

LLM Tool

A transparent hybrid pipeline for computational policy analysis

An auditable LLM and transformer workflow that turns heterogeneous policy materials into measurable structures

Antoine Claude Lemor¹, Shannon Dinan², Jérémy Gilbert³

¹Université de Sherbrooke · ²Université Laval · ³Western University



LLM Tool · GitHub

Fifteen minutes, three claims

1

Computational policy analysis

text as the common format, what it makes
measurable, and where LLMs help or threaten
validity

2

Method

LLM Tool as a local, logged hybrid workflow

3

Validation

human consensus, model comparison, speed and
traceability

PART I

Computational policy analysis

Computational policy analysis: a working definition

DEFINITION

Computational policy analysis uses machine-readable text and structured algorithms to reconstruct the objects that policy analysis has long studied: problem definitions, agendas, target populations, coalitions, framing dynamics, evaluations and their diffusion.

GRIMMER, ROBERTS & STEWART 2022 · WILKERSON & CASAS 2017 · BAUMGARTNER & JONES 1993

THEORY ASSUMES SOMETHING QUITE UNREACHABLE METHODOLOGICALLY

Policy theories describe stable typologies and long-run dynamics, but the objects they invoke live in millions of parliamentary sentences, hearings, media pieces and evaluations. Manual coding cannot follow at that scale, so the annotation infrastructure itself becomes the methodological question.

A FEW CONCRETE EXAMPLES

1 AGENDAS

Coding every parliamentary intervention over decades with the CAP typology, then tracking punctuations in policy attention.

2 TARGET PUBLICS

Detecting mentions of social groups sentence by sentence, then measuring the affective and evidentiary regimes attached to each.

3 COALITIONS AND FRAMING

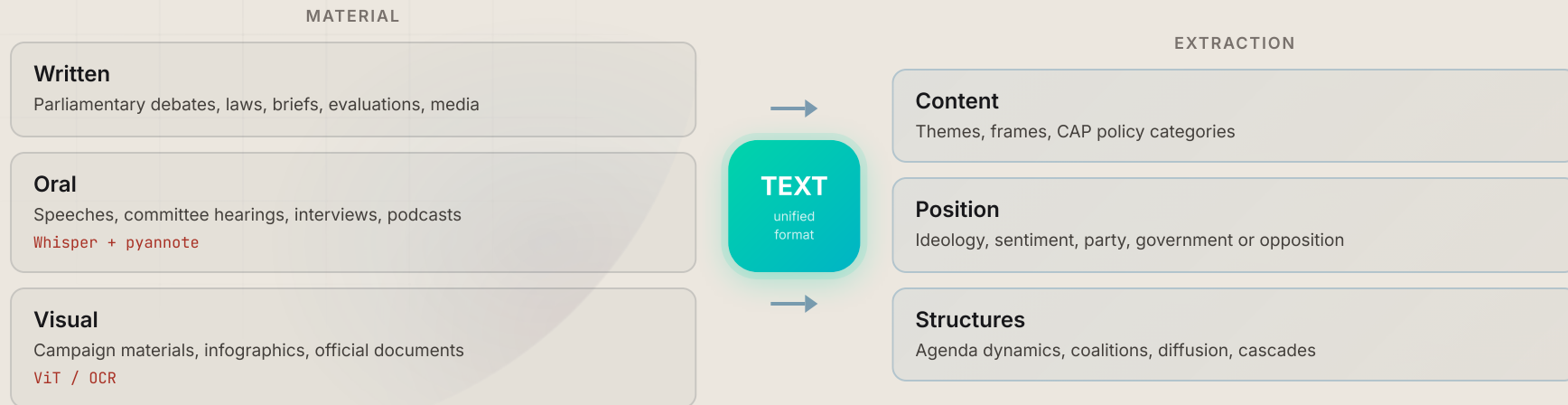
Linking actors to their framing choices in media coverage and hearings, then reconstructing coalition structures and belief diffusion.

4 EVALUATION AND DIFFUSION

Comparing evaluation reports and policy briefs across jurisdictions to trace instrument diffusion and cross-country learning.

