



# AI PHARMACOVIGILANCE WORKFLOW FOR DETECTING NEUROLOGICAL ADVERSE EVENTS FROM REPORTS, NOTES, AND LITERATURE

Jinwoo Park<sup>1\*</sup>, Minji Kim<sup>1</sup>, Seung Lee<sup>2</sup>

1. *Department of Computational Pharmaceutical Sciences, College of Pharmacy, Seoul National University, Seoul, South Korea.*
2. *Department of AI Drug Engineering, Faculty of Engineering, KAIST, Daejeon, South Korea.*

## ARTICLE INFO

### Received:

12 February 2025

### Received in revised form:

22 May 2025

### Accepted:

25 May 2025

### Available online:

28 June 2025

**Keywords:** Artificial intelligence, Pharmacovigilance, Neurological adverse events, Natural language processing, Spontaneous reports, Electronic health records

## ABSTRACT

Neurological adverse drug events are clinically consequential and may be missed when surveillance depends on a single source of evidence. Under-reporting, delayed recognition, and fragmented documentation make these events especially challenging for conventional pharmacovigilance workflows. Potential neurological safety signals may appear separately in spontaneous reports, electronic health record notes, and published biomedical literature. A unified AI workflow is needed to connect these evidence streams in a timely, traceable, and reviewer-ready manner. This article proposes an AI pharmacovigilance workflow that ingests spontaneous reports, clinical notes, and biomedical literature, then applies transformer-based NLP to identify neurological adverse event mentions. The extracted evidence is fused into a dynamic risk score with source attribution and reviewer-facing explanations. The workflow includes data ingestion and harmonization, multi-source NLP extraction, signal fusion, disproportionality analysis, confounder-aware alerting, and a human-review dashboard. Each module is designed to preserve links between computational outputs and the original source evidence. The proposed workflow would be expected to support earlier recognition of rare neurological safety concerns by cross-validating signals across heterogeneous data streams. It could reduce unsupported alerts by distinguishing consistent multi-source evidence from isolated or ambiguous mentions. A holistic, AI-driven surveillance system could transform pharmacovigilance from a fragmented, reactive process into an integrated, proactive safety intelligence function. Its value would depend on transparent evidence handling, expert oversight, and careful validation in operational settings.

This is an **open-access** article distributed under the terms of the [Creative Commons Attribution-Non Commercial-Share Alike 4.0 License](https://creativecommons.org/licenses/by-nc-sa/4.0/), which allows others to remix, and build upon the work non commercially.

**To Cite This Article:** Park J, Kim M, Lee S. AI Pharmacovigilance Workflow for Detecting Neurological Adverse Events from Reports, Notes, and Literature. *Pharmacophore*. 2025;16(3):42-52. <https://doi.org/10.51847/bGtPljJ5Yy>

## Introduction

Drug-induced neurological adverse events can have major clinical and regulatory significance because they may affect cognition, movement, sensation, autonomic function, or seizure risk in ways that alter long-term quality of life. AI-enabled pharmacovigilance has been proposed as a way to improve signal management across complex safety data, particularly when human review alone cannot keep pace with growing evidence streams [1]. Reviews of machine learning in pharmacovigilance emphasize that safety surveillance increasingly requires methods capable of handling high-volume, heterogeneous, and partially unstructured information [2]. A neurological focus is especially important because the clinical presentation of such events may be subtle, episodic, or described using variable language across reports, notes, and publications [3].

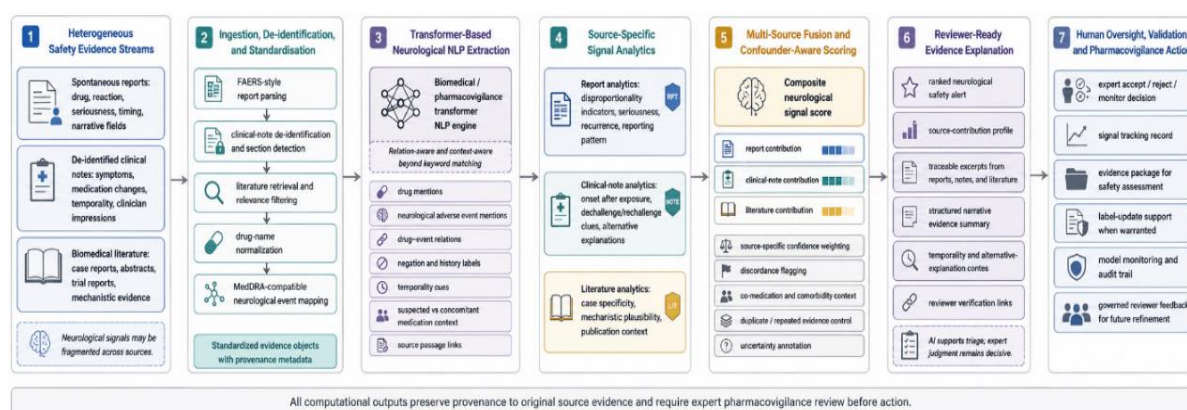
Current pharmacovigilance often operates through parallel but weakly connected processes, including structured adverse event report review, retrospective clinical-text analysis, and periodic literature surveillance. Industry and regulatory discussions of machine learning for drug and vaccine safety describe AI as most useful when it assists signal triage rather than replacing expert safety judgment [4]. Automated workflows therefore need to connect report-based evidence with narrative evidence, while preserving interpretability and auditability for human assessors [5]. Without such integration, a neurological event described in a case narrative, a discharge summary, and a case report may remain fragmented rather than forming a coherent safety signal.

**Corresponding Author:** Jinwoo Park; Department of Computational Pharmaceutical Sciences, College of Pharmacy, Seoul National University, Seoul, South Korea. E-mail: [jinwoo.park@outlook.com](mailto:jinwoo.park@outlook.com)

The maturation of biomedical NLP and transformer-based language models has made it more feasible to extract medication-event relationships across different text sources. BioBERT demonstrated the usefulness of domain-specific pretraining for biomedical text mining [6], while PubMedBERT-style pretraining further supports biomedical language understanding in downstream NLP tasks [7]. Pharmacovigilance-specific transformer frameworks show how adverse drug reaction detection can be formulated as a fine-tuning problem over safety-relevant text [8]. These advances allow a workflow to treat spontaneous-report narratives, clinical notes, and literature passages as complementary textual evidence rather than isolated information silos.

