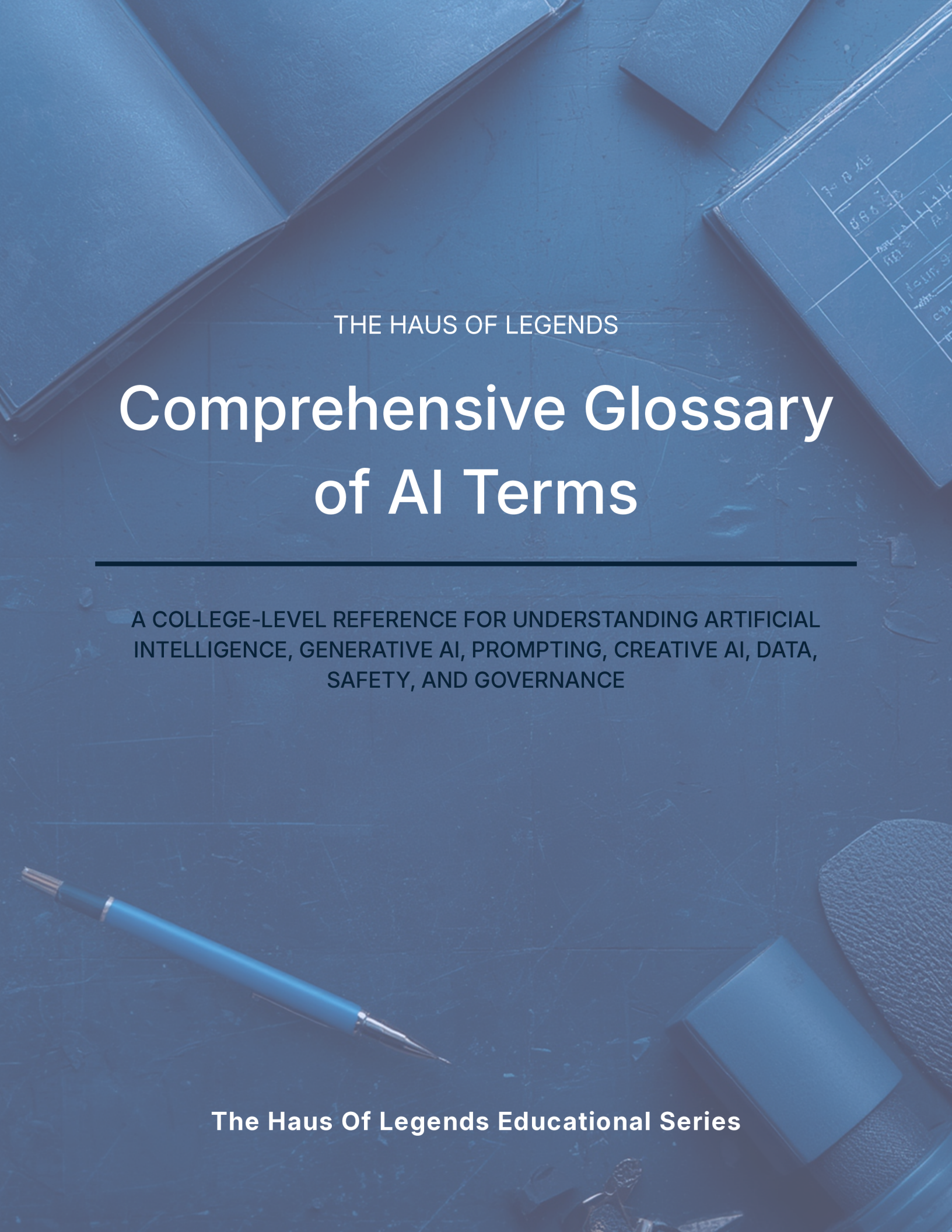




THE HAUS OF LEGENDS

Comprehensive Glossary of AI Terms

A COLLEGE-LEVEL REFERENCE FOR UNDERSTANDING ARTIFICIAL
INTELLIGENCE, GENERATIVE AI, PROMPTING, CREATIVE AI, DATA,
SAFETY, AND GOVERNANCE

The Haus Of Legends Educational Series

Purpose of This Glossary

This glossary is designed for students, creators, educators, entrepreneurs, and general readers who want to understand artificial intelligence beyond buzzwords. It is written at a college learning level: clear enough for non-engineers, but precise enough to support serious study, professional discussion, and responsible creative practice.

The Haus Of Legends frames AI as part of a longer history of artists and makers adopting new tools. Photography, printmaking, digital art, Photoshop, design software, music production tools, and now AI have all changed the creative process. The tool changes; human judgment, intention, taste, and responsibility remain central.

Use this document as a companion reference. The terms are grouped by topic so readers can build understanding progressively, then locate individual terms alphabetically in the index. Definitions intentionally emphasize practical meaning, conceptual importance, and responsible use.

How to Read the Glossary

- Start with Core AI Concepts if you are new to the field.
- Move into Machine Learning, Data, and Generative AI to understand how systems are built and used.
- Use the Safety, Ethics, Governance, Security, and Privacy sections to understand why responsible use matters.
- Use the Creative and Business Workflow section to connect AI terminology to real production, publishing, education, and brand work.

