



Software Valuation Report

Analysed Source Code: dify

Document Date: August 25, 2026

Platform:

Client / Applicant

Legal entity name:

Registered address:

Tax Identification Number:

Software name:

Existing trade mark:

Company web:

Applicant Name:

Applicant Email:

Request date and time: 25-08-2026 - 08:10:04



TABLE OF CONTENTS

1. Executive Summary

2. Platform Overview

- 2.1 Functional Description
- 2.2 Technical Architecture
- 2.3 Technology Stack
- 2.4 Third-Party Integrations

3. Production Readiness Assessment

- 3.1 Overall Score: 76/100
- 3.2 Detailed Breakdown

4. Development Investment Estimation

- 4.1 Effort Analysis
- 4.2 Team & Timeline
- 4.3 Cost Estimation
- 4.4 Codebase Metrics
- 4.5 Cloud Infrastructure & Maintenance Cost

5. Findings Summary

- 5.1 Critical Issues (Must Fix)
- 5.2 Warnings (Should Fix)
- 5.3 Recommendations (Nice to Have)
- 5.4 Strengths

6. Conclusion

- 6.1 Overall Assessment Summary
- 6.2 Readiness for Production / Scale
- 6.3 Key Areas Requiring Attention
- 6.4 Suggested Prioritisation of Improvements

Software Valuation Report

1. EXECUTIVE SUMMARY

Diffy (v1.16.1) is a mature, well-architected open-source LLM application development platform that combines AI workflow orchestration, a RAG (Retrieval-Augmented Generation) pipeline, agent capabilities, model management, and observability into a single deployable system. The platform comprises a modular monolith backend built in Python with Flask, a Next.js frontend in TypeScript, a Go-based agent runtime, and a Python agent SDK. With over 2.1 million effective lines of code across nearly 12,800 files, the codebase represents a substantial engineering investment spanning multiple languages, frameworks, and infrastructure components.

The platform achieves an overall production readiness score of 76 out of 100, rated as 'Good'. This score reflects strong engineering practices in several dimensions: comprehensive test coverage with a test-to-source ratio of approximately 1.04:1, full OpenTelemetry integration with structured JSON logging and distributed tracing, a robust CI/CD pipeline with linting, type checking, and multiple test suites enforced as required pull request checks, and well-organised dependency management with Dependabot, lockfiles, and a three-tier dependency classification system. The architecture enforces strict layer boundaries through import-linter contracts, ensuring framework neutrality for critical services.

However, production readiness is constrained by two critical security vulnerabilities: password hashing uses PBKDF2-SHA256 with only 10,000 iterations (versus OWASP's recommended 600,000+), and field encryption for sensitive data such as passwords and verification codes is merely Base64 encoding rather than cryptographic encryption. Additionally, default credentials are shipped in Docker Compose files, no SAST or DAST security scanning is integrated into CI, and no circuit breakers exist for external LLM API calls. These issues must be addressed before the platform can be considered fully production-ready for security-sensitive deployments.

The estimated development investment to date is substantial. Based on the calibrated analysis, approximately 92,300 person-hours were invested over an estimated 39-month duration with a team of 15 engineers. The estimated cost range for this development effort is between EUR 8,630,050 and EUR 11,675,950. This represents a significant capital investment reflecting the platform's complexity, breadth of integrations, and the maturity of its engineering practices.

2. PLATFORM OVERVIEW

2.1 Functional Description

Dify is an open-source platform for developing LLM-powered applications. It provides a comprehensive suite of tools for building, deploying, and managing AI-driven solutions. The platform's core capabilities include:

- **AI Workflow Orchestration:** A visual workflow builder supporting multiple node types including LLM nodes, code execution (Python, JavaScript, Jinja2), HTTP requests, knowledge retrieval, template transformation, parameter extraction, iteration loops, agent nodes, human-in-the-loop interactions, trigger events, webhooks, and datasource operations. Workflows support streaming responses, pause/resume functionality, and conversation variable persistence.
- **RAG Pipeline:** A full retrieval-augmented generation pipeline with support for multiple document extractors (PDF, Word, Excel, CSV, HTML, Markdown, Notion, Firecrawl, Watercrawl, Jina Reader, Unstructured), text splitters, indexing processors (paragraph, parent-child, QA), embedding with caching, reranking (weight-based and model-based), and multi-dataset retrieval routing.
- **Agent Capabilities:** Multiple agent runner strategies including chain-of-thought (CoT) agents for chat and completion, function-calling (FC) agents, and a newer agent_v2 architecture with human-in-the-loop support, session management, file handling, and output orchestration. The platform also provides a Go-based agent runtime and a Python agent SDK for external agent development.
- **Model Management:** A provider system supporting numerous LLM providers (OpenAI, Anthropic, and others via plugins) with load balancing, credential management, quota enforcement, and model access controls.
- **Knowledge Management:** Dataset management with document indexing, segmentation, metadata, annotations, and external knowledge API integration. Supports over 25 vector database backends including Milvus, Qdrant, Weaviate, PGVector, Chroma, Elasticsearch, OceanBase, and many cloud-native options.
- **Observability:** Full OpenTelemetry integration with distributed tracing, metrics collection, and structured logging. Additional tracing provider integrations for Langfuse, LangSmith, MLflow, Arize Phoenix, Opik, Weave, Aliyun, and Tencent.