This AIF article proposes an AI-enabled workflow that consistently extracts, links, and prioritizes neurological safety signals from spontaneous reports, clinical narratives, and biomedical literature. Signal detection from spontaneous reporting systems can support broad surveillance [9], while electronic medical record and spontaneous report integration can provide a more clinically grounded view of adverse drug reactions [10]. Clinical-note extraction studies further show that medications, adverse events, and their relations can be identified from unstructured health records using neural and knowledge-aware NLP models [11, 12]. The intended output is not an autonomous regulatory decision, but a unified, evidence-attributed safety view for drug safety teams.

**Figure 1** presents the proposed AI pharmacovigilance workflow for transforming spontaneous reports, de-identified clinical notes, and biomedical literature into source-attributed neurological safety alerts for expert review.



**Figure 1.** AI Pharmacovigilance Workflow for Detecting Neurological Adverse Events from Reports, Notes, and Literature.

### Neurological Adverse Events and the Pharmacovigilance Need

Neurological adverse events include heterogeneous syndromes such as seizures, extrapyramidal symptoms, paraesthesia, cognitive changes, neuropathic symptoms, and serotonin-toxicity-like presentations. Their detection is difficult because symptoms may overlap with comorbid disease, concomitant medication effects, or baseline neurological vulnerability, requiring careful contextual interpretation rather than simple keyword matching. AI-based pharmacovigilance reviews emphasize that machine learning can help prioritize complex safety patterns but must remain embedded in expert-led workflows [2, 3]. For neurological safety, this means the workflow should distinguish a suspected drug-related event from a historical diagnosis, differential diagnosis, or unrelated symptom mention.

### Spontaneous Reporting Systems and Structured Data Mining

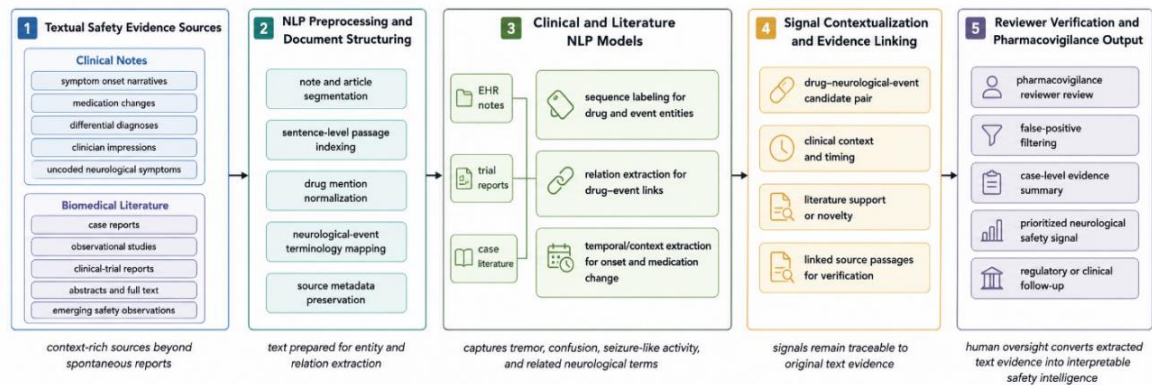
Spontaneous reporting systems such as FAERS and Vigibase are central to postmarketing safety surveillance because they provide broad coverage across medicines, populations, and reporting contexts. Machine learning studies using spontaneous reporting data show that structured case features can be used to support signal detection and triage, while also reflecting known limitations such as under-reporting, missing denominators, and reporting bias [9, 13]. Causality-oriented feature engineering in FAERS illustrates how structured report fields and narrative-derived features can be combined to inform pharmacovigilance assessment [14]. For neurological adverse events, spontaneous reports are valuable as early warning evidence but should be interpreted alongside clinical and literature context.

### Adverse Event Extraction from Clinical Text

Clinical notes provide detailed temporal and contextual information that is often missing from spontaneous reports, including symptom onset, medication changes, differential diagnoses, and clinician impressions. Deep learning approaches have been designed to extract adverse drug event information directly from electronic health record notes [12], and shared-task work has formalized medication and adverse drug event extraction as a benchmark problem in clinical NLP [15]. Later n2c2-related methods further show that entities and relations can be extracted from EHR text using ensembles, sequence labeling, and relation extraction models [16, 17]. For neurological pharmacovigilance, these methods are useful because they can capture narrative descriptions such as tremor, confusion, or seizure-like activity even when they are not coded as adverse events.

### Literature Mining for Drug Safety

Biomedical literature remains an important source for rare or emerging neurological adverse reactions, particularly when published case reports and observational studies precede routine signal recognition. Neural approaches to adverse drug event discovery from biomedical literature demonstrate how large-scale literature processing can contribute to safety surveillance [18]. Text mining of adverse events in clinical-trial reports also illustrates how deep learning can support extraction from formal biomedical documents rather than only clinical notes or spontaneous reports [19]. A literature module in the proposed workflow would therefore screen abstracts and full text for drug-neurological-event relationships while maintaining links to source passages for reviewer verification. **Figure 2** illustrates how clinical notes and biomedical literature can be integrated into a source-linked NLP workflow for detecting, contextualizing, and verifying neurological adverse drug event signals.



**Figure 2.** Text-Source Integration Workflow for Neurological Pharmacovigilance Signal Detection

### Multi-Source Pharmacovigilance and AI

Multi-source pharmacovigilance aims to reduce the weaknesses of any single data stream by combining signals from reports, clinical records, social media, and literature. Work combining electronic medical records with spontaneous reports shows that cross-source analysis can enrich signal detection beyond one surveillance system alone [10]. Other studies combining social media with FAERS suggest that heterogeneous real-world text sources can complement formal reporting channels when interpreted carefully [20]. The gap addressed here is a dedicated neurological safety workflow that aligns report data, EHR narratives, and literature evidence into a unified, source-attributed signal assessment process.