Quick Taxonomy

Area	What It Helps You Understand	Examples
Core AI	The basic language of intelligent systems and automation.	AI, model, inference, workflow
Machine Learning	How systems learn patterns and are trained.	training, loss, overfitting, fine-tuning
Data	What models learn from and how information is represented.	dataset, token, embedding, vector database
Generative AI	How systems produce text, images, audio, video, and code.	LLM, diffusion model, prompt, RAG
Evaluation	How quality, accuracy, and performance are measured.	precision, recall, evals, calibration
Governance & Safety	How risk, trust, rights, and responsibility are managed.	bias, transparency, red teaming, provenance
Security & Privacy	How systems can be attacked or misused.	prompt injection, data poisoning, PII
Creative Workflow	How AI fits into art, design, media, and brand production.	co-creation, curation, generative fill

Categories Included

- A. Core AI Concepts (18 terms)
- B. Machine Learning and Model Development (34 terms)
- C. Data, Datasets, and Representation (30 terms)
- D. Language Models and Generative AI (37 terms)
- E. Evaluation, Quality, and Performance (22 terms)
- F. AI Safety, Ethics, and Governance (26 terms)
- G. Security, Privacy, and Misuse (22 terms)
- H. Creative, Media, and Business Workflow Terms (26 terms)
- I. Deployment, Operations, and Enterprise AI (25 terms)
- J. Advanced and Emerging Concepts (25 terms)

A. Core AI Concepts

Algorithm

A formal sequence of steps used to solve a problem or complete a task. In AI, algorithms determine how data is processed, how patterns are detected, and how outputs are generated.

Artificial Intelligence (AI)

The broad field of building computational systems that perform tasks normally associated with human intelligence, such as reasoning, language use, perception, planning, classification, prediction, and decision support.

Automation

The use of software or machines to perform repeatable tasks with limited human intervention. AI-enabled automation differs from simple automation because it can adapt to patterns, language, images, or probabilistic signals.

Autonomy

The degree to which a system can act without direct human control. Higher autonomy increases efficiency but also requires stronger oversight, monitoring, and accountability.

Cognitive Computing

A business-oriented term for systems that simulate aspects of human cognition, including language processing, pattern recognition, and decision support. It is often used in enterprise contexts rather than academic AI research.

Computational Intelligence

A subfield of AI focused on adaptive, biologically inspired, or probabilistic methods such as neural networks, evolutionary algorithms, and fuzzy logic.

Decision Support System

A system that helps humans evaluate options, risks, or recommendations. In responsible AI practice, decision support should assist human judgment rather than quietly replace it in high-stakes contexts.

Expert System

Another form of AI that uses explicitly written rules and expert knowledge to make recommendations or solve problems. Unlike modern machine learning, expert systems usually do not learn patterns from large datasets.

Human-in-the-Loop

A design approach in which humans remain part of the workflow by reviewing, approving, correcting, or guiding AI outputs. This is essential when accuracy, ethics, brand identity, safety, or legal risk matters.

Inference

The process of using a trained model to produce an output from new input. When a person types a prompt into a chatbot and receives a response, the model is performing inference.

Intelligence Amplification

The use of AI to extend human capability rather than replace human agency. This framing treats AI as a partner for research, drafting, design, exploration, and decision support.

Knowledge Representation

The method by which information is structured so a system can use it. Examples include rules, ontologies, embeddings, knowledge graphs, and vector representations.

Model

A computational structure that has learned patterns from data or rules and can produce predictions, classifications, recommendations, or generated content. A model is not conscious; it is a mathematical system that maps inputs to outputs.

Narrow AI

AI designed for a specific class of tasks, such as translation, image generation, search, recommendation, or speech recognition. Most AI used today is narrow AI.

Probabilistic System

A system that produces outputs based on likelihood rather than absolute certainty. Many AI systems estimate the most probable answer, classification, word, image feature, or action given available context.

Symbolic AI

An AI approach based on explicit symbols, rules, and logic. It contrasts with neural approaches that learn statistical representations from data.

Task

A defined objective the AI system is expected to perform, such as summarization, image classification, translation, retrieval, or recommendation. Clear task definition improves evaluation and responsible use.

Workflow

A sequence of steps used to complete work. AI becomes more useful when it is placed inside a repeatable workflow rather than treated as a one-time magic answer machine.

B. Machine Learning and Model Development

Ablation Study

An experiment that removes or changes one component of a system to measure its contribution. Ablation helps researchers identify which parts of a model or workflow actually matter.

Activation Function

A mathematical function inside a neural network that determines how strongly a unit responds to input. Activation functions help neural networks learn complex, non-linear patterns.

Backpropagation

The algorithmic process used to adjust neural network weights by calculating how much each part of the model contributed to error. It is central to training deep learning systems.

Batch Size

The number of training examples processed before the model updates its parameters. Batch size affects training speed, memory use, and learning behavior.

Benchmark

A standardized test or dataset used to compare model performance. Benchmarks are useful but can be misleading if they do not match real-world use.

Classification

A task in which a model assigns an input to a category, such as spam or not spam, benign or malignant, approved or denied.

Clustering

An unsupervised method that groups similar examples together. Clustering is often used for customer segmentation, topic discovery, and pattern exploration.

Cross-Validation

An evaluation method that trains and tests a model on different splits of the data to estimate how well it may generalize to unseen examples.

