

Automated Classification of Radiation Oncology Safety Events Using Large Language Models (LLMs)

A novel approach to streamline reporting and enable retrospective analysis

Qiongge Li, PhD^{1,2} Jian Liu, PhD^{3,4,5} Xing Li, PhD⁶, Wei Nie, PhD^{1,2}, and Jiajin Fan PhD¹

¹ Department of Radiation Oncology, Inova Schar Cancer Institute, Fairfax, VA

² School of Medicine, University of Virginia, Charlottesville, VA

³ Department of Radiation Medicine, Northwell, New Hyde Park, NY

⁴ Zucker School of Medicine at Hofstra/Northwell, Hempstead, NY

⁵ Physics and Astronomy, Hofstra University, Hempstead, NY

⁶ Department of Radiation Oncology, University of Texas Southwestern Medical Center, Dallas, TX

ABSTRACT

Purpose: To evaluate the feasibility of using a large language model (LLM) to automatically classify radiation oncology safety events from free-text narratives, and to outline an AI-assisted reporting framework.

Methods: 788 incidents (Jan 2021–Sep 2025) from our institutional reporting system were classified by GPT-5 across seven dimensions via iterative prompt refinement. Output was validated by three independent, blinded expert physicists using Cohen’s κ .

Results: Model–reviewer agreement (mean $\kappa = 0.31$) was comparable to inter-reviewer agreement (mean $\kappa = 0.42$), approaching expert-level performance. Communication failures were the most common failure mode (~40%).

Conclusions: LLM-assisted classification can turn unanalyzed free-text reports into structured, analyzable data with no added reporter burden—a proof-of-concept for AI-enabled incident learning.

INTRODUCTION

Incident reporting underpins a strong safety culture, but its value depends on whether the data can be analyzed.

Our institutional platform, SafetyAlways, prioritizes ease of reporting: staff write a free-text narrative and only minimal structured fields are captured. As a result, 788 incidents accumulated over ~5 years sit essentially unanalyzed.

Structured systems (RO-ILS, SAFRON) capture rich data, but only by asking reporters to classify each event by hand, a disincentive to report in a busy clinic.

Key insight: the information needed for classification already lives in the narrative. What is missing is the extraction. We propose using an LLM to extract structured fields directly from the free text the reporter already writes.

Study goals:

- Demonstrate feasibility of multi-dimensional LLM classification from narrative alone
- Surface hidden patterns: common failure modes and workflow vulnerabilities
- Design an AI-assisted, human-in-the-loop reporting system

METHODS AND MATERIALS

Data: 788 de-identified radiation oncology safety reports (Jan 2021–Sep 2025).

Model: GPT-5 (OpenAI) via API, temperature = 0 for deterministic output. Zero-shot, instruction-based prompting; one stateless call per report.

Seven dimensions: failure mode, severity, treatment type, discoverer role, occurred & discovered workflow stage, and an exclusion flag. Four dimensions permit multiple labels.

Prompt refinement: refined on 50 incidents over 19 reasoning patterns, then frozen. The 50-incident refinement set and 80-incident validation set were strictly separate—reported agreement is held-out.

Validation: three independent, blinded expert medical physicists classified 80 incidents; pairwise Cohen’s κ per dimension. After exclusions, 650 records remained for retrospective analysis.

