

A PRACTITIONER'S FIELD GUIDE

Vibe Engineering: a guide to building *scalable* AI applications

*From the creative spark of vibe coding to the architectural
rigor professional systems require*

SEVEN PHASES · PROBLEM TO PRODUCTION



What's inside

01 Introduction: from "vibe coding" to professional vibe engineering

02 Phase 1 — Defining the problem blueprint

03 Phase 2 — The requirements document (functional & ethical)

04 Phase 3 — The design document (architecture for non-coders)

05 Phase 4 — From logic to pseudocode

06 Phase 5 — Testing, evaluation, and benchmarking

07 Phase 6 — Iteration & scalability (the fine-tuning secret)

08 Final summary

SECTION 01

Introduction: from "vibe coding" to professional vibe engineering

In the current landscape of rapid AI development, many projects begin with "vibe coding," the initial creative phase where a developer utilizes generative AI to bridge the gap between concept and prototype. However, to transform a creative spark into a scalable, professional product, one must transition to **vibe engineering**. This discipline applies software architectural rigor to ensure that the final system is reliable, trustworthy, and distinct from the "AI slop" that results from unconstrained generation. Implementing engineering rigor serves as a critical guardrail against the risks of malicious use and unintended behaviors identified in the OECD Model Openness reports.

A foundational decision in vibe engineering is the selection of the model delivery mechanism. Open-weight models are foundation models where the trained weights, the numerical results of the training process, are made publicly available for local deployment and optimization. In contrast, closed models are provided strictly as a service via Application Programming Interfaces (APIs), where the weights remain proprietary and inaccessible. Architects must understand that weights are not "source code"; they are the outputs of training. Consequently, they require specific licensing, such as the CDLA-Permissive license for weights, as opposed to Apache 2.0 or MIT licenses typically reserved for human-readable inference and training code.

PHASE 01

Defining the problem blueprint

Architectural integrity begins with a clear problem definition. In AI engineering, this is essentially a quest for **alignment**: ensuring the system's outputs strictly reflect user values and project-specific objectives. Without a defined blueprint, the system is prone to "drift," where the generative output fails to meet the strategic needs of the business.

TEMPLATE — PROBLEM DEFINITION

FIELD	DESCRIPTION
Problem Statement	Concise definition of the specific challenge.
Target Audience	The end-user profile interacting with the system.
Desired "Vibe"	The technical tone and personality (e.g., authoritative, diagnostic).
Business Impact	The measurable value and strategic justification for using AI.

INSTRUCTIONS FOR USE

The architect must provide a technical justification for why a generative AI model is required for this specific problem. If the task can be solved with a deterministic algorithm or a standard database query, the AI component should be rejected to reduce system complexity.

SAMPLE INFORMATION — SOCIAL MEDIA ANALYTICS APP

- **Problem Statement:** Small business owners are overwhelmed by engagement metrics and lack the analytical capacity to synthesize data into a content strategy.
- **Target Audience:** Non-technical entrepreneurs.
- **Desired "Vibe":** Data-driven coach; encouraging but analytically rigorous.
- **Business Impact:** Increases user retention by providing actionable "Next Post" recommendations based on historical reach.

A rigorously defined problem acts as the primary constraint, dictating the functional requirements and the necessary ethical boundaries of the system.

PHASE 02

The requirements document (functional & ethical)

Setting clear requirements is a strategic defense against "downstream impacts" and the "proliferation of risks." Because AI models can be modified or repurposed, we must define the system's operational envelope early. This includes assessing **marginal risks**, the additional risks introduced by this AI implementation compared to existing non-AI technologies.

TEMPLATE — SYSTEM REQUIREMENTS

CATEGORY	REQUIREMENT SPECIFICATION
Functional	Explicit operations (e.g., "Analyze sentiment of the last 50 data points").
Performance	Technical benchmarks (e.g., "Inference latency < 500ms").
Safety Guardrails	Strategies to mitigate vulnerabilities like attribute inference or memorization.

INSTRUCTIONS FOR USE

Every requirement must include a trade-off analysis. For example, choosing a cloud-based API may offer lower initial latency but compromises data sovereignty. Conversely, an on-device open-weight model ensures privacy but requires higher local compute resources. Architects must explicitly document these choices to prevent security oversights such as non-consensual intimate imagery (NCII) generation.

SAMPLE INFORMATION — SOCIAL MEDIA ANALYTICS APP

- **Functional Requirement:** Must ingest raw engagement metrics and identify the "Top 3" highest-performing thematic clusters.
- **Safety Requirement:** Must implement filters to detect and block malicious content and NCII, adhering to Section 3.2 of the OECD safety framework.