Decision Boundary

The conceptual line, surface, or region separating one model decision from another. In classification, it marks where an input changes from one predicted class to another.

Deep Learning

A form of machine learning based on neural networks with many layers. Deep learning is central to modern computer vision, speech recognition, translation, and generative AI.

Distillation

A process in which a smaller model is trained to imitate the behavior of a larger model. Distillation can reduce cost and latency while preserving some performance.

Epoch

One full pass through the training dataset during model training. Multiple epochs may improve learning, but too many can contribute to overfitting.

Feature

A measurable property used by a model, such as age, color, word frequency, pixel intensity, or transaction amount. In deep learning, models often learn features automatically.

Feature Engineering

The process of selecting, transforming, or constructing useful inputs for a model. Traditional ML often relies heavily on feature engineering; deep learning reduces but does not eliminate its importance.

Fine-Tuning

Additional training applied to a pre-trained model so it performs better for a specific task, style, domain, or dataset.

Generalization

A model's ability to perform well on new data, not just the examples it saw during training. Good generalization is one of the central goals of machine learning.

Gradient Descent

An optimization method that updates model parameters in the direction that reduces error. It is one of the basic engines of neural network training.

Hyperparameter

A setting chosen before or during training, such as learning rate, batch size, number of layers, or temperature at inference time. Hyperparameters shape how a model learns or behaves.

Instruction Tuning

Training that teaches a model to follow natural-language instructions more effectively. It is a major reason modern chatbots can respond to user requests conversationally.

Loss Function

A mathematical measure of how wrong a model's output is compared with the desired answer. Training tries to minimize this loss.

Machine Learning (ML)

A branch of AI in which systems learn patterns from data instead of being programmed with every rule manually. ML systems improve by identifying statistical relationships in examples.

Neural Network

A model architecture inspired loosely by networks of neurons. It consists of layers of mathematical units that transform input data into increasingly abstract representations.

Optimizer

The algorithm that controls how a model updates its weights during training. Examples include stochastic gradient descent and Adam.

Overfitting

A failure mode in which a model learns the training data too closely, including noise or accidental patterns, and performs poorly on new data.

Parameter

A learned numerical value inside a model. In large neural networks, parameters encode patterns discovered during training but do not directly store knowledge in a human-readable way.

Regression

A task in which a model predicts a numerical value, such as price, demand, temperature, or probability of churn.

Reinforcement Learning

A learning method in which an agent learns by taking actions in an environment and receiving rewards or penalties. It is used in robotics, games, optimization, and some alignment techniques.

Self-Supervised Learning

A learning method in which a model creates training signals from the structure of the data itself. Many foundation models are trained this way by predicting missing or next elements in text, images, or other data.

Semi-Supervised Learning

A training approach that uses a small amount of labeled data combined with a larger amount of unlabeled data. It is useful when labels are expensive or difficult to obtain.

Supervised Learning

A machine learning method in which a model is trained on labeled examples, such as images tagged as cat or dog, or emails labeled spam or not spam.

Training

The process of adjusting a model's parameters using data so it can perform a task. Training can require large datasets, specialized hardware, and repeated evaluation.

Transfer Learning

A technique in which a model trained for one purpose is adapted to another related purpose. Foundation models rely heavily on this principle.

Underfitting

A failure mode in which a model is too simple or insufficiently trained to capture important patterns in the data.

Unsupervised Learning

A machine learning method that detects patterns in unlabeled data, often by clustering similar examples or reducing complex data into simpler representations.

C. Data, Datasets, and Representation

Annotation

Human or automated markup added to data, such as bounding boxes around objects, sentiment tags, transcript corrections, or content categories.

Class Imbalance

A dataset condition in which one category appears much more often than another. Imbalance can cause a model to perform poorly on rare but important cases.

Concept Drift

A change in the relationship between inputs and outputs over time. For example, words, risks, customer needs, or fraud patterns may change meaning or behavior.

Context Window

The maximum amount of text, code, images, tool results, or conversation history a model can consider at once. A larger context window allows more material, but relevance and structure still matter.

Corpus

A large collection of text or other materials used for analysis or training. Language models are often trained on enormous corpora containing many genres and domains.

Data

Information used by a system, including text, images, audio, video, tables, clicks, sensor readings, or human labels. Data quality strongly influences AI behavior.

Data Augmentation

A technique that expands training data by creating modified versions of existing examples, such as rotated images, paraphrased sentences, or altered audio samples.

Data Cleaning

The process of removing errors, duplicates, inconsistencies, corrupted records, or irrelevant information from a dataset.

Data Curation

The intentional selection, organization, documentation, and refinement of data. In creative and educational AI work, curation is often as important as generation.

Data Drift

A change in real-world data over time that makes a model less accurate or less reliable. Drift often appears when user behavior, markets, language, or environments change.

Data Provenance

The record of where data came from, how it was collected, how it was transformed, and who handled it. Provenance supports accountability, licensing, reproducibility, and trust.

Dataset

A structured collection of examples used for training, testing, evaluation, or analysis. A dataset may include raw inputs, labels, metadata, and documentation.

De-Identification

The process of removing or masking personal identifiers from data. De-identification reduces privacy risk but does not always eliminate the possibility of re-identification.

