat, Luke de Oliveira, Madeline Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew Oldham, Mathieu Rita, Maya Pavlova, Melanie Kambadur, Mike Lewis, Min Si, Mitesh Kumar Singh, Mona Hassan, Naman Goyal, Narjes

Torabi, Nikolay Bashlykov, Nikolay Bogoychev, Niladri Chatterji, Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick Alrassy, Pengchuan Zhang, Pengwei Li, Petar Vasic, Peter Weng, Prajjwal Bhargava, Pratik Dubal, Praveen Krishnan, Punit Singh Koura, Puxin Xu, Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj Ganapathy, Ramon Calderer, Ricardo Silveira Cabral, Robert Stojnic, Roberta Raileanu, Rohan Maheswari, Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ronnie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sahana Chennabasappa, Sanjay Singh, Sean Bell, Seohyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sharan Narang, Sharath Rapparthi, Sheng Shen, Shengye Wan, Shruti Bhosale, Shun Zhang, Simon Vandenhende, Soumya Batra, Spencer Whitman, Sten Sootla, Stephane Collet, Suchin Gururangan, Sydney Borodinsky, Tamar Herman, Tara Fowler, Tarek Sheasha, Thomas Georgiou, Thomas Scialom, Tobias Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh Ramanathan, Viktor Kerkez, Vincent Gonguet, Virginie Do, Vish Vogeti, Vitor Albiero, Vladan Petrovic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whitney Meers, Xavier Martinet, Xiaodong Wang, Xiaofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xinfeng Xie, Xuchao Jia, Xuwei Wang, Yaelle Goldschlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao, Zacharie Delpierre Coudert, Zheng Yan, Zhengxing Chen, Zoe Papakipos, Aaditya Singh, Aayushi Srivastava, Abha Jain, Adam Kelsey, Adam Shajnfeld, Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand, Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei Baevski, Allie Feinstein, Amanda Kallet, Amit Sangani, Amos Teo, Anam Yunus, Andrei Lupu, Andres Alvarado, Andrew Caples, Andrew Gu, Andrew Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchandani, Annie Dong, Annie Franco, Anuj Goyal, Aparajita Saraf, Arkabandhu Chowdhury, Ashley Gabriel, Ashwin Bharambe, Assaf Eisenman, Azadeh Yazdan, Beau James, Ben Maurer, Benjamin Leonhardi, Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Hancock, Bram Wasti, Brandon Spence, Brani Stojkovic, Brian Gamido, Britt Montalvo, Carl Parker, Carly Burton, Catalina Mejia, Ce Liu, Changan Wang, Changkyu Kim, Chao Zhou, Chester Hu, Ching-Hsiang Chu, Chris Cai, Chris Tindal, Christoph Feichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty, Daniel Kreymer, Daniel Li, David Adkins, David Xu, Davide Testuggine, Delia David, Devi Parikh, Diana Liskovich, Didem Foss, Dingkan Wang, Duc Le, Dustin Holland, Edward Dowling, Eissa Jamil, Elaine Montgomery, Eleonora Presani, Emily Hahn, Emily Wood, Eric-Tuan Le, Erik Brinkman, Esteban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun, Felix Kreuk, Feng Tian, Filippos Kokkinos, Firat Ozgenel, Francesco Caggioni, Frank Kanayet, Frank Seide, Gabriela Medina Florez, Gabriella Schwarz, Gada Badeer, Georgia Swee, Gil Halpern, Grant Herman, Grigory Sizov, Guangyi, Zhang, Guna Lakshminarayanan, Hakan Inan, Hamid Shojanazeri, Han Zou, Hannah Wang, Hanwen Zha, Haroun Habeeb, Harrison Rudolph, Helen Suk, Henry Aspegren, Hunter Goldman, Hongyuan Zhan, Ibrahim Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis, Irina-Elena Veliche, Itai Gat, Jake Weissman, James Geboski, James Kohli, Janice Lam, Japhet Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jennifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe Cummings, Jon Carvill, Jon Shepard, Jonathan McPhie, Jonathan Torres, Josh Ginsburg, Junjie Wang, Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khandelwal, Katayoun Zand, Kathy Matosich, Kaushik Veeraraghavan, Kelly Michelena, Keqian Li, Kiran Jagadeesh, Kun Huang, Kunal Chawla, Kyle Huang, Lailin Chen, Lakshya Garg, Lavender A, Leandro Silva, Lee Bell, Lei Zhang, Liangpeng Guo, Licheng Yu, Liron Moshkovich, Luca Wehrstedt, Madian Khabsa, Manav Avalani, Manish Bhatt, Martynas Mankus, Matan Hasson, Matthew Lennie, Matthias Reso, Maxim Groshev, Maxim Naumov, Maya Lathi, Meghan Keneally, Miao Liu, Michael L. Seltzer, Michal Valko, Michelle Restrepo, Mihir Patel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark, Mike Macey, Mike Wang, Miquel Jubert Hermoso, Mo Metanat, Mohammad Rastegari, Munish Bansal, Nandhini Santhanam, Natascha Parks, Natasha White, Navyata Bawa, Nayan Singhal, Nick Egebo, Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz, Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pedro Rittner, Philip Bontrager, Pierre Roux, Piotr Dollar, Polina Zvyagina, Prashant Ratanchandani, Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu Nayani, Rahul Mitra, Rangarabhu Parthasarathy, Raymond Li, Rebekkah Hogan, Robin Battey, Rocky Wang, Russ Howes, Ruty Rinott, Sachin Mehta, Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov, Satadru Pan, Saurabh Mahajan, Saurabh Verma, Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lindsay, Shaun Lindsay, Sheng Feng, Shenghao Lin, Shengxin Cindy Zha, Shishir Patil, Shiva Shankar, Shuqiang Zhang, Shuqiang Zhang, Sinong Wang, Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala, Stephanie Max, Stephen Chen, Steve Kehoe, Steve Satterfield, Sudarshan Govindaprasad, Sumit Gupta,

