



**SUPER
CHARGE**

**G-Cloud 15 (RM1557.15)
Supercharge London Ltd.**

AI and Intelligent Automation Implementation Services

Service Definition

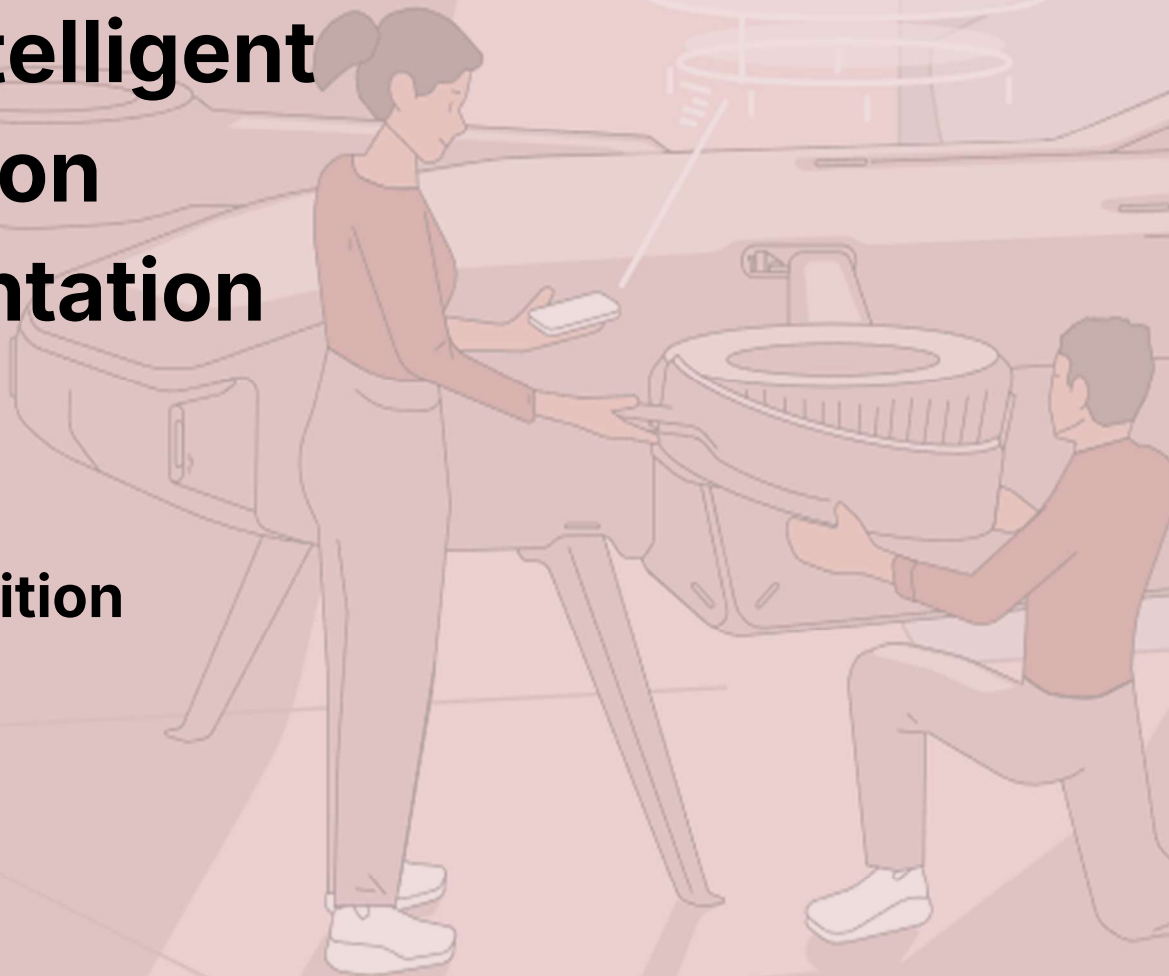


Table of contents

1 Introduction	3
2 Service description	4
2.1 About our service	4
2.2 Delivery framework.....	5
2.3 Service highlights	8
2.4 Service features	9
2.5 Service benefits	13
2.6 Service scope.....	14
2.7 User support	16
2.8 Service accessibility and usability	18
2.9 Service levels and availability	19
3. G-Cloud alignment information.....	20
3.1 Onboarding and offboarding processes.....	20
3.2 Data	24
3.3 Business continuity and disaster recovery	26
3.4 Governance	28
3.5 Security	29
3.6 Auditing, monitoring and performance	32
3.7 Training	34
3.8 Service pricing, invoicing and termination terms	35
4 Supplier information.....	37
4.1 About us.....	37
4.2 Relevant case studies	38

1 Introduction

This Service Definition document provides information about the Artificial Intelligence (AI) and Intelligent Automation Implementation Services offered by Supercharge London Ltd (Supercharge) under the G-Cloud 15 Framework. We have designed this to provide clarity regarding key aspects of our company and services, including functionality, data and security, on/offboarding and pricing, management and support.

The document comprises **3 core sections**:

- 1. Service description:** Provides an overview of our AI and Intelligent Automation Services, outlining highlights, key features, benefits for clients, delivery methodology and support measures
- 2. G-Cloud alignment information:** Details how our services are delivered under the G-Cloud 15 framework, covering purchasing, onboarding/offboarding, data and security concerns and governance
- 3. Supplier information:** Summarises our company background and how to procure our services through G-Cloud

Supporting buyers to quickly locate the information most relevant to their stage in the procurement process, the table of contents on the previous page is clickable.

2 Service description

2.1 About our service

Our AI and Intelligent Automation Implementation Services help public sector organisations move beyond AI experimentation to deliver production-grade systems that generate measurable impact. We design, build and operate governed AI and automation solutions that reduce manual effort, improve decision-making and support transparency, accountability and value for money.

The service covers the full lifecycle of AI and intelligent automation delivery, from readiness assessment and use-case prioritisation through to live operation and optimisation. This includes intelligent workflow automation, predictive machine learning models, and the safe implementation of Generative and Agentic AI with appropriate controls, human oversight and auditability. We focus on embedding AI into real business processes, particularly document-heavy and multi-step workflows, rather than delivering isolated tools or proofs of concept.

Specialising in delivering solutions for complex, regulated environments, where security, governance and reliability are critical, our AI and automation services are ideal for:

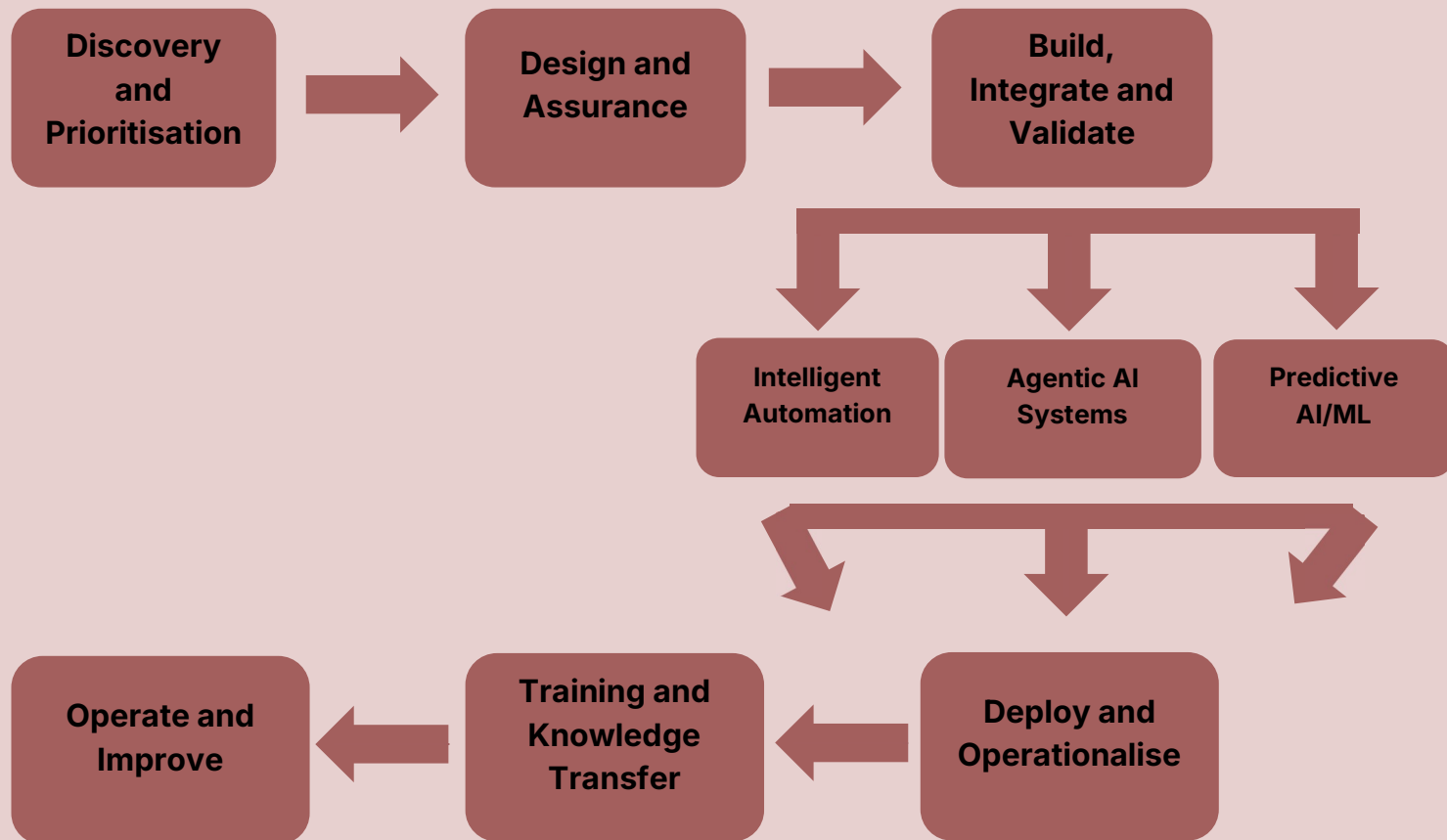
- Central government departments and arms-length bodies
- Local authorities
- NHS trusts and arms-length bodies
- Regulators
- Transport and mobility bodies
- Publicly owned utilities providers

We are part of **Siili Solutions**, a National Association of Securities Dealers Automated Quotations (NASDAQ)-traded Scandinavian group focused on digital innovation. We are proud to be working with ambitious companies from Europe, the US and Asia and helping them realise their full potential. Providing confidence of security, quality and governance throughout, delivery is underpinned by recognised standards and assurance, including:

- ISO 27001
- ISO 9001
- ISO 14001
- ISO/IEC 20000-1
- Cyber Essentials Plus
- SOC 2 Type I

2.2 Delivery framework

Ensuring a user-focused, compliant and quality solution, we deliver our services through a structured framework. Preceded by a comprehensive discovery phase, we move from prioritised use cases to production implementation, introducing solutions in Agile sprints. Our AI and Intelligent Automation Services framework is as follows:



1. Discovery and Prioritisation

To establish the current environment, desired outcomes and the steps required to move from one to the other, delivery begins with **Discovery and Prioritisation**. This phase follows onboarding, where scope, constraints and success measures are agreed. Working collaboratively with service owners, Subject Matter Experts (SMEs) and Technical Leads, we:

- Assess data availability, quality and access constraints
- Review the existing application estate and integrations
- Map critical user journeys, data flows and process bottlenecks
- Identify security, privacy and governance requirements
- Confirm organisational risk appetite for automation and AI

Prioritising the right approach (intelligent automation, Agentic AI, Predictive AI or data science, for example), we create an implementation roadmap. Focusing on delivering value quickly while managing risk, we prioritise a small number of high-impact use cases. These could include resource optimisation through Machine Learning (ML) forecasts, predicting daily staffing needs at locations, for example. As well as automation of repetitive, administration-heavy work, such as complaint management, combining intelligent automation and Generative AI models. These are selected to deliver measurable operational benefits while establishing reusable technical and governance foundations that can be scaled across our service delivery for this organisation.