Embedding

A numerical representation of an item such as a word, sentence, image, document, or user preference. Embeddings place similar items closer together in mathematical space.

Federated Learning

A training method in which models learn from data across multiple devices or institutions without centralizing all raw data in one place.

Ground Truth

The best available answer used for comparison during training or evaluation. In many social, creative, or interpretive tasks, ground truth may be contested rather than absolute.

Knowledge Graph

A structured network of entities and relationships. Knowledge graphs help systems represent facts, relationships, context, and domain knowledge.

Label

The target answer attached to a training example, such as a category, score, annotation, or correct output. Labels can be produced by humans, systems, or rules.

Label Leakage

A flaw in which training data includes information that would not be available in real use, making the model look more accurate than it actually is.

Latent Space

A compressed mathematical space where a model represents meaningful features that are not directly visible in raw data. Generative models often manipulate latent space to create new outputs.

Metadata

Data about data, such as file type, author, timestamp, location, license, camera settings, or source. Metadata can improve organization but may also expose sensitive information.

Ontology

A formal description of concepts and relationships within a domain. Ontologies support consistent classification, reasoning, and data integration.

Synthetic Data

Artificially generated data used for training, testing, privacy protection, simulation, or balancing underrepresented cases. Synthetic data can help, but it may also amplify errors if poorly designed.

Test Data

A dataset held back until evaluation so developers can estimate performance on unseen examples. Proper test data should resemble the real-world task.

Token

A unit of text processed by a language model. A token may be a word, part of a word, punctuation mark, number, or symbol depending on the tokenizer.

Tokenization

The process of breaking text into tokens so a language model can process it. Tokenization affects cost, context length, and how a model handles unusual words or formatting.

Training Data

The data used to teach a model patterns. Training data affects what a model can do, what it may ignore, and what biases it may reproduce.

Validation Data

A dataset used during development to tune settings and evaluate performance before final testing. It helps detect overfitting and guide model improvement.

Vector

An ordered list of numbers. In AI, vectors often represent meaning, visual features, user behavior, or document similarity.

Vector Database

A database designed to store and search embeddings efficiently. It is commonly used for semantic search, retrieval-augmented generation, and recommendation systems.

D. Language Models and Generative AI

Agent

An AI system that can plan, use tools, remember task state, and take multi-step actions toward a goal. Agentic systems need clear permissions, monitoring, and boundaries.

Attention Mechanism

A model component that helps determine which parts of the input are most relevant to producing the next representation or output. Attention allows models to relate distant words, pixels, or features.

Autoencoder

A neural network trained to compress input into a smaller representation and then reconstruct it. Autoencoders are used for representation learning, denoising, anomaly detection, and compression.

Autoregressive Model

A model that generates output one step at a time, where each new token or element depends on what came before. Many chat models generate text this way.

Chain-of-Thought Prompting

A prompting family that encourages stepwise reasoning for complex tasks. In many products, users should ask for concise reasoning or a final explanation rather than hidden internal reasoning traces.

Completion

The model's generated response to a prompt. In older terminology, the prompt is the input and the completion is the output.

Confabulation

A term sometimes used for AI outputs that invent details while maintaining fluent presentation. It highlights that fluent language is not the same as verified knowledge.

Decoder

A model component that generates output from internal representations. Many text-generation systems use decoder-based architectures to predict the next token.

Diffusion Model

A generative model that learns to create data by reversing a noise process. Diffusion models are widely used in image generation and increasingly in audio, video, and 3D workflows.

Encoder

A model component that converts input into internal representations. Encoders are common in classification, search, translation, and multimodal understanding.

Few-Shot Prompting

Providing a small number of examples inside the prompt so the model can imitate the pattern, format, tone, or reasoning style.

Foundation Model

A large, general-purpose model trained on broad data and adaptable to many downstream tasks. Foundation models often support chatbots, coding tools, image generators, search assistants, and agentic systems.

Function Calling

A method that allows a model to request the use of external tools or functions, such as searching a database, calculating values, creating calendar events, or retrieving files.

Generative Adversarial Network (GAN)

A generative architecture in which two networks compete: a generator creates examples and a discriminator evaluates them. GANs were important in image synthesis before diffusion models became dominant.

Generative AI

AI that produces new content, such as text, images, audio, video, code, or design variations. Its output is generated from learned patterns rather than copied mechanically from a single source.

Grounding

The practice of connecting AI outputs to specific sources, documents, data, or observable evidence. Grounding reduces unsupported claims but does not remove the need for verification.

Hallucination

A generated output that sounds plausible but is false, unsupported, fabricated, or incorrectly inferred. Hallucination is a serious issue when users rely on AI for facts, citations, legal language, health information, or technical accuracy.

JSON Schema

A formal structure describing the expected shape of JSON data. AI applications use schemas to request outputs that are easier for software to validate and process.

Large Language Model (LLM)

A large neural network trained on extensive text data to predict and generate language. LLMs can summarize, draft, translate, classify, reason probabilistically, and interact conversationally.

Latent Diffusion

A diffusion approach that performs generation in a compressed latent space rather than directly in pixel space. This can reduce computational cost while preserving expressive capability.

Multimodal AI

AI that can process or generate more than one type of data, such as text, image, audio, video, or code. Multimodal systems can interpret a screenshot, discuss a chart, generate captions, or combine media.