Text: the common format for computational policy analysis

From diverse policy materials to a single analyzable format, then to measurable structures



A concrete example: the CCF database



ccf-project.ca / observatory

266K

Articles

1978 – 2024

20

Newspapers

EN + FR

9.2M

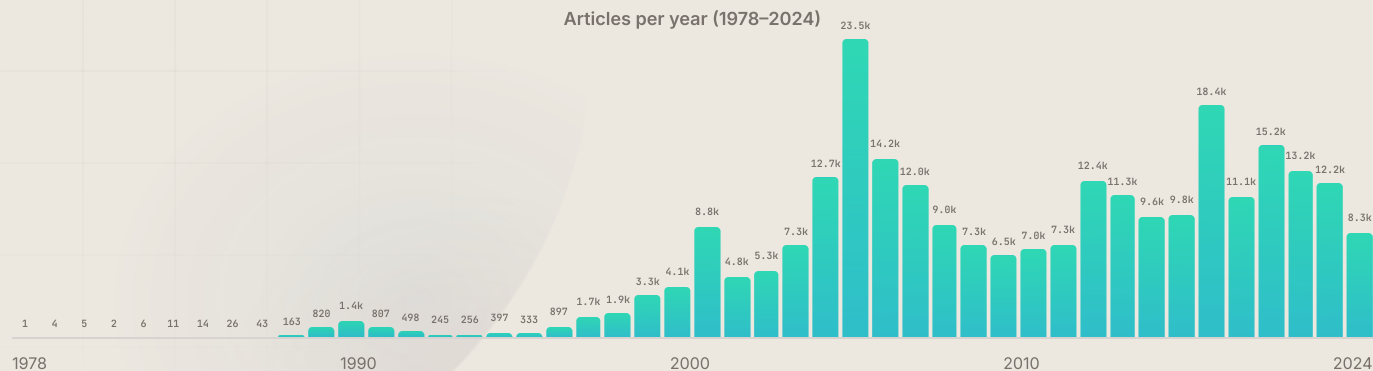
Annotated sentences

BERT / CamemBERT

65

Categories

hierarchical



The largest ML-annotated climate discourse corpus in Canada

Macro F1 : 0.866 – Lemor, Pillod, Taylor (2025)

The puzzle: annotating large corpora without losing the audit trail

Computational social science needs labels for corpora that now contain **hundreds of thousands or millions of texts**.

GRIMMER, ROBERTS & STEWART 2022 · KRIPPENDORFF 2018

Zero-shot LLMs make annotation easier, but **model opacity, prompt sensitivity, API drift and costs** weaken cumulative validation.

WEBER & REICHARDT 2023 · OLLION ET AL. 2024 · ZHAO ET AL. 2024

Hybrid distillation is promising: LLMs generate training labels, then smaller classifiers reproduce the coding logic across the corpus.

HINTON ET AL. 2015 · DI PALO ET AL. 2024 · PANGAKIS & WOLKEN 2024

1

VALIDATION

Few controlled tests compare LLM-generated training data against human-validated ground truth.

2

REPRODUCIBILITY

Published pipelines often depend on undocumented preprocessing, unstable model versions or missing implementation details.

3

ACCESS

Without an integrated local workflow, LLM-based annotation remains concentrated among teams with engineering support and API budgets.