### Workflow Architecture Overview

#### High-Level Design

The proposed workflow continuously or periodically ingests new spontaneous reports, newly generated clinical notes, and new biomedical literature records, then routes each source through a specialized preprocessing and NLP pipeline. Extracted drugs, neurological event mentions, temporality cues, and context labels are normalized to standardized drug names and MedDRA-compatible adverse event concepts, drawing on approaches that map narrative adverse reaction descriptions into MedDRA terminology [21]. A multi-source fusion layer then compares report-based patterns, clinical-note evidence, and literature evidence to generate a neurological safety score. This design follows the broader view that AI in signal management should support prioritization, traceability, and expert review rather than produce unexamined automated conclusions [1].

#### Core Data Input and Processing Modules

The FAERS module would process structured case fields and available narrative text, the clinical-note module would identify medication and neurological symptom mentions in de-identified EHR text, and the literature module would retrieve and classify biomedical publications relevant to drug-event pairs. Clinical NLP studies show that medication and adverse event extraction often requires joint recognition of entities and relations rather than separate isolated tagging [11, 22]. Relation-focused methods in clinical notes also support the distinction between a medication causing an adverse event and a medication merely appearing in the same record [23, 24]. All three modules would output standardized, provenance-preserving evidence objects that can be compared across sources.

#### Design Principles

The workflow is modular so that source-specific pipelines can be updated without redesigning the entire system, and scalable so that new drugs, neurological concepts, or literature feeds can be added over time. It is evidence-based because every alert should remain traceable to the original report field, note sentence, or literature passage that contributed to the score. Human-in-the-loop review is a central design requirement, consistent with industry perspectives that machine learning in pharmacovigilance should improve reviewer efficiency while preserving expert accountability [4, 5]. Privacy and governance controls are also required because clinical-note processing depends on de-identification, access control, and careful handling of sensitive health information.

### Data Ingestion and Standardisation

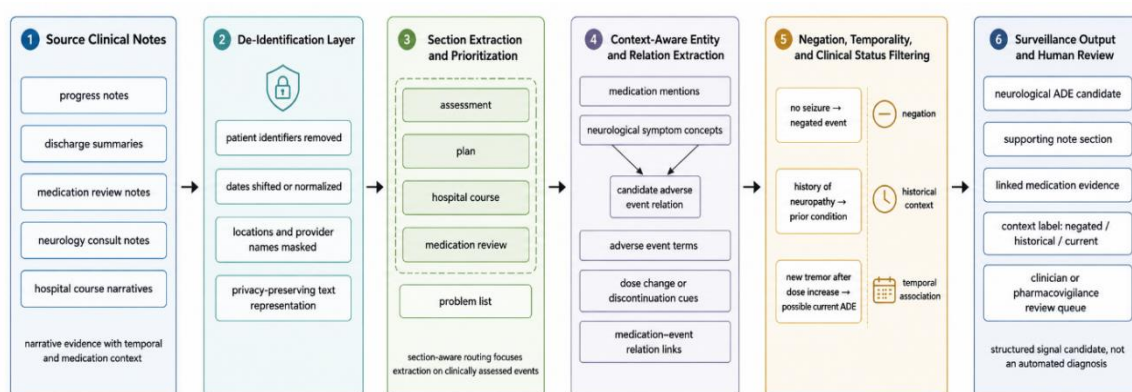
#### FAERS Ingestion and Normalization

The FAERS ingestion module would parse structured case data, drug names, reaction terms, reporting dates, seriousness indicators, and available narrative fields into a normalized safety-event schema. Machine learning work on spontaneous reporting systems shows that structured report features can support signal detection when transformed into analysis-ready representations [9]. Causality assessment research using FAERS further illustrates the importance of feature engineering for report-level interpretation, including variables that reflect timing, seriousness, co-medication, and report context [14]. Drug and reaction terms would be standardized to controlled vocabularies so that variant expressions of neurological events can be aligned before signal fusion.

#### Clinical Note De-identification and Section Extraction

The clinical-note pipeline would operate on de-identified text and use document structure to prioritize sections where adverse events are most likely to be clinically assessed, such as assessment, plan, hospital course, and medication review. End-to-end deep learning models for adverse event extraction from EHR notes demonstrate how medication mentions and adverse event concepts can be identified in narrative clinical text [12]. Shared-task systems for EHR adverse event extraction further show the need to handle abbreviations, discontinuous mentions, negation, and medication-event relations [15, 16]. For neurological surveillance, this module would emphasize context-aware extraction so that “no seizure,” “history of neuropathy,” and “new tremor after dose increase” are treated differently. Clinical-note processing should also include explicit privacy, structural, and contextual safeguards before extracted neurological events are allowed to contribute to signal generation. Benchmark de-identification tasks show that removing protected health information from longitudinal clinical narratives is a specialized NLP problem rather than a simple text-cleaning step [25], and earlier de-identification challenges established the need for systematic evaluation of PHI recognition across names, dates, locations, identifiers, and institutional references [26]. Neural de-identification methods further support the use of sequence-based models for masking sensitive clinical text while preserving downstream linguistic structure [27]. After privacy protection, section identification becomes important because clinical meaning changes depending on whether a symptom appears in the assessment, medication review, past history, or differential diagnosis section [28]. Context algorithms such as ConText, DEEPEN, and ContextD further demonstrate that negation, temporality, experienter, and historical status must be modeled directly so that neurological mentions are not misclassified as current drug-related adverse events [29–31].

**Figure 3** illustrates how de-identified clinical notes can be transformed into section-aware, context-sensitive neurological adverse event signals by combining document segmentation, medication-event extraction, negation handling, temporality assessment, and clinician-review outputs.