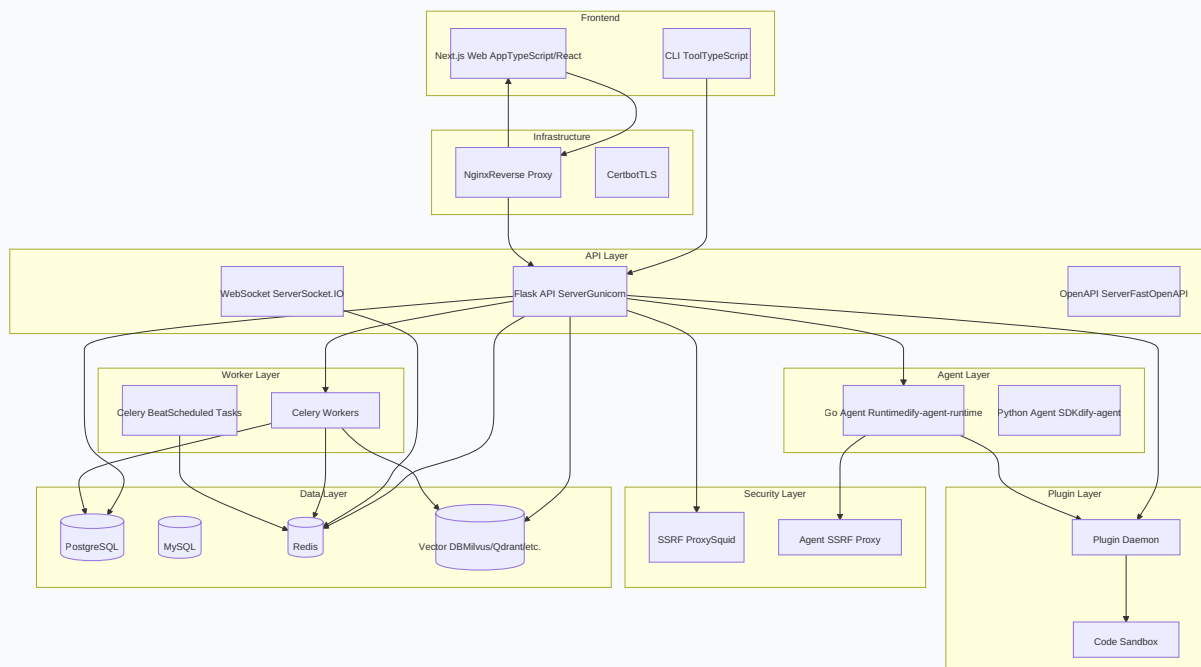


- **Plugin System:** An extensible plugin architecture supporting model providers, tools, datasources, endpoints, triggers, and OAuth flows, with a dedicated plugin daemon service.
- **Multi-tenancy and RBAC:** Workspace-based multi-tenancy with role-based access control, member management, and billing integration.
- **Deployment Options:** Self-hosted via Docker Compose, with support for multiple database backends (PostgreSQL, MySQL), multiple storage backends (AWS S3, Azure Blob, Google Cloud Storage, Aliyun OSS, and many others), and configurable vector database selection.

The target audience includes development teams building production LLM applications, enterprises requiring self-hosted AI infrastructure, and organisations needing a visual workflow builder for complex AI pipelines.

2.2 Technical Architecture

Dify employs a multi-service architecture with clear separation of concerns:



The backend follows a modular monolith pattern with strict layer boundaries enforced by import-linter. The architecture is organised into four primary layers:

1. **Controllers** (`api/controllers/`): Handle HTTP request routing, input validation, and response serialisation. Multiple controller namespaces exist for console management, service API, web app, inner API, OpenAPI, MCP, and triggers.
2. **Services** (`api/services/`): Contain business logic with framework-neutral implementations. Import-linter contracts forbid services such as machinery, workspace_query, account, and billing from importing Flask, SQLAlchemy, Werkzeug, or models directly, ensuring testability and framework independence.
3. **Core** (`api/core/`): Domain logic including workflow engine, RAG pipeline, agent runners, tools, plugins, MCP client/server, moderation, and model management. This is the largest layer with extensive sub-modules for workflow nodes, RAG extraction/indexing/retrieval, and the plugin system.
4. **Libs** (`api/libs/`): Shared utilities including authentication, encryption, pagination, rate limiting, OAuth, email, and helper functions.

Data flows from frontend requests through Nginx reverse proxy to the Flask API server, which delegates to services for business logic and to Celery workers for asynchronous task processing. The agent runtime (Go) handles agent execution with SSRF protection, while the plugin daemon manages plugin lifecycle and code sandboxing.

2.3 Technology Stack

Programming Languages (by effective lines of code):



Language	Lines of Code	Percentage
TypeScript	1,040,345	48.0%
Python	761,724	35.1%
JSON	216,017	10.0%
JavaScript	52,152	2.4%
YAML	35,543	1.6%
Markdown	34,608	1.6%
Go	11,246	0.5%
CSS	6,673	0.3%
HTML	6,207	0.3%
Shell	2,086	0.1%
PHP	173	<0.1%
SQL	1	<0.1%

Frameworks and Libraries:

- Backend: Flask, Celery, SQLAlchemy, Pydantic, Gunicorn, OpenTelemetry, Socket.IO
- Frontend: Next.js, React, Tailwind CSS
- Agent Runtime: Go with cobra, landlock LSM, modernc SQLite
- Agent SDK: Python with httpx, pydantic-ai, FastAPI, uvicorn
- Testing: pytest, Cucumber, Playwright, vitest

Databases and Data Stores:

- Primary: PostgreSQL, MySQL (configurable)
- Cache/Pub-Sub: Redis
- Vector Databases: 25+ supported including Milvus, Qdrant, Weaviate, PGVector, Chroma, Elasticsearch, OceanBase, Couchbase, Oracle, TiDB, and cloud-native options

Infrastructure and Deployment:

- Docker Compose for self-hosted deployment
- Nginx reverse proxy with TLS via Certbot
- Squid-based SSRF proxy for outbound request security
- Code sandbox for isolated execution
- Plugin daemon for plugin lifecycle management

Development and Build Tools:

- Python: uv (package manager), Ruff (linter), Pyrefly (type checker), import-linter (architecture enforcement)
- Frontend: pnpm, tsslint, Knip (dead code detection), ESLint, Oxlint
- CI/CD: GitHub Actions with path-filtered parallel jobs
- Dependency Management: Dependabot with grouped updates, lockfiles for all package managers (uv.lock, pnpm-lock.yaml, go.sum)

