

MLOps & Data Quality

Interview cheat sheet

215 questions · Data drift, class imbalance, monitoring, CI/CD for ML and everything that keeps models alive in production.

#monitoring #infrastructure #llmops #data-quality #deployment #reproducibility #serving #mlops #pipeline #experimentation

Data quality & drift

What is data drift and how do you detect it?

Data drift is a change in the input feature distribution $P(X)$ over time compared to training data. Detect it with statistical tests per feature (KS test for numeric, chi-square for categorical), Population Stability Index (PSI > 0.2 = significant drift), or distance measures (JS divergence, Wasserstein).

How is concept drift different from data drift?

Concept drift is a change in the relationship $P(y | x)$ — the mapping from features to label evolves over time (customer behavior changes, fraud patterns evolve). Data drift changes $P(X)$; label drift changes $P(y)$.

What does 'data imbalance' mean and why is it a problem?

Data imbalance means classes are unequally represented (e.g., 99% negatives, 1% positives). A naive model can hit 99% accuracy by always predicting the majority class, learning nothing about the minority.

What are common sources of data leakage in an ML pipeline?

(1) Fitting scalers/encoders on the full dataset before splitting. (2) Using future information (post-outcome features).

What should you monitor in a production ML model?

(1) Prediction distribution over time. (2) Feature distributions and drift metrics.

How do you validate data quality in an ML pipeline?

Set up automated checks with tools like Great Expectations, TFDV, or Deequ. Validate: schema (types, columns present), value ranges, allowed categories, null rates, uniqueness, referential integrity, and distributional stats (mean, std, quantiles) vs a reference.

What is schema drift and how do you detect it?

Schema drift = upstream data changes shape: column added / removed / renamed, dtype changes (int → string), enum values expanded. Silent failures because pipelines don't crash but predictions become garbage.

What is label drift and why does it matter?

Label drift = change in $P(y)$ marginal distribution — class balance shifts over time. Example: fraud rate rises from 1% to 3%.

How is Population Stability Index (PSI) computed and interpreted?

Bin feature into k quantiles based on training distribution. Compute expected % (train) and actual % (current) per bin.

Why prefer Jensen-Shannon Divergence over KL for drift?

$KL(P || Q)$ is asymmetric + unbounded + undefined when $Q=0$ for some support. $JSD = \frac{1}{2} KL(P || M) + \frac{1}{2} KL(Q || M)$ where $M = (P+Q)/2$ is symmetric + bounded to $[0, \ln 2]$ (or $[0, 1]$ with base 2) + always defined.

What is Maximum Mean Discrepancy (MMD) for drift detection?

Kernel-based two-sample test. Embed both distributions in RKHS via kernel k .

When is Wasserstein distance appropriate for drift?

Wasserstein-1 = 'earth mover's distance' — minimal work to move mass from P to Q . Advantages: (1) sensitive to shifts (unlike KL that focuses on mode overlap).

How do you monitor drift on unstructured data (text / images)?

(1) Compute embeddings from pretrained model (CLIP for images, sentence-BERT for text). (2) Reduce to low-D via PCA / UMAP for viz.

Outliers vs drift — how do you distinguish?

Outliers = individual points far from distribution; drift = the distribution itself shifts. Time-based test: (1) if score spikes then returns → outliers or anomalies.

How do you handle missing values in a production pipeline?

(1) Categorize: MCAR (random), MAR (depends on observed), MNAR (depends on missing itself). (2) Simple: mean/median (numeric) or mode/missing category (categorical).

How does imputation cause data leakage?

Fitting imputer (mean / KNN) on train + test together sees test values → optimistic evaluation. Same for standardization / target encoding.

How do you handle label noise in training data?

(1) Confident Learning (cleanlab): identifies mislabeled examples via out-of-fold predictions. (2) Symmetric losses (MAE, generalized cross-entropy) more robust than CE.

How do you measure and improve annotation quality?

(1) Inter-Annotator Agreement (IAA): Cohen's κ (2 annotators), Fleiss's κ (3+), Krippendorff's α . Below 0.6 = poor, above 0.8 = strong.

What is active learning and when should you use it?

Iteratively query most informative unlabeled examples for human labeling. Strategies: (1) Uncertainty sampling (max entropy / min margin / least confidence).

What is weak supervision (e.g., Snorkel)?

Generate noisy labels programmatically from labeling functions (LFs): heuristics, keyword patterns, existing model outputs, distant supervision from KBs. Snorkel + Data Programming: LFs vote → generative model estimates LF accuracies + correlations → produces probabilistic labels → train final model on those.

How does time-based leakage happen and how do you prevent it?

(1) Using future timestamps in features (fitting on aggregate features computed with post-event data). (2) Random shuffle instead of time-based split.

What is purged cross-validation?

López de Prado. Standard k-fold + (1) remove training samples whose labels depend on test period (purge).

SMOTE vs class weights vs threshold tuning — which to use?

(1) Class weights: cheapest, no data manipulation. Try first — often sufficient for 10-100:1 imbalance.