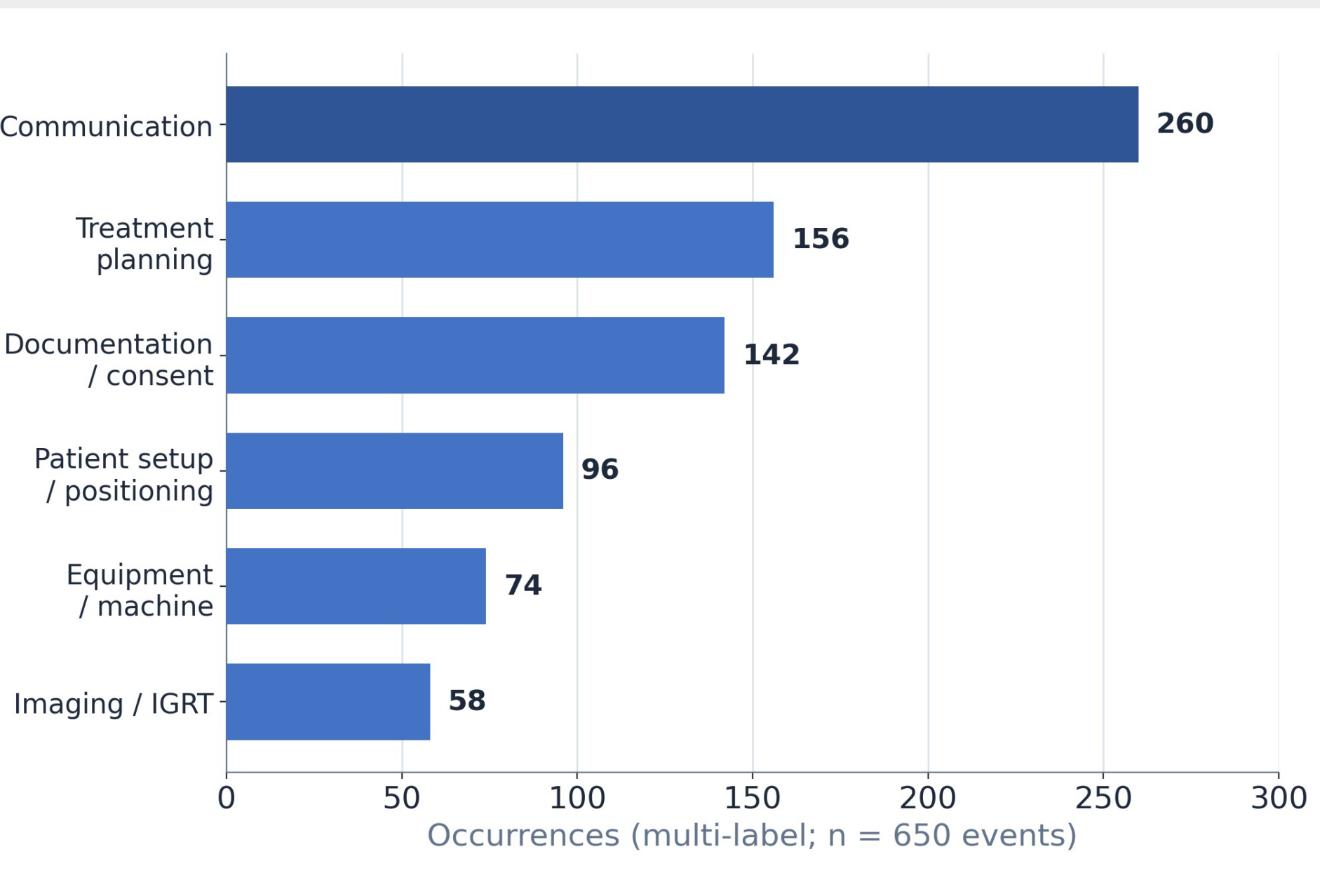


Figure 1. Failure mode distribution (n = 650; multi-label).

RESULTS

Expert-level agreement. Model–reviewer agreement (mean $\kappa = 0.31$) was of the same order as agreement among the reviewers themselves (mean $\kappa = 0.42$). On average the model’s label fell within the experts’ group-accepted set in 67.9% of comparisons. (as shown in Table 1)

Per-dimension, not pooled. Agreement ranged from 96.2% (inclusion) and 85.3% (treatment type) down to 39.0% (severity). Severity was the weakest dimension for the model and humans alike (inter-reviewer $\kappa = 0.22$).

Multi-label artifact. Exact-set scoring penalizes partial overlaps, so fields that legitimately take multiple labels score lower: occurred workflow ($\kappa \approx 0.62$) > discovered workflow (≈ 0.40) > discoverer role (≈ 0.31) > failure mode (≈ 0.24).

Communication dominates. Communication failures were the most common failure mode (260 occurrences, ~40%), followed by treatment planning (156) and documentation/consent (142). (as shown in Figure 1)

A reporting gap. 54.5% of narratives did not state who detected the event; the model abstains rather than guess—pointing to a concrete reporting-template improvement.

Table 1. Per-dimension validation agreement (Model vs. consensus; Cohen’s κ).