2.4 Third-Party Integrations

LLM Provider APIs: OpenAI, Anthropic (via plugin system with extensible provider support)

Observability and Tracing: Sentry (error tracking), Langfuse, LangSmith, MLflow, Arize Phoenix, Opik, Weave, Aliyun Trace, Tencent Trace

Cloud Storage Services: AWS S3, Azure Blob Storage, Google Cloud Storage, Aliyun OSS, Baidu OBS, Huawei OBS, Tencent COS, Oracle OCI, Supabase, Volcengine TOS, OpenDAL, ClickZetta Volume

Email Services: SendGrid, Resend, SMTP

Authentication and Security: Cloudflare Turnstile (bot protection), OAuth (multiple providers), Azure Key Vault (key management)

Databases and Vector Stores: PostgreSQL, MySQL, Redis, Milvus, Qdrant, Weaviate, Elasticsearch, and 20+ additional vector database backends

Configuration Management: Apollo, Nacos (remote settings sources)

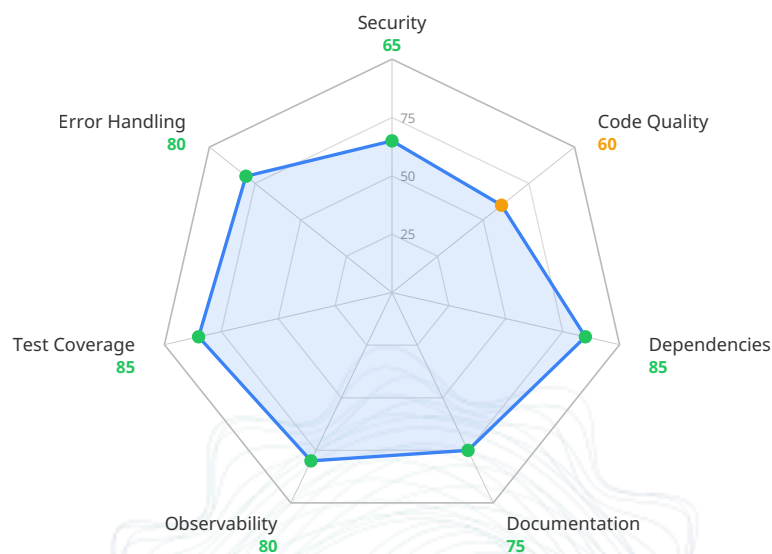
Licensing Considerations: The project uses its own licence (see LICENSE file) with a contributor agreement. Dependencies are organised in three tiers (Legacy, Stable, Emerging) with version capping to manage compatibility and security. Lockfiles are committed for reproducible builds across all package managers.



3. PRODUCTION READINESS ASSESSMENT

3.1 Overall Score: 76/100

Readiness Level: Good



The platform demonstrates strong engineering maturity across most dimensions, with particular strength in test coverage, observability, error handling, and dependency management. The overall score is pulled down by security concerns (65/100) and code quality issues (60/100) that, whilst not blocking for development or staging environments, pose risks for production deployments handling sensitive data.

3.2 Detailed Breakdown

1. Security (Score: 65/100)

Current State Analysis: The security posture is mixed. The platform demonstrates strong practices in certain areas — SSRF proxy protection with httpx, retry logic with exponential backoff, response size limits, and configurable proxy support — whilst exhibiting critical weaknesses in password hashing and field encryption.

Specific Findings:

- **Password Hashing** (`api/libs/password.py`): PBKDF2-SHA256 is used with only 10,000 iterations, significantly below OWASP's recommended 600,000+ iterations. This makes stored password hashes vulnerable to brute-force attacks. The implementation at line 20 calls `hashlib.pbkdf2_hmac("sha256", password_str.encode("utf-8"), salt_byte, 10000)`.
- **Field Encryption** (`api/libs/encryption.py`): The `FieldEncryption` class uses Base64 encoding rather than cryptographic encryption for sensitive fields including passwords and verification codes. The docstring at line 7 acknowledges this is obfuscation, not encryption. Methods `decrypt_field`, `decrypt_password`, and `decrypt_verification_code` all use `base64.b64decode`.
- **Default Credentials**: The password `difyai123456` appears in 27 files including `.env.example`, `docker-compose.yaml`, and `docker-compose-template.yaml` for Redis, PostgreSQL, MySQL, Qdrant, Chroma, OceanBase, and PGVector. This could lead to insecure deployments if operators do not change defaults.
- **Agent Secret Key** (`docker-compose.yaml`): `DIFY_AGENT_SERVER_SECRET_KEY` is hardcoded with a default value for JWE token encryption. Documentation notes it should be replaced in production, but the default is shipped in compose files.
- **SSRF Protection** (`api/core/helper/ssrf_proxy.py`): Well-implemented with `httpx`, retry logic (`BACKOFF_FACTOR=0.5`, `STATUS_FORCELIST` for 429, 500, 502, 503, 504), response size limits, configurable proxy support, and distributed tracing header injection.
- **No SAST/DAST Scanning**: Despite comprehensive CI checks for linting, type checking, and testing, no static or dynamic application security testing tools are integrated into the CI pipeline.
- **Secret Key Management** (`api/configs/secret_key.py`): The platform can auto-generate and persist a secret key if one is not provided, using `secrets.token_urlsafe(48)` and storing it in the configured storage backend. This is a reasonable approach for initial setup.

Recommendations:

- Upgrade password hashing to Argon2id or increase PBKDF2 iterations to at least 600,000 as per OWASP guidelines.
- Replace Base64 field encoding with proper cryptographic encryption such as AES-GCM for sensitive field transmission.
- Integrate SAST scanning tools such as Bandit or Semgrep into the CI pipeline.

- Remove default credentials from shipped compose files or enforce mandatory credential rotation on first boot.

2. Code Quality (Score: 60/100)

Current State Analysis: The codebase demonstrates strong architectural discipline through import-linter enforcement and a well-organised layer structure, but the score reflects the scale-related challenges of maintaining consistency across 2.1 million lines of code.