Why version data (DVC / lakeFS / Delta Lake)?

(1) Reproducibility: recreate any past experiment with exact input data. (2) Rollback: revert bad ingestion / labeling changes.

How do you monitor fairness in production?

(1) Track prediction / performance metrics per protected group (gender, race, age, geography). (2) Compute fairness metrics: demographic parity, equal opportunity, equalized odds, calibration by group.

How do you handle PII in ML pipelines?

(1) Minimize collection: only fields needed. (2) Pseudonymization: hash identifiers.

How does selection bias affect ML systems?

Training data doesn't represent deployment population. Examples: (1) recommender only trained on items shown → can't learn about unshown items (exposure bias).

What is a model feedback loop and why is it dangerous?

Model's predictions influence future training data → self-reinforcing. Example: (1) recommender pushes popular items → users click popular → more training signal for popular → runaway effect.

What is a golden test set and how do you maintain it?

Curated, high-quality, hand-labeled evaluation set that remains fixed across model versions — enables consistent comparison. Properties: (1) representative of production distribution.

What is a data contract?

Formal agreement between data producer + consumer specifying: (1) schema (columns + types + nullability). (2) semantic meaning of each field.

What is point-in-time correctness in feature stores?

When constructing training data, for each event at time t , use only feature values known BEFORE t . Prevents leakage from future data into features.

How do you guarantee online-offline feature consistency?

(1) Single source of truth for feature definition (code + config). (2) Compute both paths from same definition: batch compute (offline) + streaming compute (online) share transformation logic.

What do you monitor about features (not just models)?

(1) Freshness: is the feature updated on schedule? (2) Completeness: missing rate per feature.

How do you structure a progressive rollout with automated guardrails?

(1) Deploy to 1% traffic. (2) Monitor N minutes: error rate, latency p50/p99, business KPI, model-specific metrics (calibration).

Why do offline improvements often not translate online?

(1) Distribution shift: offline eval on historical data doesn't reflect current traffic. (2) Feedback effects: recsys changing displayed items changes user behavior.

SLO / SLI / SLA for ML services.

SLI (Service Level Indicator): metric measured (latency p99, error rate, accuracy). SLO (Objective): target for SLI (p99 < 100ms 99.9% of time).

How do you avoid alert fatigue in ML monitoring?

(1) Prioritize: only alert on things needing immediate action; log others to dashboard. (2) Escalation: warnings → email; errors → Slack; critical → page.

Three pillars of observability applied to ML.

(1) Metrics: aggregated numeric signals over time (Prometheus). ML: prediction distribution, latency, drift score, feature freshness.

How do you monitor feature / prediction freshness?

Freshness = time between event and its reflection in served feature. Monitor: (1) end-to-end lag: $\max(\text{current_time} - \text{feature_updated}_{\text{set}})$ across features.

How does shadow evaluation work?

New model runs on production requests in parallel with current model. Predictions logged, not served.

What is Canary Analysis (Kayenta / Flagger)?

Automated statistical comparison between canary + baseline metrics during progressive rollout. Netflix Kayenta: (1) fetch metrics for canary + baseline over rollout window.

How do you monitor prediction uncertainty in production?

(1) Log softmax entropy / max prob per prediction. (2) MC dropout / ensemble variance for Bayesian models.

How do you monitor model calibration in production?

Bin predictions by predicted probability. For each bin, compute observed frequency of positive outcome (needs labeled data).

Why monitor per-slice performance?

Overall accuracy can hide subgroup regressions. Monitor per (country, device, user tier, time of day, product category).

Popular ML observability tools?

(1) Evidently AI: open-source drift + performance reports. (2) WhyLabs / WhyLogs: streaming profile + monitoring.

How do you detect anomalous predictions in production?

(1) Statistical: z-score of prediction / feature vs recent history. (2) Isolation Forest / LOF on request features.

How do you set thresholds for drift alerts?

(1) Compute baseline: 1-2 weeks of production data during known-good period. (2) Set threshold: mean + 3 σ , or specific PSI > 0.2, or KS p < 0.01.

How do you monitor performance when labels are delayed?

Ground truth arrives days / weeks / months later (loan default, customer churn). Strategies: (1) proxy metrics available immediately (clicks for CTR, purchase within 1h).

How do you monitor drift when data has natural seasonality?

(1) Baseline per time period (hour-of-day, day-of-week, month) not one static. (2) Detrend via time-series decomposition (STL, Prophet) before comparing.

What is model decay and how do you measure it?

Gradual degradation of model performance over time due to drift + concept change + world evolving. Measure: (1) rolling accuracy on freshly-labeled batches.

How do you respond to a production ML incident?

(1) Alert fires → on-call paged. (2) Assess severity: user impact, revenue, safety.

How do you monitor ML infrastructure cost?

(1) Tag resources by (team, project, model, environment). (2) Cost per model per day: infra cost + prediction volume → prediction . (3) $\text{Alerts} \propto \sum \text{unexpected cost spikes} (> 2\sigma)$. (4) Budget alert to spare team . (5) FinOps dashboard : idle CPU , $\text{over-provisioned pods}$, expensive — but — unused features . (6) Right — sizing .