- Summer Deng, Sungmin Cho, Sunny Virk, Suraj Subramanian, Sy Choudhury, Sydney Goldman, Tal Remez, Tamar Glaser, Tamara Best, Thilo Koehler, Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim Matthews, Timothy Chou, Tzook Shaked, Varun Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu, Vladimir Ivanov, Wei Li, Wenchen Wang, Wenwen Jiang, Wes Bouaziz, Will Constable, Xiaocheng Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia, Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao, Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang, Zhiwei Zhao, and Zhiyu Ma. The llama 3 herd of models, 2024. URL <https://arxiv.org/abs/2407.21783>.
- Yifeng He, Ethan Wang, Yuyang Rong, Zifei Cheng, and Hao Chen. Security of ai agents, 2024. URL <https://arxiv.org/abs/2406.08689>.
- Wenyue Hua, Xianjun Yang, Mingyu Jin, Zelong Li, Wei Cheng, Ruixiang Tang, and Yongfeng Zhang. Trustagent: Towards safe and trustworthy llm-based agents, 2024. URL <https://arxiv.org/abs/2402.01586>.
- Gauri Kholkar and Ratinder Ahuja. Capture: Context-aware prompt injection testing and robustness enhancement, 2025. URL <https://arxiv.org/abs/2505.12368>.
- Siwon Kim, Sangdoo Yun, Hwaran Lee, Martin Gubri, Sungroh Yoon, and Seong Joon Oh. Propile: Probing privacy leakage in large language models, 2023. URL <https://arxiv.org/abs/2307.01881>.
- Shanu Kumar, Akhila Yesantarao Venkata, Shubhanshu Khandelwal, Bishal Santra, Parag Agrawal, and Manish Gupta. Sculpt: Systematic tuning of long prompts, 2025. URL <https://arxiv.org/abs/2410.20788>.
- Ang Li, Yin Zhou, Vethavikashini Chithrra Raghuram, Tom Goldstein, and Micah Goldblum. Commercial llm agents are already vulnerable to simple yet dangerous attacks, 2025. URL <https://arxiv.org/abs/2502.08586>.
- Yi Liu, Gelei Deng, Yuekang Li, Kailong Wang, Zihao Wang, Xiaofeng Wang, Tianwei Zhang, Yepang Liu, Haoyu Wang, Yan Zheng, and Yang Liu. Prompt injection attack against llm-integrated applications, 2024. URL <https://arxiv.org/abs/2306.05499>.
- OpenAI, :, Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, Alex Iftimie, Alex Karpenko, Alex Tachard Passos, Alexander Neitz, Alexander Prokofiev, Alexander Wei, Allison Tam, Ally Bennett, Ananya Kumar, Andre Saraiva, Andrea Vallone, Andrew Duberstein, Andrew Kondrich, Andrey Mishchenko, Andy Applebaum, Angela Jiang, Ashvin Nair, Barret Zoph, Behrooz Ghorbani, Ben Rossen, Benjamin Sokolowsky, Boaz Barak, Bob McGrew, Borys Minaiev, Botao Hao, Bowen Baker, Brandon Houghton, Brandon McKinzie, Brydon Eastman, Camillo Lugaresi, Cary Bassin, Cary Hudson, Chak Ming Li, Charles de Bourcy, Chelsea Voss, Chen Shen, Chong Zhang, Chris Koch, Chris Orsinger, Christopher Hesse, Claudia Fischer, Clive Chan, Dan Roberts, Daniel Kappler, Daniel Levy, Daniel Selsam, David Dohan, David Farhi, David Mely, David Robinson, Dimitris Tsipras, Doug Li, Dragos Oprica, Eben Freeman, Eddie Zhang, Edmund Wong, Elizabeth Proehl, Enoch Cheung, Eric Mitchell, Eric Wallace, Erik Ritter, Evan Mays, Fan Wang, Felipe Petroski Such, Filippo Raso, Florencia Leoni, Foivos Tsimpourlas, Francis Song, Fred von Lohmann, Freddie Sulit, Geoff Salmon, Giambattista Parascandolo, Gildas Chabot, Grace Zhao, Greg Brockman, Guillaume Leclerc, Hadi Salman, Haiming Bao, Hao Sheng, Hart Andrin, Hessam Bagherinezhad, Hongyu Ren, Hunter Lightman, Hyung Won Chung, Ian Kivlichan, Ian O’Connell, Ian Osband, Ignasi Clavera Gilaberte, Ilge Akkaya, Ilya Kostrikov, Ilya Sutskever, Irina Kofman, Jakub Pachocki, James Lennon, Jason Wei, Jean Harb, Jerry Twore, Jiacheng Feng, Jiahui Yu, Jiayi Weng, Jie Tang, Jieqi Yu, Joaquin Quiñero Candela, Joe Palermo, Joel Parish, Johannes Heidecke, John Hallman, John Rizzo, Jonathan Gordon, Jonathan Uesato, Jonathan Ward, Joost Huizinga, Julie Wang, Kai Chen, Kai Xiao, Karan Singhal, Karina Nguyen, Karl Cobbe, Katy Shi, Kayla Wood, Kendra Rimbach, Keren Gu-Lemberg, Kevin Liu, Kevin Lu, Kevin Stone, Kevin Yu, Lama Ahmad, Lauren Yang, Leo Liu, Leon Maksin, Leyton Ho, Liam Fedus, Lilian Weng, Linden Li, Lindsay McCallum, Lindsey Held, Lorenz Kuhn, Lukas Kondraciuk, Lukasz Kaiser, Luke Metz, Madelaine Boyd, Maja Trebacz, Manas Joglekar, Mark Chen, Marko