Positional Encoding

Information added to token representations so a transformer can understand order. Without positional information, a model would have difficulty distinguishing sequence relationships.

Prompt

The instruction, question, context, or input given to an AI model. A strong prompt defines the task, audience, constraints, source material, desired output, and evaluation criteria.

Prompt Engineering

The practice of designing effective instructions so a model consistently produces useful outputs. At a college level, it combines rhetoric, task design, constraints, examples, and evaluation.

Retrieval-Augmented Generation (RAG)

A method that retrieves relevant documents or data and supplies them to a generative model before it answers. RAG can improve factual grounding when the retrieved sources are relevant and trustworthy.

Sampling

The process of selecting the next token, pixel, note, or output element from a probability distribution. Sampling settings influence creativity, consistency, and randomness.

Self-Attention

A form of attention in which different parts of the same input sequence compare with one another. It helps a model understand relationships such as pronouns, dependencies, context, and sequence structure.

Semantic Search

Search based on meaning rather than exact keyword matching. It often uses embeddings to find conceptually similar documents, images, or passages.

Structured Output

AI output constrained into a predictable format, such as JSON, a table, a form, or a schema. Structured outputs are useful when AI results must feed into software workflows.

System Prompt

A higher-priority instruction that shapes a model's behavior, role, boundaries, or style. System prompts are typically controlled by the application rather than the end user.

Temperature

A generation setting that controls randomness. Lower temperature usually produces more predictable output; higher temperature usually produces more varied or surprising output.

Tool Calling

A broader term for model-assisted use of external software tools. Tool calling can make AI more useful because the model can access current data, perform actions, or use specialized systems.

Top-p Sampling

A generation method that restricts possible next tokens to a probability mass threshold. It can balance variety with coherence.

Transformer

A neural network architecture based on attention mechanisms that process relationships among tokens in a sequence. Transformers are foundational to modern language models and many multimodal systems.

User Prompt

The instruction provided by the person interacting with the model. Good user prompts are specific, contextual, and clear about the desired result.

Variational Autoencoder (VAE)

A generative model that learns compressed latent representations and can generate variations. VAEs are often discussed in relation to latent spaces and image-generation systems.

Zero-Shot Prompting

Asking a model to perform a task without providing examples. It works best for common tasks the model can infer from clear instructions.

E. Evaluation, Quality, and Performance

Accuracy

The proportion of predictions a model gets correct. Accuracy can be misleading when datasets are imbalanced or when some errors are more serious than others.

Benchmark Leakage

A condition in which benchmark examples appear in training data or are indirectly optimized against, making reported performance seem stronger than real generalization.

Calibration

The degree to which a model's confidence scores match real-world probabilities. A well-calibrated model that says 80 percent confidence should be correct about 80 percent of the time in similar cases.

Confusion Matrix

A table comparing predicted classes with actual classes. It shows true positives, true negatives, false positives, and false negatives.

Evals

Tests designed to measure whether an AI system performs acceptably on selected tasks, behaviors, or risk scenarios. Evals can be automated, human-scored, or hybrid.

Evaluation

The systematic assessment of model behavior, output quality, safety, reliability, or task performance. Evaluation should match the actual use case, not just generic benchmarks.

F1 Score

A metric that balances precision and recall using their harmonic mean. It is useful when both false positives and false negatives matter.

False Negative

A case in which a model fails to identify something that is present or true. For example, missing a fraudulent transaction.

False Positive

A case in which a model incorrectly identifies something as present or true. For example, flagging a legitimate email as spam.

Human Evaluation

Assessment performed by people, often used for tasks where quality is subjective, such as writing style, image aesthetics, helpfulness, or relevance.

Latency

The time it takes for a system to produce a response. In AI products, latency affects user experience, cost, and workflow efficiency.

Metric

A numerical measure used to evaluate performance, such as accuracy, precision, recall, F1 score, latency, cost, or human preference rating.

Model Card

A document that describes a model's intended use, training data, evaluation results, limitations, ethical considerations, and appropriate deployment contexts.

Observability

The ability to monitor a system's behavior, performance, failures, cost, and user interactions in operation. AI observability is central to production reliability.

Precision

The proportion of positive predictions that are actually correct. High precision matters when false positives are costly.

Recall

The proportion of actual positive cases the model successfully identifies. High recall matters when missing a case is costly.

Reliability

Consistent performance under expected operating conditions. Reliability is about whether the system behaves dependably over time, not just whether it works once.

Reproducibility

The ability to obtain the same or comparable results when the same method, data, and settings are used. Reproducibility supports scientific and operational trust.

Robustness

A model's ability to maintain acceptable performance when inputs are noisy, incomplete, adversarial, or different from ideal training examples.

Side-by-Side Evaluation

A comparison method in which reviewers judge outputs from two or more models on the same task. It is often used for preference testing.

Throughput

The amount of work a system can process in a given time, such as requests per second or documents per hour.

Validity

The extent to which a system measures or performs what it claims to measure or perform. A valid AI system must be appropriate for its intended context.

F. AI Safety, Ethics, and Governance

Accountability

The ability to identify who is responsible for an AI system's design, deployment, outputs, impacts, and correction processes.