Specific Findings:

- **Architecture Enforcement** (`api/.importlinter`): Strict layer boundaries are enforced: controllers to services to core to libs. Multiple forbidden-import contracts ensure framework neutrality for specific services. For example, machinery, workspace_query, account, and billing services are forbidden from importing Flask, SQLAlchemy, Werkzeug, models, or controllers. This is a strong architectural discipline rarely seen at this scale.
- **Code Organisation:** The codebase is well-structured with clear separation between controllers, services, core domain logic, models, and shared libraries. The `api/core/` directory contains well-organised sub-modules for workflow, RAG, agent, tools, plugins, and MCP.
- **Naming Conventions:** Consistent naming patterns are observed across Python modules (snake_case files, PascalCase classes) and TypeScript (camelCase functions, PascalCase components).
- **Technical Debt:** The presence of over 100 Alembic migration files indicates a long evolution with schema changes. Some legacy patterns may exist alongside newer implementations (e.g., agent vs agent_v2 workflow nodes).
- **Code Duplication:** The large codebase with multiple similar service implementations (e.g., numerous VDB provider integrations) likely contains some duplication, though the provider pattern mitigates this through abstraction.

Recommendations:

- Continue enforcing import-linter contracts and expand coverage to new modules.
- Consider consolidating legacy agent implementations with the agent_v2 architecture.
- Document and track technical debt items systematically.

3. Dependencies (Score: 85/100)

Current State Analysis: Dependency management is a clear strength of the platform, with a well-organised three-tier classification system, Dependabot automation, and committed lockfiles.

Specific Findings:

- **Dependency Classification** (`api/pyproject.toml`): Dependencies are organised into three tiers: Legacy (mature, widely deployed), Stable (production-proven, capped below next major), and Emerging (newer, fast-moving, compatible pins). This demonstrates thoughtful dependency management.
- **Dependabot Configuration** (`.github/dependabot.yml`): Configured with grouped updates by ecosystem (uv for Python, github-actions), targeting both main and LTS branches (lts/1.13.x). Groups include flask, google, opentelemetry, pydantic, llm, database, storage, vdb, dev, and python-packages for efficient batching.
- **Lockfiles:** Committed for all package managers — `uv.lock` (Python), `pnpm-lock.yaml` (frontend), `go.sum` (Go runtime). This ensures reproducible builds.
- **Version Pinning:** Dependencies are pinned with upper bounds to prevent unintended major version upgrades. The uv workspace pattern is used for VDB and trace provider sub-packages.
- **Python Version:** Targets Python 3.12, which is current and well-supported.

Recommendations:

- Continue monitoring for security advisories on Emerging-tier dependencies.
- Consider adding automated licence compliance checking to the CI pipeline.

4. Documentation (Score: 75/100)

Current State Analysis: Documentation is present at multiple levels but varies in completeness and depth.

Specific Findings:

- **README:** Comprehensive overview of the platform, deployment instructions, and feature descriptions. Available in multiple languages (en, zh-CN, ja-JP, ko-KR, fr-FR, de-DE, es-ES, pt-BR, and others).
- **API Documentation:** OpenAPI/Swagger specifications are generated and maintained in `api/openapi/markdown/` with separate docs for console, service, web, and open APIs. An



API Schema Guide (`api/controllers/API_SCHEMA_GUIDE.md`) documents the expected Flask-RESTX + Pydantic pattern.

- **Contributing Guide** (`CONTRIBUTING.md`): Covers bug reports, feature requests, and contribution workflow with issue prioritisation guidance.
- **Security Policy** (`SECURITY.md`): Clear vulnerability reporting process through GitHub Security Advisories with responsible disclosure guidance.
- **Inline Documentation**: Code contains docstrings and comments, though coverage varies across modules. The structured formatter and OTel instrumentation include good inline documentation.
- **Architecture Documentation**: No Architecture Decision Records (ADRs) or design documentation were found, despite the complex multi-service architecture with strict layer boundaries. This is a notable gap.
- **Setup and Deployment**: Docker Compose files with `env_file` includes for core services, databases, vectorstores, and infrastructure provide reasonable deployment guidance. The `.devcontainer/` configuration supports development environment setup.

Recommendations:

- Create Architecture Decision Records to document key architectural decisions, trade-offs, and design rationale.
- Expand inline documentation coverage for complex workflow and RAG pipeline modules.
- Add deployment guides for production scenarios beyond Docker Compose.

5. Observability (Score: 80/100)

Current State Analysis: Observability is a significant strength, with production-grade OpenTelemetry integration, structured logging, and multiple tracing provider support.

Specific Findings:

- **OpenTelemetry Integration** (`api/extensions/otel/instrumentation.py`): Full auto-instrumentation for Flask, Celery, HTTPx, Redis, and SQLAlchemy. Includes a custom `ExceptionLoggingHandler` that records exceptions on current spans, HTTP response counter metrics with status code, method, and route attributes, and SQL commenter for SQLAlchemy query tracing.
- **Structured Logging** (`api/core/logging/structured_formatter.py`): Production-grade JSON logging using orjson with a well-defined schema: timestamp (ISO 8601 UTC),

severity, service, caller (file:line), trace_id, span_id, identity (tenant_id, user_id, user_type), message, attributes, and stack_trace. Trace context is integrated into log entries.

- **Distributed Tracing:** Multiple tracing provider integrations including Langfuse, LangSmith, MLflow, Arize Phoenix, Opik, Weave, Aliyun, and Tencent, each implemented as separate provider packages with unit tests.
- **Metrics Collection:** HTTP response counters with dimensional attributes (status code, method, route). Enterprise telemetry module provides additional metric handling.
- **Health Checks:** Configured in Docker Compose for api, worker, redis, postgres, mysql, sandbox, and local_sandbox services.
- **Sentry Integration:** Error tracking via Sentry extension (`api/extensions/ext_sentry.py`).
- **Missing Elements:** No SLO (Service Level Objective) or SLI (Service Level Indicator) definitions were found. No runbooks for operational incident response are present. Alerting setup is not documented.

