

Valcon G-Cloud Lot 3 Service Definition

AI, ML, LLM Ops Accelerators

[Set Up and Migration | Service optimisation/right sizing]

Version 1.0

1. About Us

Valcon is a leading European consultancy specialising in AI and data-driven business transformation, with a growing and established presence in the UK. A trusted partner to major organisations including multiple UK Government departments and public sector agencies. Valcon connects strategy and operations, helping organisations create value in their transformation programmes.

With over 1,750 consultants across Europe, Valcon brings deep expertise across data, AI, technology and consulting. In the UK, Valcon works with clients across financial services, public sector, retail, industrials and infrastructure. Valcon is backed by private equity firm Rivean Capital.

2. Service Description

A service that accelerates the deployment, monitoring and lifecycle management of AI, ML and LLM solutions. Valcon enable organisations to move from experimentation to production by implementing automated, scalable and governed pipelines that improve reliability, auditability, collaboration and time-to-value across traditional ML models and cloud-based large language models.

Features

- Automated pipelines for AI, ML and LLM lifecycle stages
- Model and prompt registration and governance with change tracking
- Scalable training, deployment and optimisation workflows
- Reusable, repeatable components across solutions
- Versioning for data, code, models and prompts
- Cloud, on-premise or hybrid deployment flexibility
- Monitoring for drift, quality, hallucinations and anomalies
- Engineering practices for AI/ML platform operations, (CI/CD,, observability, feature stores...)
- Technology-agnostic frameworks (Azure, AWS, GCP)
- Optional managed service for continuous support

Benefits

- Faster progression from experimentation to stable production
- Improved performance and reliability of AI and LLM models
- Rapid iteration and controlled experimentation
- Scalable infrastructure supporting multiple AI workloads
- Transparent auditability and regulatory alignment
- Early detection of drift or model behaviour changes
- Stronger collaboration across data, engineering and IT
- Reduced operational burden through automation
- Measurable value from AI, ML and LLM investments
- Continuous optimisation via managed service options

3. Approach

Delivery Overview

- Our AI, ML and LLM Ops Accelerators are delivered as an engineering-led engagement, configured to the client's architecture, security requirements and operating model. We design and implement automated pipelines covering the full AI lifecycle, including build, training, evaluation, deployment, monitoring and retraining for both traditional ML models and large language models.
- An automated ML pipeline is structured around Build, Train and Release components and is designed to reflect specific constraints such as cloud or on-premise environments, security controls, performance requirements and individual use-case needs. The accelerators enable reproducible experimentation, controlled releases and continuous monitoring, supporting faster iteration while reducing operational risk and human error.
- The service is delivered through a structured sequence of assessment, design, implementation and handover. Pipelines are integrated into existing environments and designed to be operated and extended by client teams, with optional managed service support available for ongoing monitoring, optimisation and lifecycle management.

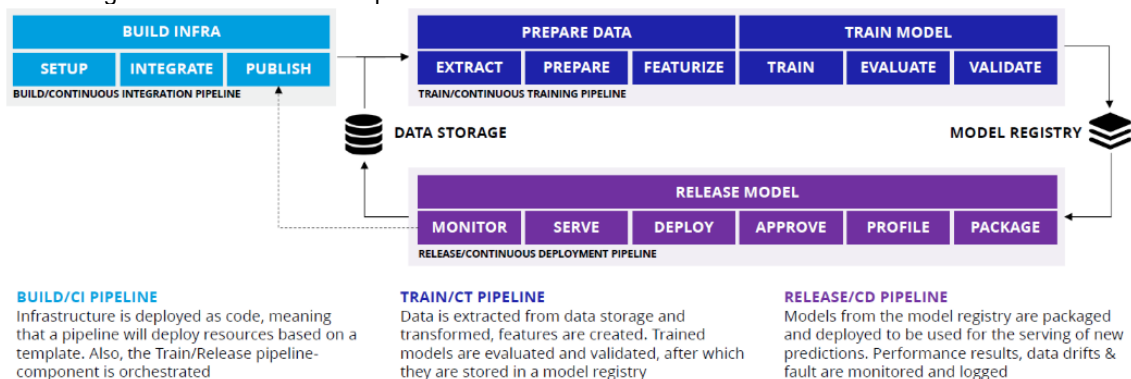
Governance, Collaboration and Quality Assurance

