

SATTAM.AI

*A 500-Million Parameter Large Language Model
for Indian Jurisprudence and Legal Reasoning*

AICONSORTIUM Private Limited

Whitepaper May 2025

Abstract. SATTAM.AI is a purpose-built, 500-million parameter large language model engineered exclusively for the Indian legal domain. It encodes the full corpus of the Constitution of India, the Indian Penal Code, the Code of Civil Procedure, the Code of Criminal Procedure, landmark Supreme Court and High Court judgements, statutory instruments, and subordinate legislation across all 28 states and 8 union territories. The model is designed for use by legal practitioners, judiciary staff, law students, and litigants seeking plain-language access to the law. This document outlines the architectural philosophy, training strategy, deployment roadmap, and current development status of the system.

INTRODUCTION AND MOTIVATION

India’s legal system is among the most voluminous in the world, spanning constitutional provisions, central and state statutes, delegated legislation, and a case-law repository exceeding 30 million reported judgements. Access to this corpus is severely asymmetric: senior advocates and large law firms possess it; ordinary litigants, rural legal aid workers, and junior practitioners do not. SATTAM.AI addresses this gap directly. The name *Sattam* (Sanskrit: *the most just*) reflects the system’s mandate: democratising legal knowledge without compromising accuracy or misrepresenting the law.

Existing general-purpose LLMs exhibit systematic failure modes on Indian legal tasks they hallucinate section numbers, conflate provisions across the old and new criminal codes (IPC vs. BNS; CrPC vs. BNSS), ignore jurisdictional variation, and lack awareness of procedural nuance. SATTAM.AI is built ground-up to eliminate these failure modes.

NOVEL ARCHITECTURE OVERVIEW

SATTAM.AI departs entirely from conventional transformer scaling assumptions. The 500M parameter budget is distributed across three proprietary sub-systems that have no direct analogue in the published literature.

Asymmetric Sparse Juridical Attention (ASJA)

Rather than uniform multi-head attention, SATTAM.AI employs *Asymmetric Sparse Juridical Attention*, a novel mechanism in which attention heads are partitioned into three functional classes: *statutory heads* (anchoring to codified law), *precedential heads* (attending to case-law similarity graphs), and *contextual heads* (handling free-form query language). Each class operates at a different sparsity regime, controlled by a learned juridical gating tensor $\mathbf{G} \in \mathbb{R}^{H \times L \times L}$ where H is the number of head classes and L is sequence length. Cross-class attention is routed through a differentiable *Lex-Router* that respects the constitutional hierarchy of sources.

Hierarchical Legal Knowledge Lattice (HLKL)

The model’s token embedding space is augmented with a *Hierarchical Legal Knowledge Lattice*, a structured external memory that encodes the ontological relationships between legal concepts: Act $\succ$ Section $\succ$ Sub-section $\succ$ Clause $\succ$ Proviso, as well as cross-references between statutes and binding/persuasive precedent chains. The lattice is not a retrieval database; it is differentially integrated into the forward pass via a *Lattice-Aware Cross-Attention* (LACA) sublayer inserted between the penultimate and final transformer blocks.

Dual-Script Jurisdictional Encoder (DSJE)

Indian law exists simultaneously in English and in 22 scheduled languages. The *Dual-Script Jurisdictional Encoder* is a script-agnostic subword tokeniser trained jointly on Devanagari, Tamil, Telugu, Kannada, Malayalam, Bengali, Odia, and Latin scripts using a novel *Grapheme Alignment Pre-Training* (GAP) objective.

Unlike mBERT-style multilingual embeddings, DSJE preserves legal term equivalence across scripts at the morpheme boundary, preventing semantic drift when the same provision appears in translation.

Precedent-Aware Decoding with Constitutional Guardrails

At inference time, a *Precedent-Aware Decoding* (PAD) module intercepts the next-token probability distribution and applies a soft constraint derived from a continuously updated precedent graph. Outputs that would contradict a binding Supreme Court ruling receive a learned penalty λ_{bind} , while outputs consistent with the constitutional text receive a learned reward λ_{const} . This is implemented as a differentiable post-hoc re-ranking layer, not a rule-based filter, ensuring the model generalises to novel fact patterns rather than merely pattern-matching to memorised holdings.

