

DEPLOYMENT REQUIREMENTS

On-Premises Hardware and Environment Requirements

EmbodiFlow Embodied-AI Data Platform · Deployed in the customer's own data centre

DOCUMENT NUMBER	VERSION	ISSUE DATE	LENGTH
EF-DR-2026.09	V3.0	2026-09-09	10 pages

Overview

Procurement and budgeting

Section 2 to choose a version, Section 8 for equipment and localised alternatives, Section 9 for the full list of deliverables.

IT and infrastructure

Sections 3 to 7 are the itemised specifications: architecture, servers, storage, network and ports, operating system and software.

Project and implementation

The checklist, responsibility matrix and acceptance tests in Section 10 can be used directly for scheduling and sign-off. Print that page separately.

Contents

01	Deployment Requirements	Cover · Audience · Contents	1
02	Choosing a Version	Two versions compared · Common deviations	2
03	Reference Architecture	Components · Responsibility boundaries	3
04	Workstation Version	Minimum · Recommended	4
05	Cluster Version	Node roles · High availability and failure domains	5
06	Storage and Capacity Planning	Data model · Data protection · Growth · Sizing	6
07	Network and Ports	Topology · Firewall · Bandwidth and latency	7
08	Runtime Environment and Software	OS · Open-source components · Localisation	8
09	Delivery and Operations	Deliverables · Upgrades · Backup · Monitoring	9
10	Prerequisites and Acceptance	Checklist · Responsibility matrix · Acceptance tests	10

02

Choosing a Version

Table 1 · The two versions compared

Item	Workstation version	Cluster version
Reference model	 Dell Pro Precision 9 T6	 Dell PowerEdge R770
Intended use	Single team, centralised capture	Multiple teams in parallel, shared across departments
Robots capturing at once	5 or more	50 or more
Concurrent users	50 or more	500 or more
Server count	1 tower workstation	2 application + 2 storage + 1 training, 5 in total
CPU / memory	From 16 cores / 64 GB, 32 cores / 128 GB recommended	Allocated by role; see Table 5
GPU memory	From 24 GB, 48 GB recommended	From 48 GB, 96 GB recommended
Capture data disks	From 2 × 16 TB, 8 × 32 TB recommended	12 × 32 TB per storage node, 2 nodes
Data protection	Disk-level redundancy; tolerates 1–2 disk failures	Node-level redundancy; tolerates the loss of 1 whole server
Planned net capacity ^①	13 – 146 TB	From 307 TB, scales linearly per node
Deployment method	Docker Compose	Kubernetes

^① Net capacity already excludes data-protection overhead and the 80% planning watermark. See Section 6 for the model.

Common deviations

Table 2 · Common deviations and their effect

Change	Effect	Feasibility
Using existing S3 object storage	Removes the storage-node investment	Supported; see Section 3
Adding a GPU later	Requires slot, power and cooling headroom	Supported; reserve capacity at selection
Grow to 4 or more storage nodes	Enables erasure coding, adding about 43% net capacity on the same disks	Supported; see Section 6
Workstation version to cluster version	Requires a data migration and recabling; business configuration is preserved	Supported, assisted by a delivery engineer
One cluster spanning two data centres	Inter-node latency exceeds database and storage requirements	Not supported

IMPORTANT

The workstation version has no whole-server redundancy

Redundancy protects disks, not servers. A mainboard, power-supply or memory failure takes the platform offline, and recovery typically takes several hours to a working day. If that window is unacceptable, choose the cluster version.

NOTE

Two storage nodes are enough for disaster recovery

Each node holds a complete copy of the data, so either one can fail outright without interrupting reads or writes — and it is the cheapest way to start. At four or more nodes you can switch to erasure coding for about 43% more net capacity from the same disks.

03

Reference Architecture

Every component runs inside the customer's data centre. Data, annotations and model weights all reside on customer-owned storage. The platform needs no outbound internet access and sends nothing off site.