RESEARCH QUESTION

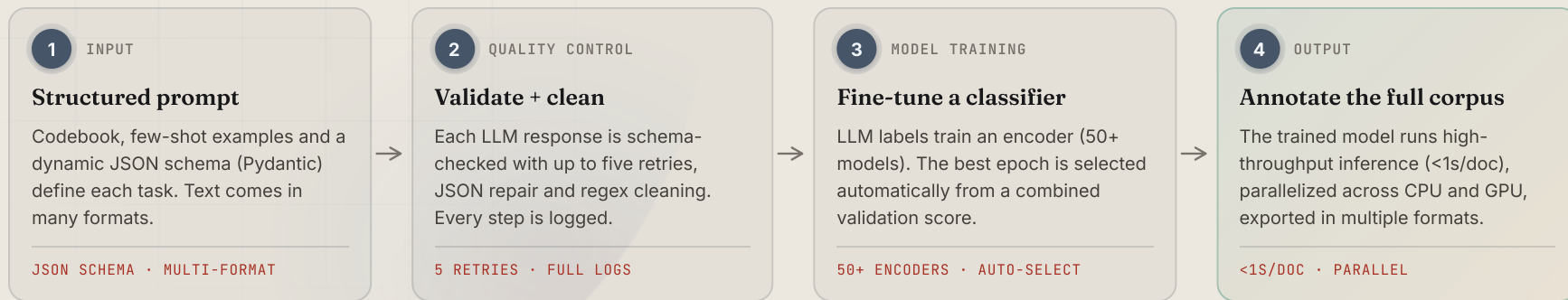
Can a locally executable encoder-decoder pipeline produce valid annotations for large corpora while documenting enough of the workflow to make replication and auditing possible?

PART II

Method and tool

LLM Tool: a local hybrid workflow, not just another prompt

From a structured codebook to a reusable classifier, with automatic documentation at each stage



Every parameter, prompt, model version and environment setting is documented automatically. The pipeline runs entirely on a local machine, no commercial API required, and each run is auditable under documented conditions.

What makes the tool transparent: modes plus persistent records

1 The Annotator

Zero-shot LLM annotation of a raw text file, exported to Label Studio or Doccano.

2 The Annotator Factory

End-to-end pipeline: chains LLM annotation, data preparation and BERT fine-tuning in one session.

3 Training Arena

Dedicated training and benchmarking workspace: an 18-step wizard from dataset to evaluation.

4 BERT Annotation Studio

High-throughput inference with a trained classifier, parallel GPU/CPU, resumable sessions.

5 Validation Lab

Quality-control workspace: ingests annotator exports and computes Cohen's κ and Krippendorff's α .

6 Resume Center

Session registry: lists partially completed runs so they resume from the last checkpoint.

Every mode is driven from an interactive command line, no code required. Each session persists its configuration, progress and outputs, so any run can be resumed or reproduced later.

RECORDED INPUTS

prompts, concept definitions, labels, examples, JSON schema, provider and model version.

RECORDED EXECUTION

temperature, seeds, retries, raw outputs, cleaned outputs, validation failures, runtime and cost logs.

MINIMUM REPRODUCIBILITY

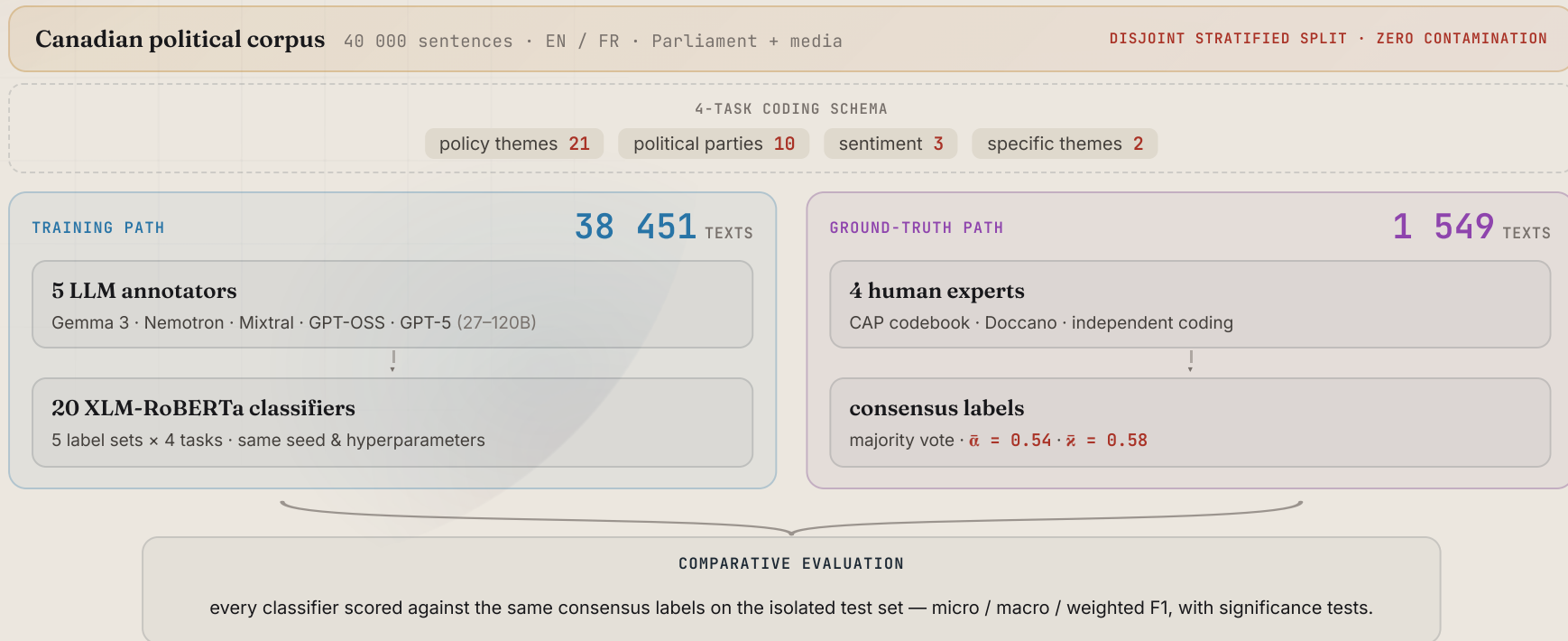
Even when inference is stochastic, the analytical choices are auditable and the run can be repeated under documented conditions.