**Figure 3.** Clinical-Note Processing Architecture for Context-Aware Neurological Adverse Event Surveillance

#### Literature Retrieval and Relevance Filtering

The literature module would run scheduled searches for drug-neurological-event combinations, classify abstracts for relevance, and process eligible full-text passages through the same conceptual extraction layer used for reports and notes. Literature-focused adverse event discovery has shown that biomedical publications can be processed with neural methods to surface safety-relevant drug-event associations [18]. Deep learning approaches for extracting adverse events from clinical trial text further support the feasibility of automated safety-oriented document screening [19]. Transformer-based biomedical language models would help identify whether a publication describes a suspected adverse reaction, background pharmacology, confounding disease, or unrelated neurological outcome [6, 7].

**Table 1** defines the evidence architecture required to convert heterogeneous pharmacovigilance sources into standardized, provenance-preserving neurological safety evidence objects.

**Table 1.** Evidence Architecture for Multi-Source Neurological Adverse Event Detection

| Evidence stream                          | Primary neurological safety value                                                                         | Key data elements to extract                                                                                                                      | NLP / analytic task                                                                                                                               | Source-specific limitations                                                                                                 | Contribution to composite signal score                                                                    | Reviewer-facing traceability requirement                                                                     |
|------------------------------------------|-----------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------|
| <b>Spontaneous adverse event reports</b> | Broad postmarketing coverage and early warning for rare or unexpected neurological events                 | Suspected drug, concomitant medicines, reaction terms, seriousness, reporter type, onset timing, narrative description, outcome                   | Drug–event extraction, MedDRA-compatible event normalization, seriousness classification, duplicate detection, disproportionality support         | Under-reporting, stimulated reporting, missing denominator, incomplete timing, variable narrative detail                    | Provides report-based signal strength, reporting pattern, seriousness burden, and recurrence across cases | Link every score contribution to original case fields, reaction terms, and narrative excerpts                |
| <b>De-identified clinical notes</b>      | Rich clinical context for temporality, medication changes, differential diagnosis, and symptom evolution  | Medication exposure, dose change, neurological symptom, onset phrase, negation, history, clinician impression, comorbidity, competing explanation | Context-aware entity recognition, relation extraction, negation detection, temporality classification, suspected versus historical event labeling | Local data access constraints, documentation variation, incomplete medication reconciliation, institutional population bias | Provides clinical plausibility, temporal coherence, dechallenge/rechallenge clues, and confounder context | Link extracted evidence to note section, sentence-level excerpt, document date, and de-identification status |
| <b>Biomedical literature</b>             | External scientific context for rare, emerging, or mechanistically plausible neurological safety concerns | Drug name, neurological event, publication type, case specificity, trial context, mechanism, comparator, author conclusion                        | Literature retrieval, relevance filtering, drug–event relation extraction, case-versus-background classification                                  | Publication bias, delayed indexing, variable full-text availability, hypothesis-generating evidence quality                 | Provides external corroboration, mechanistic support, early case evidence, and signal plausibility        | Link evidence to abstract/full-text passage, publication metadata, and classified relevance rationale        |
| <b>Standardized terminology layer</b>    | Harmonizes heterogeneous language across sources into comparable safety concepts                          | Drug vocabularies, MedDRA-compatible neurological terms, synonyms, abbreviations, spelling variants, severity modifiers                           | Concept normalization, synonym expansion, terminology mapping, hierarchy-aware grouping                                                           | Over-mapping may collapse clinically meaningful distinctions; under-mapping may fragment related events                     | Enables cross-source aggregation of related neurological events without losing source wording             | Preserve both original text and mapped concept so reviewers can inspect translation fidelity                 |
| <b>Cross-source evidence object</b>      | Converts raw evidence fragments into comparable, provenance-preserving units                              | Source type, drug, event, relation type, temporality, confidence score, excerpt, metadata, uncertainty label                                      | Evidence object construction, provenance tagging, confidence assignment                                                                           | Evidence objects may inherit upstream extraction errors or missing metadata                                                 | Serves as the unit of fusion for source weighting, discordance assessment, and alert generation           | Each object must remain traceable to a report field, note sentence, or literature passage                    |
| <b>Confounder and context layer</b>      | Reduces unsupported neurological alerts by identifying                                                    | Comorbid neurological disease, concomitant neuroactive                                                                                            | Confounder recognition, context classification, alternative-                                                                                      | Confounders may be under-documented or incompletely structured                                                              | Modifies signal priority by distinguishing plausible drug-related events from ambiguous                   | Show which confounders were detected and whether they                                                        |

|                                        | alternative<br>explanations                                                        | medicines,<br>indication,<br>baseline<br>symptom<br>history,<br>disease<br>progression,<br>overdose<br>context                                     | explanation<br>flagging                                                                        | or competing<br>explanations                                                                                           | lowered,<br>raised, or<br>qualified the<br>alert                                                                    |
|----------------------------------------|------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------|
| <b>Human-<br/>review<br/>interface</b> | Converts<br>computational<br>evidence into<br>assessable<br>safety<br>intelligence | Ranked alert,<br>source-<br>contribution<br>profile,<br>excerpts,<br>evidence<br>narrative,<br>uncertainty<br>notes,<br>reviewer<br>action options | Explainable alert<br>presentation,<br>evidence<br>summarization,<br>triage workflow<br>support | Poor interface<br>design may<br>increase<br>reviewer<br>burden or<br>encourage<br>over-reliance<br>on model<br>outputs | Determines whether the<br>signal becomes<br>accepted, rejected,<br>monitored, or escalated                          |
|                                        |                                                                                    |                                                                                                                                                    |                                                                                                |                                                                                                                        | Maintain<br>audit trail of<br>reviewer<br>decision,<br>rationale,<br>evidence<br>inspected,<br>and model<br>version |

### *Ai Detection and Signal Generation Engine*

#### *Multi-Source NLP Extraction*

The NLP engine would use a fine-tuned biomedical transformer to extract neurological adverse event mentions such as akathisia, status epilepticus, paraesthesia, encephalopathy, dystonia, and cognitive impairment from spontaneous-report narratives, clinical notes, and literature passages. Pharmacovigilance transformer work shows how BERT-style models can be adapted to adverse drug reaction detection tasks [8], while biomedical language model pretraining supports domain-specific recognition of biomedical entities and relations [6]. Medication and adverse event extraction systems developed for clinical notes demonstrate the importance of modeling entities and relations together rather than assuming that co-occurrence implies causality [22, 32]. A separate context classifier would label the drug as suspected, concomitant, historical, contraindicated, or merely mentioned, preserving the uncertainty needed for safety review.