How do you handle biased ground truth collection?

Feedback often depends on model output. Recsys: only see clicks on shown items — no signal for unshown.

How do you explain individual predictions in production?

(1) SHAP: computes Shapley values per feature contribution. TreeSHAP fast for trees.

How do you compute SHAP at production scale?

SHAP is expensive (2^N feature subsets). Optimizations: (1) TreeSHAP: exact + $O(\text{TLD}^2)$ for tree ensembles.

Features & pipelines

What problem does a feature store solve?

A feature store centralizes definitions, computation, storage, and serving of features. Online store (low-latency, key-value) serves predictions in real time; offline store (parquet / warehouse) provides consistent historical features for training.

The four golden signals — Google SRE.

(1) Latency: time to serve request (p50, p99). (2) Traffic: requests per second.

RED vs USE vs Golden Signals — which methodology?

RED (Requests / Errors / Duration): request-driven services (API, ML inference). USE (Utilization / Saturation / Errors): resource-focused (CPU, disk, GPU).

Why use latency heatmaps instead of averages?

Average / median hides bimodal or long-tail distributions. Heatmap: histogram over time (Y = latency bucket, X = time).

Common bad-alert patterns to avoid.

(1) Alert on threshold that fires daily (noise). (2) Alert without runbook (no action).

How do you monitor GPU utilization?

(1) `nvidia-smi` real-time. (2) `dcpm-exporter` for Prometheus (per-GPU metrics: SM utilization, memory, temp, power).

How do you evaluate LLM outputs in production?

(1) LLM-as-judge: strong model rates outputs on rubric (accuracy, helpfulness, safety) — cheap + fast, but bias toward similar-style output. (2) Golden set: hand-labeled Q&A / rubric — expensive but reliable.

What are the biases of LLM-as-judge?

(1) Position bias: prefers first / second response systematically → randomize order + swap. (2) Verbosity bias: prefers longer answers → normalize length.

How do you monitor a RAG pipeline?

(1) Retrieval metrics: recall@k, precision@k, MRR — needs labeled query-doc pairs. (2) Query embedding drift: distribution shift over time.

How do you observe LLM agents?

Multi-step, tool-using agents: (1) Trace each step: user request → LLM call → tool call → response → LLM call → ... final answer. (2) Log tokens per step, latency per step, tool errors.

LangSmith / Langfuse — what do they provide?

LangChain's LangSmith + open-source Langfuse: LLM-specific observability + eval. (1) Trace every LLM call, chain, agent step.

You inherit a model in production with no documentation.

What do you check in your first week?

Establish what it does before touching anything. Find the inference path and confirm which artefact is actually being served, since the deployed version is often not the one in the repository.

Accuracy has probably dropped but labels arrive 60 days late.

What can you monitor now?

Everything upstream of the label. Input drift per feature, using a population stability index or a Kolmogorov-Smirnov test against a fixed training reference, catches the changes most likely to matter.

Your drift monitoring fires 40 alerts a day and everyone ignores it. How do you fix it?

Alert on impact, not on statistics. A statistical test on a large sample detects differences too small to matter, so add a minimum effect size and require the drift to persist over several windows before firing.

How do you prove there is no training-serving skew?

Compare the actual feature vectors, not the code that computes them. Log the features used at inference time, then take the same entities and recompute their features through the training pipeline for the same timestamp, and diff the values.

What is point-in-time correctness, and how does violating it look in practice?

It means every feature value used for a training example reflects only what was knowable at that example's timestamp. Violating it produces leakage that is invisible offline: the model looks excellent in validation and disappoints in production, because at serving time the future information it relied on does not exist yet.

Scheduled retraining or triggered retraining?

Scheduled is the sane default because it is predictable, testable, and forces the retraining path to stay working, which is the failure nobody notices until an emergency. Its weakness is that it retrains when nothing changed, wasting compute, and does not react between runs.

What do you monitor for an LLM feature that you would not monitor for a classifier?

Output-side quality signals, because there is no single correct label to compare against. Track refusal and empty-response rates, which spike when a prompt template or provider changes.

How do you handle high-cardinality categorical features?

One-hot explodes memory; label encoding assumes ordinal (bad). Better: (1) Target encoding with proper CV to avoid leakage.

How do you handle cold-start users / items in production?

(1) Content-based fallback: use features (demographics, item metadata) via a separate model. (2) Global popular items.

Online vs offline feature store — architecture.

Offline store: columnar format (Parquet / Delta / Iceberg) on object storage. Cheap, high throughput, for training + batch inference.

How do you version features?

(1) Each feature has definition (name + computation code + owner + description). (2) Version bumps on compute change (eg, `user_purchase_amount_v2`).

How do you compute real-time features (streaming aggregates)?

(1) Streaming engine (Flink / Spark Streaming / Materialize) consumes events (Kafka). (2) Windowed aggregation: tumbling / sliding / session windows.

Orchestration tools for ML pipelines — Airflow vs Prefect vs Dagster.