PART III

Validation results

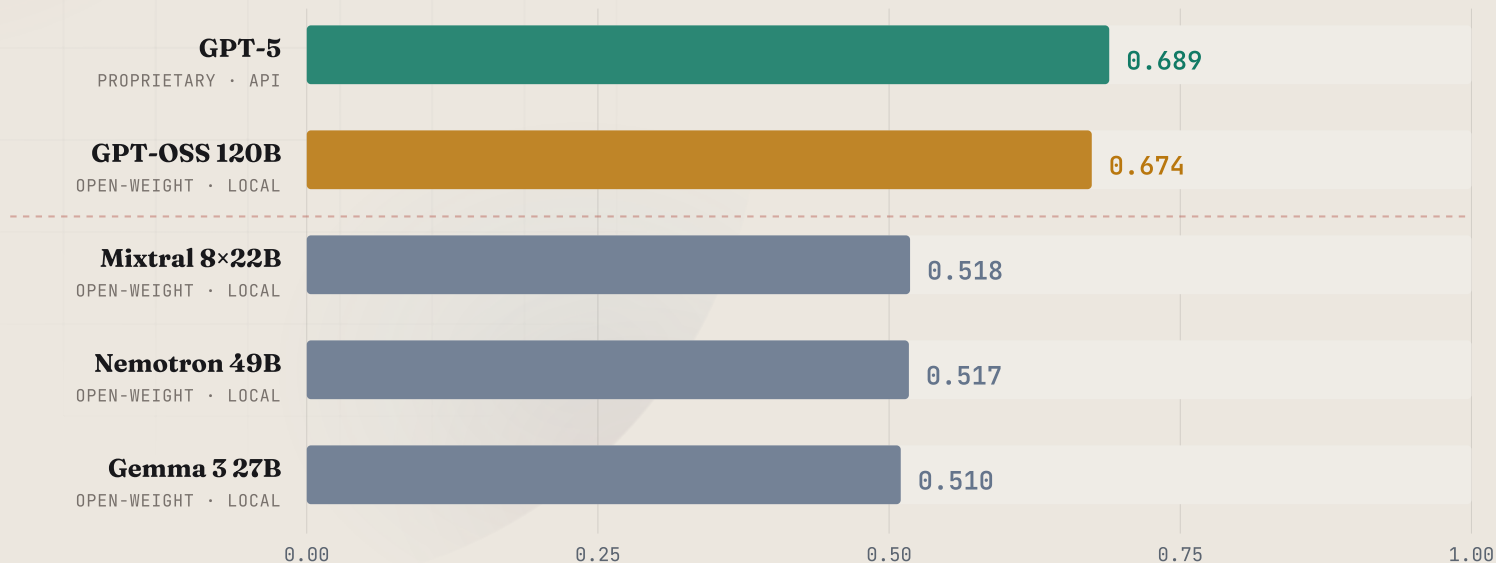
Validation design: isolate the human test set