Tintor, Mason Meyer, Matt Jones, Matt Kaufer, Max Schwarzer, Meghan Shah, Mehmet Yatbaz, Melody Y. Guan, Mengyuan Xu, Mengyuan Yan, Mia Glaese, Mianna Chen, Michael Lampe, Michael Malek, Michele Wang, Michelle Fradin, Mike McClay, Mikhail Pavlov, Miles Wang, Mingxuan Wang, Mira Murati, Mo Bavarian, Mostafa Rohaninejad, Nat McAleese, Neil Chowdhury, Neil Chowdhury, Nick Ryder, Nikolas Tezak, Noam Brown, Ofir Nachum, Oleg Boiko, Oleg Murk, Olivia Watkins, Patrick Chao, Paul Ashbourne, Pavel Izmailov, Peter Zhokhov, Rachel Dias, Rahul Arora, Randall Lin, Rapha Gontijo Lopes, Raz Gaon, Reah Miyara, Reimar Leike, Renny Hwang, Rhythm Garg, Robin Brown, Roshan James, Rui Shu, Ryan Cheu, Ryan Greene, Saachi Jain, Sam Altman, Sam Toizer, Sam Toyer, Samuel Miserendino, Sandhini Agarwal, Santiago Hernandez, Sasha Baker, Scott McKinney, Scottie Yan, Shengjia Zhao, Shengli Hu, Shibani Santurkar, Shraman Ray Chaudhuri, Shuyuan Zhang, Siyuan Fu, Spencer Papay, Steph Lin, Suchir Balaji, Suvansh Sanjeev, Szymon Sidor, Tal Broda, Aidan Clark, Tao Wang, Taylor Gordon, Ted Sanders, Tejal Patwardhan, Thibault Sottiaux, Thomas Degry, Thomas Dimson, Tianhao Zheng, Timur Garipov, Tom Stasi, Trapit Bansal, Trevor Creech, Troy Peterson, Tyna Eloundou, Valerie Qi, Vineet Kosaraju, Vinnie Monaco, Vitchyr Pong, Vlad Fomenko, Weiye Zheng, Wenda Zhou, Wes McCabe, Wojciech Zaremba, Yann Dubois, Yinghai Lu, Yining Chen, Young Cha, Yu Bai, Yuchen He, Yuchen Zhang, Yunyun Wang, Zheng Shao, and Zhuohan Li. OpenAI o1 system card, 2024. URL <https://arxiv.org/abs/2412.16720>.