Recommendations:

- Add SLO and SLI definitions along with runbooks for operational incident response.
- Implement alerting rules based on collected metrics (error rates, latency, queue depth).
- Consider adding Grafana dashboard templates for visualising OTel metrics.

6. Test Coverage (Score: 85/100)

Current State Analysis: Test coverage is exceptional, with a test-to-source ratio of approximately 1.04:1 and multiple testing layers from unit to end-to-end.

Specific Findings:

- **Test Layers:** The platform includes unit tests (`api/tests/unit_tests/`), integration tests (`api/tests/integration_tests/`), container integration tests (`api/tests/test_containers_integration_tests/`), and end-to-end tests using Cucumber and Playwright (`e2e/`).
- **Unit Test Scope:** Extensive unit tests covering controllers, services, core domain logic, models, extensions, factories, fields, libs, repositories, tasks, and tools. Tests are organised to mirror the source directory structure.
- **Container Integration Tests:** Tests that run against actual containerised services (PostgreSQL, Redis) for realistic integration testing of database operations, service workflows, and controller endpoints.



- **E2E Tests:** Cucumber feature files covering app creation, authentication, agent configuration, accessibility (keyboard navigation, axe), and workflow execution. Playwright is used for browser automation.
- **Coverage Configuration** (`api/pytest.ini`): Coverage collection enabled with `--cov=./api`, branch coverage (`--cov-branch`), and JSON/XML report generation. Mock environment variables are configured for all LLM providers with clearly fake test keys.
- **CLI Tests:** The CLI tool (`cli/`) includes both unit tests and E2E test suites with mock servers and fixture-based testing.
- **VDB Tests:** Dedicated test suites for vector database providers with both unit and integration tests for each supported backend.
- **CI Integration:** All test suites are run as required checks on pull requests, with path-filtered parallel jobs for efficiency.

Recommendations:

- Continue maintaining the high test-to-source ratio as the codebase grows.
- Consider adding property-based testing for complex workflow node transformations.
- Expand E2E test coverage for agent_v2 workflows and trigger-based scenarios.

7. Error Handling (Score: 80/100)

Current State Analysis: Error handling is well-structured with an extensive custom exception hierarchy and consistent error response patterns.

Specific Findings:

- **Exception Hierarchy** (`api/core/errors/error.py`): Over 100 custom exception classes providing structured error handling. Key exceptions include `LLMError`, `LLMBadRequestError`, `ProviderTokenNotInitError`, `QuotaExceededError`, `ModelCurrentlyNotSupportError`, `InvokeRateLimitError`, and many domain-specific errors. Each exception follows consistent description patterns.
- **Error Responses:** Structured error responses that do not leak internal details. Controller error handlers (`api/controllers/common/errors.py`, `api/controllers/console/error.py`, and others) translate domain exceptions to appropriate HTTP responses.
- **Retry Logic:** Exponential backoff with configurable parameters is implemented for outbound HTTP requests via the SSRF proxy (`api/core/helper/ssrf_proxy.py` with `BACKOFF_FACTOR=0.5` and `STATUS_FORCELIST` for 429, 500, 502, 503, 504).

- **Graceful Degradation:** The workflow engine supports pause/resume functionality, allowing workflows to be interrupted and resumed. Human-in-the-loop nodes provide graceful handling of user input timeouts.
- **Missing Circuit Breakers:** No circuit breakers are implemented for external LLM API calls. Whilst retry logic with exponential backoff exists, the absence of circuit breaker patterns means cascading failures during provider outages are possible.
- **Warm Shutdown:** The worker extension (`api/extensions/workflow_warm_shutdown.py`) implements graceful shutdown for Celery workers, allowing in-progress tasks to complete.

Recommendations:

- Implement circuit breakers for external LLM provider calls to improve resilience against provider outages and prevent cascading failures.
- Consider adding bulkhead patterns for resource isolation in high-concurrency scenarios.
- Document error handling patterns and conventions for new contributors.

8. Economic (Narrative — No Numerical Score)

The economic standing of the Dify platform reflects a substantial development investment commensurate with its scope and complexity. The estimated 92,300 person-hours of development effort, valued at between EUR 8,630,050 and EUR 11,675,950, represent a significant capital outlay spread across an estimated 39-month development period with a team of 15 engineers. This investment has produced a platform with broad market applicability in the rapidly growing LLM application development space.

The ongoing maintenance and hosting cost burden is estimated at between EUR 1,523,000 and EUR 2,031,000 annually, reflecting the multi-service architecture with 13 compute services, 3 database systems, and 2 additional managed services. The platform's support for multiple cloud providers and vector databases provides flexibility in optimising infrastructure costs, but the breadth of integrations also increases the maintenance surface area.

Overall cost efficiency is reasonable given the platform's capabilities. The strong test coverage (1.04:1 test-to-source ratio) and automated CI/CD pipeline reduce the marginal cost of ongoing development and maintenance. The modular architecture with import-linter enforcement helps contain the cost of changes by preventing cross-layer dependencies. However, the security remediation work required (password hashing, field encryption, SAST integration) represents an additional near-term investment that should be factored into any acquisition or investment decision.



4. DEVELOPMENT INVESTMENT ESTIMATION

4.1 Effort Analysis

Base Hours Calculation:

The codebase contains 2,166,791 effective lines of code (non-blank, non-comment) across 12,797 files. Using industry-standard productivity benchmarks for complex multi-language systems, the base effort is derived from the effective LOC with adjustments for the platform's complexity profile.