Capture side

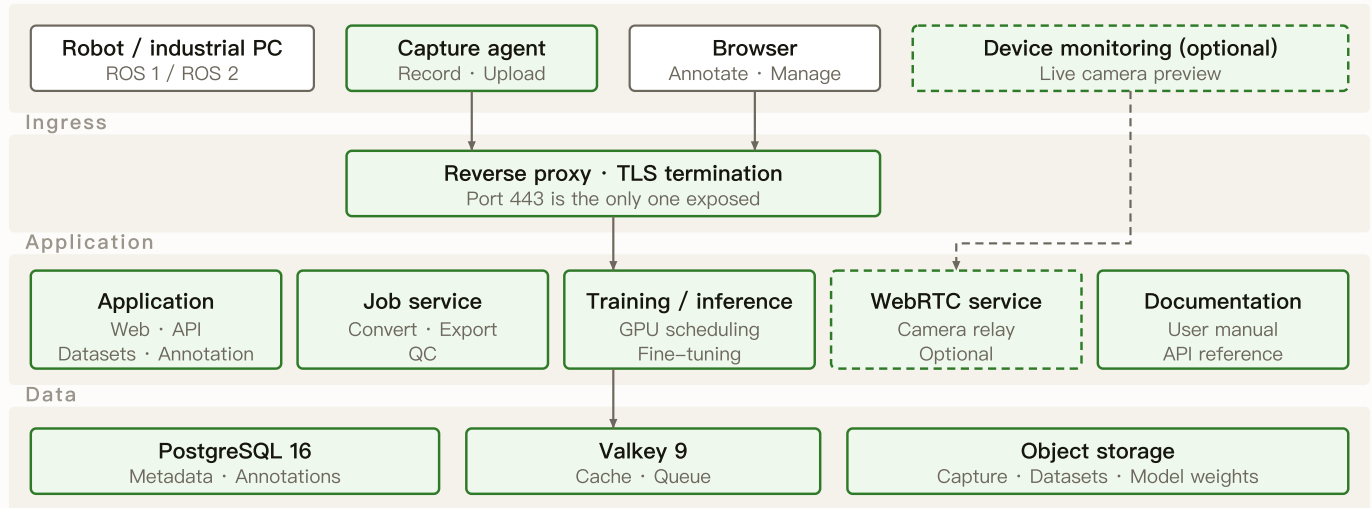


Figure 1 · Reference architecture

■ Delivered by EmbodiFlow
 □ Provided by customer
 - - - Optional

Table 3 · Component responsibilities

Component	Responsibility	Primary resource demand	GPU required	Scaling method
Application service	Web interface, REST API, permissions and quotas, dataset and annotation management	CPU, memory	No	Horizontal — add replicas
Job service	Asynchronous work: format conversion, dataset export, QC, scheduled tasks	CPU, scratch disk	No	Add replicas by queue depth
Training and inference	GPU scheduling, model fine-tuning, inference orchestration	GPU memory, sequential read bandwidth	Yes	Add GPUs or training nodes
PostgreSQL	Structured data: users, projects, dataset catalogue, annotations, jobs, audit logs	Memory, NVMe random IO	No	Primary/standby replication, read-write split
Valkey	Session cache and job queue	Memory	No	Sentinel primary/replica
Object storage	Raw capture, exported datasets, model weights	Disk capacity, sequential throughput	No	Add storage nodes
Capture agent	Runs on the robot: topic recording, integrity verification, resumable upload	Local disk, uplink bandwidth	No	One per robot

RECOMMENDED

Object storage can be the customer's existing S3 service

The platform accesses object data over the standard S3 protocol. If Tencent COS, Alibaba OSS, Huawei OBS, AWS S3, Azure Blob, Volcengine TOS or self-hosted Ceph is already in place, it can be connected directly and the storage nodes listed here are not required — only the capacity and throughput checks in Section 6 still apply.

04

Workstation Version

A single tower workstation runs every platform component and supports 5 or more robots capturing at once. Both tiers are functionally identical; they differ only in compute and capacity.