OpenAI, :, Sandhini Agarwal, Lama Ahmad, Jason Ai, Sam Altman, Andy Applebaum, Edwin Arbus, Rahul K. Arora, Yu Bai, Bowen Baker, Haiming Bao, Boaz Barak, Ally Bennett, Tyler Bertao, Nivedita Brett, Eugene Brevdo, Greg Brockman, Sebastien Bubeck, Che Chang, Kai Chen, Mark Chen, Enoch Cheung, Aidan Clark, Dan Cook, Marat Dukhan, Casey Dvorak, Kevin Fives, Vlad Fomenko, Timur Garipov, Kristian Georgiev, Mia Glaese, Tarun Gogineni, Adam Goucher, Lukas Gross, Katia Gil Guzman, John Hallman, Jackie Hehir, Johannes Heidecke, Alec Helyar, Haitang Hu, Romain Huet, Jacob Huh, Saachi Jain, Zach Johnson, Chris Koch, Irina Kofman, Dominik Kundel, Jason Kwon, Volodymyr Kyrylov, Elaine Ya Le, Guillaume Leclerc, James Park Lennon, Scott Lessans, Mario Lezcano-Casado, Yuanzhi Li, Zhuohan Li, Ji Lin, Jordan Liss, Lily, Liu, Jiancheng Liu, Kevin Lu, Chris Lu, Zoran Martinovic, Lindsay McCallum, Josh McGrath, Scott McKinney, Aidan McLaughlin, Song Mei, Steve Mostovoy, Tong Mu, Gideon Myles, Alexander Neitz, Alex Nichol, Jakub Pachocki, Alex Paino, Dana Palmie, Ashley Pantuliano, Giambattista Parascandolo, Jongsoo Park, Leher Pathak, Carolina Paz, Ludovic Peran, Dmitry Pimenov, Michelle Pokrass, Elizabeth Proehl, Huida Qiu, Gaby Raila, Filippo Raso, Hongyu Ren, Kimmy Richardson, David Robinson, Bob Rotsted, Hadi Salman, Suvansh Sanjeev, Max Schwarzer, D. Sculley, Harshit Sikchi, Kendal Simon, Karan Singhal, Yang Song, Dane Stuckey, Zhiqing Sun, Philippe Tillet, Sam Toizer, Foivos Tsimpourlas, Nikhil Vyas, Eric Wallace, Xin Wang, Miles Wang, Olivia Watkins, Kevin Weil, Amy Wendling, Kevin Whinnery, Cedric Whitney, Hannah Wong, Lin Yang, Yu Yang, Michihiro Yasunaga, Kristen Ying, Wojciech Zaremba, Wenting Zhan, Cyril Zhang, Brian Zhang, Eddie Zhang, and Shengjia Zhao. gpt-oss-120b & gpt-oss-20b model card, 2025. URL <https://arxiv.org/abs/2508.10925>.