Complexity Multiplier Breakdown:

Factor	Rating	Impact
Architectural Complexity	4/5	High — multi-service architecture with strict layer boundaries, plugin system, and agent runtime
Domain Complexity	4/5	High — LLM workflow orchestration, RAG pipeline, multi-tenancy, RBAC
Integration Complexity	5/5	Very High — 25+ vector databases, 10+ storage backends, 8+ tracing providers, multiple LLM providers
Security Surface	4/5	High — authentication, encryption, SSRF protection, multi-tenancy isolation

The combined complexity profile yields a 'very_high' complexity classification, reflecting the platform's extensive integration surface, multi-language codebase, and sophisticated workflow engine.

Quality Adjustment: The high test-to-source ratio (1.04:1), comprehensive CI/CD pipeline, and structured architecture enforcement indicate that the codebase was developed with strong engineering practices, which is reflected in the effort estimation.

Final Estimated Hours: 92,300 person-hours

Complexity Classification: Very High

4.2 Team & Timeline

Estimated Team Size: 15 engineers

Team Composition:

Role	Count
Tech Lead	1
Backend Developer	5
Frontend Developer	3
DevOps / SRE	2
QA Engineer	2
AI/ML Engineer	2

Estimated Project Duration: 39 months

Assumptions:

- The team composition reflects the multi-disciplinary nature of the platform, with backend developers forming the largest group due to the extensive Python codebase and service layer.
- Frontend developers handle the Next.js web application and CLI tool.
- DevOps/SRE engineers manage the Docker Compose infrastructure, CI/CD pipeline, and multi-service deployment.
- QA engineers maintain the comprehensive test suite across unit, integration, container, and E2E layers.
- AI/ML engineers develop the LLM workflow engine, RAG pipeline, and agent runtime.
- The 39-month duration is consistent with the platform's version number (1.16.1), migration history spanning from 2024 to 2026, and the breadth of features implemented.

4.3 Cost Estimation

Cost Range: EUR 8,630,050 — EUR 11,675,950



This range is calculated as 92,300 hours multiplied by a blended rate of EUR 75-150 per hour, reflecting the mix of senior and mid-level engineering roles across multiple geographies and specialisations.

Confidence Level: Medium

The confidence level is medium because the estimation is based on codebase metrics and complexity signals rather than direct access to development team records. The large codebase with clear architectural patterns and comprehensive test coverage provides reasonable confidence in the effort estimation, but the actual team composition and development timeline may differ from the estimated values.

4.4 Codebase Metrics

Metric	Value
Total files analysed	12,797
Total effective lines of code	2,166,791
Primary language (TypeScript)	1,040,345 LOC (48.0%)
Secondary language (Python)	761,724 LOC (35.1%)
JSON configuration/data	216,017 LOC (10.0%)
JavaScript	52,152 LOC (2.4%)
YAML	35,543 LOC (1.6%)
Markdown	34,608 LOC (1.6%)
Go	11,246 LOC (0.5%)
Other (CSS, HTML, Shell, PHP, SQL)	15,140 LOC (0.7%)

4.5 Cloud Infrastructure & Maintenance Cost

Detected Infrastructure Components:



Component Type	Count	Details
Compute Services	13	API, WebSocket, Worker, Beat, Web, Sandbox, Local Sandbox, Plugin Daemon, Agent Backend, Agent SSRF Proxy, SSRF Proxy, Nginx, Certbot
Databases	3	PostgreSQL, MySQL, Redis
Message Queues	0	(Celery uses Redis as broker)
Storage Buckets	0	(Configured via external cloud storage providers)
CDN Endpoints	0	(Nginx serves as reverse proxy)
ML/GPU Services	0	(LLM inference via external provider APIs)
Other Managed Services	2	Plugin Daemon, Agent Backend

Detected Cloud Provider: No single cloud provider is mandated. The platform is cloud-agnostic with support for AWS, Azure, Google Cloud, Aliyun, and others through configurable storage and vector database backends. Docker Compose is the primary deployment mechanism for self-hosted installations.

Suggested Managed Services Mapping:

- PostgreSQL: Amazon RDS, Azure Database for PostgreSQL, or Google Cloud SQL
- Redis: Amazon ElastiCache, Azure Cache for Redis, or Google Cloud Memorystore
- Vector Database: Managed Milvus (Zilliz Cloud), Qdrant Cloud, or Weaviate Cloud
- Object Storage: AWS S3, Azure Blob Storage, or Google Cloud Storage
- Compute: Amazon ECS/Fargate, Azure Container Apps, or Google Cloud Run for containerised services
- Monitoring: Managed OpenTelemetry collectors with Datadog, New Relic, or Grafana Cloud

Estimated Monthly Hosting Cost Range: EUR 4,000 — EUR 12,000

This range assumes a production deployment with moderate traffic (thousands of daily active users), single-region deployment, and managed database services. The lower bound reflects a minimal deployment with self-managed databases; the upper bound reflects a production deployment with managed services, redundancy, and monitoring.

Key Assumptions:

- Traffic: Moderate production load with thousands of daily active users
 - Redundancy: Single-region with no multi-AZ replication (lower bound); single-region with managed service redundancy (upper bound)
 - Vector database: Single instance (lower bound); clustered with replication (upper bound)
 - LLM inference costs are excluded as they are usage-dependent and billed by external providers
 - Storage costs are minimal for the application itself; document storage scales with usage
-

5. FINDINGS SUMMARY

5.1 Critical Issues (Must Fix)

1. **Weak Password Hashing** (`api/libs/password.py`): PBKDF2-SHA256 with only 10,000 iterations is used, significantly below OWASP's recommended 600,000+ iterations. This makes stored password hashes vulnerable to brute-force attacks and poses an immediate risk to production deployments handling user authentication.
2. **Non-Cryptographic Field Encoding** (`api/libs/encryption.py`): The `FieldEncryption` class uses Base64 encoding rather than cryptographic encryption for sensitive fields including passwords and verification codes during transmission. Base64 is reversible encoding, not encryption, and provides no confidentiality protection. The docstring acknowledges this is obfuscation, not encryption.