Five LLMs annotate the training pool; five identical XLM-RoBERTa classifiers are scored against human consensus



Validation results: two tiers, one local competitor

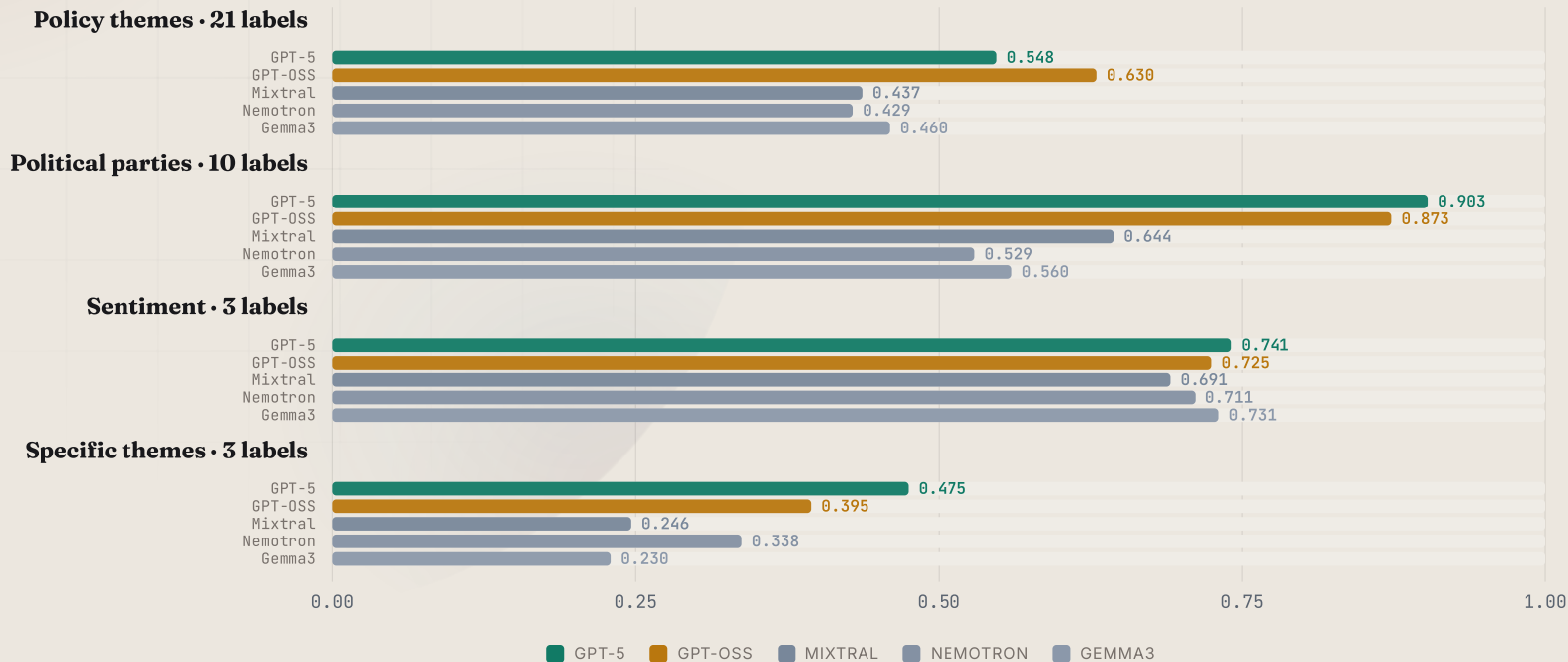
Mean Micro F1 across four annotation tasks, evaluated on the isolated human-consensus test set (N = 1,549)



A 120B open-weight model running locally is statistically indistinguishable from GPT-5 ($p = 0.327$). The annotator's parameter count does not predict the gap; codebook adherence does.