Once the requirements are set, they dictate the physical arrangement and interaction of the system components.

PHASE 03

The design document (architecture for non-coders)

Design is the map of how system components, data, inference code, and model weights, interact. Professional vibe engineering utilizes the Model Openness Framework (MOF) to categorize these components into three distinct buckets: code, data, and documentation. A complete architecture considers all 17 components of the MOF to ensure transparency and reproducibility.

TEMPLATE — HIGH-LEVEL DESIGN

MOF CATEGORY	COMPONENT SPECIFICATION
Code	Inference code, training/fine-tuning code, preprocessing scripts, and evaluation tools.
Data	Training datasets (where applicable), evaluation data, and sample model outputs.
Documentation	Data cards, model cards, research papers, and technical reports.
Processing Layer	Selection: open-weight (local inference) vs. closed (cloud API).

INSTRUCTIONS FOR USE

Architects must evaluate the necessity of on-device solutions. On-device deployment using open-weight models is preferred for applications involving sensitive business data to maintain data sovereignty, whereas cloud APIs are selected when peak foundation model performance outweighs privacy concerns.

SAMPLE INFORMATION — SOCIAL MEDIA ANALYTICS APP

- **Architecture:** Utilizes an open-weight model to ensure that proprietary engagement data remains within the user's local environment.
- **Inference:** Local model weights optimized for text-to-text synthesis of brand-specific slang.

The architectural map is the skeleton; the logic flow provides the nervous system that drives model behavior.

PHASE 04

From logic to pseudocode

Before generating functional code, we must develop a "logic flow." Pseudocode ensures the sequence of operations is architecturally sound and that data ingress points lead logically to the desired inference.

TEMPLATE — LOGIC FLOW

- 1

Data Ingress: Fetch the last 30 days of engagement metrics via Platform API.
- 2

Filter: Identify the "Top 3" posts based on the reach-to-follower ratio.
- 3

Model Inference: Pass top-performing post content to the AI with the prompt: "Synthesize common thematic elements."
- 4

Trend Comparison: Cross-reference identified themes against current industry trend data.
- 5

Output Generation: Format results into a "Post Recommendation" report for the user dashboard.

Architectural logic must be validated against objective performance metrics to ensure the system meets its benchmarks.

PHASE 05

Testing, evaluation, and benchmarking

Subjective "vibes" must be replaced by objective evaluation. We utilize external evaluation and industry-standard benchmarks, such as MMLU (Massive Multitask Language Understanding) for general knowledge or Arena ELO for comparative quality, to quantify model performance.

TEMPLATE — EVALUATION CRITERIA

TEST SCENARIO	EXPECTED OUTCOME	ACTUAL RESULT	PASS/FAIL	REMEDATION PATH
PR Crisis Context	Cautious, brand-safe advice.	Sarcastic joke.	Fail	Implement stricter system prompt constraints.
Reach Summary	Accurate list of themes.	Accurate list.	Pass	N/A

This phase marks the transition from initial testing to a lifecycle of continuous improvement and scaling.

PHASE 06

Iteration & scalability: the fine-tuning secret

In professional engineering, a system is never "finished." We use **marginal benefit** assessments to decide if updates are worth the compute cost. Fine-tuning allows us to adjust model weights with niche-specific data, evolving a prototype into a high-performance asset.

STRATEGIC WARNING

Architects must guard against the "boiling frog" effect. This occurs when a model is only compared to its immediate predecessor, leading to a gradual and dangerous increase in risk tolerance over time. Always evaluate the current system against the original safety baseline.

TEMPLATE — ITERATION LOG

FEEDBACK	IDENTIFIED VULNERABILITY	REFINEMENT ACTION
Advice is too generic.	Lack of niche context.	Fine-tune weights on local historical successful post data.
Safeguard bypass detected.	Jailbreaking vulnerability.	Update output guardrails and prompt-filtering logic.

SAMPLE INFORMATION — SOCIAL MEDIA ANALYTICS APP

- Observation:** The model fails to capture the brand's unique "technical-but-accessible" voice.
- Action:** Execute a low-rank adaptation (LoRA) fine-tuning session using a curated dataset of the brand's best-performing historical content.

IN CLOSING

Final summary

By adhering to these architectural phases, we transform "AI slop" into a high-value, professional asset. This process mirrors the evolution of the broader ecosystem, where the surge in open-weight models provides the foundation for robust, scalable engineering. True vibe engineering is not just about the spark of the AI; it is about the rigor of the guardrails that ensure that spark provides lasting, reliable value.