Airflow: mature, ubiquitous, DAG-based scheduler; Python operators. Downside: verbose, poor local dev experience, weak lineage.

DAG-based vs imperative ML pipelines — tradeoffs.

DAG-based (Airflow, Kubeflow): explicit dependency graph; scheduler manages execution + parallelism; better observability + retries. Downside: harder to develop locally + parametrize.

Why must ML pipelines be idempotent?

Re-running same pipeline on same input must produce same output. Enables: (1) safe retries after failures (no double-writes).

Why partition ML training tables by date?

(1) Prune scans: only read partitions needed for training window (e.g., last 90 days) — 10-100x faster. (2) Incremental compute: only recompute affected date's features.

How do you safely backfill features / labels historically?

(1) Isolate backfill to non-production tables (`backfill_YYYYMMDD`). (2) Validate output matches sample of production for overlap dates.

What tests should a training pipeline have?

(1) Data validation: schema + ranges + expected volume. (2) Feature parity: features computed in training match serve-time features.

Deployment & experimentation

What is training-serving skew?

Training-serving skew is a mismatch between how features are computed at training and at serving time — different preprocessing code, different data sources, different schemas, different aggregations. Result: production performance is worse than offline eval.

What is a shadow deployment and how does it differ from A/B testing?

In shadow deployment, the new model runs alongside the current one on real production traffic, but its predictions are logged and not served to users. You compare its outputs to the production model without user impact.

What is a canary release for ML models?

Canary release routes a small percentage of traffic (e.g., 1-5%) to the new model, monitors metrics and errors closely, and gradually increases traffic if healthy. If a regression is detected, you roll back quickly.

Batch vs online inference — how do you choose?

Batch inference precomputes predictions on a schedule (nightly scoring) — cheap, high-throughput, tolerates seconds/minutes of latency. Online inference serves predictions in real time via HTTP/gRPC — needed for user-facing systems (recommendations, fraud checks).

How do you A/B test a model rigorously?

Randomly assign users to control (old model) and treatment (new model). Define the primary business metric before starting.

MLflow model registry — what does it provide?

Centralized store for model artifacts + metadata. Features: (1) staging → production lifecycle.

What is a model signature and why does it matter?

Explicit schema of model inputs + outputs: field names, types, shapes. Enables: (1) validation at serve time (reject invalid inputs).

What is a blue-green deployment for ML models?

Maintain two identical production environments: 'blue' (current) + 'green' (new). Deploy new model to green, run smoke + integration tests, then flip load balancer to green.

Canary vs blue-green — when to use each?

Canary: gradual traffic shift (1% → 10% → 50% → 100%); great for ML where model quality is uncertain + wanting real-user metrics before full rollout. Blue-green: instant flip; great when new model was thoroughly tested offline + rollback speed matters more than gradual verification.

What is an 'asset' in Dagster and why is it useful for ML?

Asset = data artifact (dataset, model, feature) with computation defined declaratively. Dagster tracks: which upstream assets it depends on, when it was last computed, whether it's stale.

Push vs pull materialization for features.

Push: compute feature eagerly, store in online store, serve reads from cache. Low serve latency, higher storage cost, potential staleness.

ETL vs ELT for ML data — which pattern wins?

ETL: extract → transform → load. Traditional; heavy ETL layer transforms before storage.

What is the medallion architecture (bronze/silver/gold)?

Data lakehouse pattern popularized by Databricks. Bronze: raw ingested data, immutable, append-only, minimal parsing.

Common ML pipeline stages in Airflow / Kubeflow.

(1) Data ingestion / snapshot. (2) Validation (schema + drift check).

Metaflow — what and when to use it?

Netflix ML workflow framework. Decorator-based Python (`@step`, `@resources`) — writes like normal Python but runs as DAG on cloud.

How does Hydra help ML config management?

Facebook's config framework. Compose configs from YAML files: `defaults` list + overrides.

Continuous training pipeline — what triggers retraining?

(1) Schedule (nightly / weekly). (2) New data threshold (X labeled examples).

TFX (TensorFlow Extended) — what is it?

Google's production ML platform for TensorFlow. Pipeline of components: ExampleGen → StatisticsGen → SchemaGen → ExampleValidator → Transform → Trainer → Evaluator → InfraValidator → Pusher.

ZenML — how is it different from Kubeflow?

ZenML: framework-agnostic, Python-first pipelines running on any backend (Airflow, Kubeflow, Vertex, Sagemaker, local). Decorator-based (`@step`, `@pipeline`).

When is a feature store worth the operational cost?

When several models share features and you are already paying the skew tax. The genuine problems a feature store solves are duplicated feature logic across teams, training-serving skew from two implementations, and point-in-time correct joins for training data, which are tedious and easy to get wrong.

What is champion-challenger pattern?

Current production model = 'champion'. New candidate = 'challenger'.

A/B test vs multi-armed bandit — which to use?

A/B: fixed allocation, wait for significance, use when you need clean causal comparison + all variants matter equally. Bandit (Thompson sampling / UCB): dynamically allocate more to winning variants — minimizes regret.