5.2 Warnings (Should Fix)

1. **Default Credentials in Deployment Files**: The password `difyai123456` is present in 27 files including `.env.example`, `docker-compose.yaml`, and `docker-compose-template.yaml` for Redis, PostgreSQL, MySQL, Qdrant, Chroma, OceanBase, and PGVector. This could lead to insecure deployments if operators do not change defaults.
2. **No SAST or DAST Security Scanning in CI**: Despite comprehensive CI checks including linting (Ruff, ESLint, Oxlint, tsslint), type checking (Pyrefly), dead code detection (Knip), import boundary enforcement, and multiple test suites, no static or dynamic application security testing tools are integrated.

3. **No Circuit Breakers for External LLM API Calls:** Retry logic with exponential backoff exists in the SSRF proxy, but no circuit breaker patterns are implemented for external LLM provider calls. This could lead to cascading failures during provider outages.
4. **No Architecture Decision Records:** No ADRs or design documentation were found, despite the complex multi-service architecture with strict layer boundaries enforced by import-linter. This makes it difficult for new contributors and evaluators to understand architectural rationale.
5. **Hardcoded Agent Secret Key** (`docker-compose.yml`): `DIFY_AGENT_SERVER_SECRET_KEY` is hardcoded with a default value for JWE token encryption. Documentation notes it should be replaced in production, but the default is shipped in compose files, creating a risk of insecure deployments.
6. **No SLO/SLI Definitions or Runbooks:** Whilst observability infrastructure is production-grade, no service level objectives, service level indicators, or operational runbooks are defined for incident response.

5.3 Recommendations (Nice to Have)

1. **Upgrade Password Hashing:** Migrate to Argon2id or increase PBKDF2 iterations to at least 600,000 as per OWASP guidelines. Consider supporting password re-hashing on login for seamless migration.
2. **Replace Base64 Field Encoding:** Implement proper cryptographic encryption such as AES-GCM for sensitive field transmission, with key management via Azure Key Vault or an equivalent KMS.
3. **Integrate SAST Scanning:** Add tools such as Bandit or Semgrep to the CI pipeline to catch security vulnerabilities early in development.
4. **Implement Circuit Breakers:** Add circuit breaker patterns for external LLM provider calls to improve resilience against provider outages and prevent cascading failures.
5. **Create Architecture Decision Records:** Document key architectural decisions, trade-offs, and design rationale for the complex multi-service architecture, including the layer boundary enforcement strategy and plugin system design.
6. **Add SLO/SLI Definitions and Runbooks:** Define service level objectives and indicators along with operational runbooks for incident response, covering common scenarios such as LLM provider outages, database failures, and queue backlogs.

5.4 Strengths

1. **Well-Architected Modular Monolith:** Strict layer boundaries enforced by import-linter (controllers to services to core to libs), with multiple forbidden-import contracts ensuring framework neutrality for specific services. This is a strong architectural discipline rarely seen at this scale.
2. **Comprehensive Test Coverage:** Test-to-source ratio of approximately 1.04:1, including unit tests, integration tests, container integration tests, and E2E tests using Cucumber and Playwright. All test suites are enforced as required checks on pull requests.
3. **Full OpenTelemetry Integration:** Structured JSON logging with orjson, distributed tracing, metrics collection, and auto-instrumentation for Flask, Celery, HTTPx, Redis, and SQLAlchemy. Custom exception logging handler and SQL commenter enhance observability.
4. **Extensive Custom Exception Hierarchy:** Over 100 exception classes providing consistent error handling patterns and structured error responses that do not leak internal details.
5. **Comprehensive CI/CD Pipeline:** Linting (Ruff, ESLint, Oxlint, tsslint), type checking (Pyrefly), dead code detection (Knip), import boundary enforcement, response contract linting, and multiple test suites all enforced as required checks on pull requests.
6. **SSRF Proxy Protection:** Well-implemented with httpx, retry logic with exponential backoff, response size limits, configurable proxy support, and distributed tracing header injection.
7. **Dependabot with Grouped Updates:** Configured with grouped updates by ecosystem, targeting both main and LTS branches, with lockfiles committed for all package managers (uv.lock, pnpm-lock.yaml, go.sum).

6. CONCLUSION

6.1 Overall Assessment Summary

Dify (v1.16.1) is a mature, well-architected open-source LLM application development platform that demonstrates strong engineering practices across most dimensions of production readiness. With an overall score of 76/100 (Good), the platform excels in test coverage (85/100), dependency management (85/100), observability (80/100), and error

handling (80/100). The modular monolith architecture with import-linter enforcement, comprehensive CI/CD pipeline, and full OpenTelemetry integration indicate a team that values engineering discipline and operational excellence.

The platform's two critical security vulnerabilities — weak password hashing (PBKDF2 with 10,000 iterations) and non-cryptographic field encoding (Base64 instead of encryption) — are the primary factors constraining production readiness. These issues, combined with the absence of SAST/DAST scanning in CI, default credentials in deployment files, and missing circuit breakers for external API calls, prevent the platform from achieving a higher readiness score despite its strengths in other areas.

The estimated development investment of 92,300 person-hours (EUR 8,630,050 — EUR 11,675,950) over 39 months with a 15-person team reflects the substantial engineering effort required to build a platform of this scope and complexity. The ongoing maintenance cost of EUR 1,523,000 — EUR 2,031,000 annually is commensurate with the multi-service architecture and broad integration surface.

6.2 Readiness for Production / Scale

The platform is conditionally ready for production deployment. It can be deployed in production environments where:

- User authentication data is not stored (e.g., delegated to external identity providers), or where the password hashing weakness is mitigated through infrastructure controls.
- Sensitive field transmission does not rely on the built-in field encryption (e.g., TLS provides transport-layer protection).
- Default credentials are changed before deployment.
- LLM provider outages are handled at the infrastructure level (e.g., via external load balancers or API gateways with circuit breaking).

The platform is not ready for production deployment in environments handling sensitive user data without first addressing the critical security vulnerabilities. For scaling, the multi-service architecture with Redis-based Celery workers, configurable vector databases, and multi-tenant workspace isolation provides a reasonable foundation, though horizontal scaling of the Flask API server would require session affinity or stateless session configuration.