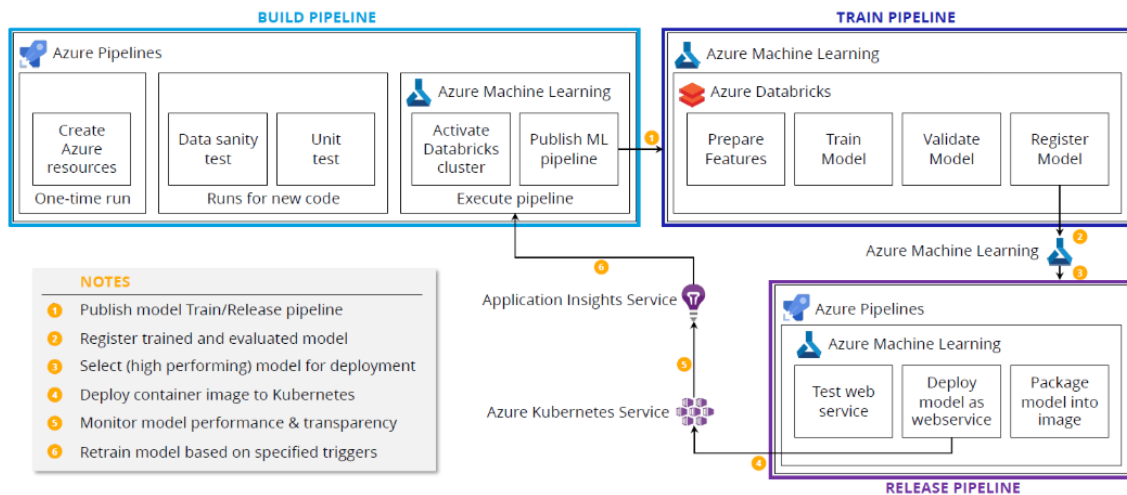
- We work with your data scientists, ML engineers, IT, architecture and security teams to implement controlled and compliant operations. Governance includes version control, approval workflows, audit trails, prompt and model risk assessments, and performance monitoring. Our accelerators include safeguards to detect drift, degradation or unexpected model behaviour (e.g., LLM hallucinations).
- Quality assurance spans automated testing, evaluation metrics, reproducibility checks and adherence to Responsible AI principles. Where applicable, delivery aligns with recognised standards and guidance, such as ISO/IEC 23894 (AI risk management), ISO/IEC 27001 (information security), ISO/IEC 42001 (AI management systems), the UK Government AI Framework, and relevant cloud provider best practices. All processes are documented and aligned with organisational and regulatory standards.

Change Management, Stakeholder Engagement and Methodology

- We collaborate with stakeholders throughout discovery, design, implementation and rollout phases. Our approach aligns with DevOps, MLOps and emerging LLMOps best practices, ensuring teams can sustain and evolve AI solutions. Training and handover activities support long-term ownership and capability uplift.
- The service may be delivered as a standalone accelerator or combined with managed service support for continuous monitoring, retraining, improvement, incident handling and lifecycle management.

The diagram below illustrates a representative example of how the accelerator structures an end-to-end automated ML pipeline across build, train and release stages, showing how data, models and deployments are managed in a controlled and repeatable manner:





4. Effort and Cost

Effort will be determined collaboratively between the Valcon and the client based on the skills mix and scope required.

Please refer to the Pricing Document for current rates.

- Rate card prices are exclusive of VAT.
- Expenses are recharged at cost and supported by receipts.
- Expenses may include reasonable travel, accommodation, and subsistence costs.
- The payment schedule will follow milestone completions as agreed in the Statement of Work.

5. Assumptions

The successful delivery of the engagement depends on assumptions and dependencies jointly agreed between the Valcon and the client.

These may include (but are not limited to):

- Access to required stakeholders, systems, or information
- Timely decisions and approvals
- Availability of resources and infrastructure
- Defined governance and communication processes

All assumptions and dependencies will be recorded and validated as part of the Statement of Work (SoW).