What is interleaving in recsys A/B testing?

Alternative to standard A/B for ranking systems. Instead of splitting users, merge rankings from both models into one presented list (team-draft or probabilistic interleaving).

What are guardrail metrics in an A/B test?

Metrics you don't want to regress even if primary metric improves. Examples: (1) latency p99 (should not degrade).

What is a sample ratio mismatch (SRM) alert?

Actual traffic split (e.g., 51/49) differs significantly from designed (50/50). Chi-square test detects.

What is CUPED and why use it?

Controlled Using Pre-Experiment Data. Adjust post-experiment metric Y by subtracting $\beta \times (X - E[X])$ where X is pre-experiment metric for same user, $\beta = \text{Cov}(X,Y)/\text{Var}(X)$.

How does stratified randomization help experiments?

Randomize within strata (segments) — country, device, user tier. Guarantees balance on covariates.

Long-term holdout — what and why?

Reserve 1-10% of users as permanent control who never receive new model / experiment. Measure long-term effects of everything shipped combined.

How do you meet latency budgets in production ML?

Budget = p50 or p99 target (e.g., 50ms). Techniques: (1) Model compression: quantization (int8/int4), pruning, distillation.

How do you control ML serving cost?

(1) Right-size infrastructure: autoscale on utilization. (2) Batch requests where latency allows.

When is batch inference dramatically cheaper than online?

Batch: (1) high GPU utilization via full batches (100-1000). (2) commodity hardware fine.

How do you version model artifacts in production?

(1) Semantic versioning (v1.2.3): major = incompatible schema, minor = new feature, patch = bug fix. (2) Store immutable artifacts in registry (S3 + metadata DB).

What's a good rollback strategy?

(1) Previous model version always available in registry. (2) Deployment supports instant revert (K8s rolling update revert, feature flag toggle).

How do you handle model errors in production?

(1) Timeout: enforce max inference time; fallback if exceeded. (2) OOM protection: reject oversized inputs.

What is an A/A test and why run it?

Both variants receive the SAME model. Expected result: no significant difference (~5% false-positive rate).

How do network effects complicate A/B tests?

User-level randomization assumes independence between users. Broken by: (1) social networks: friend sees new feed, tells friend in control.

Switchback experiments — when to use?

Assign entire population to A or B alternately over time windows (e.g., every hour). Analyze paired periods.

Difference-in-differences for ML experiments — when?

When randomization impossible (e.g., can't ethically deny users, city-wide rollout). Compare treatment vs control unit change over pre-post period: $DiD = (Y_{treat,post} - Y_{treat,pre}) - (Y_{ctrl,post} - Y_{ctrl,pre})$.

How do you compute sample size for A/B?

$n \approx 16 \times \sigma^2 / (MDE)^2$ per arm for two-sided test at $\alpha=0.05$ and 80% power. Inputs: (1) baseline metric variance σ^2 .

How do you serve multiple models efficiently on shared infrastructure?

(1) Multi-model server (TF Serving / Triton): load N models per container; route by request. (2) Model swap on GPU: pre-load, unload cold.

Why make inference requests idempotent?

Same request produces same result (or same side effects). Enables: (1) safe retry on network / server errors.

How to trade model quality vs latency in production?

(1) Distill large model into smaller. (2) Ensemble → single model.

How do you configure autoscaling for ML serving?

(1) Metric: QPS + GPU util + p95 latency (not just CPU). (2) HPA (K8s Horizontal Pod Autoscaler) with custom metrics via Prometheus adapter.

How do you serve ML models across multiple regions?

(1) Deploy model to each region: latency to nearest user. (2) Global load balancer routes by geography.

Shadow deployment, canary, or A/B test — which do you use when?

They answer different questions, so the sequence matters more than the choice. Shadow mode sends real traffic to the new model without using its output, which validates that it runs, that latency is acceptable and that predictions are sane, with zero user risk but no information about business impact.

What does a real rollback plan for a model require?

The previous model artefact still deployable, which means versioned artefacts with their exact dependency set, not just a file in object storage. It also requires the previous feature pipeline, because a rollback that feeds the old model new-format features is not a rollback.

Monitoring & observability

Distributed tracing for ML inference — why?

Single request spans: feature store fetch → preprocess → model call → postprocess → API response. Tracing (Jaeger, OpenTelemetry) shows per-span latency + errors + attributes.

Why sample prediction logs and how?

High-volume services (millions QPS) can't log every prediction. Sample: (1) uniform 1% for baseline monitoring.

What goes into an ML alert runbook?

(1) Alert description + severity. (2) Impact assessment: what's affected?

How do you design useful ML dashboards?

(1) Top-of-funnel: at-a-glance health (traffic + errors + latency + drift). (2) Drill-down per model / feature / slice.

How long should you retain ML prediction logs?

Depends on: (1) regulatory (banking / medical: 7-10 years). (2) debug window (recent logs hot 30 days; older cold storage).

Why is high-cardinality label bad for Prometheus?

