

G-Cloud 15

Service Definition Document

Lot 3: Cloud Support

Cloud ML Platform Implementation and Support (AWS & Azure)

Provider:	Zartis
Framework:	G-Cloud 15 – Lot 3: Cloud Support
Document Version:	1.0 – January 2026
Contact:	ricky.hill@zartis.com

Table of Contents

1. Executive Summary	3
2. Service Overview	4
3. Delivery Methodology	6
4. Team Structure and Roles	8
5. Case Studies	9
6. Tools and Technologies	11
7. Service Capabilities	11
8. Pricing	12
9. Support Arrangements	12
10. Security and Compliance	13
11. Engagement Governance	14
12. Social Value	16
13. About Zartis	17

1. Executive Summary

Zartis provides Cloud ML Platform Implementation and Support (AWS & Azure) to public sector organisations through the G-Cloud 15 framework. Our machine learning platform expertise enables organisations to operationalise AI and ML capabilities in secure, scalable cloud environments. We help teams move from experimentation to production, establishing the infrastructure and practices needed for responsible AI deployment.

Our service is delivered by experienced consultants who combine deep technical expertise with proven delivery experience across regulated industries. We hold ISO 27001 and Cyber Essentials certifications, and our methodologies align with NCSC Cloud Security Principles.

Key service highlights:

- Proven delivery methodology with defined phases and deliverables
- Experienced team with regulated industry track record
- Flexible engagement models including Time & Materials and fixed-price options
- Comprehensive knowledge transfer ensuring client self-sufficiency
- ISO 27001 and Cyber Essentials certified delivery

This document provides a complete description of the service, including our delivery methodology, team structure, case studies demonstrating relevant experience, support arrangements, and governance framework.

2. Service Overview

2.1 Service Description

Zartis provides Cloud ML Platform Implementation and Support (AWS & Azure) to public sector organisations seeking to operationalise machine learning and AI capabilities. Our service covers the infrastructure and practices needed to move from ML experimentation to production deployment.

We help organisations implement ML platforms on AWS and Azure, establishing the compute infrastructure, data pipelines, and MLOps practices needed for sustainable AI operations. Our consultants bring experience with ML workflows, model deployment, and the governance frameworks required for responsible AI.

The service is designed for organisations ready to scale their ML capabilities beyond experimentation. We help teams build the platforms and practices needed for production AI.

Our consultants bring experience across multiple sectors including healthcare, financial services, education, utilities, and logistics. This cross-sector experience enables us to apply proven patterns and avoid common pitfalls, while respecting the specific regulatory and operational requirements of each industry.

2.2 Target Audience

This service is designed for:

- Central government departments
- Local authorities and councils
- NHS trusts and healthcare organisations
- Educational institutions
- Arm's length bodies and agencies
- Any public sector organisation seeking cloud capabilities

2.3 Key Features

The service encompasses the following capabilities:

ML platform design and implementation (SageMaker, Azure ML)

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

MLOps pipeline and workflow automation

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

Model deployment and serving infrastructure

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

Model monitoring and performance tracking

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

Feature store implementation

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

Responsible AI framework and governance

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

Team training on ML platform usage

This capability is delivered by our experienced consultants as part of the overall service engagement. Our approach ensures alignment with industry best practices and client-specific requirements.

2.4 Benefits

Organisations engaging this service can expect:

Production-ready ML infrastructure

Organisations engaging this service can expect to realise this benefit through our proven delivery approach and experienced team.

Faster model deployment cycles via MLOps

Organisations engaging this service can expect to realise this benefit through our proven delivery approach and experienced team.

Scalable inference for demanding workloads

Organisations engaging this service can expect to realise this benefit through our proven delivery approach and experienced team.

Reduced time from experimentation to production

Organisations engaging this service can expect to realise this benefit through our proven delivery approach and experienced team.

3. Delivery Methodology

3.1 Overview

Our delivery follows a structured MLOps methodology designed for production machine learning platforms. This approach ensures model quality, reproducibility, and operational excellence.

3.2 Phase 1: ML Discovery

The discovery phase assesses ML readiness and defines platform requirements.

Activities

- ML use case identification and prioritisation
- Data readiness assessment for ML workloads
- Current ML practices and tooling evaluation
- Infrastructure and compute requirements analysis
- Model governance and compliance requirements identification

Gate Criteria

Phase 1 concludes with the ML Readiness Report presented to stakeholders for approval.

3.3 Phase 2: Platform Design

The design phase creates the MLOps platform architecture and operational workflows.