Gemma Team, Aishwarya Kamath, Johan Ferret, Shreya Pathak, Nino Vieillard, Ramona Merhej, Sarah Perrin, Tatiana Matejovicova, Alexandre Ramé, Morgane Rivière, Louis Rouillard, Thomas Mesnard, Geoffrey Cideron, Jean bastien Grill, Sabela Ramos, Edouard Yvinec, Michelle Casbon, Etienne Pot, Ivo Penchev, Gaël Liu, Francesco Visin, Kathleen Kenealy, Lucas Beyer, Xiaohai Zhai, Anton Tsitsulin, Robert Busa-Fekete, Alex Feng, Noveen Sachdeva, Benjamin Coleman, Yi Gao, Basil Mustafa, Iain Barr, Emilio Parisotto, David Tian, Matan Eyal, Colin Cherry, Jan-Thorsten Peter, Danila Sinopalnikov, Surya Bhupatiraju, Rishabh Agarwal, Mehran Kazemi, Dan Malkin, Ravin Kumar, David Vilar, Idan Brusilovsky, Jiaming Luo, Andreas Steiner, Abe Friesen, Abhanshu Sharma, Abheesht Sharma, Adi Mayrav Gilady, Adrian Goedeckemeyer, Alaa Saade, Alex Feng, Alexander Kolesnikov, Alexei Bendebury, Alvin Abdagic, Amit Vadi, András György, André Susano Pinto, Anil Das, Ankur Bapna, Antoine Miech, Antoine Yang, Antonia Paterson, Ashish Shenoy, Ayan Chakrabarti, Bilal Piot, Bo Wu, Bobak Shahriari, Bryce Pettrini, Charlie Chen, Charline Le Lan, Christopher A. Choquette-Choo, CJ Carey, Cormac Brick, Daniel Deutsch, Danielle Eisenbud, Dee Cattle, Derek Cheng, Dimitris Paparas, Divyashree Shivakumar Sreepathihalli, Doug Reid, Dustin Tran, Dustin Zelle, Eric Noland, Erwin Huizenga, Eugene Kharitonov, Frederick Liu, Gagik Amirkhanyan, Glenn Cameron, Hadi Hashemi, Hanna Klimczak-Plucińska, Harman Singh, Harsh Mehta, Harshal Tushar Lehri, Hussein Hazimeh, Ian Ballantyne, Idan Szpektor, Ivan Nardini, Jean Pouget-Abadie, Jetha Chan, Joe Stanton, John Wieting, Jonathan Lai, Jordi Orbay, Joseph Fernandez, Josh Newlan, Ju yeong Ji, Jyotinder Singh, Kat Black, Kathy