#### *Disproportionality and Temporal Analysis*

For spontaneous reports, the workflow would compute established signal-detection indicators conceptually, while avoiding treating any statistical association as sufficient evidence of causality. For clinical notes, it would evaluate whether the neurological event appears after treatment initiation, improves after withdrawal, recurs after re-exposure, or is better explained by disease progression or another medication. Studies integrating electronic medical records with spontaneous reports show why temporal and source-specific evidence should be interpreted together rather than as independent fragments [10]. Literature-derived signals would be weighted by publication context, clinical plausibility, and the specificity of the reported drug-event relationship, consistent with literature-mining approaches for adverse event discovery [18].

#### *Multi-Source Signal Fusion and Scoring*

The fusion core would combine report-based disproportionality evidence, clinical-note temporal evidence, and literature-based mechanistic or case-level evidence into a composite neurological signal score. Machine learning models such as random forests or gradient-boosted classifiers are conceptually appropriate for combining heterogeneous features, as shown by pharmacovigilance work that uses engineered report features and multiple analytical methods for adverse reaction signal detection [13, 14]. Relation extraction and ensemble methods from EHR adverse event tasks also support the idea that multiple model views can improve robustness when interpreting complex clinical text [17, 33]. The score would be accompanied by a transparent source-contribution profile so that reviewers can see whether the alert is driven mainly by FAERS, EHR notes, literature, or convergent evidence across all three.

### *Evidence Fusion and Causal Assessment*

#### *Source-Specific Confidence Weighting*

The evidence fusion layer would assign different confidence weights to spontaneous reports, clinical notes, and literature according to their evidentiary role, internal consistency, and clinical specificity. A literature signal from an isolated case report would be treated as hypothesis-generating, whereas a cluster of spontaneous reports with consistent timing and compatible clinical narratives would receive stronger signal weight [9, 14]. Clinical-note evidence would be weighted for temporal coherence, medication-event linkage, and documentation context, drawing on adverse event extraction approaches that identify both entities and relations in EHR text [11, 34]. Literature evidence could receive additional priority when it appears before structured reporting signals, because publication-based safety observations may precede formal signal escalation [18].

#### *Handling Conflicting Evidence*

Conflicting evidence would not be automatically suppressed, because disagreement between sources may itself be informative for pharmacovigilance review. For example, a neurological event may appear frequently in FAERS but be absent from a local EHR corpus because of population differences, prescribing patterns, coding practices, or limited clinical-note availability. Multi-source studies combining spontaneous reports with other real-world data sources show that heterogeneous evidence streams should be interpreted as complementary rather than interchangeable [10, 20]. The workflow would therefore flag discordance for human assessors, preserving source-level explanations instead of collapsing uncertainty into a single opaque score.

#### *Narrative Evidence Summary*

A narrative evidence-summary module would use a large language model to generate a concise, reviewer-facing synthesis of the safety concern, including suspected drug, neurological event, temporality, alternative explanations, and supporting source excerpts. Biomedical transformer models provide a foundation for extracting and organizing safety-relevant text [6, 7], while pharmacovigilance-specific transformer approaches show how adverse reaction detection can be adapted to drug safety tasks [8]. The narrative would not replace expert assessment; instead, it would organize evidence into a structured safety story that can be checked against the underlying reports, notes, and publications. The summary should remain traceable to source passages so that reviewers can verify whether the model has accurately represented the evidence.

#### *Human-In-The-Loop Review and Signal Validation*

##### *Alert Dashboard and Triage*

The alert dashboard would present neurological signals in a priority queue, with each entry showing the composite signal score, the source-contribution profile, a confidence indicator, and a concise evidence narrative. Industry perspectives on machine learning in drug and vaccine safety emphasize that AI tools are most appropriate when they support expert triage and review rather than act as independent decision-makers [4, 5]. The reviewer would be able to inspect the original FAERS report fields, clinical-note excerpts, and literature passages before accepting, rejecting, or requesting additional information. This design makes the workflow a safety-intelligence assistant rather than a substitute for pharmacovigilance judgment.

##### *Feedback Loop and Model Refinement*

Reviewer decisions would be logged as structured feedback and used to refine the signal fusion model over time. This feedback could help the workflow learn which evidence patterns pharmacovigilance scientists consider persuasive, ambiguous, or unsupported, while still requiring governance controls to avoid reinforcing historical review bias. Studies of adverse event extraction from clinical notes show that model behavior depends heavily on task definitions, annotation standards, and relation-label quality [15, 16, 22]. Continuous refinement should therefore be managed through version control, model monitoring, and periodic expert review rather than uncontrolled automatic learning.

#### *Integration Into Pharmacovigilance Operations*

##### *Periodic Signal Management Process*

The workflow would be integrated into periodic signal management by running on a scheduled basis and producing structured outputs for company or agency safety systems. AI signal-management discussions emphasize the need for automation that supports real-world pharmacovigilance operations while maintaining traceability, auditability, and human accountability [1, 4]. Accepted neurological signals could automatically create or update internal tracking records, while unresolved signals could remain in a monitored queue for future evidence accumulation. The operational goal would be to reduce manual fragmentation across reports, notes, and literature without weakening the formal signal assessment process.

##### *Label Update Support and Regulatory Reporting*

When a strong neurological signal emerges, the workflow would assemble an evidence package containing source excerpts, mapped terminology, temporality assessment, literature context, and a reviewer-edited narrative. Automated mapping from narrative adverse reaction descriptions to MedDRA concepts supports the creation of standardized safety summaries that can be compared across cases and sources [21]. Systematic review work on extracting adverse drug events from clinical notes highlights that documentation quality and extraction reliability remain central concerns when using EHR text in safety workflows [35]. The evidence package could support internal safety assessment and regulatory discussion, but final label-related decisions would remain grounded in expert review and applicable regulatory standards.