Activities

- ML platform architecture design (SageMaker, Azure ML, or custom)
- Feature store and data pipeline design
- Model training and experiment tracking workflow design

- Model deployment and serving architecture
- Model monitoring and drift detection design

Gate Criteria

Phase 2 concludes with platform architecture approval before implementation begins.

3.4 Phase 3: Implementation

The implementation phase builds the ML platform and establishes MLOps workflows.

Activities

- ML platform infrastructure provisioning
- Feature engineering pipeline development
- Model training pipeline and experiment tracking implementation
- Model registry and deployment pipeline setup
- Monitoring and alerting configuration

Gate Criteria

Implementation milestones include pipeline validation checkpoints. Phase completion requires end-to-end workflow validation.

3.5 Phase 4: Validation and Handover

The validation phase ensures platform reliability and transfers MLOps knowledge to client teams.

Activities

- End-to-end ML pipeline validation
- Model deployment and rollback testing
- Performance and scalability validation
- Knowledge transfer on MLOps practices and platform operations
- Documentation and runbook handover

Gate Criteria

Phase 4 concludes with platform acceptance and operational handover sign-off.

3.6 Deliverables

The engagement produces the following tangible deliverables:

Documentation and Governance

- **MLOps Platform Architecture:** Detailed technical blueprints of the implemented platform (e.g., using AWS SageMaker or Azure ML components), including data flow and security controls.
- **ML Pipeline Operational Runbooks:** Comprehensive guides for MLOps practices, including model retraining, deployment, and performance monitoring.
- **Responsible AI Governance Report:** Documentation detailing model governance, bias detection setup, and compliance with ethical AI principles.

Technical Assets

- **Production-Ready MLOps Pipeline:** Fully functional CI/CD/CT (Continuous Training) pipelines for automated model building, testing, and deployment (e.g., using SageMaker Pipelines or Azure ML Pipelines).
- **Codified ML Infrastructure (IaC):** Version-controlled Infrastructure as Code (IaC) templates (Terraform/CloudFormation/Bicep) for the ML environment, feature store, and compute resources.
- **Observability Dashboards:** Custom monitoring dashboards (e.g., in CloudWatch or Azure Monitor) to track platform health, model drift, and prediction latency in real-time.

4. Team Structure and Roles

4.1 Engagement Team

Engagements are staffed with experienced consultants in defined roles. Team composition scales based on engagement size and complexity.

ML Engineer

Implements technical solutions and activities. Expertise in relevant platforms and tooling. Executes implementation plans and resolves technical issues.

Data Scientist

Provides expertise in data engineering, analytics, and data platform implementation. Ensures data quality and governance requirements are met.

Solution Architect

Leads technical design and ensures architectural decisions align with business requirements and best practices. Responsible for technical quality assurance.

DevOps Engineer

Implements technical solutions and activities. Expertise in relevant platforms and tooling. Executes implementation plans and resolves technical issues.

Cloud Engineer

Implements technical solutions and activities. Expertise in relevant platforms and tooling. Executes implementation plans and resolves technical issues.

4.2 Client Team Requirements

Successful delivery requires client participation. We typically request:

- Executive Sponsor with authority to make decisions and remove blockers
- Technical Lead with knowledge of current systems
- Subject Matter Experts for relevant business areas
- Operations representatives who will support systems post-implementation
- Security and compliance representatives for reviews and approvals

5. Case Studies

5.1 Risk Intelligence: AI Platform Development

Client Context

An AI company providing news and risk intelligence to enterprises required development of their machine learning platforms for collecting, analysing, and understanding vast amounts of unstructured content.

Challenge

The client needed to develop and maintain production AI platforms processing large volumes of news and content data, implementing natural language processing, entity extraction, and risk scoring at enterprise scale.

Our Approach

Zartis engineers worked on the core news intelligence platform and RADAR risk intelligence product. The team was involved across the full development lifecycle

including planning, design, implementation, testing, deployment, and maintenance. We implemented ML pipelines using Python and Scala, with Kubernetes for orchestration.

Outcomes

- Production AI platforms processing enterprise-scale data volumes
- Risk intelligence product successfully launched
- ML pipelines delivering accurate entity extraction and risk scoring
- Platform contributed to successful acquisition by Decision Intelligence leader

5.2 Digital Health: Production ML Models

Client Context

A digital health company required enhancement of their machine learning capabilities to support clinical decision-making and patient insights.

Challenge

The client needed to deploy, monitor, and maintain production machine learning models supporting clinical operations, with appropriate governance and monitoring to ensure model performance.

Our Approach

Zartis ML engineers implemented production machine learning model deployment pipelines, established model monitoring and performance tracking, and integrated ML outputs with the broader data platform using streaming pipelines for real-time inference.