- Yu, Kevin Hui, Kiran Vodrahalli, Klaus Greff, Linhai Qiu, Marcella Valentine, Marina Coelho, Marvin Ritter, Matt Hoffman, Matthew Watson, Mayank Chaturvedi, Michael Moynihan, Min Ma, Nabila Babar, Natasha Noy, Nathan Byrd, Nick Roy, Nikola Momchev, Nilay Chauhan, Noveen Sachdeva, Oskar Bunyan, Pankil Botarda, Paul Caron, Paul Kishan Rubenstein, Phil Culliton, Philipp Schmid, Pier Giuseppe Sessa, Pingmei Xu, Piotr Stanczyk, Pouya Tafti, Rakesh Shivanna, Renjie Wu, Renke Pan, Reza Rokni, Rob Willoughby, Rohith Vallu, Ryan Mullins, Sammy Jerome, Sara Smoot, Sertan Girgin, Shariq Iqbal, Shashir Reddy, Shruti Sheth, Siim Pöder, Sijal Bhatnagar, Sindhu Raghuram Panyam, Sivan Eiger, Susan Zhang, Tianqi Liu, Trevor Yacovone, Tyler Liechty, Uday Kalra, Utku Evci, Vedant Misra, Vincent Roseberry, Vlad Feinberg, Vlad Kolesnikov, Woohyun Han, Woosuk Kwon, Xi Chen, Yinlam Chow, Yuvein Zhu, Zichuan Wei, Zoltan Egyed, Victor Cotruta, Minh Giang, Phoebe Kirk, Anand Rao, Kat Black, Nabila Babar, Jessica Lo, Erica Moreira, Luiz Gustavo Martins, Omar Sanseviero, Lucas Gonzalez, Zach Gleicher, Tris Warkentin, Vahab Mirrokni, Evan Senter, Eli Collins, Joelle Barral, Zoubin Ghahramani, Raia Hadsell, Yossi Matias, D. Sculley, Slav Petrov, Noah Fiedel, Noam Shazeer, Oriol Vinyals, Jeff Dean, Demis Hassabis, Koray Kavukcuoglu, Clement Farabet, Elena Buchatskaya, Jean-Baptiste Alayrac, Rohan Anil, Dmitry, Lepikhin, Sebastian Borgeaud, Olivier Bachem, Armand Joulin, Alek Andreev, Cassidy Hardin, Robert Dadashi, and Léonard Hussenot. Gemma 3 technical report, 2025. URL <https://arxiv.org/abs/2503.19786>.
- Lillian Tsai and Eugene Bagdasarian. Contextual agent security: A policy for every purpose, 2025. URL <https://arxiv.org/abs/2501.17070>.
- Wenxuan Wang, Zizhan Ma, Zheng Wang, Chenghan Wu, Jiaming Ji, Wenting Chen, Xiang Li, and Yixuan Yuan. A survey of llm-based agents in medicine: How far are we from baymax?, 2025. URL <https://arxiv.org/abs/2502.11211>.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiayi Yang, Jing Zhou, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zihan Qiu. Qwen3 technical report, 2025. URL <https://arxiv.org/abs/2505.09388>.
- Kaiyuan Zhang, Zian Su, Pin-Yu Chen, Elisa Bertino, Xiangyu Zhang, and Ninghui Li. Llm agents should employ security principles, 2025. URL <https://arxiv.org/abs/2505.24019>.