2. Design and Assurance

Transforming our findings into clear improvement plans, we create and consolidate the target solution and operating model during **Design and Assurance**. This includes establishing:

- Integrations
- Workflow orchestration
- Model approach
- Monitoring procedures
- Support processes
- Governance Framework

Where Generative or Agentic AI will feature, this phase includes explicit assurance design. We define control points requiring human oversight, verification or approval, and agree on which decisions can be automated versus those that must remain human-led. This ensures solutions are aligned to governance, regulatory and ethical requirements before build begins.

3. Build, Integrate and Validate

Build, Integrate and Validate follows a comprehensive, client-approved design. We begin to build automations, implement them in Agile sprints and validate their compatibility with the client's workflows through sprint reviews. To ensure solutions meet real operational needs before progressing, delivery is iterative and collaborative, with continuous validation against agreed acceptance criteria. Following User-Centred Design principles, we involve client-side end-users, SMEs and future customers (citizens or patients, for example) to provide continuous validation and feedback, ensuring a perfect match between the solution and needs as the software evolves.

Intelligent automations

For intelligent automation use cases, we build the end-to-end workflow, including orchestration, integrations and document processing where required. Confirming accuracy, effective exception handling and full auditability, we validate each automation against real process scenarios. Ensuring the solution operates reliably and consistently in live conditions, we test performance under expected volumes.

Agentic AI systems

Where agentic AI is introduced, we implement both the agentic workflow steps and the surrounding controls. Ensuring accountability, we design clear Human-in-the-Loop (HITL) verification points, define structured hand-offs to human users and implement comprehensive logging. We run formal evaluations covering both typical tasks and edge cases, measuring quality, consistency and failure modes against agreed acceptance criteria.

To manage risk, we implement guardrails such as input and output constraints, policy checks, grounded responses where required, and defined escalation paths for sensitive actions. We establish observability through detailed trace logs of agent steps, tool calls and key decisions, supported by monitoring and alerting. This allows us to detect degraded performance or risky behaviour early, both during and following delivery.

Predictive AI/ML

For predictive AI and ML use cases, we prepare and quality-check training data before defining a baseline model approach. During the build phase, we train or fine-tune models using agreed metrics and acceptance thresholds, then validate performance using holdout data and realistic operational scenarios. Before deployment, ensuring transparency, assurance and readiness for production use,

G-Cloud 15 (RM1557.15) – Service Definition

we document model behaviour, limitations and monitoring requirements. We build with sustainable MLOps practices in mind, including automated pipelines, version control and continuous monitoring.

4. Deploy and Operationalise

Following successful validation, we **Deploy and Operationalise** solutions through an agreed release process into the client's environments. We configure access controls, logging, monitoring and operational runbooks, integrating with the client's existing incident and change management processes. Following go-live, we provide a structured hypercare period of 1 month, with 2 months of additional support, resolving issues, stabilising performance and confirming the service is operating as expected. This includes testing of AI outputs and workflow behaviour in practice.

5. Training and Knowledge Transfer

For successful adoption and seamless transition to new systems for clients and end-users, we can provide training as an optional add-on. Depending on client requirements and the complexity of systems implemented, this could be a comprehensive, intense programme of live sessions, recorded materials and tailored guides. Or, for clients with internal training teams and good AI literacy, we may provide simple guidance documents and no tailored training. The **Training and Knowledge Transfer** phase is led by Innovation Consultants and the AI Expert who fronted the implementation. This can include bespoke sessions/materials for all roles who will engage with the solution, including users, administrators, support teams and technical staff.

Enabling client teams to confidently operate the solution, we also provide training materials, clear documentation and operational runbooks as a part of offboarding. This minimises ongoing support requirements and facilitates knowledge transfer for new users.

6. Operate and Improve

Once live, we offer an optional **Operate and Improve** service. We collaboratively agree on a level of operational support aligned with the demands of the solution, which could be a complete handover with no ongoing support, or hands-on, 24/7/365 L1-L3 support.

If the client chooses for their support package to include optimisation, we proactively improve the service. Driving continuous improvement, we refine our solution based on user feedback and operational data, implementing updates or tweaks to automations and AI systems. For stability, accessibility and scalability, we prioritise the backlog, then turn to performance tuning, cost management, security strengthening and other enhancements.

If applicable, this includes ML-specific activities such as monitoring for data and concept drift, recalibrating thresholds, retraining models on fresh data, and revalidating performance against agreed metrics. Depending on whether the client opts for it in the scope, this can also include Agentic AI-specific activities such as:

- Aligning agent workflows with evolving business processes
- Expanding tool access and scope where appropriate
- Refining routing and orchestration as new tasks are introduced
- Systematically updating to newer model versions when beneficial
- Maintaining robust, repeatable evaluations to ensure quality, safety and consistency

2.3 Service highlights

Designed to maximise productivity in complex, regulated environments: We specialise in delivering AI and automation where governance, security and assurance matter most. Avoiding the common trade-off between control and usability, our services embed performance, accessibility and user experience alongside security and compliance from the outset.

End-to-end delivery from a single, accountable partner: From AI readiness assessment through to live operation, we act as a single point of responsibility, even when managing several platforms or systems.

Provider-neutral AI and cloud engineering: Avoiding vendor lock-in and supporting long-term flexibility as AI capabilities evolve, design and deliver cloud solutions across cloud and AI platforms. Selecting AI platforms and models based on security, compliance, data residency and cost requirements, we engage with a range of providers, including Microsoft Azure, AI Foundry, AWS Bedrock and SageMaker, Google Vertex AI and third-party LLM services.

AI-specific security and risk controls: We apply targeted controls for agentic and predictive AI, mitigating risks such as prompt injection, data leakage, misuse of tool-enabled agents and model drift. This allows organisations to adopt advanced AI techniques without increasing operational or reputational risk.

Long-term value without over-engineering: We prioritise high-value use cases, right-size infrastructure and seek to build reusable foundations. Supporting sustainable cost management and defensible value for money, clients retain IP ownership and control over roadmaps.

Ethical, responsible AI solutions: We design AI systems with explicit control points, human oversight and escalation paths. Clearly defining what can be automated and what must remain a human decision, we enable confident, ethical and auditable use of AI in public services.

2.4 Service features

1. AI capability assessment

We assess organisational readiness for AI by evaluating data availability and quality, existing infrastructure, security and governance controls and internal skills. Through structured workshops with Service Owners, Subject Matter Experts (SMEs) and Technical Leads, we identify constraints, risks (and risk appetite) and opportunities.



This assessment informs proportionate decisions about where solutions are optimal and whether we implement intelligent automation, Agentic AI, Predictive AI or data science solutions.

2. Use case prioritisation

To deliver value quickly and create a scalable model, we identify, prioritise and implement use cases where AI or intelligent automation can deliver instant value and build reusable foundations. Working collaboratively with client teams, we assess processes against feasibility, risk, data readiness and operational impact.



By selecting a small number of high-value, low-risk use cases, initially, we enable early validation, faster outcomes and a clear pathway for scaling AI capability over time. This could include introducing ML to domains with strong data availability (such as staffing and triage) for resource optimisation. Or, implementing AI in areas where significant administrative work can be automated, such as for document intake and data entry.

3. Secure, compliant AI architecture/platform design

Approaching the engagement strategically through our structured framework, we design the target AI architecture and operating model on the chosen cloud platform. Key deliverables include:

- An implementation roadmap
- Target AI architecture plans
- Workflow and integration designs
- Design data pipelines
- Model environments
- Deployment layers and monitoring arrangements

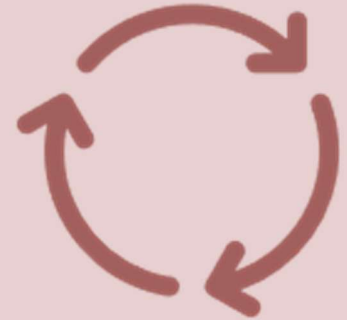
Upholding compliance, designs are aligned to the client's security, governance and data handling requirements. Adhering to Responsible AI Practices, where Generative or Agentic AI is used, we design explicit control points, with HITL verification steps, logging and human oversight.



4. Intelligent workflow implementation with automation and document processing

We build end-to-end intelligent workflows, including orchestration, system integrations and document processing where required. Implementing in structured phases aligned to Government Digital Service (GDS) methods, we deliver in Agile sprints, with continuous validation against real process scenarios and agreed acceptance criteria.

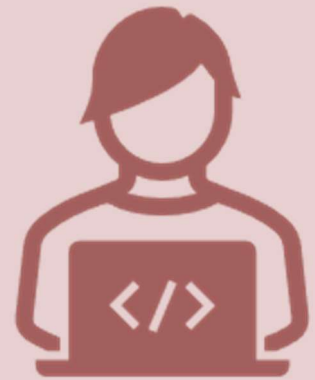
Confirming automations are robust, reliable and suitable for live operational use, we test workflows for accuracy, exception handling, auditability and performance under expected volumes as a part of the Build, Integrate and Validate phase.



5. Agentic AI workflow implementation

To streamline tasks, we implement bespoke agentic workflows that integrate Large Language Models (LLM)-driven steps with clear human handoffs and safeguards. We design and run structured evaluations covering typical tasks and edge cases, measuring quality, consistency and failure modes. We also establish observability, including trace logs of agent steps, tool calls and decisions, with monitoring and alerting to detect degraded or risky behaviour.