Table 4 · Hardware specifications

Item	Minimum	Recommended
Reference model ^①	Dell Pro Precision 9 T4 9 storage bays, 4 PCIe slots	Dell Pro Precision 9 T6 21 storage bays, 15 PCIe slots, supports 2 × 600 W GPU
CPU	16 cores / 32 threads, base clock ≥ 2.5 GHz	From 32 cores, 60 recommended Intel Xeon w7 / w9 series or equivalent
Memory	64 GB DDR5 ECC	From 128 GB, 256 GB DDR5 ECC recommended
GPU	Optional. 24 GB required if training is enabled e.g. NVIDIA RTX PRO 4000 24GB	From 32 GB, 48 GB recommended RTX PRO 4500 32GB → RTX PRO 5000 48GB
System disk	1 × 2 TB NVMe	2 × 1.92 TB NVMe Mirrored, to avoid a single point on the system disk
Capture data disks	2 × 16 TB Enterprise SATA / SAS hard drives	8 × 32 TB Enterprise hard drives, identical model and capacity
Data protection	Two replicas. 50% usable, tolerates 1 disk failure	Erasure coding 10+4. 71.4% usable, tolerates 2 simultaneous disk failures
Capacity arithmetic	32 TB raw → 16 TB usable → 13 TB net	256 TB raw → 182.9 TB usable → 146 TB net
Network	1 GbE	10 GbE Ten robots uploading at once need about 900 Mbps
Supported load	About 50 concurrent users, 5 robots	About 200 concurrent users, 10 or more robots
Power and cooling ^②	Single-phase 220 V, peak about 1.0 kW	Single-phase 220 V, peak about 1.8 kW including GPU
Deployment method	Docker Engine 24+ with Docker Compose 2.20+	

① Reference models illustrate the specification level and do not designate a brand. Equivalent HPE, Lenovo or domestic Chinese systems may be substituted; see Section 8 for localised equivalents.

② Whole-chassis peak estimate, for circuit and UPS planning. We recommend a data centre or equipment room with dedicated power and stable temperature.

Chassis and footprint

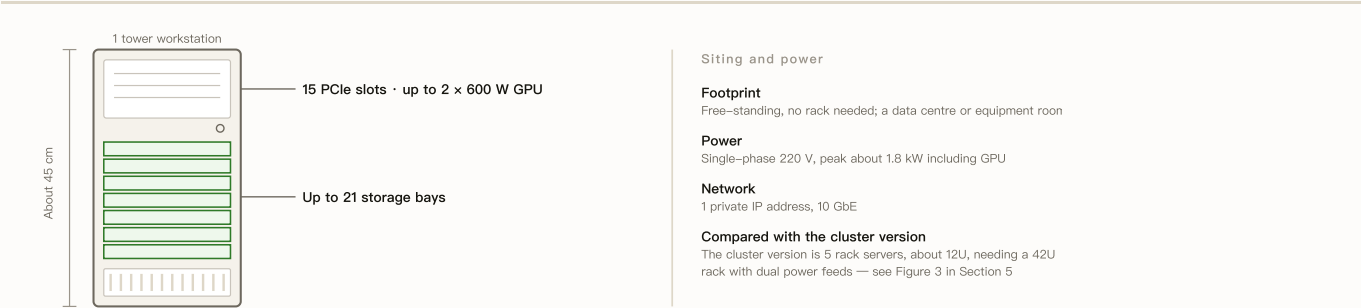


Figure 2 · Chassis and siting requirements for the workstation version

NOTE

Check drive bays and slots before you buy

Memory, GPUs and disks can all be added later, but once the drive bays and PCIe slots are full the only option is a new chassis. Sizing bay count against three years of expected capture is cheaper than filling every bay up front.

05

Cluster Version

Built on Kubernetes, with multiple application replicas, automatic database failover and object storage redundant across nodes. Loss of any single node interrupts neither service nor data.

Table 5 · Node roles and specifications

Role	Qty	Minimum	Recommended
Application node	2	32 cores / 128 GB · 2 × 960 GB NVMe · 10 GbE Dell PowerEdge R770 or equivalent 2U dual-socket	64 cores / 256 GB · 2 × 1.92 TB NVMe · 25 GbE Hosts application, job service and database primary/standby
Storage node	2	16 cores / 64 GB · 12 × 32 TB · 25 GbE Dell PowerEdge R7725xd or equivalent dense storage chassis	32 cores / 128 GB · 24 × 32 TB · 25 GbE More bays now means cheaper expansion later
Training node	1+	32 cores / 256 GB · GPU 48 GB NVIDIA L40S 48GB or RTX PRO 5000 48GB	64 cores / 512 GB · GPU 96 GB NVIDIA RTX PRO 6000 Blackwell 96GB · additional cards as needed

On node counts: automatic leader election needs three members, and the training node doubles as the third, so “two plus one” is the cheapest highly available arrangement. Two storage nodes are enough for whole-server disaster recovery.