Validation results by task: Figure 5

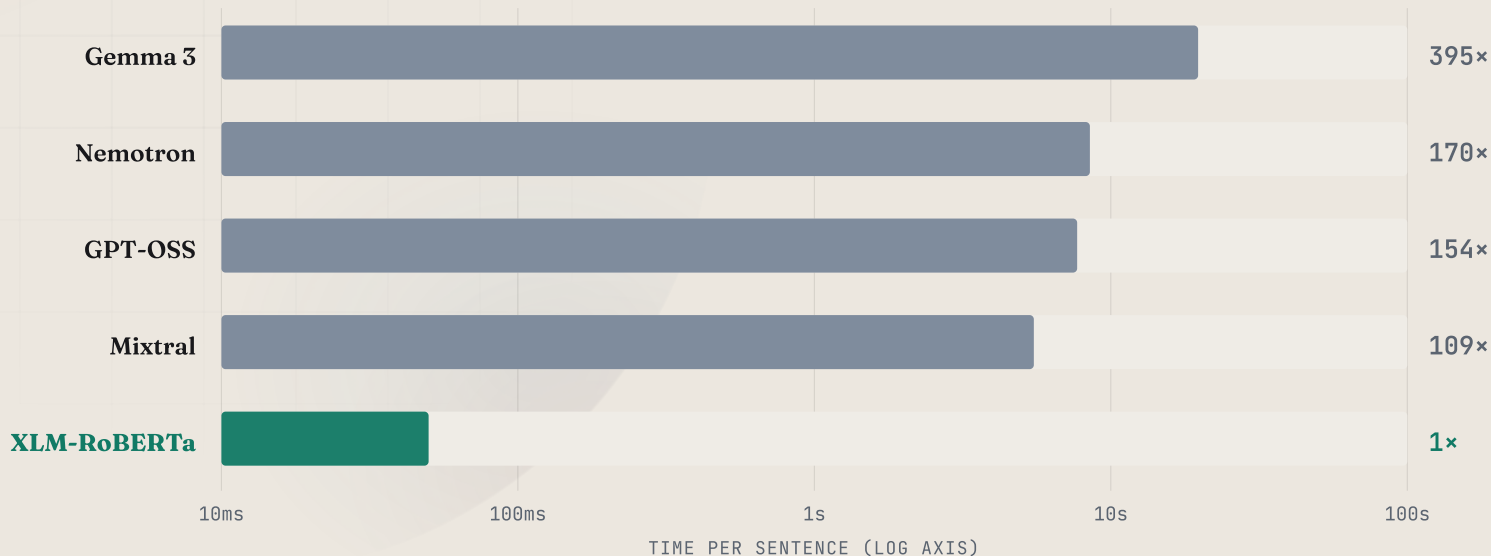
Micro F1 by annotation task and source model, evaluated on the same human-consensus test set



The task matters more than parameter count: GPT-OSS leads on policy themes, GPT-5 leads on parties and sentiment, and specific themes remain difficult because of sparse positive cases.

Validation throughput: 109-395× faster after distillation

Per-sentence inference time after distillation, compared with source-model inference



Source-model inference for 38,451 sentences takes 2.4 to 8.8 days. The distilled classifier processes the same corpus in 32 minutes, a 109× to 395× gain. The initial annotation cost is amortized across every later run.

LLM TOOL

Thank you

1

Validated

The pipeline is evaluated against a human-consensus test set, not assumed valid because an LLM produced labels.

2

Transparent

Prompts, schemas, model versions, hyperparameters, failures and outputs are recorded as part of the method.

3

Open-source

The code, data and documentation are available so the workflow can be inspected, adapted and reused.

Antoine Claude Lemor · Shannon Dinan · Jérémy Gilbert

IPPA IWPP5 · Workshop T09W02 · Ottawa · 7 July 2026



LLM Tool · GitHub