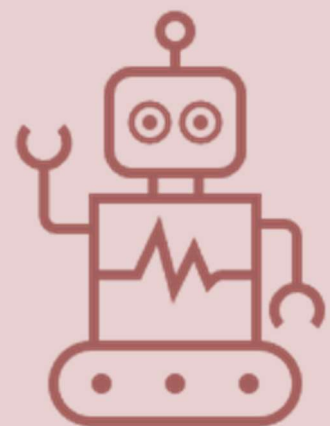
Ensuring accountability, we restrict agent actions through allowlists and least-privilege access, enforcing guardrails and defining escalation paths for sensitive actions.



6. Predictive AI and ML development

Improving decision-making, we develop predictive AI and machine learning models for forecasting, anomaly detection, resource optimisation and decision-support use cases. This includes preparing and quality-checking training data, defining baseline model approaches and training or fine-tuning models using agreed metrics. We validate performance on holdout datasets and in realistic operational scenarios, documenting model behaviour, limitations and monitoring requirements before Deployment.

Enabling early intervention, we establish ML Operations (MLOps) practices to support the reliable operation of AI systems in production. This includes model versioning, deployment controls, automated testing and approval gates, and monitoring for performance degradation or drift where relevant.



7. Embedded security and quality testing

Providing assurance throughout delivery, we integrate security and quality testing into every stage of the lifecycle. Security is treated as a vital design consideration rather than a late-stage check, with requirements defined in early stages, controls aligned to secure-by-design principles and ISO27001-certified processes. Recognising the additional risks AI poses, we maintain a suite of AI-specific security controls.

Similarly, automated and manual quality testing provides confidence that services are robust, compliant and ready for live use, reducing the risk of vulnerabilities, defects or instability. This includes rigorous testing of AI outputs and workflow behaviour as well as end-user accessibility testing.



8. Training and Knowledge Transfer

Equipping staff to operate, maintain and improve AI systems, we can provide a comprehensive Training and Knowledge Transfer phase for users, administrators, support teams and technical staff. This could include live sessions, recorded materials, tailored guides and materials, such as operational runbooks.

We enable client teams to operate, maintain and improve AI systems confidently and independently after delivery by handing over documentation and practical knowledge. This includes involving SMEs in every stage of delivery, fostering internal expertise and supporting the adoption and resolution of queries.



9. End-to-end support and ongoing maintenance

Ensuring services remain stable, secure and operational throughout their lifecycle, we provide structured end-to-end support and maintenance during and, if desired, after delivery. Designed to suit and scale with service criticality, our support is underpinned by clearly defined, ITIL 4-aligned escalation paths and response processes.

Ensuring support is available when it is most needed, during ongoing delivery, we maintain a 24/7/365 incident reporting channel, with hypercare for the Deploy and Operationalise phase. We can provide ongoing L1, L2 and L3 support with appropriately skilled Agents triaging and resolving cases in line with agreed Service Level Agreements (SLAs).



10. Post delivery optimisation

Beyond initial delivery, if contracted, we support continuous optimisation of live services, refining based on end-user feedback, performance data, usage patterns and operational insight. This includes:

- Fine-tuning performance
- Managing drift or degradation (towards bias, for example)
- Optimising cloud spend
- Strengthening security controls
- Enhancing systems to support evolving business operations
- Adjusting in line with emerging AI/ML legislation or developments



By adjusting based on real-world evidence, we ensure solutions remain efficient, resilient and cost-effective as demand and requirements change

2.5 Service benefits

Enabling clients to use AI and intelligent automations safely, proportionately and with appropriate governance, our AI and Intelligent Automation Services deliver measurable, operational benefits.

Specifically designed for complex, regulated environments, the AI systems we implement deliver benefits including:

✓ **Faster delivery of value with reduced delivery risk**

We prioritise a small number of high-value, low-risk use cases early, enabling clients to see tangible outcomes quickly while building reusable foundations for future automation. Early discovery, validation and continuous testing reduce the risk of misaligned solutions, rework or failed implementations.

✓ **Improved efficiency and productivity**

Intelligent automation streamlines manual, repetitive and document-heavy workflows, reducing operational effort and freeing staff to focus on higher-value activities. By orchestrating workflows end-to-end and integrating directly with existing systems, we remove bottlenecks and reduce manual tasks.

✓ **Safer adoption of AI in regulated environments**

Designing AI systems with explicit controls, human oversight and auditability built in, we ensure safe and secure use suitable for government environments. By clearly defining what can be automated and what must remain a human decision, we enable organisations to adopt AI confidently without undermining accountability, governance or public trust.

✓ **Better user outcomes and adoption**

Ensuring AI systems align with real workflows and user needs, we continuously validate with users and SMEs from Design to Operate. To facilitate sustained adoption, we provide comprehensive training and ongoing support.

✓ **Flexibility without vendor lock-in**

We design solutions using open standards, client-owned infrastructure and configurable platforms. This gives organisations flexibility to evolve their AI capability over time, change providers if required, and avoid dependency on proprietary tooling or single-vendor ecosystems.

✓ **Scalable foundations and cost benefits**

By focusing on reusable architecture, shared services and repeatable patterns, we enable clients to scale AI use safely across additional processes and services. This means that early investments continue to deliver value beyond the initial use cases.

2.6 Service scope

Service constraints

Providing buyers with transparency upfront, we require certain practical conditions to operate the service. As our service is delivered through close collaboration with client teams, for delivery to proceed effectively, we need:

- Timely access to client stakeholders and subject-matter experts, including business, technical and governance representatives, to support discovery, decision-making and validation activities
- Access to relevant systems, environments and documentation, as agreed in the Statement of Work, including legacy platforms, interfaces, operational processes and historical data of preexisting operations conducted
- Access to data assets and training data
- Appropriate permissions to deliver tooling and environments, such as cloud accounts, code repositories, CI/CD pipelines and monitoring or logging tools, where development, testing or deployment activities are part of the scope

Where access or availability is delayed, delivery timelines and sequencing may be impacted. We mitigate the risk of this through collaborative early planning, outlining clear dependency points with clients.

Minimising disruption, we plan maintenance windows to align with customer needs, scheduling work for the lowest-traffic times.

System requirements

Our AI and Intelligent Automation Services do not require specialist hardware or prerequisites. Our extensive in-house capabilities enable us to develop and modernise applications across all major platforms, including:

- **Mobile:** Native iOS and Android development, plus cross-platform expertise in Kotlin Multiplatform, React Native, Flutter and other modern frameworks
- **Desktop:** Windows and macOS applications
- **Server:** Windows Server and Linux-based systems

AI solutions can include a mix of components, including APIs, web interfaces and data pipelines, with system requirements depending on the client's needs. Reasonable baselines for common implementations are outlined below

- **Simple API service:**
For cloud-hosted API services
CPU: 250m (millicores) request, 500m limit
Memory: 512MB request, 1GB limit
Storage: 1-2GB ephemeral storage
- **Web application-fronted:**
Windows 10+, macOS, or Linux
Modern browsers (Chrome, Edge, Firefox, Safari - latest 2 versions)
Minimum 4GB RAM, dual-core processor
Internet: 2 Mbps minimum, 5+ Mbps recommended
- **Mobile application:**

G-Cloud 15 (RM1557.15) – Service Definition

iOS 15+ or Android 10+ Minimum 2GB RAM
Internet: 2 Mbps minimum, 5+ Mbps recommended

Where Graphics Processing Unit (GPU) acceleration is required, such as for computer vision, we size the cloud infrastructure based on the model and throughput requirements. We align this with cost and security constraints.

Integration options

Creating solutions tailored to the client's security and governance needs, we select services, platforms, hosting region and data-handling controls carefully. Depending on requirements and the sensitivity of the information, we can use third-party LLM and multimodal LLM services, for example, OpenAI APIs. Alternatively, we can deploy models through managed cloud platforms such as Azure AI Foundry and AWS Bedrock.

Designed to seamlessly integrate with existing systems, and using secure, standards-based approaches, integration options commonly include:

- **Model hosting and AI platforms:**
 - Microsoft Azure AI Foundry
 - AWS Bedrock and SageMaker
 - Google Vertex AI
 - Third-party LLM services, such as OpenAI, Claude and Google Gemini (subject to approved connectivity, data-handling and governance controls)
- **Authentication and identity:**
 - Azure Active Directory (Entra ID)
 - Amazon Cognito and Auth0
 - SAML 2.0-based single sign-on (SSO) providers
 - OAuth 2.0 and OpenID Connect
 - Multi-Factor Authentication systems
- **Enterprise and line-of-business systems:**
 - CRM platforms such as Salesforce and Dynamics 365
 - CMS platforms including DatoCMS, Strapi and WordPress
 - SAP and other finance, ecommerce and accounting systems
- **Data and analytics:**
 - Data warehouses such as Snowflake, BigQuery and Redshift
 - Business intelligence tools including Power BI and Tableau
 - Analytics platforms such as Google Analytics
- **Payments and financial services:**
 - Payment gateways including Stripe, GoCardless and Worldpay
 - Open Banking APIs
- **Monitoring and observability:**
 - Grafana, Prometheus and Loki
 - Sentry and Firebase

Where the modernised service integrates with buyer-owned third-party systems, the buyer is responsible for maintaining any required licences or subscriptions. For example, Microsoft 365 for Dynamics 365 integrations or OpenAI licences.

2.7 User support

For efficient and effective user support, we provide a structured, responsive service aligned to ITIL4 best practices. **Supercharge Support Team** comprises:

- **23 L1 Technical Support Agents**
- **6 L2 Reliability Engineers**
- **L3 Application Engineers, assigned per project**

These agents are responsible for providing **24/7/365 service**. Additionally, each client is assigned a dedicated **Service Operations Manager**, who acts as their key escalation contact. All of our Operated Services are overseen by our **Lead Service Operations Manager** and by our **Head of Operations**.

Offering a tailored approach, support is designed to scale based on buyer needs. Both during delivery and for ongoing assistance, our model can be adjusted from standard (for example, 9 am – 5 pm, Monday-Friday availability) to full 24/7 coverage.

As the level of ongoing support required is best understood once AI workflows and automations are live with integrations in place, we typically contract this separately. This way, we can provide an initial estimate, then agree on the support scope and commercials based on operational needs and usage.