Each unique label combination = separate time series. Explosive: `latency{user_id=..., request_id=...}` → millions of series → memory blowup.

A model caused a costly wrong decision. What does a useful post-mortem produce?

A timeline anchored in evidence, which is only possible if predictions, inputs and versions were logged, and a first finding of a missing log is itself the most valuable output. Then the actual causal chain, which for model incidents is usually upstream: a schema change, a feature pipeline default, or drift nobody was alerted on, rather than the model being wrong in an interesting way.

Reproducibility & versioning

What is data / model lineage and why does it matter?

Lineage = trace of every data + code + config + upstream dependency that produced an artifact. Model lineage: dataset version → feature spec → training code commit → hyperparameters → resulting weights → deployed endpoint.

What audit trail is required for regulated ML?

(1) Training data source + version + snapshot. (2) Preprocessing code + config.

How do you manage ML config (hyperparameters, thresholds)?

(1) YAML / Hydra: declarative + composable + versioned in git. (2) Separate infra config from experiment config.

MLflow vs Weights & Biases vs Neptune — which to pick?

MLflow: open-source, self-hosted, most common, minimal features. W&B: SaaS, best UI + collaboration + reports, sweeps, integration with everything.

Why containerize ML training + serving?

(1) Reproducibility: exact Python + CUDA + system libs pinned. (2) Portability: same image on laptop / cluster / prod.

Why isn't setting a seed enough for full reproducibility?

Even with fixed seed, non-determinism from: (1) parallel GPU ops (cuDNN atomic ops, reduction order varies). (2) multi-threaded data loading order.

How does DVC work?

Data Version Control: git-like for large files. Stores metadata (.dvc files) in git; actual data in remote storage (S3, GCS, Azure). `dvc add data.csv` → creates data.csv.dvc with hash + points to remote. `git commit metadata` + `dvc push` data.

What is lakeFS?

Git-like versioning for object storage (S3, GCS, Azure Blob). Full branch / commit / merge for data lakes.

What does Delta Lake add over plain Parquet?

(1) ACID transactions: atomic multi-file writes, no partial reads. (2) Schema enforcement + evolution (add columns safely).

Weights & Biases vs MLflow — key differences.

MLflow: open-source, self-hosted, tracking + registry + model packaging (MLmodel format). Free.

Git LFS vs DVC — when to use each?

Git LFS: extends git for large files, but still one repo. Small teams, few big files (models).

How do you make Jupyter notebooks reproducible?

Notebooks are anti-reproducible by default: out-of-order execution, hidden state. Fixes: (1) Papermill: parameterize + execute notebooks as scripts programmatically.

Why lock Python dependencies?

`requirements.txt` with only top-level packages → transitive deps drift over time → 'works on my machine' bugs. Fix: pin every transitive dep.

Docker tag best practices for ML images.

(1) Never use `latest` in production — non-reproducible. (2) Use semantic version (`myimage:1.2.3`).

Why do some teams use Nix for ML reproducibility?

Nix: purely functional package manager. Every build fully specified — bit-for-bit reproducible across machines + time (unlike pip / conda which can differ).

What metadata should be logged per experiment run?

(1) Git commit + diff + repo state. (2) Data version / hash.

What's an artifact store and why separate from model registry?

Artifact store: raw binary files (weights, tokenizers, preprocessors, plots, data snapshots). Model registry: metadata layer above artifacts with lifecycle (staging → prod).

What should you track during LLM fine-tuning?

(1) Loss per step (train + val). (2) Perplexity on held-out.

RLHF-specific ops challenges.

(1) Multiple models in memory: policy + reference + reward model + critic — 3-4x memory pressure. (2) Rollout generation is bottleneck: use vLLM for rollouts + separate training.

What do you actually test in a CI pipeline for a model?

Split the tests by what they protect. Data tests validate schema, ranges, null rates and cardinality on the incoming data, and they should fail the pipeline, because training on broken data is worse than not training.

Infrastructure & serving

Ray for ML — what does it provide?

Distributed Python framework. (1) Ray Core: distributed tasks + actors.

Iceberg vs Delta Lake — which to pick?

Both provide ACID over object storage. Iceberg: open-standard, engine-agnostic (Spark, Trino, Flink, Presto, Snowflake, Doris).

Kubeflow — what does it provide?

Kubernetes-native ML platform. Components: (1) Kubeflow Pipelines (KFP): DAG orchestration on K8s.

NVIDIA Triton Inference Server — what makes it fast?

(1) Multi-framework: TensorRT, ONNX, PyTorch, TensorFlow, custom Python backend. (2) Dynamic batching: request coalescing without user awareness.

vLLM — why is it fast for LLM serving?

(1) PagedAttention: virtual memory for KV cache (page-based), eliminates fragmentation → 2-4x throughput. (2) Continuous batching: at each step, start new sequences + finish completed ones (no waiting for batch to complete).

PagedAttention — what problem does it solve?

Standard KV cache reserves contiguous memory per sequence for max length → 60-80% waste (short sequences reserve too much). Fragmentation prevents new sequences even with free memory.