Dimension	Within set	Model κ	Inter κ
Treatment type	85.3%	0.46	0.38
Occurred workflow	76.6%	0.62	0.63
Failure mode	79.5%	0.23	0.51
Discoverer role	50.0%	0.31	0.49
Discovered workflow	48.7%	0.40	0.41
Severity	39.0%	0.14	0.22

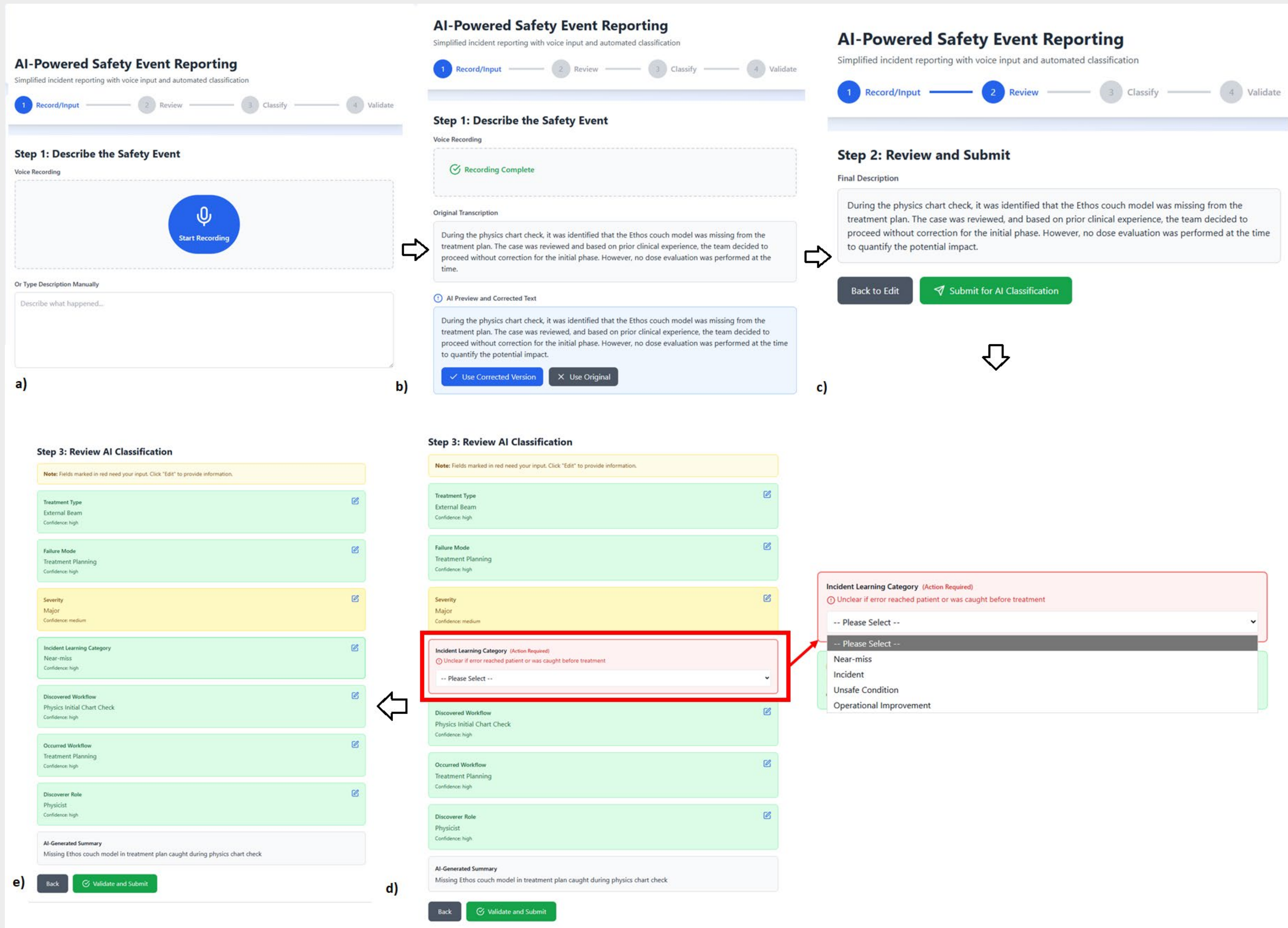


Figure 2. Future AI assisted safety reporting system

CONTACT

Qiongge Li, PhD
Department of Radiation Oncology, Inova
Schar Cancer Institute, Fairfax, VA
Email: Qiongge.li@inova.org

This work is under review of Practical
Radiation Oncology (PRO) since Feb 2026