Support channels

To manage customer interactions, we use Zendesk's telephony and web chat management software. Leveraging our omnichannel tooling and capabilities, buyers can contact the Supercharge Support Team through:

- Email
- Online ticketing (tied to emailing services)
- Telephone, using Zendesk's built-in telephony and call centre tooling
- Web chat, through Zendesk, with information provided by our Service Management Toolkit

Ensuring visibility and auditability, we log, prioritise and track all support requests through the ticketing platform. Promoting agency, clients can manage the priority status of tickets by raising them as a priority with their Service Operations Manager.

Support availability

We can provide up to 24/7/365 support for L1 and L2, with standard working hour coverage for L3. However, response and resolution times may vary depending on customer demand and the agreed support package. Supercharge Support Team's maximum availability is as follows:

- **L1 (Technical Support):** 24/7/365 availability in standard 8-hour shift patterns
- **L2 (Advanced Infrastructure Support):** 24/7/365 on-call rotation
- **L3 (Application Support):** 9am-5pm CET, Monday to Friday, Hungarian Business Days

Response times

We operate a priority-based support model with response and resolution times defined for Priority 1 to Priority 4 (P1–P4) incidents and requests. Our standard service level commitments are as follows:

L1 – Technical Support:

Technical Support Agents are available 24/7/365, with staff working in 8-hour shift rotations. P1-4 Technical Support queries are **responded to immediately** and **resolved at the port of call**. If immediate resolution is not possible, the Agent passes the query to the L2 or L3 team. This is consistent across weekdays and weekends.

L2 – Advanced Infrastructure Support:

- **P1:** Response within 1 hour, resolution within 8 hours
- **P2:** Response within 4 hours, resolution within 24 hours
- **P3:** Response within 24 hours, resolution within 5 days
- **P4:** Response within 2 days, resolution within the next scheduled release

These response and resolution timescales are consistent across weekdays and weekends.

L3 – Application Support:

- **P1:** Response within 4 working hours, resolution within 2 working days
- **P2:** Response within 1 working day, resolution within 3 working days
- **P3:** Response within 2 working days, resolution within 5 working days
- **P4:** Response within 5 working days, resolution within the next scheduled release

L3 support is primarily available on weekdays during working hours CET, Hungarian working days. However, providing backup for emergencies, we maintain a wider resource pool of Developers who can be reallocated to resolve urgent L3 issues.

Escalation and management

Our Service Operations Managers function as key escalation contacts. We apply ITIL4 processes across incident management, major incident handling, request fulfilment and problem management. For ongoing engagements, we manage and refine our support approach through:

- **Weekly ticket review calls** with an operational focus
- **Monthly service review meetings**, with a tactical focus
- **Quarterly service reviews**, with a strategic focus

2.8 Service accessibility and usability

We embed accessibility and usability throughout the design, delivery and operation of our Service. Considering accessibility from the outset, we develop user-centred designs during **Discovery and Prioritisation**. Through workshops with Service Owners, SMEs and Technical Leads, we gain insight into current processes and what users and operators like and dislike about them, including any accessibility/usability issues.

During the **Design and Assurance** phase, we co-create solutions with our client, considering user journeys and real-person feedback from Discovery and Prioritisation. For Generative AI and Agentic AI integrations, we confirm with our client which processes can be automated and which require human decisions.

We perform accessibility testing throughout **Build, Integrate and Validate**, with automated and manual accessibility checks forming a part of our quality assurance process. Validating usability and compatibility continuously, we implement in Agile Sprints, using test cases and agreed acceptance criteria. Any potential barriers are identified and removed proactively.

Deploy and Operationalise typically includes an intensive hypercare period, where we monitor the solution closely, seek to stabilise performance issues and resolve problems quickly. To improve usability, we confirm whether the service is operating as expected and make refinements if not.

When opted for, our **Training and Knowledge Transfer Phase** is user-centred. We coach relevant roles through jargon-free and accessible live/remote training sessions. Additionally, to enable effective day-to-day adoption and seamless use of our services, we produce user-friendly, printable handover and operation packs. This supports useability of the products we design.

During **Operate and Improve**, if contracted, we continue to improve accessibility and usability based on user feedback and operational insight, feeding improvements into the prioritised backlog for ongoing enhancement.

Ensuring users can engage with support in a way that best meets their needs, all service interfaces and user-facing support channels, including online ticketing and web chat, meet accessibility standards. We aim for anyone accessing with assistive technology, such as screen readers/keyboard navigation, to have a successful and positive experience.

Our provider, Zendesk, maintains a robust Product Accessibility Policy that **exceeds the WCAG 2.2 AA industry standard**. Built on Zendesk's accessible Garden design system, our support functions have undergone both manual and automatic testing. Web chat accessibility testing, conducted on our behalf, includes:

- Formally consulting with agents, admins and end users who rely on assistive technology to collect feedback
- Accepting customer/end-user feedback through a variety of channels
- Using all feedback to inform improvements
- Retesting products following remediation and new feature updates
- Engaging third-party auditors to conduct regular compliance audits of their products

2.9 Service levels and availability

Providing a tailored approach, service levels and availability are agreed in collaboration with the client following Go-Live, to ensure these are bespoke to the final product and client needs. Service performance is monitored continuously by the Service Operations Manager and reported in an agreed upon format at an agreed upon frequency.

Fostering transparency, we present our performance during weekly ticketing review calls, monthly service review calls and quarterly service reviews.

A breakdown of key service level commitments we measure as a minimum can be viewed below:

Metric	Service level/target	Measuring and reporting
Service availability	Agreed per engagement, typically 24/7/365 for L1-2 and 9 am-5 pm CET, Hungarian Working Days for L3	We monitor tickets and report response and resolution times during weekly reviews with clients
Incident response times	Incidents are prioritised from P1-P4, with response targets defined for each support level (L1-L3)	All incidents are logged using our omnichannel ticketing platform. Response and resolution times are tracked automatically through Zendesk and reported on during review calls and meetings
Incident resolution and request fulfilment times	Resolution/fulfilment times are defined by and measured against the hourly/daily targets we set for each priority (P1-P4) as part of standard service level commitments	Ticket lifecycle data is captured within Zendesk, resolutions are logged manually, and performance is assessed during review meetings
Downtime	99% uptime	Major incidents are tracked end-to-end, with status updates, post-incident reviews and corrective actions documented and shared for future prevention

Preventing slippage, the Service Operations Manager proactively monitors performance against these service levels. If they identify risks, such as a rising number of incidents responded to outside of agreed-upon timescales, they will take action immediately. For example, assigning additional Support Desk Agents to the contract.

Reacting appropriately to any breaches of SLA targets, slippage or failure to meet agreements triggers internal escalation to the Service Operations Manager. They will manage escalations on a case-by-case basis, developing and implementing action plans for rectification with minimal disruption. This process is informed by ITIL4 best practices.

Driving continuous improvement, we will use slippage or breaches as a learning experience, dissecting root causes, the effectiveness of our resolution and outcomes to inform future delivery.

3. G-Cloud alignment information

This section explains how the service is delivered, managed and supported once a buyer has purchased it through the G-Cloud framework. This includes how we onboard and support buyers, manage data securely, ensure service resilience, monitor performance, and provide training, support and controlled offboarding. Together, they set out how client and supplier teams work in partnership to deliver outcomes safely, efficiently and in line with G-Cloud 15 requirements.

3.1 Onboarding and offboarding processes

Onboarding

To align expectations, each engagement begins with a collaborative onboarding process, where we discuss all the engagement details to align with the client's expectations. By clarifying the objectives, scope, governance and approach before activity begins, we reduce risk and establish a collaborative foundation for successful delivery.

Our onboarding process includes:

Internal and external alignments

1. **Scope of Works (SoW) confirmation:** Meeting with key representatives from the client's side, we establish scope, deliverables, objectives, constraints and success indicators (KPIs and SLAs). This provides a shared reference point for measuring progress and performance, aligning expectations.
2. **Commercial confirmation:** As our pricing model is tailored to each engagement, based on the scope and required outcomes, we cement the commercial terms following SoW confirmation. This includes defining the exact Fixed Price or Time and Materials (T&M) quote and invoicing preferences. At this point, we draw up and co-sign a contract with the client.
3. **Indicative timeline & team allocation:** As part of our pre-sales discussions, we share a prospective timeline to confirm how quickly we could deliver the project based on the scope. Though our initial timeline is indicative, it can estimate project length and volume. We plan resourcing internally, selecting the best team for the project based on experience and skill sets and securing them for the duration, and sync internally to prepare for kick off.

Project kick-off

4. **Kick-off session:** After internal preparation and onboarding the team, we book a project kick-off with the client. The goal of this session is to meet, present the detailed project approach, ways of working and planning, and align on next steps for efficient progress. It includes:
 - a. **Team introductions:** Based on the SoW, we create a dedicated team comprising a Project and Account Manager and all the Digital Experts (with project-specific specialities required for the project). We introduce our key personnel and ask our clients to assign and introduce key people, including:
 - A named product owner or point of contact to manage day-to-day alignment, unblock issues and give official approvals on the milestones
 - Optional: Business and technical stakeholders for decisions, prioritisation and approvals (if they will be closely involved)

G-Cloud 15 (RM1557.15) – Service Definition

- Subject Matter Experts (SMEs) for discovery interviews, requirements clarification and regular testing (who can also be assigned without joining the kick off)
 - Optional: End users for user research, usability testing and accessibility testing
- b. **Access and setup:** The client provides access and appropriate permissions for documentation, environments, assets and tools. During the kick-off, we align on which accesses are needed for this specific project. When we are creating new products, we typically require access to previous research materials and related tools, as well as data sources and data sets to be used for training, inference and automation. For AI/automation improvements to existing solutions, we require access to:
 - Current architecture
 - APIs
 - Data models and data quality documentation
 - Data sets to be used for training, inference and automation
 - Security policies and key sensitivities
 - Operational runbooks
 - Cloud accounts
 - Repositories
 - CI/CD
 - Logging/monitoring software
- c. **Ways of working:** We establish communication preferences and details for key points of contact on all sides, outlining how the project milestone reviews and approvals will be handled. For robust planning, we also outline client input needed and establish their availability for meetings and ad hoc communication/queries.
- d. **Delivery planning:** We agree on indicative timelines and establish key milestones and review points in line with the agreed delivery lifecycle. For example, outlining indicative dates for the completion of each phase.

Project opening: Discovery

5. **Discovery:** Completing onboarding, we begin the Discovery process, co-creating requirements and outlining in collaboration all the product features and user journeys to create the desired solution. During Discovery, we apply our Change Control Procedure for any required amendments to the SoW, so we can ensure all is fully detailed, planned and quoted before starting the Design and Development phases. The Discovery phase length varies depending on the scope, size of the solution and client's previous preparation work: it can last anywhere from a week to several months.