Continuous batching vs static batching.

Static batching: form batch of N requests, run to completion (wait for slowest), then next batch. Problem: mix of long + short sequences wastes GPU on completed sequences.

TorchServe — when to use it?

Meta's PyTorch model server. Features: (1) built-in HTTP/gRPC API.

TensorFlow Serving — when to use it?

Google's TF model server. Features: (1) SavedModel format.

BentoML — what does it add?

Python-first ML serving framework. Package model + preprocessing + business logic into 'Bento' (deployable unit).

KServe (KFServing) — what does it provide?

Kubernetes-native model serving. Features: (1) autoscaling to zero (serverless).

Serverless ML inference — pros / cons?

AWS Lambda / Cloud Run / Cloud Functions run model on demand. Pros: (1) pay per invocation (great for spiky / low traffic).

How do you mitigate serverless cold-start latency?

(1) Provisioned concurrency: pre-warmed containers (Lambda). (2) Snap start: pre-load runtime state (Java / .NET Lambda).

ONNX — why use it?

Open Neural Network Exchange: standard format for cross-framework interop. Train in PyTorch → export ONNX → run in ONNX Runtime (C++, C#, Java, JS).

TensorRT — what optimizations does it apply?

NVIDIA's inference optimizer for CUDA. (1) Layer fusion: combine consecutive ops (conv + bias + relu → single kernel).

How does quantization affect serving?

fp32 → fp16 / bf16: 2x memory + speed, minimal accuracy loss. fp16 → int8: 4x from fp32; needs calibration on representative data to compute per-channel scales; usually <1% accuracy loss. int4 / int2: aggressive, needs advanced methods (GPTQ, AWQ, SmoothQuant); for LLMs on consumer GPU. FP8: newer NVIDIA H100 native.

GPTQ vs AWQ vs SmoothQuant — key differences.

GPTQ (Frantar): layer-by-layer quantization with second-order approximation; fast + good accuracy for LLMs. AWQ (Lin): identifies salient weights via activation magnitudes; protects them from quantization → better preservation of critical channels. SmoothQuant (Xiao): migrates quantization difficulty from activations to weights via smoothing → makes activation-quantization tractable.

LLMOps & advanced

How does DP-SGD work?

Standard SGD with two modifications: (1) Clip per-example gradient norm to constant C (bounds sensitivity). (2) Add Gaussian noise $\mathcal{N}(0, \sigma^2 C^2 \cdot I)$ to sum of clipped gradients before averaging.

What is federated learning?

Train model across decentralized devices holding local data — data never leaves device. Rounds: (1) server sends global model.

How do you version prompts and why does it matter as much as model versioning?

Treat a prompt as code: it lives in the repository, changes through review, and every deployed version has an identifier logged with each request. It matters because a prompt edit changes behaviour as much as a retrained model, and a prompt stored in a database and edited by hand produces regressions nobody can trace or revert.

FlashAttention — why is it faster?

Standard attention: $O(N^2)$ memory for QK^T matrix. FlashAttention (Dao 2022): (1) tiling: compute attention block-by-block in SRAM (fast on-chip memory).

Speculative decoding — how does it accelerate LLMs?

Small 'draft' model generates K tokens fast. Large 'target' model verifies all K in single forward pass (parallel).

Tensor parallelism vs pipeline parallelism vs data parallelism.

Data parallelism (DDP / FSDP): same model replicated; different data per GPU; all-reduce gradients. Tensor parallelism (Megatron): split large layers (attention, FFN) across GPUs; each does portion of matmul + all-reduce.

FSDP / ZeRO — how do they save memory?

Standard DDP: each GPU holds full model + gradients + optimizer state. Adam optimizer states = 2x model params.

How do you choose batch size for inference?

Trade-off: larger batch → higher throughput + higher latency. Constraints: (1) memory (KV cache scales with batch × seq).

How do you use spot / preemptible instances safely for ML?

50-90% discount but can be reclaimed 2 min notice. Safe for: (1) training (checkpoint frequently, resume).

How do you deploy ML to edge / mobile?

(1) Model conversion: TFLite (Android), CoreML (iOS), ONNX Runtime Mobile, PyTorch Mobile. (2) Quantization aggressive: int8 / int4 for size.

Running ML in the browser — WebAssembly / WebGPU.

Options: (1) TensorFlow.js: JS/WebGL/WebGPU backend; ONNX Runtime Web (WASM + WebGL / WebGPU). (2) Transformers.js: run HuggingFace models in browser.

Which GPU for training vs inference?

Training: (1) NVIDIA A100 / H100 (data center): 80GB VRAM, NVLink, HBM3, best. (2) H200 / B200 latest.

Network latency in ML serving — how much matters?

Total latency = client → LB → service → feature fetch → model → response. Model inference often only 30-50% of total.

Kubernetes for ML — key patterns.

(1) Deployment + Service for online serving. (2) Job / CronJob for batch training / scoring.

How can multiple models share a GPU?