Table 6 · High availability and failure domains

Layer	Redundancy mechanism	Failure tolerated
Ingress	Load balancing with health checks	Loss of any application node
Application and job service	Multiple stateless replicas, queue re-dispatch	Loss of any replica; no jobs lost
Database	Primary/standby replication with automatic failover	Primary fails, standby takes over
Cache and queue	Sentinel primary/replica, at least 3 sentinels	Primary fails, new primary elected automatically
Object storage	Two replicas — each node holds a complete copy	Loss of any storage node
Rack	Both replicas placed in separate racks	Loss of one rack when split across two racks

Table 7 · Capacity arithmetic

Item	Minimum	Recommended
Data disks per node	12 × 32 TB	24 × 32 TB
Cluster raw capacity	768 TB	1 536 TB
Usable capacity after protection	384 TB	768 TB
Planned net capacity	307 TB	614 TB

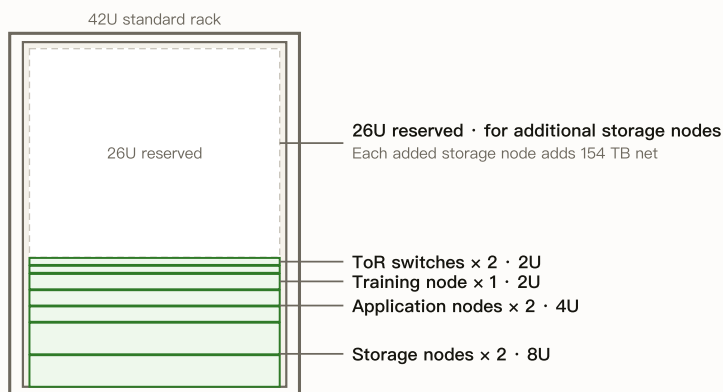
Each additional 12 × 32 TB node adds about 154 TB net. At four or more nodes, erasure coding yields about 43% more from the same disks. Expansion is online.

RECOMMENDED

Plan the rack and power for 5 servers

The minimum tier is about 12U and peaks near 9 kW; the recommended tier is about 16U and 14 kW. A single 42U rack holds either, with room for more storage nodes.

Rack planning

**Figure 3** · Rack elevation and facility requirements

Rack and facility requirements

Rack specification

EIA-310, ≥ 42U, width ≥ 600 mm, depth ≥ 1100 mm

Space occupied

12U minimum; 16U recommended (24-bay chassis is 4U)

Power

Dual independent feeds; peak 9 kW (14 kW recommended)

Airflow

About 1,350–2,100 CFM, estimated at 150 CFM per kW

Loaded weight

About 280–380 kg; confirm floor and rack ratings

Environment

18–27 °C, 40–60% relative humidity

06

Storage and Capacity Planning

Storage is the main cost item in an on-premises deployment and the only one that keeps growing. This section gives the capacity model and every parameter behind it, so capacity can be calculated against your own capture plan.

Data volume assumptions

Table 8 · Capture data parameters

Raw capture rate	40 GB / robot-hour
Capture hours per day	8 hours
Capture days per year	250 days
Growth per robot-year	80 TB
Training-format export share	About 1/6 of raw

Basis: 40 GB per robot-hour corresponds to 3–6 streams of 720p or 1080p video plus pose, torque and tactile data; the public RoboMIND dataset measures 40.3 GB per hour. Writing uncompressed images at the capture agent raises the rate by an order of magnitude, so enable hardware encoding.

Capacity model

```
// Demand side
Required capacity = robots × 80 TB × retention years × 1.17
// 1.17 = raw 1 + export 1/6

// Supply side
Usable capacity = raw capacity × usable ratio (Table 9)
Net capacity = usable capacity × 0.80

// Check
Net capacity ≥ Required capacity
```

NOTE

What the 80% planning watermark is for

The 20% held back covers rebuilds after a drive failure and the lead time on buying more. Plan to full capacity and the first rebuild exhausts the space.

Data protection

Table 9 · Data protection options

Scheme	Usable %	Fault tolerance	Applies to
Single disk, unprotected	100%	None	Not for production use
Two replicas	50%	1 replica lost	Two storage nodes, or fewer than 8 disks in one chassis
Erasure coding 10+4	71.4%	Any 4 shards lost	8 or more disks in one chassis, or 4 or more storage nodes

Erasure coding splits data into 10 data shards plus 4 parity shards; any 4 can be lost and the data still reads back in full, at only 1.4× storage overhead — well below the cost of replication. Tolerating a whole-node failure needs the shards spread over at least 4 nodes, so a 2-node cluster uses two replicas instead.