Offboarding

After the initial phase (including Discovery, Design, and Development of the solution agreed upon), we reconvene with the client to discuss further support for lifecycle product enhancements or maintenance/troubleshooting.

If we are to continue partnering, for optimal collaboration, we discuss a new support agreement, which may include:

G-Cloud 15 (RM1557.15) – Service Definition

- Production environments support and monitoring to ensure good performance (with our L1, L2 and L3 services)
- Product maintenance and enhancements, with a dedicated or ad-hoc team of digital experts depending on the client's demands

If support, maintenance and optimisation functions will be provided in-house, with the client's own teams, we conclude the engagement with a structured handover and offboarding process to:

- Facilitate a smooth transition into operation
- Enable sustained, effective day-to-day use
- Confirm we have met all deliverables and targets
- Transfer full ownership to the client
- Ensure organisation-wide knowledge, understanding and adoption
- Revoke our access to systems and data

Offboarding typically comprises the following steps:

1. **Project closure review:** After official Go Live, we arrange a project closure review with all key stakeholders to confirm all deliverables have been provided/objectives met and assess performance against agreed success criteria (KPIs and SLAs).
2. **Optional training:** Supporting adoption, offboarding could include an **intensive training programme** of knowledge-sharing sessions. Enabling client teams to use the platforms without external support, our experts deliver live remote and in-person sessions, with recorded materials to support long-term continuity.
3. **Documentation:** We provide a central Handover Pack containing comprehensive technical documentation, such as:
 - Architecture diagrams
 - Infrastructure-as-code
 - Runbooks
 - Operational procedures
 - All credentials, integration points and dependencies
 - AI-specific artefacts, such as model documentation, evaluation criteria and monitoring dashboards
 - Training materials

Exact documentation and training materials to be provided are agreed with clients.

4. **Data return:** Following our **ISO27001-accredited procedures**, we identify all storage locations, extract data in agreed formats and validate completeness. We either transfer ownership of cloud resources or securely destroy data, verifying through destruction certificates
5. **Account closure:** Revoking our access to your systems, we remove service accounts and integrations, decommission temporary infrastructure and provide formal, co-signed confirmation of access withdrawal and data deletion. We retain audit logs in accordance with our Document Retention Policy, which is outlined in the original contract, keeping anonymised client data for a required period.
1. **Intensive post-Go-Live care:** Whether the client opts for support or not, we offer our a hypercare month after Go-Live, ensuring we keep our team available for any unexpected issues.

- 6. Ongoing support assessment:** Ensuring post-engagement assistance proportionate to operational usage and client needs, we assess and agree on a continued optimisation and advisory support package through a new SoW if the client wishes to continue the collaboration. Our post-go-live agreements may include:
- **Production support**, with our L1 and L2 teams available on a 24/7 window and our L3 team, during working hours
 - **Maintenance and Optimisation:** we make changes and enhancements after Go-Live based on user feedback and operational data, either with a dedicated team or ad-hoc

3.2 Data

Data importing, exporting and deletion

Data is shared as part of the engagement. However, rather than being uploaded or provided through a portal, our clients provide access to data sources, systems and datasets required to deliver the project, with the necessary permissions.

Clients may request data extraction at any point during the engagement or at contract end through a written request to the provided contact email. To minimise access and risk, all data handling follows a least-privilege approach. By default, data is provided as database exports or CSV files, using open, non-proprietary formats that support portability and reuse. Where required, to meet client preferences, export formats can be adjusted. To ensure data confidentiality and integrity, all exports are delivered through encrypted transfer.

Following ISO27001-accredited procedures, we use certified secure wiping methods for deleting media and cloud-stored data, offering physical destruction with certificates where applicable. Additionally, **active systems are purged within 30 days of contract end** or deletion request, with **backups deleted 90 days post-contract** (unless legally required to be longer).

Data at rest

Protecting data at rest, we ensure secure data storage and location cloud infrastructure through AWS, Google Cloud Platform and Microsoft Azure. For flexibility, Regions can be configured based on client requirements. Sensitive data at rest is encrypted under A3/A4 classification, and we store business data exclusively on approved corporate systems, such as Google Drive and approved SaaS platforms. Backups to cloud storage are securely encrypted and tested annually by our Data Protection Officer.

To maintain data sovereignty, meet regulatory requirements and reduce exposure risk, we configure AI training and inference to run in the agreed geographic region where supported. We avoid cross-region processing unless explicitly approved. Ensuring consistent control over where sensitive information is stored and processed, operational logs, monitoring data and backups are also retained in-region by default. Where third-party model APIs are required, we only use approved providers and configure them to comply with the client's data-handling, retention and security policies.

Securing assets, we implement access controls, including Okta Multi-Factor Authentication (MFA), which is enforced across all systems. Our Password Policy requires all passwords to:

- Contain a lower-case letter and an upper-case letter
- Contain a number (0-9) and symbol (e.g., !@#\$%^&*)
- Not contain part of the username, first or last name
- Not be a 'common password'
- Have a minimum of 8 characters
- Not be one of the user's last 4 passwords
- Be changed every 90 days
- Lock after 20 unsuccessful attempts (brute force protection)

G-Cloud 15 (RM1557.15) – Service Definition

We apply the **Principle of Least Privilege** for access to all assets and store privileged credentials in a secure vault. Minimum necessary access is granted per project, per developer, and logical segregation is maintained between client environments. Preventing cross-client exposure, we segregate client data by storing it under the given cloud provider (AWS/GCP/Azure) in a separate account dedicated and owned by the client. We maintain detailed logs highlighting who has access to which client data, which are available for auditing and review by senior staff.

Our Network Security includes:

- Network segmentation, separating development, testing, and production environments
- Border firewalls to disable unused ports and protocols
- Real-time malware protection with automatic updates
- Vulnerability patching, occurring within 30 days of disclosure

Client data is prohibited on personal devices or personal cloud storage, and remote access on authorised devices requires VPN connection through Tailscale, our trusted partner.

Data in transit

To protect data in transit, we use mandatory TLS/SSL encryption. All confidential email attachments and sensitive data transfers are encrypted, with passwords sent by a separate email or communication. Minimising risk, we only use mock or anonymised data in non-production environments.

Sub-processors

Owing to the nature of the services, we may use sub-processors for effective delivery. We engage these entities in line with our ISO 27001-accredited ISMS and GDPR obligations, and only where necessary.

Our primary sub-processors typically include:

- **Cloud infrastructure providers:** AWS, Google Cloud Platform and Microsoft Azure. Hosting regions (UK, EU or global) are configurable to meet client data residency, security and compliance requirements
- **Development and collaboration tools:** Used where required for project delivery, such as secure communication, design and code management tools (for example, Slack, Figma, or code repositories and CI/CD platforms). Where possible, we use buyer-provided tools and environments
- **AI platforms and model providers (where applicable):** Managed AI and machine learning services within the selected solution. Only where explicitly approved by the client, third-party model APIs may be configured to meet agreed data-handling, retention and security requirements

We apply a strict data minimisation approach, ensuring sub-processors only process the minimum data necessary to perform their function and only for the duration required. Reducing risk, no sub-processor is permitted to use client data for any purpose beyond the contracted service.

3.3 Business continuity and disaster recovery

The service is designed to be highly available, resilient and recoverable. We ensure continuity of service even in the event of infrastructure failures, security incidents or other adverse circumstances as specified in our Operational Ways of Working document. Our 24/7 teams, who are always available to investigate and resolve issues when flagged, further support our continuity approach. Evidencing the robustness of our services and policies, we have never encountered a long-lasting major service disruption or business continuity issue.

Because delivery of the service does not rely on a single physical location or system, business availability is unlikely to be compromised. Work is delivered using secure, cloud-based environments provided by major platforms with regions configurable to UK or EU requirements. We maintain backup resources, allowing capacity to be flexed and reallocated quickly if individuals, locations or systems are unavailable.

Strong security controls, including encryption, access management, continuous monitoring and incident response procedures, further reduce the likelihood of disruption caused by security incidents. This combination of decentralised resources, cloud infrastructure, resilient system design and mature operational controls significantly lowers the risk of continuity incidents and supports rapid recovery.

For business continuity and resilience, we implement the following measures:

- **Incident response:** For rapid rectification of issues, we maintain documented incident response procedures. This includes our Business Continuity Plan (BCP) and Supercharge Operational Procedures. For ongoing preparedness, our Head of Operations reviews our BCP annually, and it is tested at least annually. Ensuring consistent knowledge across our teams, we provide incident training upon induction, refresh learning annually and confirm agreement with these policies. Additionally, our Service Operation Teams (including all service levels) have a recurring tabletop exercise system on a half-yearly basis.
- **Redundancy and recovery:** Enabling systems to continue to function despite failures, we practice redundant system design. This includes load sharing, using multiple cloud regions/availability zones and backup systems. For rapid recovery, we maintain documented and formalised system restore and reboot processes.
- **Offsite storage:** Encrypted cloud backups are stored geographically separate from primary systems.
- **Crisis management:** We maintain a 24/7 incident reporting channel, with the Supercharge Support Team aiming to rectify issues within 1-4 hours (P1) – 2-5 days (P4). This includes L1, L2 and L3 live service capability from expert Agents and Engineers.

Disaster recovery and rectification

For timely restoration of service following an incident, we maintain formal disaster recovery targets, these are as follows:

- **Recovery Point Objective (RPO):** Up to 24 hours using standard daily backups (configurable to shorter intervals)
- **Recovery Time Objective (RTO):** System-dependent and verified through annual recovery testing

G-Cloud 15 (RM1557.15) – Service Definition

Upon client request, critical systems can be configured with shorter RPO/RTO targets through more frequent backups and redundant architecture, which would incur additional costs. Ensuring predictable and repeatable recovery, our system restore and reboot procedures, including timescales, are documented, tested and formalised.

3.4 Governance

Accreditations and certifications

Demonstrating the robustness of our management system and alignment with industry standards, the following accreditations, certifications and frameworks underpin our services and delivery:

- **ISO 9001:** Quality Management Systems
- **ISO 14001:** Environmental Management Systems
- **ISO/IES 20000-1:** IT Service Management Systems
- **ISO 27001:** Information Security Management Systems (ISMS)
- **Cyber Essentials**
- **Cyber Essentials Plus**

Adding value, we uphold secure management of customer data in our delivery processes by aligning with SOC 2. Applying secure-by-design principles and SOC Trust Services Criteria, we protect data with resilient systems, particularly suited to public security needs. Evidencing compliance with this framework, we completed and passed a **SOC type I audit** in 2022.