AI Governance

The policies, roles, processes, documentation, and oversight mechanisms used to control how AI is developed and used. Governance turns ethics into operational practice.

Alignment

The effort to make AI behavior consistent with human intentions, values, policies, and safety boundaries. Alignment is difficult because human goals can be ambiguous, conflicting, or context-dependent.

Bias

A systematic tendency in data, model behavior, or decision processes that can produce unfair, inaccurate, or harmful outcomes. Bias may originate from historical data, sampling choices, labels, objectives, or deployment context.

Black Box

A system whose internal decision process is difficult to inspect or understand. Many deep learning systems are considered black boxes because their learned representations are complex.

Consent

Permission from a person or rights holder for data use, likeness use, voice cloning, training, publication, or transformation. Consent is a central issue in generative media.

Content Credentials

A content authenticity approach that attaches metadata about creation, editing, and provenance to digital media. It is often discussed in relation to AI-generated or AI-edited content.

Content Provenance

Information about where content came from, how it was created, and whether AI was involved. Provenance helps audiences evaluate authenticity, authorship, and manipulation.

Disclosure

A clear statement that AI was used in a process or output. In creative and educational contexts, disclosure can build trust when it is honest, specific, and not exaggerated.

Disparate Impact

A pattern in which a system affects groups differently, even if the system does not explicitly use protected characteristics. It is a major concern in employment, housing, credit, education, and public services.

Explainability

The ability to describe how a system reached a result in terms understandable to its audience. Explainability is especially important in high-stakes or contested decisions.

Fairness

The effort to identify, reduce, and monitor harmful differences in AI behavior across people, groups, contexts, or outcomes. Fairness requires both technical analysis and social judgment.

Guardrails

Rules, filters, policies, classifiers, prompts, or system designs that limit harmful, unsafe, off-brand, illegal, or low-quality outputs.

Harmful Bias

Bias that produces unjust, unsafe, exclusionary, discriminatory, or materially damaging outcomes. Responsible AI focuses on detecting and managing harmful bias rather than pretending all bias can be eliminated.

Impact Assessment

A structured analysis of how an AI system may affect people, organizations, rights, safety, opportunity, labor, culture, or the environment.

Incident Disclosure

The process of reporting significant AI failures, harms, vulnerabilities, or misuse events. Disclosure supports accountability and learning across organizations.

Interpretability

The degree to which a human can understand the internal logic or relationships used by a model. Interpretability is related to but not identical with explainability.

Red Teaming

A structured testing process in which evaluators deliberately try to make an AI system fail, behave unsafely, leak information, or violate policy. Red teaming helps uncover risks before deployment.

Responsible AI

The discipline of designing, deploying, and monitoring AI systems in ways that account for human rights, safety, fairness, transparency, accountability, privacy, security, and social impact.

Risk Management

The process of identifying, measuring, prioritizing, reducing, and monitoring risks. In AI, risk can arise from design, data, deployment, misuse, outputs, and organizational context.

Safety

The property that an AI system does not create unacceptable risk to people, property, rights, institutions, or the environment in its intended use.

Safety Filter

A mechanism that detects and blocks certain categories of input or output, such as violent instructions, private data exposure, malware assistance, or explicit harmful content.

Transparency

The degree to which users, stakeholders, or regulators can understand that AI is being used, what it is being used for, and what its limitations are.

Trustworthy AI

AI that can be reasonably relied upon in its intended context because it is valid, reliable, safe, secure, resilient, accountable, transparent, explainable, privacy-conscious, and fair with harmful bias managed.

Value Alignment

A subset of alignment focused on making systems behave according to specified human values or institutional principles. It requires translating abstract values into operational constraints.

Watermarking

A technique for embedding detectable signals into content to indicate origin or generation method. AI watermarking can support provenance but is not always reliable or tamper-proof.

G. Security, Privacy, and Misuse

Access Control

Rules that determine who can view, use, modify, export, or delete data and systems. AI workflows require strong access control because prompts and outputs may contain sensitive material.

Adversarial Example

An input intentionally modified to cause a model to make a mistake. In image systems, the change may be nearly invisible to humans but influential to the model.

Adversarial Machine Learning

The study of attacks and defenses involving machine learning systems. It includes adversarial examples, poisoning, evasion, extraction, and prompt-based attacks.

Context Poisoning

The insertion of misleading, malicious, or irrelevant information into the context window so the model produces compromised outputs.

Data Poisoning

An attack in which malicious or corrupted data is inserted into training or retrieval sources to change system behavior.

Data Retention

Policies governing how long data is stored and when it is deleted. Users should understand whether prompts, files, or outputs may be retained by a tool provider.

Deepfake

Synthetic media that convincingly depicts a person saying or doing something they did not actually say or do. Deepfakes raise issues of consent, fraud, reputation, and evidence integrity.

Differential Privacy

A privacy framework that adds mathematical noise or constraints so aggregate analysis reveals less about any individual record.

Jailbreak

A prompt or interaction designed to bypass model safety rules, system instructions, or usage policies. Jailbreaks are a form of adversarial behavior.

Malware Assistance

AI-enabled help with creating, modifying, deploying, or evading malicious software. Responsible AI systems restrict this behavior.

Membership Inference

An attack that tries to determine whether a specific record or person's data was included in a model's training dataset.