#### *Evaluation Strategy*

**Table 2** provides a governance and validation framework for determining whether the proposed AI neurological pharmacovigilance workflow is operationally reliable, auditable, and suitable for expert-led signal management.

**Table 2.** Governance, Validation, and Operational Readiness Framework for AI Neurological Pharmacovigilance

| Operational domain | Core risk addressed | Required design safeguard | Evaluation question | Practical metric or review criterion | Pharmacovigilance decision-use implication |
|--------------------|---------------------|---------------------------|---------------------|--------------------------------------|--------------------------------------------|
|--------------------|---------------------|---------------------------|---------------------|--------------------------------------|--------------------------------------------|

|                                      |                                                                                                          |                                                                                                                                    |                                                                                                                                                |                                                                                                                                                                     |                                                                                                            |
|--------------------------------------|----------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------|
| <b>Neurological event extraction</b> | Clinically important events may be missed because symptoms are described indirectly or variably          | Use domain-tuned transformer extraction with neurological synonym handling, negation detection, and section-aware clinical parsing | Does the system identify neurological adverse event mentions across reports, notes, and literature without relying on simple keyword matching? | Event-level precision and recall; error review by event class; performance on subtle terms such as tremor, confusion, paraesthesia, dystonia, seizure-like activity | Determines whether the workflow can support early detection of rare or subtle neurological safety concerns |
| <b>Drug–event relation modeling</b>  | Co-occurrence may be mistaken for causality or suspected association                                     | Model medication–event relations, suspected versus concomitant context, temporality, and alternative explanations                  | Does the system distinguish suspected adverse reactions from historical diagnoses, background disease, or unrelated medication mentions?       | Relation-level accuracy; temporality classification accuracy; proportion of alerts with documented suspected-drug linkage                                           | Reduces unsupported alerts and improves reviewer confidence in prioritized signals                         |
| <b>Source attribution</b>            | AI-generated alerts may become difficult to audit or verify                                              | Preserve original report fields, clinical-note excerpts, literature passages, and source metadata                                  | Can reviewers trace every alert component back to the evidence that generated it?                                                              | Percentage of score components with inspectable source links; reviewer-rated evidence traceability                                                                  | Makes the workflow suitable for expert-led pharmacovigilance and regulatory discussion                     |
| <b>Multi-source fusion</b>           | A single noisy source may dominate the score and generate misleading priority                            | Apply source-specific weighting, discordance labeling, and contribution profiles rather than opaque averaging                      | Does the composite score show whether evidence is convergent, isolated, or conflicting across sources?                                         | Source-contribution profile completeness; concordance/discordance flag accuracy; calibration of high-priority alerts                                                | Helps safety teams distinguish robust multi-source signals from weak single-source alerts                  |
| <b>Confounder-aware alerting</b>     | Neurological symptoms may reflect comorbidity, disease progression, overdose, or concomitant medications | Include comorbidity, co-medication, indication, baseline symptom, and documentation-context features                               | Does the system identify plausible alternative explanations before escalating a neurological signal?                                           | Reviewer agreement with confounder flags; false-alert reduction after context adjustment; audit of suppressed or downgraded alerts                                  | Supports more clinically credible signal triage and avoids overstating causality                           |
| <b>Narrative evidence summary</b>    | Large language model summaries may omit uncertainty or overstate causal interpretation                   | Require extractive evidence grounding, uncertainty labels, reviewer verification, and source-linked summaries                      | Does the narrative accurately represent the underlying evidence without adding unsupported causal claims?                                      | Summary faithfulness review; unsupported statement rate; reviewer correction frequency                                                                              | Allows efficient review while maintaining accountability for final safety interpretation                   |
| <b>Human-in-the-loop review</b>      | Automation may be misinterpreted as replacing                                                            | Require expert accept/reject/monitor actions, reviewer                                                                             | Are safety professionals able to                                                                                                               | Reviewer usability rating; time-to-triage; proportion of alerts                                                                                                     | Positions the system as decision support rather than                                                       |

|                                                | pharmacovigilance judgment                                                                      | rationale capture, and override mechanisms                                                                | inspect, challenge, and revise AI-prioritized signals?                                                                       | modified after expert review                                                                               | autonomous signal adjudication                                                                                |
|------------------------------------------------|-------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------|
| <b>Model monitoring and update control</b>     | Model drift, terminology changes, and evolving prescribing patterns may degrade performance     | Maintain model versioning, periodic validation, drift monitoring, and locked update cycles                | Does workflow performance remain stable as new reports, notes, literature, and drug-event terminology accumulate?            | Drift indicators; periodic benchmark performance; model-version audit trail                                | Enables long-term operational reliability and defensible safety governance                                    |
| <b>Bias and representativeness</b>             | Signals may be under-detected in populations or settings with poorer documentation or reporting | Evaluate performance across demographic groups, care settings, drug classes, and documentation patterns   | Does the workflow detect neurological safety evidence equitably across populations and data environments?                    | Stratified performance review; missingness analysis; subgroup alert-rate comparison                        | Prevents the safety system from amplifying inequities in reporting, care access, or documentation             |
| <b>Regulatory and organizational readiness</b> | AI outputs may not be accepted without transparent validation and documented oversight          | Produce audit-ready evidence packages, validation documentation, reviewer logs, and governance procedures | Can accepted signals be converted into defensible evidence packages for internal safety assessment or regulatory discussion? | Completeness of evidence package; auditability of decision pathway; documentation of final expert judgment | Supports responsible integration into periodic signal management, label-review support, and safety governance |

#### *Retrospective Detection Performance*

Retrospective evaluation would test whether the workflow could identify known neurological safety concerns from historical spontaneous reports, clinical notes, and literature without reporting artificial performance claims in this conceptual article. Prior studies combining electronic medical records with spontaneous reports provide a useful model for comparing multi-source signal detection against single-source approaches [10]. Spontaneous-report machine learning methods and report-level causality feature engineering also provide design precedents for evaluating structured report evidence in pharmacovigilance settings [9, 14]. The evaluation should focus on whether the workflow retrieves clinically plausible evidence, preserves source attribution, and supports expert interpretation.