Outcomes

- Production ML models deployed with appropriate governance
- Model monitoring established tracking performance metrics
- Real-time inference capabilities integrated with data platform
- Continued model iteration improving clinical insights

6. Tools and Technologies

We utilise industry-standard tools and machine learning platforms native services:

We deploy robust MLOps workflows using SageMaker Pipelines and Azure ML Pipelines for automation, integrated with MLflow for experiment tracking and model versioning. For specialized containerized requirements, we utilize Kubeflow on Kubernetes (EKS/AKS) to orchestrate complex model training and deployment cycles.

7. Service Capabilities

This section provides detailed answers to common questions about our service capabilities, technical approach, and delivery practices.

What MLOps tools do you use? (MLflow, Kubeflow, SageMaker Pipelines)

We deploy robust MLOps workflows using SageMaker Pipelines and Azure ML Pipelines for automation, integrated with MLflow for experiment tracking and model versioning. For specialized containerized requirements, we utilize Kubeflow on Kubernetes (EKS/AKS) to orchestrate complex model training and deployment cycles.

What LLM/foundation model services do you integrate? (Bedrock, Azure OpenAI, Claude)

We specialize in integrating foundation models through Amazon Bedrock and Azure OpenAI Service, with a focus on high-performance models like Anthropic's Claude Sonnet. Our approach prioritizes secure Retrieval-Augmented Generation (RAG) architectures to ground model outputs in client-specific data while maintaining strict privacy.

Describe your Anthropic partnership and Claude expertise

As an Anthropic partner, Zartis provides specialized expertise in the Claude model family, focusing on responsible AI deployment and prompt engineering for enterprise document processing. We leverage Claude's high reasoning capabilities and large context windows to build secure AI applications that integrate with internal systems via the Model Context Protocol (MCP).

Do you implement feature stores? Which ones?

Yes, we implement managed feature stores to ensure data consistency between model training and real-time inference, primarily using the Amazon SageMaker Feature Store

and Azure Machine Learning Managed Feature Store. These implementations allow for high-scale feature reuse and low-latency access for production models.

8. Pricing

Services are priced on a Time and Materials basis using SFIA-aligned day rates. Fixed-price engagements are available for well-defined projects with clear scope and acceptance criteria. All rates are quoted in GBP excluding VAT.

Please refer to the separate Pricing Document for the full rate card and engagement pricing details.

9. Support Arrangements

9.1 Delivery Model

Services are delivered remotely by UK and EU based teams. Our consultants follow ISO 27001 standards and connect to client environments via secure VPN with multi-factor authentication.

9.2 Working Hours

Standard delivery is conducted during UK business hours: 09:00 to 17:30 GMT/BST, Monday to Friday, excluding UK public holidays. Extended support including 24/7 coverage is available by prior agreement subject to commercial terms.

9.3 Support Channels

Channel	Availability	Use Case
Email	Yes – included	General queries, non-urgent issues, documentation requests
Ticketing System	Yes – included	Issue tracking, change requests, formal support requests
Phone	No	–
Web Chat	No	–

9.4 Support Levels

Standard support includes Priority 1–5 incident management with response times from 1 business hour (P1 Critical) to 5 business days (P5 Service Requests). Enhanced support with accelerated response times and dedicated resources available at additional cost. Support operates via monthly retainer; unused hours roll over monthly but forfeit at engagement end.

9.5 Hypercare

Following go-live, we provide a standard hypercare period. During hypercare, the delivery team remains engaged to monitor systems, resolve issues, optimise performance, complete knowledge transfer, and support transition to client operations.

10. Security and Compliance

10.1 Our Certifications

Zartis maintains the following certifications relevant to secure service delivery:

Certification	Status	Reference
ISO 27001:2013	Certified	34/5700/19/10070
Cyber Essentials	Certified	7f243da7-4b18-402f-973d-90fe65086dc2
Microsoft Solutions Partner	Active	Data & AI, Infrastructure, Digital & App Innovation

10.2 Compliance Frameworks

Our delivery approach aligns with:

- NCSC Cloud Security Principles
- ISO 27001 Information Security Management
- Cyber Essentials
- UK GDPR and Data Protection Act 2018

10.3 Data Handling

During delivery:

- All client data remains within client-controlled environments
- We do not extract or retain client data beyond engagement requirements
- Any data used for testing is anonymised or synthetic
- Data handling procedures are documented and agreed at engagement start
- Data is encrypted in transit and at rest

10.4 Insurance

Zartis maintains appropriate insurance coverage:

Insurance Type	Cover Level	Provider
Professional Indemnity	£10,000,000	Hiscox SA
Public Liability	£10,000,000	Hiscox SA
Employer's Liability	£13,000,000	Hiscox SA
Cyber Insurance	£5,000,000	Hiscox SA

11. Engagement Governance

11.1 Onboarding

New engagements begin with a structured onboarding process:

- Kick-off meeting to confirm scope, objectives, and ways of working
- Access provisioning for required systems and environments
- Introduction to key stakeholders on both sides
- Communication channels and meeting cadence established
- Project governance framework agreed

11.2 Governance and Reporting

Throughout delivery, we maintain governance through:

Weekly Status Reporting

Written status reports covering progress against plan, risks and issues, upcoming activities, and decisions required.