Memorization

A model behavior in which specific training examples are retained closely enough to be reproduced or inferred. Memorization raises privacy and copyright concerns.

Misuse

The use of an AI system for harmful, deceptive, illegal, or unintended purposes. Misuse can come from insiders, outsiders, customers, or automated agents.

Model Extraction

An attack that attempts to copy or approximate a model by repeatedly querying it and analyzing outputs.

Model Inversion

An attack that attempts to infer sensitive training data or input features from model outputs or parameters.

Personally Identifiable Information (PII)

Information that can identify a person directly or indirectly, such as name, address, Social Security number, email, phone number, biometric data, or combinations of details.

Privacy-Enhancing Technology

A technical method designed to reduce privacy risk, such as de-identification, differential privacy, federated learning, encryption, or access controls.

Prompt Injection

An attack in which a user or document includes instructions that try to override the intended behavior of an AI system. It is especially relevant when models read web pages, emails, files, or tool outputs.

Secure and Resilient AI

AI designed to withstand attacks, failures, misuse, and changing conditions while continuing to operate appropriately or fail safely.

Secure Deployment

The process of releasing an AI system with appropriate access controls, monitoring, testing, documentation, incident response, and rollback procedures.

Sensitive Data

Data that requires special protection because exposure could cause harm. Examples include health, financial, legal, biometric, precise location, personal identity, and confidential business information.

Voice Cloning

The generation or imitation of a person's voice using AI. It can be useful for accessibility and production, but risky when used without consent or for deception.

H. Creative, Media, and Business Workflow Terms

AI-Assisted Creation

A creative process in which a human uses AI as one tool among others for ideation, drafting, image generation, editing, variation, or refinement. The human remains responsible for taste, direction, curation, and final use.

Brand Voice

The consistent tone, vocabulary, rhythm, and point of view used by a brand. AI can help draft in a brand voice, but humans should define and approve that voice.

Co-Creation

A workflow in which human and AI contributions are iteratively combined. Co-creation can include prompting, selecting, editing, remixing, revising, and contextualizing outputs.

Commercial Use

Use connected to business, promotion, sales, paid services, products, or monetized distribution. AI tools may have different rules for personal, educational, and commercial use.

Creative Direction

The human decision-making process that defines concept, audience, emotional tone, visual language, story, constraints, and final judgment. AI can assist execution but does not replace direction.

Curation

The act of selecting, arranging, editing, and presenting work according to a concept or standard. In AI workflows, curation separates intentional work from mass output.

Derivative Work

Work based on or adapted from existing protected material. AI-assisted creation can raise derivative-work questions when outputs closely follow source material or protected characters, images, voices, or text.

Editorial Review

Human review of AI-assisted output for accuracy, tone, ethics, legal risk, grammar, originality, and brand alignment.

Generative Fill

An image-editing workflow that fills selected areas with AI-generated content that attempts to match surrounding context.

Image-to-Image Generation

A workflow that uses an existing image as input to guide the generation of a new or modified image.

Inpainting

Editing inside a selected region of an image, often to remove, replace, or repair visual elements.

Iteration

Repeated cycles of generation, critique, revision, and refinement. Iteration is one of the most practical advantages of AI-assisted creative work.

Licensing

The legal permission structure governing how data, models, images, music, fonts, or outputs may be used. AI creators should check tool terms, source rights, and commercial permissions.

Negative Prompt

An instruction, common in image generation, describing what should be avoided, such as blur, extra fingers, distorted text, low resolution, or unwanted objects.

Noncommercial Use

Use that is not primarily intended for commercial advantage or monetary compensation. Noncommercial permission does not automatically permit resale, advertising, or product creation.

Outpainting

Extending an image beyond its original boundaries by generating new surrounding content.

Prompt Library

A saved collection of reusable prompts organized by task, style, audience, format, or workflow. Prompt libraries support consistency and efficiency.

Seed

A numerical value used to influence random generation. Reusing a seed can help reproduce or vary outputs in some generative systems.

Speech-to-Text

The transcription of spoken audio into written text. It is used for captions, interviews, meetings, accessibility, and research.

Style Prompt

A prompt component that describes the desired tone, visual style, genre, medium, mood, or reference vocabulary. Style prompts should avoid infringing protected living artists' styles where tool policies or ethics require caution.

Style Transfer

A technique that applies the visual or aesthetic characteristics of one image, genre, or medium to another. It raises questions of authorship, originality, and influence.

Synthetic Voice

A generated voice that may be fictional or based on a real speaker. Synthetic voices require careful consent, labeling, and licensing practices.

Text-to-Image Generation

The creation of images from written prompts. The quality of output depends on prompt clarity, model capability, style constraints, and post-processing.

Text-to-Speech

The generation of spoken audio from written text. AI voices can support accessibility, narration, education, and production when consent and licensing are handled correctly.

Text-to-Video Generation

The creation of video clips from written prompts. Current workflows often require strong human editing, continuity management, and post-production.

Upscaling

Increasing image, video, or audio resolution or quality using computational methods. AI upscalers infer missing detail but can also introduce artifacts.

I. Deployment, Operations, and Enterprise AI

API

An application programming interface: a structured way for software systems to communicate. AI APIs allow applications to send prompts, retrieve outputs, use tools, or process data.