#### *Prospective Signal Monitoring and Alert Burden*

Prospective monitoring would evaluate how the workflow behaves when new reports, notes, and publications are added over time, with particular attention to alert burden and reviewer usability. AI pharmacovigilance reviews emphasize that practical value depends not only on detection capability but also on whether alerts can be prioritized, explained, and acted on by safety teams [2, 3]. Methods that combine heterogeneous sources, such as FAERS with social media or EHR data, show why prospective monitoring must assess both convergent evidence and noisy single-source signals [10, 20]. The workflow should therefore be evaluated for whether it produces manageable, interpretable, and clinically meaningful neurological alerts.

#### *User Satisfaction and Workflow Efficiency*

User evaluation would assess whether pharmacovigilance professionals find the dashboard, source attribution, evidence summaries, and signal scores useful in routine safety review. Work on predicting drug-side effect relationships from biomedical BERT models illustrates how language-model representations may support safety reasoning, but such outputs must be made understandable to human users [36]. Clinical NLP studies also show that extracted adverse event information is only useful when the relation between medication, symptom, context, and temporality is clear enough for review [24, 32, 33]. The evaluation should therefore examine trust, interpretability, and workflow fit rather than only model-centered technical performance.

#### *Limitations*

*Data Availability and Quality*

The proposed workflow depends on access to timely, representative, and sufficiently detailed data across spontaneous reports, clinical notes, and literature. Clinical notes may be unavailable for some drugs, institutions, or populations, and even available notes may incompletely document neurological symptoms, medication timing, or alternative explanations. Systematic review evidence on clinical-note adverse event extraction shows that NLP performance and usefulness are shaped by documentation practices, annotation definitions, and local data quality [35]. These limitations mean the workflow should be viewed as decision support rather than a comprehensive source of neurological safety truth.

*Regulatory Acceptance and Interpretability*

Before a multi-source AI workflow could be used as a primary signal-detection method in regulatory submissions, robust validation, governance, and explainability expectations would need to be defined. AI pharmacovigilance literature repeatedly emphasizes that automation must remain auditable, interpretable, and compatible with expert-led signal management [1, 4, 5]. Transformer-based models may improve extraction from complex biomedical language, but their internal representations do not automatically provide the causal explanations needed for regulatory confidence [6-8]. The workflow would therefore require transparent source attribution, documented model updates, bias monitoring, and reviewer override mechanisms.

**Conclusion**

An AI pharmacovigilance workflow for neurological adverse event detection could connect spontaneous reporting systems, electronic health record notes, and biomedical literature into a single evidence-oriented process. By extracting neurological event mentions, linking them to suspected medicines, and preserving source provenance, the system could help safety teams move from isolated signal fragments to a more coherent safety view.

The major strength of the proposed workflow is its integrated evidence picture. Source-specific weighting, narrative summarization, and continuous human feedback would allow the system to support expert review while maintaining transparency about where each signal originated and why it was prioritized.

Important challenges remain in data access, clinical-note quality, regulatory acceptance, and equitable signal detection across populations, medicines, and care settings. The workflow would also need careful governance to ensure that AI-generated summaries remain faithful to source evidence and do not overstate causality.

Collaborative pilots involving pharmaceutical companies, health systems, regulators, and informatics researchers would be an appropriate next step. Such pilots could test whether the workflow can be implemented responsibly, evaluated transparently, and scaled into a practical safety intelligence function.

**Acknowledgments:** None

**Conflict of interest:** None

**Financial support:** None

**Ethics statement:** None

**References**

1. Warner J, Prada Jardim A, Albera C. Artificial intelligence: applications in pharmacovigilance signal management. *Pharm Med.* 2025;39(3):183-98.
2. Salas M, Petracek J, Yalamanchili P, Aimer O, Kasthuril D, Dhingra S, et al. The use of artificial intelligence in pharmacovigilance: a systematic review of the literature. *Pharm Med.* 2022;36(5):295-306.
3. Kompa B, Hakim JB, Palepu A, Kompa KG, Smith M, Bain PA, et al. Artificial intelligence based on machine learning in pharmacovigilance: a scoping review. *Drug Saf.* 2022;45(5):477-91.
4. Bate A, Hobbiger SF. Artificial intelligence, real-world automation and the safety of medicines. *Drug Saf.* 2021;44(2):125-32.
5. Painter JL, Kassekert R, Bate A. An industry perspective on the use of machine learning in drug and vaccine safety. *Front Drug Saf Regul.* 2023;3:1110498.
6. Lee J, Yoon W, Kim S, Kim D, Kim S, So CH, et al. BioBERT: a pre-trained biomedical language representation model for biomedical text mining. *Bioinformatics.* 2020;36(4):1234-40.
7. Gu Y, Tinn R, Cheng H, Lucas M, Usuyama N, Liu X, et al. Domain-specific language model pretraining for biomedical natural language processing. *ACM Trans Comput Healthc.* 2021;3(1):1-23.
8. Hussain S, Afzal H, Saeed R, Iltaf N, Umair MY. Pharmacovigilance with transformers: a framework to detect adverse drug reactions using BERT fine-tuned with FARM. *Comput Math Methods Med.* 2021;2021(1):5589829.
9. Bae JH, Baek YH, Lee JE, Song I, Lee JH, Shin JY. Machine learning for detection of safety signals from spontaneous reporting system data: example of nivolumab and docetaxel. *Front Pharmacol.* 2021;11:602365.