6.3 Key Areas Requiring Attention

The most important technical areas requiring investment in the short term are: first, the remediation of the two critical security vulnerabilities — upgrading password hashing to



Argon2id or increasing PBKDF2 iterations to OWASP-recommended levels, and replacing Base64 field encoding with proper cryptographic encryption such as AES-GCM. Second, the integration of SAST scanning tools into the CI pipeline to provide ongoing security vulnerability detection. Third, the removal or enforcement of default credential changes in deployment files. Fourth, the implementation of circuit breakers for external LLM provider calls to prevent cascading failures. Fifth, the creation of Architecture Decision Records and operational runbooks to support long-term maintainability and operational readiness.

6.4 Suggested Prioritisation of Improvements

The improvements should be tackled in the following order. First, the two critical security issues (password hashing and field encryption) must be addressed immediately as they pose direct risks to user data confidentiality. These are blocking issues for any production deployment handling user authentication. Second, SAST scanning integration should be added to the CI pipeline to provide continuous security vulnerability detection and prevent regressions. Third, default credentials should be removed from deployment files or replaced with mandatory first-boot credential generation. Fourth, circuit breakers should be implemented for external LLM provider calls to improve resilience and prevent cascading failures during provider outages. Fifth, Architecture Decision Records should be created to document the architectural rationale and support onboarding of new contributors. Sixth, SLO/SLI definitions and operational runbooks should be developed to support production operations and incident response. The first three items should be completed before any production deployment; the remaining items can be addressed in parallel with ongoing development.

ANNEX — METHODOLOGY

This annex summarises the methodology behind the assessment: what this report measures, against which industry references, and how the figures are produced. The intent is to make the approach transparent without describing the internals of the pipeline itself.

The assessment has two complementary sides:

1. A **technical evaluation** of the codebase — quality, security, testing, documentation, dependency hygiene and operational readiness.
 2. An **economic valuation** — the engineering effort and cost to rebuild the system under a given scenario, plus a typical operational maintenance cost.
-

1. Technical evaluation

The technical assessment combines two complementary layers:

- **Static analysis.** Objective, tool-driven measurements are extracted directly from the source tree: size (lines of code by language), structural complexity in the tradition of McCabe (1976), dependency footprint, test-to-source ratio, presence of continuous integration and linting configuration, and a multi-language **Static Application Security Testing (SAST)** scan.
- **AI-assisted code review.** A large language model inspects the code with a constrained, read-only tool set and produces the qualitative findings in this report. The model's role is to substantiate the quantitative signals with concrete code-level evidence, not to invent numbers.

Reference frameworks

Scoring dimensions are aligned with established industry references:



Dimension	Reference
Code Quality & Maintainability	ISO/IEC 25010 <i>Maintainability</i> characteristic; ISO/IEC 5055 automated source-code quality measures; McCabe cyclomatic complexity.
Test Coverage & Quality	The test pyramid (Cohn) and ISO/IEC 25010 <i>Reliability</i> .
Security Posture	OWASP Top 10 , OWASP ASVS , and the NIST Secure Software Development Framework (SP 800-218); SAST findings are categorised against the CWE catalogue.
Documentation	SWEBOK v4 recommendations on software documentation.
Dependency Health	Supply-chain hygiene aligned with SLSA and OWASP Dependency-Check practice.
Error Handling & Resilience	Site Reliability Engineering (SRE) literature and ISO/IEC 25010 <i>Reliability</i> .

Each dimension is scored on a 0–100 scale. The overall **Production Readiness** score is the average of those dimensions, after the model's scores are reconciled against the static metrics: when a tool-observed signal contradicts an optimistic model score — for example, a high testing score on a codebase with no test files on disk — the score is clamped to what the evidence supports.

2. Economic valuation

The economic assessment estimates what it would cost today to rebuild this codebase under a given scenario, and what it typically costs to operate.

Build effort

Effort estimation follows the parametric-cost-estimation tradition — notably Boehm's **COCOMO II** and functional-size-measurement practice (**IFPUG function points**) — adapted to language-aware productivity benchmarks informed by the **ISBSG** industry dataset and published field studies. The estimate is anchored in observable properties of the codebase

(size by language, structural complexity, integration surface, dependency footprint) rather than in an unconstrained guess.

Three scenario factors refine that baseline:

- **Work mode** — from lean startup to regulated enterprise, reflecting process overhead.
- **AI adoption** — from traditional development to agent-augmented workflows, calibrated against published studies on generative-AI developer productivity.
- **Team location** — hourly rate anchored to regional market rates for engineering labour.

Cost range

Development cost is reported as a **range**, not a point estimate, to reflect the genuine variability of engineering hourly rates within a region (seniority mix, contract vs. employee, engagement type).

Maintenance cost

Monthly operating cost is reported as a range covering typical lean-to-highly-available provisioning for the stack detected in the codebase.

3. What the numbers mean — and don't mean

- **Scores** summarise what is *observable* at the analysed commit. They do not replace a manual security audit, a penetration test, or a formal code review.
- **Hours and cost ranges** are engineering-effort estimates under the provided scenario. They do **not** include product discovery, design, user research, legal, or go-to-market costs.
- **Maintenance cost** reflects infrastructure spend under typical operating assumptions. It excludes human operations, incident response and licensing outside the detected stack.

4. Reproducibility

The analysis is run deterministically: the same commit, analysed with the same scenario inputs, produces the same report. Variability in the AI layer is suppressed through zero-



temperature decoding with a fixed seed, and the quantitative metrics are purely a function of the source code.

This report was generated by CODEEGO LTD.

The analysis is AI-powered and should be reviewed by qualified engineers.

CODEEGO LTD · Company number 17056638

Building 3 Chiswick Park, 566 Chiswick High Road, W4 5YA

© 2026 CODEEGO LTD. All rights reserved.