Ways to expand

Table 10 · Three ways to expand

Method	Net capacity added	Downtime
Fill the empty drive bays	About +18 TB per 32 TB disk	Power down to install disks
Add a storage node	+154 TB per 12 × 32 TB node	Online
Switch to erasure coding at 4 nodes	+43% on existing capacity	Online, rebalanced in the background

There is no need to buy everything up front: fit disks for the first year of capture, then add disks or nodes as the real growth rate becomes clear. The platform keeps reading and writing throughout a cluster expansion, and data is rebalanced automatically.

Table 11 · Storage configuration by fleet size

Robots	Raw growth per year	Net capacity required	Suggested storage configuration	Net capacity provided	Version
1	80 TB	94 TB	One chassis, 8 × 32 TB, erasure coding	146 TB	Workstation version
3	240 TB	281 TB	2 storage nodes × 12 × 32 TB, two replicas	307 TB	Cluster version
6	480 TB	562 TB	2 storage nodes × 24 × 32 TB, two replicas	614 TB	Cluster version
10	800 TB	936 TB	5 storage nodes × 12 × 32 TB, erasure coding	1 097 TB	Cluster version
20	1 600 TB	1 872 TB	5 storage nodes × 24 × 32 TB, erasure coding	2 194 TB	Cluster version
50	4 000 TB	4 680 TB	11 storage nodes × 24 × 32 TB, erasure coding	4 827 TB	Cluster version

How to use this table: figures assume one year of retention and already include the export share and the planning watermark. For a different retention period, scale the required capacity by the same factor — keep data for six months instead and both the requirement and the disk spend halve.

07

Network and Ports

Only one HTTPS port needs to be open externally. Database, cache and object storage are reachable on the LAN only, and the platform needs no outbound connectivity.

Table 12 · Ports to open

Direction	Source	Destination	Port	Protocol	Purpose
Inbound	Client browsers (LAN subnets)	Application entry point	443	TCP	Web interface and REST API
Inbound	Capture agents / workstations	Application entry point	443	TCP	Capture upload, chunked and resumable
Inbound	Operations terminals (restrict source IP)	All servers	22	TCP	SSH administration
Inbound	Client browsers (optional)	Application entry point	443	TCP	Device monitoring preview, reuses the same port
Outbound	All platform components	Internet	None required	—	Offline deployment; no call-home and no telemetry

Table 13 · Internal service ports

Service	Port	Permitted access
PostgreSQL	5432	Application and job service only
Valkey	6379	Application and job service only
Object storage S3 endpoint	8333	Application, job and training services
Object storage master	9333	Storage cluster internal only
Object storage filer	8888	Storage cluster internal only
Object storage volume server	8080	Storage cluster internal only
Training service API	8088	Application service only
Device monitoring signalling (optional)	8890	Via the ingress reverse proxy

IMPORTANT

Internal services must not be exposed to the internet

The ports in Table 13 are not designed for internet-facing authentication. Restrict them to the platform node subnets and confirm the default cloud security group does not permit them.

Table 14 · Bandwidth model

Traffic type	Per unit	Aggregate estimate
Capture upload	About 90 Mbps per robot	About 900 Mbps for 10 robots; 4.5 Gbps for 50
Web and annotation access	1–2 Mbps per active session	About 0.5–1 Gbps for 500 users
Dataset export download	Fills available client bandwidth	2 concurrent export jobs by default
Inter-node storage rebuild	Rises sharply during disk rebuild	Dedicated 25 GbE recommended

Uplink: from 1 Gbps for the workstation version, 10 Gbps preferred; 25 GbE between cluster nodes. Where capture and platform share a data centre, traffic stays on the LAN. The 90 Mbps per robot is an average; bulk upload creates peaks, so reserve about twice that.

Table 15 · Network constraints

Application node ↔ database	< 5 ms
Between storage nodes	< 2 ms
Deployment span	Within a single data centre
Private IP addresses	1 for a single host; 5 plus a virtual IP for a cluster
Hostname resolution	All nodes must resolve each other
Time synchronisation	NTP or Chrony, inter-node skew < 1 s
HTTPS	Mandatory in production, TLS 1.2 or above

Certificates: use the customer's own domain certificate, or a *.telexdata.com subdomain certificate provided by EmbodiFlow. Private keys stay with the customer. With a self-signed certificate, the root must be trusted on each client.

08

Runtime Environment and Software

The operating system, network and certificates are prepared by the customer; container images and runtimes ship with the delivery package. Every third-party component is under a permissive OSI-approved licence.