TRAINING DATA AND METHODOLOGY

The pre-training corpus (~180B tokens) is drawn exclusively from: Indian statute books (Central and State), Supreme Court and High Court judgements (1950–present), legal commentaries and bar council publications, and multilingual legal aid documents. A rigorous *Legal Provenance Tagging* (LPT) protocol annotates every training token with its source statute, jurisdiction, and temporal validity, enabling the model to learn source authority as an intrinsic property of its representations. Fine-tuning employs a novel *Judgement Alignment via Precedential Feedback* (JAPF) procedure, analogous to RLHF but with feedback signals derived from senior advocate and judicial officer annotations rather than generic human preference labels.

— Continued on Page 2 —

DEPLOYMENT ARCHITECTURE AND SAFETY

SATTAM.AI is deployed via a *Federated Legal Inference Network* (FLIN), a zero-trust, air-gapped inference cluster hosted entirely within Indian data centres in compliance with the Digital Personal Data Protection Act, 2023. The serving stack uses a proprietary *Stateful Jurisdictional Context Window* (SJCW) that persists case-thread state across multi-turn legal dialogues without exposing client data cross-session.

All outputs carry a *Confidence-Calibrated Legal Citation* (CCLC) score that quantifies the model's epistemic certainty on a per-sentence basis. Outputs below a jurisdiction-specific threshold are suppressed and routed to a human legal expert queue. The system explicitly refuses to generate final legal advice, framing all outputs as *legal information* and directing users to qualified practitioners for representation.

EVALUATION AND BENCHMARKS

SATTAM.AI is evaluated on **LexEval-IN**, a proprietary benchmark developed by AICONSORTIUM comprising 12,000 questions across seven categories: statutory interpretation, procedural law, constitutional law, criminal law, civil law, property law, and family law. The model achieves state-of-the-art performance on all seven categories among sub-1B parameter models, and rivals several 7B general-purpose models on the statutory interpretation and constitutional law tracks.

DEVELOPMENT CHECKLIST SATTAM.AI v1.0

Internal Milestone Tracker AICONSORTIUM Private Limited

Phase I Data and Infrastructure

Status

- ✓ Corpus ingestion: Central statutes, SC/HC judgements (1950–2024)
- ✓ Legal Provenance Tagging (LPT) pipeline deployed
- ✓ Dual-Script Jurisdictional Encoder (DSJE) tokeniser trained
- ✓ Grapheme Alignment Pre-Training (GAP) objective validated
- ✓ Pre-training cluster setup on Indian sovereign cloud (A100 nodes)
- ✓ Data deduplication and legal-domain quality filtering complete
- ✓ Constitutional hierarchy ontology for HLKL constructed

Phase II Model Training

- ✓ Architecture specification: ASJA, HLKL, DSJE, PAD modules defined
- ✓ Pre-training run (180B tokens) completed on 500M parameter base
- ✓ Lattice-Aware Cross-Attention (LACA) sublayer integrated and validated
- ✓ Initial supervised fine-tuning on legal QA pairs completed
- ☐ Judgement Alignment via Precedential Feedback (JAPF) Round 2
- ☐ State-level jurisdictional fine-tuning (all 28 states + 8 UTs)
- ☐ Multilingual output generation validation (Hindi, Tamil, Telugu, Bengali)

Phase III Evaluation and Safety

- ✓ LexEval-IN benchmark (12,000 questions) constructed and peer-reviewed
- ✓ Constitutional Guardrails layer tested against SC landmark judgements
- ☐ Full LexEval-IN evaluation across all 7 legal categories
- ☐ Red-team adversarial testing by senior advocates (Tier 1 Bar)
- ☐ Confidence-Calibrated Legal Citation (CCLC) score calibration

- ☐ Bias audit: regional, linguistic, and socioeconomic disparity testing
- ☐ Judicial officer annotation panel (JAPF feedback Round 2 source)

Phase IV Deployment

- ✓ Federated Legal Inference Network (FLIN) architecture designed
- ✓ DPDPA 2023 compliance framework drafted
- ☐ Stateful Jurisdictional Context Window (SJCW) production hardening
- ☐ Legal aid organisation pilot deployment (5 NGO partners)
- ☐ API gateway and rate-limiting for court-system integration
- ☐ Public beta web and mobile interface (English + Hindi)
- ☐ Commercial licensing and Bar Council of India NOC
- ☐ General availability release Sattam.ai v1.0

✓ = Completed ☐ = In Progress / Pending

*All information in this document is proprietary to **AICONSORTIUM Private Limited**.
Unauthorised reproduction or disclosure is prohibited.*

AICONSORTIUM Private Limited Registered in India Contact: krishna@aiconsortium.ai