



Vertesia Generative AI Platform - Service Definition

Service Overview

SynApps Solutions is a leading provider of enterprise software solutions and associated consultancy, implementation and support services. SynApps is 100% focused on delivering content centric enterprise software solutions and is generally acknowledged by vendors, SIs and its customers as an organisation that employs the foremost subject matter experts in the UK and provides the highest level of customer service.

Vertesia addresses the fundamental disconnect in the enterprise AI landscape: while organisations are excited about generative AI's potential, many fail to move beyond pilots to production deployment with measurable ROI. Vertesia's strategy is rooted in a clear market observation—enterprises are currently forced to choose between building custom infrastructure, adopting narrow point solutions, or accepting the constraints of vendor ecosystems.

Vertesia bridges the gap between intuitive low-code tools and powerful pro-code flexibility, empowering teams to rapidly build and launch secure AI apps and agents using the best of both approaches. Vertesia's AI Studio makes it easy for users to quickly turn ideas into enterprise-grade AI solutions without getting bogged down by technical complexity, while the Vertesia Platform API, SDK, CLI, and open-source repositories provide developers with the in-code access they need for deeper customisation and integration of AI solutions. This architecture is powered by systemic composability, allowing enterprises to treat every AI component as a modular building block that can be rearranged and reused across different business units without starting from scratch.

Vertesia's platform strategy centers on three foundational pillars designed for rapid application generation:

1. **Intelligent content preparation for more accurate model output.** A critical differentiator for developers is Vertesia's proprietary semantic document preparation layer and Agentic RAG pipeline. The platform provides built-in tools and a visual workflow builder that handles complex multimodal content (PDFs, images, video) that typically causes hallucinations in standard LLM implementations. This production-ready data foundation ensures that the results are accurate and usable from day one.



2. **A unified enterprise foundation for scalable AI.** Vertesia eliminates the "DIY burden" of stitching together fragmented tools by providing a pre-integrated environment where production-grade requirements are native properties, not manual configurations.
 - **Continuous observability and governance:** Vertesia provides full-spectrum transparency by capturing the entire lifecycle of an AI interaction—from raw data inputs and retrieved context to the model's reasoning and final outputs. By logging every autonomous decision alongside human interventions, we create a native audit trail that transforms "black box" AI into a compliant, observable workflow.
 - **Performance, scalability, and resilience:** Designed to overcome the inherent constraints of current LLM providers, Vertesia features a high-concurrency architecture with automated failover and intelligent load balancing. If a primary model hits a rate limit or experiences latency spikes, the platform dynamically reroutes requests to ensure zero downtime and consistent performance, allowing your AI operations to scale horizontally without manual infrastructure tuning.
 - **Robust agentic underpinnings:** Beyond simple chat, Vertesia provides the underlying orchestration layer required for long-running, multi-step autonomous workflows. This includes persistent state management, secure tool execution, and the ability to manage complex "skills" across disparate systems. By decoupling the execution logic from the underlying model, your agents remain stable and functional even as the external LLM landscape shifts.

3. **Strategic model optionality is designed for automatic failover and peak performance.** By supporting the Model Context Protocol and a diverse range of leading providers—including OpenAI, Anthropic, Google, Amazon, and Microsoft—Vertesia offers a broader and more versatile model library than any single hyperscaler. This vendor-agnostic architecture prevents lock-in and allows developers to swap models based on performance, cost, or residency requirements without rewriting application code, and the ability to reuse prompts and embeddings, ensuring the highest degree of agility in a rapidly evolving AI landscape.
 - **Universal tool & skill portability:** Vertesia standardises the integration of external Tools (executable functions like APIs and databases) and specialised Skills (packaged domain expertise and procedural logic). By decoupling these capabilities from the underlying model, any agent or app you build retains its full functionality—including custom workflows and third-party integrations—even if you swap the core inference provider.



Vertesia enables rapid prototyping that transitions to production within 2-4 weeks, compared to the 6-12 month timelines typical of custom development.

By unifying intuitive low-code interfaces with robust pro-code extensibility, Vertesia empowers business users to innovate while providing developers with the deep programmatic control required for complex integrations. This dual-track approach enables Vertesia to be the enterprise-approved choice for the next generation of autonomous AI software.

Service and Product Features

Enterprise Content Management in the Cloud provides Alfresco's enterprise content management platform for development and delivery of configured solutions.