Maintaining robust governance and ensuring AI systems are used ethically, we uphold Responsible AI Practices. These include:

- **Explainability:** Providing clear explanations of how results are generated to support transparency, auditability and informed decision-making. Where models influence outcomes, we document decision logic, inputs, limitations and confidence thresholds
- **Fairness and bias monitoring:** We assess potential bias risks during design and validate model behaviour using representative data and agreed evaluation criteria. Where applicable, we monitor outcomes over time to identify unintended bias or performance drift and implement corrective actions
- **Appropriate human oversight:** For both client-facing systems and AI used within our own delivery processes, we ensure human oversight and controls. We define which actions can be automated and which require human approval, implement escalation paths for edge cases or sensitive decisions and ensure ultimate accountability remains with named individuals, rather than automated systems
- **Privacy and security controls:** Our solutions and methodology include data minimisation, least-privilege access controls, secure model hosting, encrypted data handling and alignment with client security policies and regulatory requirements. Third-party AI services require explicit approval, and may be configured to meet agreed data protection and retention standards

3.5 Security

Protecting customer systems and data and supporting regulatory compliance, we build security into every aspect of our delivery. Our ISO27001-accredited ISMS, Cyber Essentials Plus certification and SOC Type I assurance evidence robust security management across people, processes and technology.

Combining preventative controls, active monitoring and defined response procedures, we instil confidence that security risks are managed proactively, with incidents managed efficiently and effectively.

Operational security

Minimising the likelihood of security incidents, we embed operational security controls into day-to-day delivery and service management. Live services are protected through structured practices, including:

- **Secure change and release management:** Reducing the risk of vulnerabilities or service instability during updates, we manage changes to systems and configurations through controlled processes. Our methodology is supported by peer review, testing and approval workflows
- **Vulnerability and patch management:** We protect systems against emerging threats and reduce exposure to known weaknesses by actively monitoring for newly disclosed vulnerabilities, applying security patches within defined timeframes. Additionally, we partner with external vendors, such as Silent Signal or the client's preferred vendor, for annual penetration testing
- **Monitoring, logging and auditability:** Enabling early detection of unusual activity, we centrally log and manage security-relevant events across environments, including access activity and system changes. Logs are protected from tampering, retained in line with our Data Retention policy, and made available for audit or investigation on request.
- **Incident detection and response:** We maintain a 24/7/265 incident reporting channel, accessible through email or phone, with L1, L2 and L3 live service capability. Adhering to our crisis management strategy, we follow the defined incident response procedures outlined in our Supercharge Operational Procedures and BCP
- **Tightly controlled accesses:** We secure our accesses via MFA on all platforms, and we use time-based access levels (Just-in-Time Access) for production-level systems. In addition to technical measures, we review our access control and our access policy on a half-yearly basis.

AI-specific security controls

To protect AI-enabled services in production, we apply security controls specifically designed to address the unique risks associated with generative AI, agentic systems and predictive machine-learning models. Reducing exposure while enabling safe, effective use, these controls are embedded throughout Design, Build, Deploy and Operate stages.

For agentic and generative AI solutions, we actively mitigate threats such as:

- Prompt injection, including indirect injection through documents or web content

G-Cloud 15 (RM1557.15) – Service Definition

- Data poisoning
- Model inversion
- Misuse of tool-enabled agents

We treat all external content as untrusted by default and restrict agent capabilities through explicit allowlists and least-privilege credentials. To ensure sensitive actions receive appropriate checks and human approval, we enforce guardrails, verification steps and defined escalation paths.

To reduce the risk of data leakage, we apply context minimisation and redaction rules and use grounded responses where required, meaning models only act on approved, relevant information. We validate these controls through structured testing, including red-teaming exercises that simulate real-world misuse.

For predictive AI and ML solutions, we protect training data integrity through:

- Clear provenance, versioning and controlled access
- Secure MLOps pipelines with automated testing, approval gates and deployment controls
- Monitor for model drift and anomalous inputs

Across all AI-enabled services, to provide transparency and support investigations, we maintain comprehensive, auditable logs covering agent steps, tool calls, model versions and deployments.

Staff security

Ensuring only suitable, trusted individuals can access customer systems and data, we apply strict personnel security controls from pre-employment to live work. This includes:

- **Pre-employment screening and confidentiality:** For security from day 1, we conduct comprehensive background screening for all staff before employment. Upon induction, we require them to sign confidentiality agreements confirming intent to comply with our policies
- **Security-cleared personnel:** For public sector engagements, we confirm/provide any additional security screening requirements. For example, UK-based experts screened in line with BS7858:2019 and up to and including Developed Vetting (DV), subject to authority requirements
- **Security training and awareness:** Fostering company-wide awareness of security legislation, controls and responsibilities, we mandate information security and GDPR training as a part of induction. For ongoing alignment, we deliver annual refresher training
- **Role-based, least-privilege access:** Minimising risks, we grant access to client systems strictly on a minimum-necessary basis, aligned to each employee's role and project responsibilities. We review access regularly and remove it promptly when no longer required. For any production-level access, we use Just-in-Time access methods
- **Segregation of client environments:** We prevent potential errors and misuse by limiting cross-client access, maintaining logical separation between customer environments and teams

Secure development

To reduce security risk from the outset, we design and build systems using secure-by-design principles and a structured secure development lifecycle. We:

G-Cloud 15 (RM1557.15) – Service Definition

- **Embed security from the outset**, with security requirements and sensitivities gauged early and data protection, access control and risk mitigation measures built into solutions. For example, establishing which systems contain sensitive information which cannot be shared with LLMs
- **Apply secure coding and review practices**, applying secure coding standards, peer review and automated testing to identify defects, vulnerabilities and unintended behaviours early. For AI and agentic components, this includes validating inputs, outputs and failure modes against agreed acceptance criteria
- **Conduct independent testing**, with regular security audits and annual penetration testing, assuring that effective, best practice measures are in place

Identity and authentication

Preventing unauthorised access to systems and data, we enforce strong identity and authentication controls across all environments. Our approach ensures that only verified users can access services, that access is appropriately restricted and that all authentication activity is auditable. We protect identity and access through:

- **Multi-Factor Authentication (MFA)**: Enforced across corporate systems using centrally managed identity services
- **Strong password and credential management**: Applying password complexity requirements, minimum length controls and regular credential rotation, with privileged credentials stored securely in a managed vault to prevent exposure or misuse
- **Role-based access control**: Assigning system permissions based on defined roles and responsibilities with least-privilege access, adjusted when roles change or engagements end
- **Just-in-Time Access control**: For production-level environments, we distribute access based on a request policy, and we grant access for a certain time period (e.g. 3 hours)
- **Controlled administrative access**: Protecting sensitive systems, we tightly restrict administrative privileges and require additional authentication controls for elevated access. All privileged actions are logged and monitored to support auditability and accountability
- **Regular access reviews**: Maintaining ongoing assurance, we conduct periodic access reviews to confirm that permissions remain appropriate, removing dormant or unnecessary access to minimise risk

3.6 Auditing, monitoring and performance

Providing transparency, accountability and regulatory assurance, we maintain comprehensive audit and monitoring capabilities across all environments where client data is processed or accessed. These controls ensure our/user activity is traceable, reviewable and available to support governance, compliance and investigation requirements.

Audit logs

Enabling all significant activity to be traced, reviewed and evidenced, we maintain clear audit logs in a searchable index. For a tailored approach, we confirm the audit logs to be maintained and processes to be followed with the client during mobilisation, adjusting our approach as necessary. Typically, our audit logs capture:

- **User access and authentication activity:** Enabling visibility of who accessed systems and when, we log authentication events and access attempts to support security monitoring
- **System and configuration changes:** Reducing risk and supporting accountability, we log changes to system configurations, deployments and environments, ensuring that operational actions are traceable to authorised individuals or processes
- **Data processing and handling activity:** Supporting GDPR compliance and data governance, we log data processing activities and access to client data, allowing data owners to audit usage and respond to data subject requests where required
- **AI actions and versions:** To provide transparency and support investigations and compliance, we maintain comprehensive, auditable logs covering agent steps, tool calls, model versions and deployments
- **Supplier and administrative activity:** Providing full transparency, all supplier access to client environments is logged and can be audited, reviewed and extracted on request

Protection, retention and access

Protecting the integrity and confidentiality of audit data, logs are centrally managed and protected from tampering under our ISMS. Ensuring consistent control over where sensitive information is stored and processed, logs are retained in-region by default, avoiding cross-region processing unless explicitly approved.

Retention is managed in line with our Data Retention Policy, contractual agreements and GDPR requirements, with retention periods defined and enforced for all data categories. We anonymise customer data following engagement completion.

Fostering transparency, we extract and provide audit logs to clients on request to support compliance reviews, investigations or assurance activities.

Performance monitoring and metrics

Proactively managing our performance, progress and any issues/risks, we leverage our centralised monitoring capabilities and report results to clients. Providing accountability, Service Operations Managers are responsible for monitoring, measuring and reporting on the service, providing updates at regular check-ins. The frequency and format of check-ins/reviews, as well as communication channels and protocols, are agreed collaboratively with clients during mobilisation. For immediate

G-Cloud 15 (RM1557.15) – Service Definition

visibility of unforeseen changes, they also provide ad hoc updates using the client's preferred channel, enabling real-time awareness and resolution, where necessary.

Maintaining visibility of operations and tracking performance data, we monitor:

- System behaviour
- Access activity
- Operational events across environments
- Incidents and response times
- Milestones achieved
- Anomalies and issues
- Slippage risks

To report on performance and metrics for each engagement, we hold regular review meetings/calls. This typically comprises:

- **Weekly ticketing review calls**, with an operational focus
- **Monthly service review calls**, focusing on key issues, action plans and tactical improvements
- **Quarterly service reviews**, with a tactical/strategic and improvement focus

At these reviews, the Service Operations Manager presents our reports, highlights progress, addresses concerns and identifies scope for improvement. We are receptive to client feedback, using their comments to inform the next delivery period.

3.7 Training

Supporting effective adoption and long-term value, we offer training and knowledge transfer as an optional add-on, sitting within our Go Live and offboarding process. We quote training separately during onboarding or following development, allowing clients the flexibility to choose to add or remove this element based on budget and needs.

To ensure client teams understand the solution and can operate it confidently, we offer tailored training and support options. Depending on the client's requirements, training could be a 2-month intensive programme delivered through live sessions with comprehensive materials. Or a simple, basic guidance sheet, if the client decides to take training, maintenance and support in-house, for example.

Training approach

Delivering relevant, usable and proportionate training, we design training activities around real workflows, user roles and operational responsibilities.