Steering Committee

Regular steering meetings with senior stakeholders to review progress, address escalated issues, and make strategic decisions.

Risk Management

RAID (Risks, Assumptions, Issues, Dependencies) log maintained throughout delivery. Risks are assessed for probability and impact, with mitigation actions tracked.

11.3 Change Control

Changes to agreed scope are managed through a formal Change Request process:

- Change Requests documented with description, rationale, and impact
- Impact assessment covering timeline, cost, risk, and dependencies
- Client approval required before change implementation
- Changes tracked and reflected in project documentation

11.4 Acceptance and Sign-Off

Deliverables are accepted through formal sign-off at defined milestones:

- Phase gate reviews with documented acceptance criteria
- Technical deliverables validated against requirements
- Sign-off required before proceeding to next phase
- Final acceptance upon successful hypercare completion

11.5 Exit and Termination

Engagement conclusion includes:

- Knowledge transfer to client teams or successor provider
- Documentation handover
- Access revocation and security credential rotation
- Project closure meeting and lessons learned

12. Social Value

Zartis is committed to delivering social value through our G-Cloud engagements. Our approach addresses the five social value themes defined in the Social Value Model.

12.1 Tackling Climate Change and Reducing Waste

Cloud adoption inherently supports environmental sustainability by enabling organisations to leverage hyperscale providers with strong renewable energy commitments.

Our delivery model further reduces environmental impact through:

- Remote-first delivery reducing travel-related emissions
- Digital-first documentation and collaboration eliminating paper waste
- Optimisation services reducing compute resource consumption
- Guidance on sustainable architecture patterns that minimise energy usage

12.2 Creating Jobs, Skills and Growth

We contribute to skills development and employment through:

- Knowledge transfer programmes that build client team capabilities
- Training and coaching as part of engagement delivery

12.3 Tackling Economic Inequality

As an SME ourselves, we understand the importance of supporting smaller businesses:

- We actively engage SME and VCSE subcontractors where appropriate
- We support local economies through distributed team employment across UK and Europe
- Our services help public sector organisations deliver more efficient services to citizens

12.4 Equal Opportunity

Zartis is committed to diversity and inclusion:

- We maintain inclusive hiring practices and actively work to increase diversity in technology roles

- Our remote-first model enables participation from individuals who may face barriers to traditional office-based work
- We provide equal opportunities regardless of background, disability, or personal circumstances

12.5 Wellbeing

We prioritise wellbeing:

- Flexible working arrangements supporting work-life balance
- Mental health support and employee assistance programmes
- Public sector digital transformation enables better citizen access to essential services

13. About Zartis

13.1 Company Overview

Zartis is a technology consultancy established in 2009, with over 280 professionals delivering cloud transformation, software development, and digital services to organisations across Europe and North America. We combine deep technical expertise with practical delivery experience to help organisations achieve their technology objectives.

Zartis UK Limited is our UK-registered entity (Company Number: 11207666), through which we deliver services to UK clients.

13.2 Our Expertise

We have particular depth in:

- Cloud platforms: AWS, Microsoft Azure
- Cloud-native development: Kubernetes, containers, microservices
- Data platforms: Data engineering, analytics, machine learning
- DevOps and platform engineering: CI/CD, Infrastructure as Code, SRE
- Application development: Full-stack, mobile, API development

13.3 UK Client Experience

We serve a range of UK-based clients including:

- Kaluza (OVO Group) – Energy technology platform development
- Kooth – NHS-partnered digital mental health services
- SMS Group – Smart metering and energy services
- EcoOnline – Environmental health and safety software
- Nourish Care – Digital care planning software
- EML Payments – Payments technology platform

13.4 Contact Information

Legal Entity	Zartis UK Limited
Company Number	11207666
Registered Address	128 City Road, London, EC1V 2NX
Website	www.zartis.com
G-Cloud Enquiries	ricky.hill@zartis.com

We welcome the opportunity to discuss how Zartis can support your requirements. Please contact us to arrange an initial consultation.