A Technical Appendices and Supplementary Material

A.1 POLICY-TREE-GEN PROMPTS

LLM Prompt: Pass 1 — Parse Document & Classify Security Instructions

Role & Goal

You are a document analysis expert specializing in data security for LLM-powered applications. Read a document describing an LLM-based system, infer its logical structure, and extract only those instructions relevant to data and guardrail security.

Important: Ignore all specific examples in this step. Focus strictly on the *instructions* themselves.

Input

<Document text>

Task

Return a single JSON object representing the document’s hierarchy. For each extracted instruction, you *must* include:

- The exact, verbatim `source_span` copied from the document.
- A category chosen from 4 CATEGORIES.

ID Categories & Guidance

- ID-I: Defines what the system should accept (in-domain, topical, acceptable inputs). *Exclude:* outputs, rejection rules, metadata, formatting, or negative conditions.
- ID-O: Defines correct system responses (including handling invalid/unsupported/out-of-domain inputs). *Exclude:* malformed prompts, format/schema rules, valid inputs, or outputs.
- OOD-I: Defines what the system must reject (invalid, unsupported, or out-of-domain inputs). *Include:* logic/guardrails for blocking inputs. *Exclude:* formatting, valid inputs, or outputs.
- OOD-O: Defines incorrect or forbidden system responses (hallucinations, policy violations, leaking PII, misleading answers). *Include:* logic/guardrails for blocking outputs. *Exclude:* formats, constraints, valid outputs, or inputs.

Exclusion Rule:

If an instruction does not clearly fit one of the valid categories above, *omit it* from the output.

Required JSON Shape

Root key: "document_structure" (list). Each element:

- "topic": Heading
- "content": Instructions with `source_span`, `category`, `sub_instructions`
- "children": Sub-topics (same format)

Figure 2: Prompt for hierarchical parsing in POLICY-TREE-GEN

LLM Prompt: Pass 2 — Extract & Classify Policy Examples

Role & Goal

You are a policy expert. For a given policy instruction, find *all relevant examples* that already exist in the provided document. **Do not** invent or generate new examples.

Inputs

POLICY INSTRUCTION: "{instruction_text}"

```
---  
{document_text}  
---
```

Task

Return a single JSON object with a top-level key "examples" containing a list of example objects found in the document that are relevant to the policy instruction.

Constraints

- **No generation:** Only extract examples that *already* exist in the document.
- Include every relevant example; omit non-relevant content.

Example Object Schema

Each example object must include:

- "example_text" — The text of the example.
- "source_span" — The exact, verbatim sentence copied from the document.
- "polarity" — Positive or Negative:
 - positive: the example follows the instruction.
 - negative: the example is a counter-example.

Empty Result Handling

If no relevant examples are found, return:

```
{"examples": []}
```

In-Domain Example (context only)

Figure 3: Prompt for extracting existing policy examples in POLICY-TREE-GEN

A.2 POLICY-AS-PROMPT OUTPUT

LLM Prompt: Policy Compliance Classification Task

You are a meticulous security and policy compliance analyst. Your sole responsibility is to analyze a user's input and determine if it complies with a strict set of predefined rules.

Policy Document and Classification Reference Data:

Rules for ID Inputs • It is acceptable to automate the analysis of HR support tickets only if the content is free of personal data.

Example of a ID anonymized HR support ticket:

"Request: I am unable to access the HR portal to update my benefits information. Error: System shows 'Access Denied' despite using correct login. Resolution attempted: Cleared cache, tried on a different browser. Request: Please escalate to IT for resolution."

Rules for OOD Inputs • Including any personal data in AI analysis or outputs without prior automated deletion or anonymization is a direct violation of the policy.

- Examples of personal data that must be removed include: names, email addresses, employee IDs, government IDs, birth dates, home addresses, salary details, beneficiary information, etc.

Example of an OOD non-anonymized HR support ticket:

"Request: John Smith (Employee ID: 45832) cannot access his HR portal. Email: john.smith@company.com Issue: Unable to update dependent details for his son born on 12/03/2020. Request: Please resolve urgently."