10. Wang L, Rastegar-Mojarad M, Ji Z, Liu S, Liu K, Moon S, et al. Detecting pharmacovigilance signals combining electronic medical records with spontaneous reports: a case study of conventional disease-modifying antirheumatic drugs for rheumatoid arthritis. *Front Pharmacol*. 2018;9:875.
11. Dandala B, Joopudi V, Tsou CH, Liang JJ, Suryanarayanan P. Extraction of information related to drug safety surveillance from electronic health record notes: joint modeling of entities and relations using knowledge-aware neural attentive models. *JMIR Med Inform*. 2020;8(7):e18417.
12. Li F, Liu W, Yu H. Extraction of information related to adverse drug events from electronic health record notes: design of an end-to-end model based on deep learning. *JMIR Med Inform*. 2018;6(4):e12159.
13. Jeong E, Park N, Choi Y, Park RW, Yoon D. Machine learning model combining features from algorithms with different analytical methodologies to detect laboratory-event-related adverse drug reaction signals. *PLoS One*. 2018;13(11):e0207749.
14. Kreimeyer K, Dang O, Spiker J, Muñoz MA, Rosner G, Ball R, et al. Feature engineering and machine learning for causality assessment in pharmacovigilance: lessons learned from application to the FDA Adverse Event Reporting System. *Comput Biol Med*. 2021;135:104517.
15. Henry S, Buchan K, Filannino M, Stubbs A, Uzuner O. 2018 n2c2 shared task on adverse drug events and medication extraction in electronic health records. *J Am Med Inform Assoc*. 2020;27(1):3-12.
16. Wei Q, Ji Z, Li Z, Du J, Wang J, Xu J, et al. A study of deep learning approaches for medication and adverse drug event extraction from clinical text. *J Am Med Inform Assoc*. 2020;27(1):13-21.
17. Ju M, Nguyen NT, Miwa M, Ananiadou S. An ensemble of neural models for nested adverse drug events and medication extraction with subwords. *J Am Med Inform Assoc*. 2020;27(1):22-30.
18. P Tafti A, Badger J, LaRose E, Shirzadi E, Mahnke A, Mayer J, et al. Adverse drug event discovery using biomedical literature: a big data neural network adventure. *JMIR Med Inform*. 2017;5(4):e51.
19. Chopard D, Treder MS, Corcoran P, Ahmed N, Johnson C, Busse M, et al. Text mining of adverse events in clinical trials: deep learning approach. *JMIR Med Inform*. 2021;9(12):e28632.
20. Li Y, Yepes AJ, Xiao C. Combining social media and FDA adverse event reporting system to detect adverse drug reactions. *Drug Saf*. 2020;43(9):893.
21. Combi C, Zorzi M, Pozzani G, Moretti U, Arzenton E. From narrative descriptions to MedDRA: automagically encoding adverse drug reactions. *J Biomed Inform*. 2018;84:184-99.
22. Christopoulou F, Tran TT, Sahu SK, Miwa M, Ananiadou S. Adverse drug events and medication relation extraction in electronic health records with ensemble deep learning methods. *J Am Med Inform Assoc*. 2020;27(1):39-46.
23. Luo Y. Recurrent neural networks for classifying relations in clinical notes. *J Biomed Inform*. 2017;72:85-95.
24. Yang X, Bian J, Fang R, Bjarnadottir RI, Hogan WR, Wu Y. Identifying relations of medications with adverse drug events using recurrent convolutional neural networks and gradient boosting. *J Am Med Inform Assoc*. 2020;27(1):65-72.
25. Stubbs A, Filannino M, Uzuner Ö. Automated systems for the de-identification of longitudinal clinical narratives: overview of 2014 i2b2/UTHealth shared task Track 1. *J Biomed Inform*. 2015;58 Suppl:S11-9.
26. Uzuner Ö, Luo Y, Szolovits P. Evaluating the state-of-the-art in automatic de-identification. *J Am Med Inform Assoc*. 2007;14(5):550-63.
27. Dernoncourt F, Lee JY, Uzuner Ö, Szolovits P. De-identification of patient notes with recurrent neural networks. *J Am Med Inform Assoc*. 2017;24(3):596-606.
28. Denny JC, Spickard A, Johnson KB, Peterson NB, Peterson JF, Miller RA. Evaluation of a method to identify and categorize section headers in clinical documents. *J Am Med Inform Assoc*. 2009;16(6):806-15.
29. Harkema H, Dowling JN, Thornblade T, Chapman WW. ConText: an algorithm for determining negation, experienter, and temporal status from clinical reports. *J Biomed Inform*. 2009;42(5):839-51.
30. Mehrabi S, Krishnan A, Sohn S, Roch AM, Schmidt H, Kesterson J, et al. DEEPEN: a negation detection system for clinical text incorporating dependency relation into NegEx. *J Biomed Inform*. 2015;54:213-9.
31. Afzal Z, Schuemie MJ, van Blijderveen JC, Sen EF, Sturkenboom MCJM, Kors JA. ContextD: an algorithm to identify contextual properties of medical terms in a Dutch clinical corpus. *BMC Bioinformatics*. 2014;15:373.
32. Dai HJ, Su CH, Wu CS. Adverse drug event and medication extraction in electronic health records via a cascading architecture with different sequence labeling models and word embeddings. *J Am Med Inform Assoc*. 2020;27(1):47-55.
33. Chen L, Gu Y, Ji X, Sun Z, Li H, Gao Y, et al. Extracting medications and associated adverse drug events using a natural language processing system combining knowledge base and deep learning. *J Am Med Inform Assoc*. 2020;27(1):56-64.
34. Narayanan S, Mannam K, Achan P, Ramesh MV, Rangan PV, Rajan SP. A contextual multi-task neural approach to medication and adverse events identification from clinical text. *J Biomed Inform*. 2022;125:103960.
35. Modi S, Kasmiran KA, Sharef NM, Sharum MY. Extracting adverse drug events from clinical notes: a systematic review of approaches used. *J Biomed Inform*. 2024;151:104603.
36. Jeon W, Park M, An D, Nam W, Shin JY, Lee S, et al. Predicting drug-side effect relationships from parametric knowledge embedded in biomedical BERT models: methodological study with a natural language processing approach. *JMIR Med Inform*. 2025;13:e67513.