Training is typically delivered through a combination of collaborative sessions and supporting materials, including:

- **Role-based training and walkthroughs:** Helping users understand how the service supports their day-to-day responsibilities, we provide guided walkthroughs of functionality, workflows and operational processes. Sessions can be remote or in-person, and are tailored for different audiences, such as operational users, administrators and technical stakeholders
- **Knowledge transfer during delivery:** Providing a strong foundation for successful adoption, we co-create solutions with the client, minimising and streamlining knowledge-transfer requirements. Our discovery workshops, design reviews, acceptance testing and handover sessions ensure client teams gain understanding incrementally, rather than training beginning during the Live phase
- **Technical and operational enablement:** Supporting client technical teams, we provide structured handover sessions covering system architecture, CI/CD pipeline usage, integrations, data flows, operational scenarios and support/escalation processes
- **Recorded materials and accessible guides:** Enabling access to information as and when required and continuity as teams change, we provide recorded materials for ongoing consultation
- **Internal champions:** By involving key personnel in ongoing user research and testing, supplemented by comprehensive training, we create in-house internal champions who can support adoption, train colleagues and answer queries

Reinforcing learning and supporting self-service, we provide clear, accessible documentation as part of training and handover. Tailored to the solution and audience, materials may include:

- User guides and operational documentation
- Process flows and configuration guidance
- API documentation and integration guides (where applicable)
- Runbooks and support handover materials for live services
- AI-specific artefacts such as model documentation, evaluation criteria and monitoring dashboards

3.8 Service pricing, invoicing and termination terms

Pricing

Providing transparency and flexibility for public sector buyers, our commercial approach includes both Fixed Price and Time and Materials (T&M) models under G-Cloud. Full pricing details are published separately in the **Pricing Document** provided as part of our submission.

As a services-led provider, we rarely offer fixed, off-the-shelf packages. We tailor each engagement to the client's needs, agreeing on the scope, deliverables and commercials based on the required outcomes. Our models are as follows:

- **Time & Materials (T&M):** Suitable where requirements are evolving or where discovery and iteration are required. Providing flexibility to reprioritise as needs change, clients pay only for the time and roles used
- **Fixed Price:** Suitable where scope and deliverables are well-defined, we agree on pricing upfront based on outcomes and deliverables. Providing cost certainty for clearly bounded work, we define clear acceptance criteria, requirements and deliverables which we use for the fixed price quote

Typical UK day rates vary depending on role and seniority. Supporting value for money, we sometimes offer volume-based discounts for larger teams or longer-running engagements, subject to agreement.

Providing all-inclusive, transparent quotations, we capture all client requirements upfront, which minimises surprises and unplanned costs. If additional requirements arise during delivery, we apply our formal Change Control Process, costing all variations and mutually agreeing the charge where they fall outside of the original scope and quote.

As the level of ongoing support required is best understood once the solution is live, we typically contract this separately following Deployment. This enables us to provide a package and price aligned to the client's actual needs, whether they choose to receive ongoing support, maintenance or optimisations.

Invoicing process

For strong financial management and compliance with public sector requirements, we agree on invoicing arrangements upfront. We take a flexible approach to invoicing cycles, operating according to the client's chosen commercial model and preferences. Typically:

- T&M engagements invoiced **monthly** in arrears
- Fixed Price engagements are invoiced **against agreed milestones**

In line with UK government requirements, standard payment terms are 30 days. We typically send invoices against a Purchase Order (PO) and expect transfers to the bank account detailed on our invoices.

Contract duration and termination terms

Reducing commercial risk for buyers, there is no minimum contract term. Our engagements may span weeks or years.

Designing termination agreements to be fair, transparent and proportionate, we maintain procedures for natural and early termination:

- **Natural termination:** Where contracts end at the end of their natural term, we carry out our formal off-boarding process, revoking access to systems and sending all final deliverables to the client, or performing any training if the client has opted for this
- **Early termination:** Buyers may terminate an engagement early, and specific terms for each engagement will be agreed and outlined during mobilisation. No additional early termination fees apply. However, we require payment for work completed to date and throughout the notice period
- **Early termination notice:** We request notice proportionate to the team size and project scope, typically between 1 and 3 months for scale down, with larger teams requiring a longer notice period. This period will be agreed at the kick-off

4 Supplier information

4.1 About us

Recognised by the **Financial Times (FT1000)** and **Deloitte (Fast50)** as one of the fastest-growing tech companies in Europe, we are a next-generation product innovation agency driving growth, efficiency and resiliency for our partners. Cutting through business, tech and regulatory complexity, we fuse behaviour-based design with unique data competence to propel businesses, offering:

- Software engineering
- Product design
- Data and AI services
- Strategy services
- Managed support services

Turning some of the most challenging digital initiatives into success stories, we believe in the power of technology to enable human achievement, drive progress and be a force for positive change. Every solution we deliver is shaped by clients who choose to innovate and redefine the future, and we're proud to support them on that journey.

Since our inception in 2011, when Supercharge was founded by 4 Engineers in Budapest, we have expanded to operate internationally with **200+ in-house digital experts**. We are part of **Siili Solutions**, a NASDAQ-traded Scandinavian group focused on digital innovation, employing **1,000+ staff**. With offices in **London, Amsterdam, Helsinki** and **New York**, we are fortunate to partner with the most ambitious companies across **Europe**, the **United States** and **Asia**.

Demonstrated through our achievement of the following accreditations, our work is underpinned by quality, security and sustainability:

- **Cyber Essentials**
- **Cyber Essentials Plus**
- **ISO 9001:** Quality Management Systems
- **ISO 14001:** Environmental Management Systems
- **ISO/IES 20000-1 2018:** IT Service Management
- **ISO 27001:** Security Management Systems
- **SOC 2 Type I**

4.2 Relevant case studies

Our end-to-end delivery and ongoing support service has generated measurable impact for companies across a range of sectors. Evidencing our successful implementation of AI and Intelligent Automation services and production-grade systems, we have provided the following case studies:



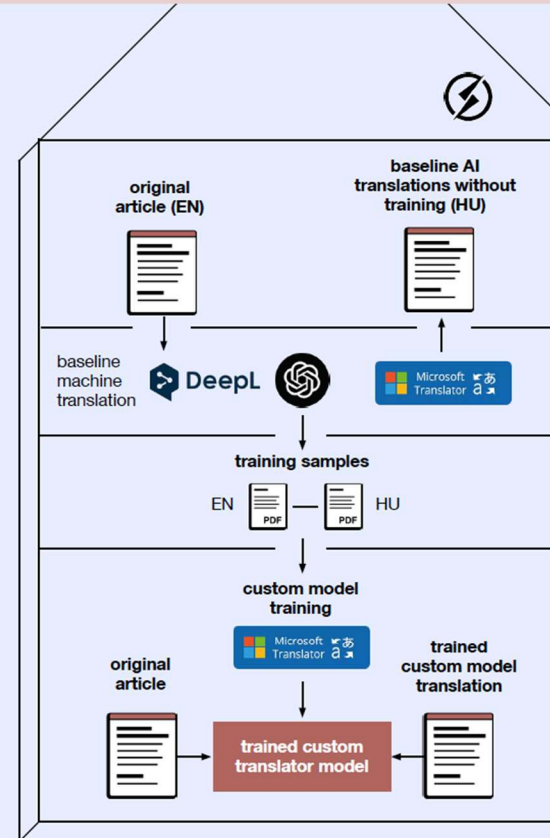
machine translation following specific domain nomenclature

General purpose language model can't produce medical translations that meet academic standards. We built a **purpose built generative AI translation system** for a **global pharmaceutical client** to translate medical research articles from English to Hungarian while preserving strict medical terminology and nomenclature.

The **goal** was to create an **exact translation** of a pdf article from **English to Hungarian** using the **specific domain nomenclature** that Pfizer provided through reference documents

For the best possible result, we selected a **flexible AI** model as the base, which we further **trained on 10,000 example sentences** for the translator model to follow.

Through several iteration cycles, we **achieved a high quality of translation** for these documents. With our custom training solution, we **improved the baseline model accuracy by 40%** (measured using the BLEU score).

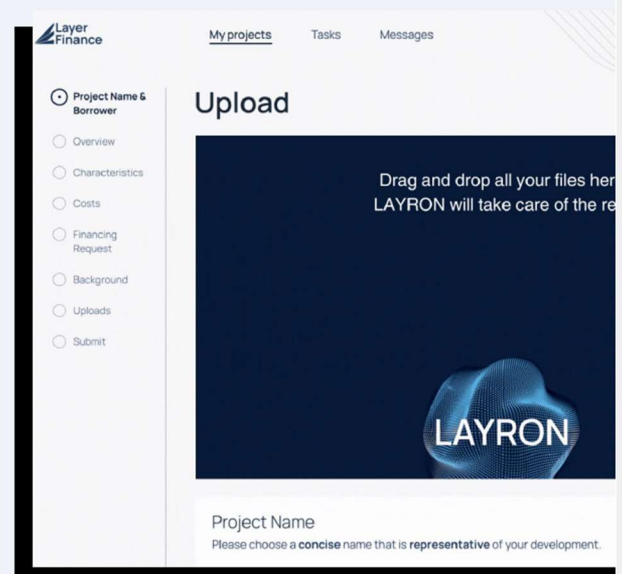


making real-estate deal flows more efficient with AI

Austrian company Layer Finance needed a new platform to digitalise their B2B workflow and establish a solid foundation for future growth through a white label solution. The platform is an ambitious first step in their digital strategy, **connecting real estate developers seeking various forms of capital with professional investors**.

We digitalised and **streamlined their operational processes**, tailored to real estate developers, investors.

We removed a key adoption barrier by adding a **purpose built, LLM driven project onboarding AI-system**. Developers upload their existing project documents, and the AI extracts and summarises the key details to populate the structured deal view. This speeds up onboarding and improves data quality.




VYRD

helping insurance agents to tackle any policy-related questions within seconds

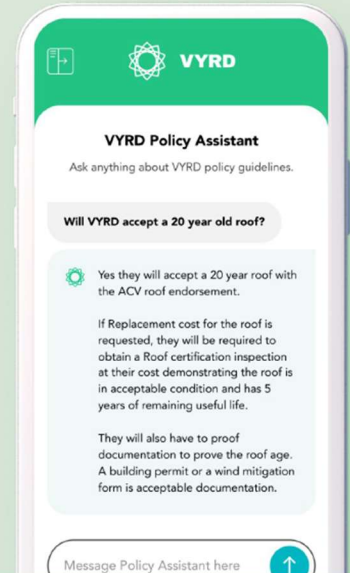
VYRD is a Florida-based, innovative home insurer that enhances its coverage with smart water-leak sensors, helping customers prevent damage before it happens.