- Deliver value from AI quickly, simply and securely
- A single platform to deliver AI applications across the enterprise with highly flexible integration options
- Deploy AI applications up to 10x faster than custom or toolkit development
- Avoid vendor lock-in and future proof AI investments
- Reduce risk from inaccurate or unreliable AI outputs
- Automate complex end to end workflows, not just simple tasks
- Meet security and compliance expectations from day one
- Lower total cost of ownership and operational effort by centralising AI application management
- Unlock value from your existing data and documents
- Full security, auditability and observability across the entire platform

Example Use Cases

- **Intelligent Content Generation:** Create, summarise, edit and enhance structured and unstructured content at a massive scale, automating tasks that were previously time-consuming and manual.
- **Advanced Data Analysis and Retrieval:** Process, analyse, and extract insights from vast amounts of unstructured data, such as documents, reports, and user/customer/consumer feedback.
- **Ensure regulatory compliance:** Simplify audit activities, enforce policy adherence, and maintain full transparency and traceability across all AI-driven decision-making.
- **Grant Processing (Philanthropic Nature Foundation):** Automates a highly manual, 8-month grant proposal review process, achieving 7x



efficiency gains and faster funding decisions through AI-powered summarisation and scoring.

- As a second step, Vertesia replaced their document-heavy review process with an intelligent AI workflow that synthesises hundreds of pages into comprehensive impact reports, unlocking the ability to manage larger portfolios while significantly strengthening donor confidence.
- **Global E-Commerce Leader:** Optimises high-volume marketing investments by analysing performance across diverse channels and vendors. The platform streamlines creative operations through automated digital asset enrichment, replacing manual tagging with intelligent metadata. This integrated approach drives data-driven decision-making and significantly accelerates the sales-to-delivery lifecycle.

Technical Features

- Scalable, serverless architecture with multi-cloud support
- Unified Low-Code Agent Builder
- Multi-model AI support and resilience
- Semantic DocPrep - Precision Document Intelligence with built in RAG pipelines
- Durable AI Orchestration for Long-Running Workflows
- Centralised Governance and Transparency of AI agent operations
- API-First integration
- Enterprise grade security and governance by design
- Ready to Use AI Assistant and UI components
- Intelligent cost optimisation and resource management

An overview of the G-Cloud Service (functional, non-functional)

Generative AI platform technology for organisations to deliver AI applications, including Agentic workflows, rapidly and efficiently.

Live service delivery of configured enterprise AI applications on a subscription model.

Information assurance

All datacentres (locales) are highly resilient Tier3, UK sovereign status.

Level of backup/restore and disaster recovery that will be provided



Organisations can choose from a range of protection levels. The default level is Local Protection - data is held in a single named UK locale and protected using server virtualisation and cloud storage which improves data durability. If organisations cannot tolerate the loss of data, they are recommended to consider the optional Remote Protection level.

Optional: Remote Protection can be chosen where virtual servers and data are stored in two UK sovereign locales, with a copy maintained in a primary named UK locale and copied to a geographically remote UK locale.

Organisations can also purchase additional backup capacity for the Cloud platform on which secondary copies of data can be automatically created and retained. Such data can also be replicated for even higher data durability.

On-boarding and Off-boarding processes/scope etc.

An organisation will need to engage in definition of their AI system requirements and the configuration / composition of the system, as well as the more detailed activities for the system to integrate to the organisation's existing business processes.

SynApps provides comprehensive on-boarding and off-boarding AI Cloud Support services through its submission to G-Cloud Framework.

Service Profile

Item	Description
Minimum Term	12 Months
Invoicing Profile	Annually in advance
Managed Service Support	Service Management 9.00 - 17.30; Mon-Fri; Excl UK Bank Holidays and including:
	Annual Disaster Recover (DR) Test
	At least three (3) Major Release Patch Deployed Annually;
	At least three (3) Minor Release Patch Deployed Annually
	Hotfixes and Security Fixes applied as required.
	Please note that End-User Account management is the responsibility of the Customer.
Out of Hours Support	Available as part of an on-call service, provided via Lot 3 AI Cloud Support Services.
System Environments	As needed

Service Management Details

SynApps Solutions provides a flexible Service Management facility, delivering the most appropriate level of support depending on the requirements of our customers and tailored to their needs.



SynApps Customer Support methodology is based on ITIL methodology, which aims to maintain a positive relationship with customers by identifying the needs of existing and potential customers and ensuring that appropriate services are developed to meet those needs. The service management is designed to be flexible and tailored to deliver the most appropriate level of support depending on the requirements of our customers.

Each customer is assigned a Server Manager who is responsible for managing the on-going customer relationship and is accountable for the success of the client's support application. The Service Manager is assigned prior to go live and works closely with the project team (customer and internal) and with project manager to ensure a smooth transition from project to business-as-usual support. The Service Manager is accountable to the head of Support and account teams for ensuring customer success.