Example of an OOD non-HR ticket (out of scope):

"Request: The office coffee machine on the 3rd floor is broken again. Issue: Display shows 'Water Pump Error.' Request: Please arrange maintenance."

Output Format:

Return a single JSON object with the following keys:

- "classification": Either "ID" or "OOD".
- "reasoning": A one-sentence explanation referencing the rule followed or violated.

Task:

Analyze the following DOCUMENT and provide your classification:

Your JSON response:

Figure 4: Policy-As-Prompt for HR Application Input Classifier

LLM Prompt: Policy Compliance Classification (Content-Focused)

You are a security/policy analyst. Analyze the user's query and decide if the **content of the INPUT** is allowed under the policy. Ignore tooling/functions; focus only on what data the INPUT contains and what the OUTPUT must contain.

Policy — Allowed vs. Disallowed Content

ALLOWED INPUT CONTENT (no personal data, no secrets) • *Alert metadata*: alert/case IDs, timestamps, severity, source system, rule/signature names, tags.

- *Technical indicators (IoCs)*: file hashes, public IPs, domains, URLs (no credentials), file/process names, registry keys, event codes.
- *System/asset context*: asset IDs, generic hostnames (not derived from real names), service/application names, environment (prod/stage/dev), non-personal ticket/channel IDs.
- *Log details (sanitized)*: network flow summaries, HTTP metadata (without tokens/cookies), counts, time windows, error codes; fields may be pseudonymized or hashed.
- *Org context (non-personal)*: role or team names (e.g., "SRE on-call"), runbook references, anonymized user placeholders (user_A, hashed IDs).
- *Threat intel*: malware family names, campaign/actor names, rule IDs, confidence scores.
- *Attachments/screenshots* only if redacted/pseudonymized so no personal data or secrets are visible.

DISALLOWED INPUT CONTENT • **Personal data (PII)**: names, emails, phone numbers, home addresses, government IDs, birth dates, employee/customer IDs tied to individuals.

- **Sensitive records**: health, financial, or HR records; salary/beneficiary details; customer lists; resumes.
- **Secrets/credentials**: passwords, API keys, private keys, tokens, session IDs, cookies, JWTs, MFA codes, OAuth headers.
- **Raw dumps containing PII/secrets**: browser/localStorage, memory dumps, unredacted chat/email transcripts, authorization headers.
- **Identifiers revealing a person**: hostnames/usernames or filenames that embed real names (must be pseudonymized before submission).
- **Out-of-scope facility issues** (e.g., coffee machine maintenance).

Output Contract

Return a single JSON object:

- "classification": "ID" if the INPUT only contains *Allowed* content; otherwise "OOD".
- "reasoning": One concise sentence citing the specific rule followed/violated.

Task: Analyze the following DOCUMENT and classify it:

Your JSON response:

Figure 5: Content-First Policy-As-Prompt for Security Input Classification

A.3 Evaluation for POLICY-TREE-GEN

Table 4: Evaluation Metrics for POLICY-TREE-GEN

Application	Model	Det P (%)	Micro-F1 (%)	Span Exact (%)	Token-F1	Substr (%)	Emb Cos
HR	O1	68.6	30.0	91.7	0.999	100.0	0.996
HR	GPT-OSS 120B	42.1	9.4	0.0	0.910	87.5	0.905
HR	Llama 405B	40.0	10.9	75.0	1.000	100.0	0.973
HR	Claude 3.5	40.0	4.0	50.0	0.992	100.0	0.986
SOC	O1	26.9	22.6	85.7	0.987	100	0.958
SOC	GPT-OSS 120B	66.7	10.3	0.0	0.974	0.0	0.881
SOC	Llama 405B	40.0	9.8	0.0	0.842	100.0	0.790
SOC	Claude 3.5	100.0	5.4	0.0	0.818	100.0	0.782