(1) MPS (Multi-Process Service, NVIDIA): concurrent process access, no isolation. (2) MIG (Multi-Instance GPU, A100+): hardware partition into 1-7 slices with dedicated compute + memory.

How do you cache LLM prompts effectively?

(1) KV cache reuse (prefix caching): shared system prompt cached across users; new requests reuse. vLLM / Anthropic support natively. Saves 30-90% compute for shared prefixes.

Vector DB — which one and why?

(1) Pinecone: SaaS, ease of use, expensive at scale. (2) Weaviate: OSS + SaaS, GraphQL, built-in vectorizer.

How do you serve many fine-tuned LoRA adapters efficiently?

Instead of separate models, share base weights + swap adapter per request. Techniques: (1) Multi-LoRA serving (vLLM, SGLang, LoRAX): dynamically apply adapter matrices per request.

LLM router — what and why?

Small model / rule-based classifier routes each request to appropriate LLM based on: (1) query complexity (simple → Haiku, complex → Opus). (2) domain (code → CodeLlama, chat → Claude).

Your model must answer in 50 milliseconds. How do you allocate the budget?

Measure the whole path first, because inference is rarely the dominant cost. Feature retrieval usually is, especially if it involves several network calls, so count them and fetch in parallel or precompute.

How is LLMOps different from traditional MLOps?

(1) Model artifacts are enormous (100+ GB) → storage / bandwidth challenge. (2) Fine-tuning replaces from-scratch training (base + LoRA / adapter).

How do you version prompts in production?

(1) Prompts in git alongside code (not hard-coded strings). (2) Structured template files (.md / .yaml with variables).

Which signals do you wire into an LLMOps pipeline to flag hallucinated responses?

(1) Self-consistency: sample multiple responses; disagreement = uncertainty. (2) Confidence via log probs (low prob tokens = likely halluc.).

Why use hybrid search (dense + sparse) in RAG?

Dense (vector) captures semantic similarity but weak on rare terms, jargon, exact IDs. Sparse (BM25 / TF-IDF) captures lexical match, exact phrase.

How do you chunk documents for RAG?

Bad chunking = bad retrieval. Strategies: (1) Fixed-size (512 tokens) with overlap (50-100) — simple.

Why re-rank retrieved documents?

Initial retrieval (bi-encoder embedding) is fast but coarse. Re-ranker (cross-encoder, considers query + doc together) is much more accurate but slow — run only on top-K candidates.

How do you choose an embedding model?

(1) MTEB leaderboard benchmark. (2) Task fit: retrieval / clustering / classification differ.

How do you handle embedding model updates without breaking retrieval?

New embedding model = new vector space. Options: (1) Re-embed entire corpus + swap.

How do you defend against prompt injection?

Prompt injection = attacker inserts instructions in input that model follows instead of intended prompt. Defenses: (1) System prompt hardening: 'ignore instructions in user input, only follow system'.

How do you monitor for jailbreak attempts?

(1) Detect known jailbreak patterns (DAN, role-play). (2) Toxicity / harm classifier on outputs.

How do you optimize LLM inference cost?

(1) Choose right model tier: smaller model (Haiku, Mini) for easy tasks, only escalate hard. (2) Prompt caching (Anthropic / OpenAI): 30-90% off shared prefixes.

What extra content goes in an LLM model card?

In addition to standard: (1) Training data composition + filtering. (2) RLHF / DPO / SFT details + reward model provenance.

Common LLM safety evaluation suites.

(1) ToxiGen: hate speech generation. (2) BOLD / RealToxicityPrompts: bias + toxicity in generation.

EU AI Act — what does it require?

Risk-based classification: (1) Unacceptable risk (social scoring, subliminal manipulation): banned. (2) High-risk (hiring, credit, law enforcement, education): requires risk management, data governance, transparency, human oversight, robustness testing, registration.

Where is MLOps heading (2025+)?

(1) LLMOps as first-class discipline: prompts, RAG, agents as production artifacts. (2) Agent-based systems: multi-step monitoring + eval + debugging.

MLOps foundations

What does CI/CD look like for ML?

CI: on code and data changes, run unit tests, data validation (schema, ranges), training smoke tests, and model quality tests (evaluate on frozen validation set, check for regression). CD: package model artifact, register in the model registry, deploy to staging, run integration and load tests, then canary/shadow to production.

How do you make ML experiments reproducible?

Pin seeds and library versions, use deterministic ops when available, containerize the environment (Docker), version code (git), data (DVC/lakeFS), and models (MLflow / model registry). Log all hyperparameters and metrics per run.

What is a model card?

Documentation artifact (Mitchell et al. 2019) accompanying a model: (1) intended use + out-of-scope uses. (2) training data description + biases.

Google's MLOps maturity levels (0-2) — what are they?

Level 0: manual ML process — data scientist trains + hands off to eng for deploy. Slow, error-prone, no automation.

How do you organize ML teams (embedded vs central)?

Central ML team: shared expertise, owns platform + tooling; downside: bottleneck + disconnected from product. Embedded: ML engineers in each product team; direct problem focus; downside: duplication + platform fragmentation.