Table 16 · Operating system and runtime

Item	Requirement
Operating system	Ubuntu 22.04 LTS or 24.04 LTS
CPU architecture	x86_64 (amd64)
File system	EXT4 on the system disk; XFS on data disks, each mounted separately
Disk controller	HBA in pass-through mode; no hardware RAID
Container runtime	Docker Engine 24 or above, Compose 2.20 or above
Cluster orchestration	Kubernetes 1.28 or above (cluster version)
GPU driver	NVIDIA Driver 550 or above, CUDA 12.4 or above
GPU container support	NVIDIA Container Toolkit
System settings	Swap disabled; raised file-descriptor limits; time synchronisation configured

Table 17 · Open-source components and licences

Component	Used for	Licence
PostgreSQL 16	Primary database	PostgreSQL License
Valkey 9	Cache and job queue	BSD-3-Clause
SeaweedFS	Object storage	Apache-2.0
Docker Engine / Compose	Container runtime	Apache-2.0
Kubernetes	Cluster orchestration	Apache-2.0
NVIDIA Container Toolkit	GPU access inside containers	Apache-2.0
Nginx	Reverse proxy and TLS termination	BSD-2-Clause
Node.js / Python	Application and training runtimes	MIT / PSF
PyTorch / JAX	Training frameworks	BSD-3-Clause / Apache-2.0

COMPLIANCE STATEMENT

Every component above is free for commercial use and none is licensed under AGPL, SSPL or any other source-available licence, so no customer-owned code is drawn into an open-source obligation. GPU drivers are installed by the customer and licensed separately by NVIDIA.

Table 18 · Model memory requirements

Model	Inference	LoRA fine-tune	Full fine-tune
SmoVLA	8 GB	16 GB	—
π 0-FAST	4 GB	8 GB	—
π 0 / π 0.5	8 GB	22.5 GB	> 70 GB
OCTO	16 GB	24 GB	—
OpenVLA-7B	15 GB	27 GB	80 GB × 8
GR00T N1.5	—	≥ 40 GB	≥ 80 GB

How to read this table: anything at or below 32 GB is covered by the recommended configuration; above that needs 48 GB or 96 GB. Figures come from each model's official repository and vary with batch size and sequence length.

Localisation

Table 19 · Domestic Chinese equivalents

Role	Domestic alternatives
Application node	Inspur NF5280M7 · H3C UniServer R4900 G7 · xFusion FusionServer 2288H V7
Storage node	xFusion FusionServer 5288 series · Inspur NF5466 series (dense storage chassis)
Training node	Same chassis as the application node, with the dual-width GPU option

NOTE

How the CPU choice affects deployment

Hygon C86 is an x86 line and runs the container stack described here directly, with no image rebuild — the lowest-change option. ARM lines such as Kunpeng and Phytium need every image rebuilt and revalidated, which is custom delivery, so settle the CPU line during selection. Kylin V10 or UOS are available as the operating system.

Classified protection: the platform provides role-based and module-level permissions, login and operation logs, and permission-change auditing. Logs are held in the database and retained long-term, supporting the customer's classified-protection assessment.

09

Delivery and Operations

Delivery includes images, scripts, credentials, certificates and documentation. Deployment, upgrades and version migration are carried out with on-site or remote assistance from an IO-AI TECH delivery engineer; the customer's IT engineers can also run them independently using the delivered scripts.

Table 20 · Deliverables

Category	Contents
Installation media	Docker image tar archives for every service, supporting fully offline import
Deployment scripts	Installation script, upgrade script, orchestration configuration templates
Credentials	Image registry account for obtaining later versions, object storage account, initial platform administrator password
Certificates	Domain certificate and server root certificate, where an EmbodiFlow subdomain is used
Documentation	User manual, API reference, functional test checklist, performance test report, security test report, operations manual
Developer interfaces	Full REST API documentation and a Node.js SDK

Upgrades

Table 21 · Upgrade methods

Deployment form	Downtime	Method
Workstation version	Planned maintenance window required	Run the upgrade script, including database migration
Cluster version	No application downtime	Rolling upgrade, replica by replica

Before upgrading: the script migrates the database schema. The delivery engineer takes a backup and retains the previous version before the upgrade, then verifies the platform afterwards.