We've created VYRD's Policy Assistant to **support VYRD's insurance agents by providing instant, accurate answers to policy and guideline questions**. VYRD's Homeowners Program Manual is a dense 40+ page document - challenging for agents to search quickly, but an ideal knowledge source for an LLM to retrieve relevant details on demand.

The chatbot **handles both basic and more complex questions**, and when it lacks high confidence in an answer, it prompts the user for clarification and/or directs them to customer support.

The chatbot was developed using a **Retrieval-Augmented Generation (RAG) approach**, grounding its answers in the **Homeowners Program Manual** to ensure accuracy and consistency with VYRD's policies.

To **ensure usefulness** in real-world scenarios, the tool **was co-tested with VYRD agents and refined based on their feedback**. Future versions will extend its coverage to additional stakeholder groups, such as adjusters and customers.


Kodak alaris

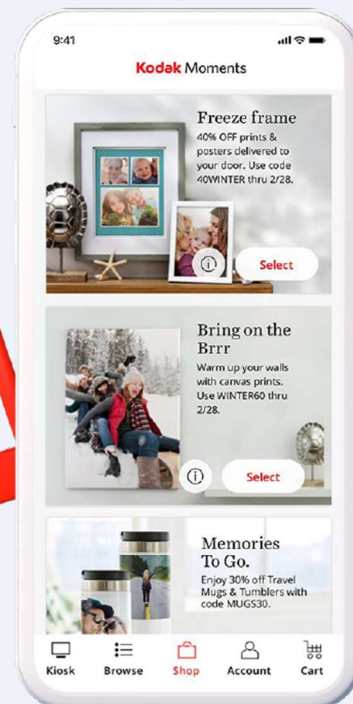
finding Kodak Moments with a single tap

We trained a **computer vision ML-model** in the Kodak Moments app to classify photos on the user's device and automatically select the best ones for an album, creating true "Kodak Moments."

The main photo editor is supplemented by a smart photo picker, a core feature which is driven by **on-device Machine Learning**

We **trained a model** on thousands of labelled images to identify the best photos based on criteria like number of faces, open eyes, smiles, and strong landscapes. These are the "Kodak Moments."

Training was done server side using **Keras** and **TensorFlow**, while **CoreML** and **MLKit** were used on the client to run the models.



ML-based forecasting

gategroup

data-driven operations transformation

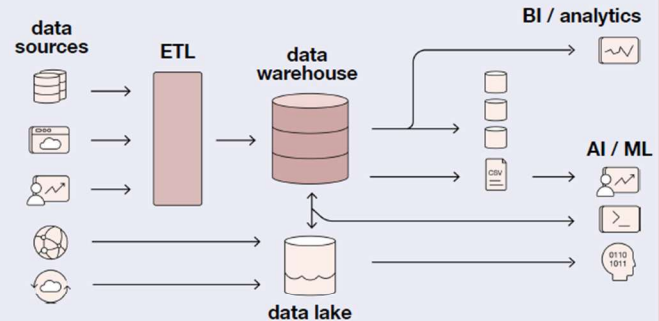
GateRetail embarked on a journey to transform their retail and supply chain operations with a data-driven approach. The project's goal was to create real impact by funneling back data insights and forecast results into the everyday decision-making process to improve revenue, and operational efficiency. Besides setting the direction for the data strategy, the Supercharge team is developing automated solutions with machine learning at the core to create efficiency in how GateRetail supplies 100,000s of flights.

Step 1: Created a sales demand analyses to understand the customer behaviour, reveal hidden patterns including seasonality, pricing elasticity, promotion effectiveness.

Step 2: Built the right data architecture and consolidating all important data points to have a feature store that serves as the **foundation for our forecasting solution**.

Step 3: Developed the ML driven forecasting capability that both supply chain and operations teams can embed into their internal processes to save cost and drive revenue.

Modern data platform as a solid foundation



8

FFD

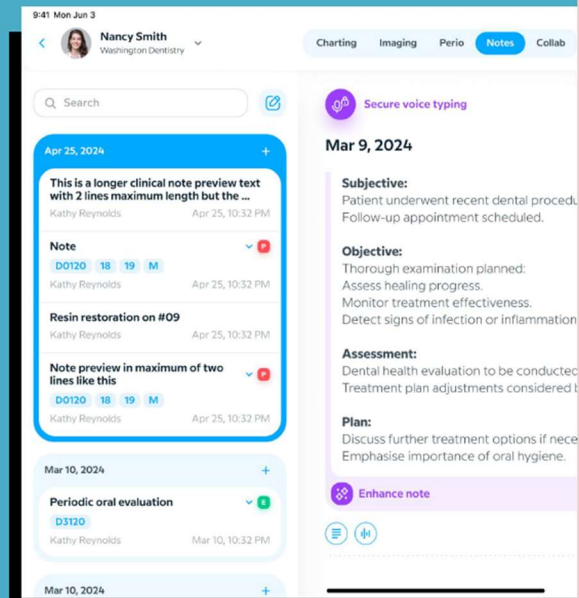
AI-powered dental note-taking

In the FFD dental application, we used Whisper AI together with **LLM based extraction and structured summarisation** to help dental practitioners **turn voice recordings into well structured transcripts** and generate summaries in a standard format.

Revolutionizes clinical note-taking with **secure, high-accuracy AI voice typing** that captures medical details while **filtering out other conversations** (operates reliably in a noisy background)

The captured details of the visits are **post processed by a purpose built LLM system** that corrects terminology, removes filler words (um, uh, er), and produces a **structured summary in a predefined clinical report format**.

The dental collaboration platform **smartly integrates** with current dental management systems that lack modern APIs, powered by a unique AI Agent based automation solution,





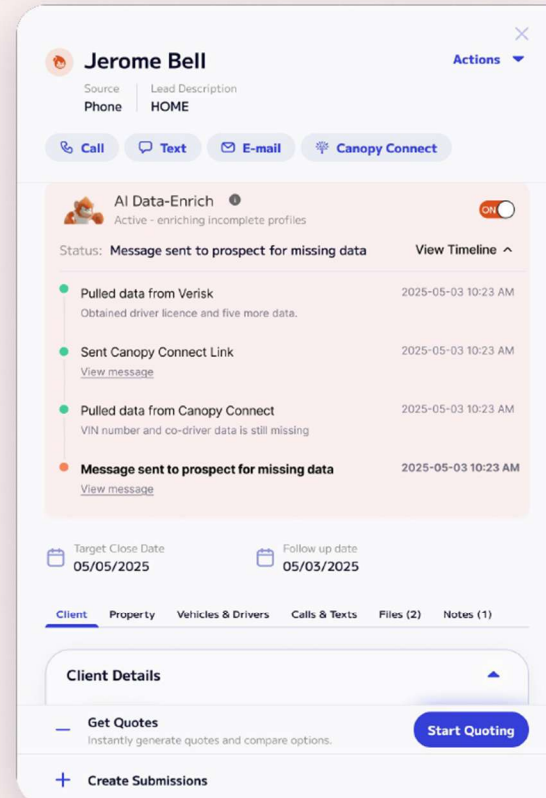
transforming insurance agency workflows with Agentic AI

Quote Monkey focuses on transforming how **insurance producers** work in the US by **empowering them with AI "superpowers"** to convert more leads and provide better policies for their clients. The platform achieves this through various **AI Agents that act as digital assistants for the user**, conducting research to target leads more effectively while simultaneously automating repetitive, time-consuming tasks to allow for greater focus on core responsibilities.

↖ The **Sales Support AI Agent** empowers insurance producers by **conducting online market research and analyzing customer context** and CRM data to generate tailored sales strategies and conversion-optimized call scripts.

↖ The **Lead Data Enrichment AI workflow** **autonomously contacts prospects via text or voice to gather missing information** and extracts data from their responses, while simultaneously **pulling from third-party sources** to ensure producers can offer the best possible policies.

↖ **AI Agents specializing in document generation and parsing** enable insurance producers to **complete complex documentation tasks** in minutes while ensuring their CRM data remains precise and up-to-date.



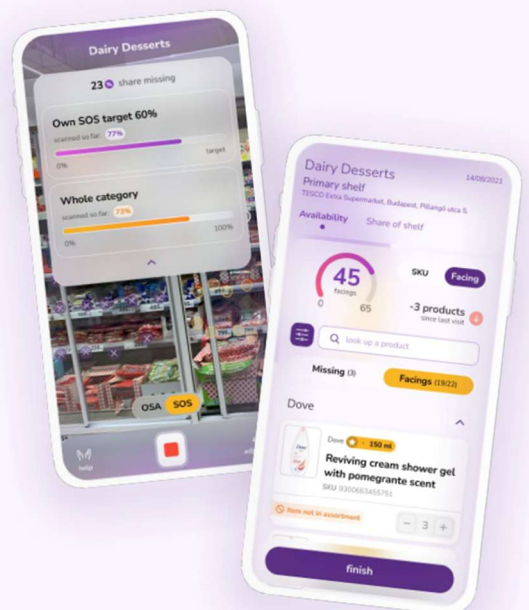
real-time AR-based shelf analytics

The TraxNow AR application allows **brand ambassadors in the retail industry to quickly audit their category**. Through the clean and simple interface ambassadors gain **real-time product placement insights** using the **on-device data processing** of the application. This allows them to get their jobs done quickly regardless of network coverage. The immersive onboarding process enables an effortless and speedy roll-out to staff members globally.

↖ **At the core** of this is a **robust AI model** developed by Supercharge, **trained to accurately recognise goods** such as chocolates and other products across varied shelf conditions.

↖ **Real-time scanning** accommodates the ever-changing environment of stores and it also **reduces the time spent on auditing** tasks.

↖ Supercharge's custom **object recognition engine** enables precise and responsive shelf analytics, significantly reducing the time spent on auditing tasks.



4.3 How to Buy

Supercharge's services can be procured directly through the Crown Commercial Service (CCS) G-Cloud 15 Digital Marketplace.

Buyers can search for Supercharge on the Digital Marketplace and procure the service through direct award, selecting the 'AI and Intelligent Automation Implementation Services' listing.

Once selected, the buyer and Supercharge will:

- Raise a Call-Off Contract under the G-Cloud framework terms
- Agree a Statement of Work (SOW) detailing scope, deliverables, timescales and pricing
- Issue a Purchase Order (PO) to confirm commencement

Ensuring transparency and compliance with public-sector procurement regulations, all engagements are delivered in line with G-Cloud framework conditions.

Enquiries or pre-contract discussions can be arranged through our central contact point:

Contact: Laura Rodriguez, Managing Director

Email: laura.rodriquez@supercharge.io

Telephone: 07425 248371