

1. HOW IT WORKS

The Prompt Intelligence Engine (PIE) Workbench is a production-oriented AI workflow for generating compliant, brand-consistent pharmaceutical web content. Built for the University of Liverpool CXI+AI Hackathon in March 2026. It addresses the challenge of scaling high-quality, regulated content across global sites in multiple countries while maintaining regulatory compliance, brand voice, and accessibility, turning a slow, friction-heavy approval process into a clear, governed flow.

Step-by-step Flow	
1. Register	Capture requester details for full traceability.
2. Ideation	Input prompt, audience, jurisdiction, product/therapy area + optional document uploads.
3. Brief Generation	PIE performs pre-classification (risk, tone, readability) → RAG-injected structured brief. Deterministic pre-classification ensures high-risk or jurisdiction-specific requests are handled with extra scrutiny before any LLM generation occurs.
4. Enriched Content	LLM expands brief into full webpage copy (hero, sections, CTA) while preserving facts from uploads. Uploaded images/documents bypass model scoring and require explicit human approval. All decisions are logged immutably.
5. Review	Structured compliance review with scored findings (accept/decline). Human gate required for media and high-risk content. Every generation and review step receives curated regulatory and brand context, directly addressing hallucination and compliance risks
6. Submit	Export HTML page + PDF audit report + reviewer notification.

2. RISKS & CONSIDERATIONS

Risk 1: Regulatory non-compliance & Hallucination

- **Mitigation:** Constrain the LLM to a strict RAG pattern as it can only reference context retrieved by the agents, not its parametric memory. Every claim in the generated content must be tagged with a source and timestamp. Human sign-off before any action removes the last-mile failure mode.

Risk 2: Data Leakage & Privacy

Why it matters: Proprietary product data, clinical claims and pre-launch material are commercially sensitive and subject to GDPR and pharmaceutical confidentiality obligations. Routing prompts through a public LLM API could expose proprietary or patient data outside the organisation's control.

- **Mitigation:** Deploy the LLM on-premises or in a private-cloud tenant with no outbound internet from the inference endpoint. Apply role-based access controls so the assistant only retrieves data the requesting user is authorised to see. No proprietary or patient data should transit an external API at any point.

Risk 3: Automation Bias & Over-reliance

- **Why it matters:** Content reviewers may begin treating AI scores as authoritative, reducing their own scrutiny, which is risky because the model cannot grasp full clinical or regulatory nuance (for example, an off-label implication or a market-specific disclaimer).
- **Mitigation:** Design the UI to surface confidence caveats and source counts alongside every score. Require an explicit, timed acknowledgement step and not a passive approve button. Run a quarterly calibration review: compare AI-recommended actions against outcomes to detect drift and adjust scoring weights accordingly.

3. TOOLS / APPROACH

Application of Technology and its Solutions	
React 18 TypeScript Vite (Frontend)	What it does: Powers the 6-step guided interface (Register → Submit). Real-time PIE scoring feedback and structured review panels reduce cognitive load for content requesters. Why useful here: Modern, type-safe frontend for complex multi-step workflows with real-time validation and accessibility.
Supabase Edge Functions (Deno/TypeScript)	What it does: Hosts the core agentic pipeline (pie-classify, generate-brief, enrich-content, review-content). Enables hybrid rule plus LLM execution with structured JSON outputs. Why useful here: Serverless, secure backend with built-in auth, RLS, and low-latency global deployment.
Prompt Intelligence Engine (PIE)	What it does: The first layer. Detects jurisdiction, scores regulatory risk (keyword plus context), analyses brand tone, and predicts readability. Outputs structured JSON that conditions every downstream RAG and generation step. Why useful here: A custom hybrid engine combining deterministic rules with lightweight ML signals.
Custom RAG (pharmaRag.ts)	What it does: Domain-specific retrieval for pharma compliance, brand voice, accessibility, and regulatory codes Why useful here: Retrieves and injects relevant context (FDA/EMA/MHRA/GDPR/WCAG) into every LLM prompt. Ensures the content generated is grounded rather than hallucinated.
Document Handling (mammoth, pdfjs-dist, jsPDF)	What it does: Extracts facts from uploaded DOCX/PDFs and preserves them in generated briefs and content. Produces auditable HTML and PDF deliverables with embedded images and decision logs. Why useful here: Robust extraction and export for real-world pharma content workflows.
Vercel + Supabase (Auth/Storage/RLS)	What it does: Frontend on Vercel, backend + data on Supabase with Row Level Security and immutable audit trails. Supports the mandatory human review gate Why useful here: Production-grade hosting with strong security and observability.