Backup and recovery

Table 22 · Data to include in backups

Object	Contents	Suggested policy
Database volume	Users, projects, dataset catalogue, annotations, jobs, audit logs	Daily full backup, 14-day retention
Object storage volume	Raw capture, exported datasets, model weights	Off-site or tape archive
Cache and queue volume	Job queue state	Daily snapshot
Certificates	TLS certificate and private key	Archive on change
Environment configuration	Environment variable files and orchestration secrets	Archive on change

IMPORTANT

A database backup alone will not restore a working system

The database holds file locations and metadata; the files live in object storage. Both must be backed up together and restored to the same point in time, or the catalogue will list files that cannot be read. The operations manual gives step-by-step instructions.

Monitoring and alerts

Table 23 · Alert thresholds

Trigger condition	Response
Free net object storage below 25%	Begin expansion procurement
Redundancy degraded — one more failure would lose data	Replace the failed disk or node immediately
Job queue backlog > 500 for 5 minutes	Add job service replicas
Database CPU > 85% for 5 minutes	Add database resources
Capture agent offline > 15 minutes	Check agent and network
Capture upload failure rate above 1%	Check capture-side network and platform ingress bandwidth

The platform includes built-in log and metric views. In the cluster version, node and container metrics can be fed into the customer's own Prometheus, Grafana or cloud monitoring.

10

Prerequisites and Acceptance

This page can be printed separately for implementation scheduling. Every item on the checklist should be complete before the installation window opens; any one missing will block deployment.

Table 24 · Customer prerequisites checklist

Item	Owner	Date completed
HARDWARE AND FACILITY		
☐ Servers delivered, racked and powered		
☐ Rack units, power and cooling meet Section 5		
☐ Disk controller set to HBA pass-through mode		
OPERATING SYSTEM		
☐ Operating system installed and patched to current		
☐ Data disks mounted separately with XFS, mount points written to fstab		
☐ Time synchronisation configured, inter-node skew under 1 second		
☐ Hostnames and internal DNS resolution configured		
☐ Operations accounts and sudo rights provisioned		
NETWORK AND CERTIFICATES		
☐ Internal IPs allocated (1 for single host / 5 plus virtual IP for cluster)		
☐ Port 443 inbound permitted, inter-node connectivity verified		
☐ DNS record for the access domain points to the entry address		
☐ TLS certificate and private key in place, or subdomain certificate confirmed		
GPU (WHERE TRAINING IS ENABLED)		
☐ NVIDIA driver and container runtime installed and verified		

Table 25 · Responsibility matrix

Activity	EmbodiFlow	Customer
Hardware procurement and racking	Supports	Owens
Operating system installation and hardening	Supports	Owens
Network, firewall and certificates	Supports	Owens
Platform installation and initialisation	Owens	Supports
Version migration (workstation ↔ cluster)	Owens	Supports
Capture agent onboarding and format adaptation	Owens	Supports
Acceptance test execution	Supports	Owens
Version upgrades	Owens	Supports
Routine backup execution	Supports	Owens
Day-to-day monitoring and capacity review	—	Owens

Note: “Owens” means leading the activity and being accountable for the outcome; “Supports” means providing documentation, scripts and technical guidance. Version upgrades and migration are led by IO-AI TECH delivery engineers; the customer’s IT engineers can also run them independently using the delivered scripts.

Table 26 · Acceptance tests

Test	Pass criterion
Service health	Health checks return healthy for every platform service
Login and permissions	Administrator login succeeds and module permissions apply by role
Object storage read/write	Uploaded and downloaded files match on checksum; capacity accounting is correct
Capture agent onboarding	Agent comes online; recorded data uploads completely and is catalogued
Format recognition and playback	MCAP and ROS Bag files are recognised, verified and replayed in the viewer
Dataset export	LeRobot, HDF5, MCAP and RLDS all export successfully and download
GPU visibility	GPUs enumerate inside the training container with correct memory and driver version
Training job	A fine-tuning job is submitted and completes normally
Disk-level fault tolerance	Data reads back intact after removing 1 data disk
Node-level fault tolerance	Service remains available after shutting down any node, with recovery within 30 minutes
Cluster version only	

Notice: this document describes the hardware and environment conditions required to deploy EmbodiFlow and serves as the specification of record. Reference server models illustrate the specification level only and do not designate a brand or model; actual models should follow what can be supplied at quotation. Third-party names and trademarks belong to their respective owners. Functional scope and service levels are governed by the executed contract.