Success is typically achieved through a variety of mechanisms a summary of which can be found below: -

- ❑ **Maintaining Customer Relationships** - SynApps continues to understand the needs of customers and establishes relationships with potential new customers. This is conducted as part of regular review meetings and reporting.
- ❑ **Identify Service Requirements** - SynApps will understand and document the desired outcome of a service and decide if the customer's need can be fulfilled using an existing service offering or if a new or changed service must be created.
- ❑ **Customer Satisfaction Survey** - SynApps will plan, carry out and evaluate regular customer satisfaction surveys. The principal aim of this process is to learn about areas where customer expectations are not being met before customers are lost to alternative service providers.
- ❑ **Handle Customer Complaints** - SynApps will continuously monitor and records customer complaints and compliments, to assess the complaints and to instigate corrective action if required.

The Service manager is backed up by our in-house support team and infrastructure project team. The service manager becomes the primary point of contact for the customer to ensure that smooth transition from the build and implementation phases of a deployment into the live operation goes smoothly. The Service Manager will be available to address and answer queries and ensure that any technical issues are handled with the highest priority.

Financial recompense model for not meeting service levels

If the service level falls below the stated availability percentage (excluding Planned and Emergency maintenance periods), the Client will be eligible for Service Credits. Service Credits will be calculated as a percentage (5%) of the



infrastructure fees for the monthly billing period during which the failure occurred (to be applied at the end of the billing cycle).

Personnel Checks

All our resources are UK based and have either undergone Baseline Personnel Security Standard (BPSS) checks or are Security Cleared (SC).

Service Constraints

No built-in service constraints.

Service Maintenance

SynApps will adhere to the following in terms of maintenance windows:

“Planned Maintenance” means any pre-planned maintenance of any infrastructure relating to the Services. SynApps shall provide the Client with at least twenty four (24) hours’ advance notice of any such planned maintenance:

- Planned maintenance of the infrastructure relating to the Services shall happen between the hours of 00:00 and 06:00 (UK local time) Monday to Sunday and/or between the hours of 08:00 and 12:00 (UK local time) on a Saturday and/or Sunday;
- Planned Maintenance shall be excluded from any availability calculation in regard to Service Credits but shall be included in the monthly service reporting;

“Emergency Maintenance” means any emergency maintenance of any of the infrastructure relating to the Services. Whenever possible, SynApps shall provide the Client with at least six (6) hours’ advance notice:

- Whenever possible Emergency Maintenance of the infrastructure will happen between the hours of 00:00 and 06:00 (UK local time) Monday to Sunday and/or between the hours of 08:00 and 12:00 (UK local time) on Saturday and/or Sunday unless there is an identified and demonstrable immediate risk to a Clients environment;

Emergency Maintenance shall be excluded from any availability calculation in regard to service credits but shall be included in the monthly service reporting.

Service Levels - Support Hours

Our service desk is operated on a 9am-17.30 (UK Time), Monday to Friday as standard. With an option for 24/7 on-call support which can be negotiated as required.

Tickets can be raised via email, our portal or via a dedicated telephone number.



Service Levels - Platform availability

SynApps offers a 99.95% availability SLA based on the availability of both the portal and the virtual server environment itself.

Service Levels - Severity Definitions

Our Service Level Agreements definitions are as follows: -

- S1 - Severe problem preventing customer or workgroup from performing critical business functions.
- S2 - Customer or workgroup able to perform job function, but performance of job function degraded or severely limited.
- S3 - Customer or workgroup performance of job function is largely unaffected.
- S4 - Minimal system impact; includes feature requests and other non-critical questions.

Service Levels - SLA Summary

Our standard service levels are applied as follows: -

- S1 - 1 Hour response time, 4 hours resolution relief/target.
- S2 - 2 Hour response time, 8 hours resolution relief/target.
- S3 - 8 Hour response time, 5 business days resolution relief/target.
- S4 - 24 Hours response time, 20 business days resolution relief/target.

Training

A range of user and administrator training options are available as standard or custom training modules can be created. These are provided by our Lot 3 - Artificial Intelligence (AI) Cloud Support Services.

Ordering and invoicing process

Billing for the service is monthly in advance. Payment can be via the following methods: Direct Debit or purchase order.

Service Lead Time

New Clients (organisations) will typically be deployed within 1 month from order. Shorter deployment times may be available and prioritised upon request.

Existing Organisations have instant access to additional usage with no notice period required.

Data restoration / service migration



Customer has ability to extract their information. A migration service is offered separately and is recommended to maintain integrity of information under retention.

Consumer responsibilities

- Definition of and configuration / composition of document management system.
- The control and management of access and responsibilities for end users.
- Organisations must also be aware of the variable nature of the billing based on usage.
- The consumer is also responsible for ensuring it applies suitable controls to its sensitive data/application.

Details of any trial service available

Demonstration service is available by agreement.