Application Layer

The user-facing software built around an AI model, including interface, data connections, permissions, prompts, workflows, logging, and business rules.

Audit Log

A record of actions, requests, decisions, changes, or access events. Audit logs help organizations investigate errors, misuse, security incidents, and compliance questions.

Caching

Storing previous results so repeated requests can be answered faster or more cheaply. Caching can reduce latency and cost but must respect privacy and freshness requirements.

Closed Model

A model whose weights and internal training details are controlled by a provider. Users typically access it through an application or API.

Cloud AI

AI services delivered through cloud infrastructure. Cloud AI provides scalability and managed services but requires attention to data handling, cost, and vendor dependence.

Edge AI

AI processing that happens on a device or local environment rather than a central cloud server. Edge AI can reduce latency and improve privacy in some contexts.

Endpoint

A network-accessible location where software can send requests to an AI service. In API usage, prompts and data are sent to endpoints for processing.

Fallback

A backup behavior used when an AI system fails, times out, produces low confidence, or encounters restricted content. Fallbacks may include human review, simpler models, or refusal.

Human Escalation

A process that routes uncertain, sensitive, or high-impact AI cases to a human reviewer or specialist.

LLMOps

Operational practices for deploying and managing large language model applications. It includes prompt management, evaluation, retrieval systems, cost control, monitoring, and safety controls.

MLOps

Machine learning operations: the discipline of deploying, monitoring, updating, and managing ML systems in production environments.

Model Deployment

The process of making a model available for use in an application, product, or internal workflow. Deployment requires performance, security, monitoring, and fallback planning.

Model Monitoring

Continuous observation of model performance, errors, drift, cost, latency, and user feedback after deployment.

On-Premises Deployment

Running a system on an organization's own infrastructure. This may support security or compliance goals but requires more technical capacity.

Open-Weight Model

A model whose trained weights are publicly available under specific license terms. Open-weight does not always mean open-source, unrestricted, or free of obligations.

Orchestration

The coordination of models, tools, data sources, prompts, business logic, and user permissions across a workflow.

Pipeline

A sequence of processing steps, such as ingesting data, cleaning it, embedding it, retrieving relevant content, sending it to a model, and formatting the result.

Rate Limit

A restriction on how many requests or tokens a user or application can send in a given time. Rate limits manage capacity, cost, and abuse.

Rollback

Returning to a previous model, prompt, dataset, or system version after a new release causes problems.

Scalability

The ability of a system to handle increased users, data, requests, or complexity without unacceptable degradation.

Service Level Agreement (SLA)

A formal expectation for availability, latency, support, or reliability. AI vendors and enterprise deployments often define service levels contractually.

Token Cost

The cost associated with processing input and output tokens. Token cost affects budgeting, application design, and prompt efficiency.

Vendor Lock-In

A dependency risk that occurs when an organization becomes difficult to move away from a specific provider, model, API, format, or ecosystem.

Versioning

The practice of tracking changes to models, prompts, datasets, code, or workflows. Versioning allows teams to reproduce behavior and roll back problems.

J. Advanced and Emerging Concepts

Active Learning

A training strategy in which the model identifies examples that would be most useful for humans to label. This can reduce labeling cost and improve model performance efficiently.

Adapter

A small trainable component added to a model to adapt it for a task or domain without retraining the full model.

Agentic Workflow

A workflow where an AI system plans steps, calls tools, reviews intermediate results, and continues toward a goal with limited human direction. The more agentic the workflow, the more important boundaries and review become.

Artificial General Intelligence (AGI)

A debated concept referring to AI with broad, flexible capability across many domains at or beyond human-level performance. There is no universally accepted technical threshold for AGI.

Autonomous Agent

An agent that can pursue goals over multiple steps with less human intervention. Autonomy increases risk when actions affect money, safety, privacy, reputation, or legal obligations.

Catastrophic Forgetting

A problem where a model loses previously learned capabilities when trained on new data or tasks.

Causal Inference

The study of cause-and-effect relationships rather than mere correlation. Causal reasoning is important when decisions require understanding what will happen if an action is taken.

Continual Learning

The ability of a model to keep learning from new data over time without forgetting previous knowledge. Continual learning is difficult because new learning can disrupt old capabilities.

Counterfactual Explanation

An explanation showing how an input would need to change to produce a different outcome. For example, a loan system might explain what factors would have changed a denial to an approval.

Dense Model

A model in which most or all parameters are used for each input. Dense models are simpler to route but may require more computation per request.

Embodied AI

AI integrated with a physical body or robot that perceives and acts in the real world. Embodiment introduces safety, control, and environmental complexity.

Emergent Behavior

Capabilities or behaviors that appear when models reach sufficient scale or are used in new ways. The term is debated because some behaviors may reflect measurement artifacts or gradual improvement.

Low-Rank Adaptation (LoRA)

A parameter-efficient fine-tuning method that adapts a model using smaller trainable matrices rather than updating all model weights.

Mixture of Experts (MoE)

A model architecture that routes different inputs or tokens to specialized sub-networks called experts. MoE can increase capacity without activating every parameter on every request.

Model Collapse

A degradation scenario in which models trained heavily on AI-generated data lose diversity, accuracy, or connection to real-world distributions.

Multimodal Reasoning

The ability to combine information across modalities, such as reading text in an image, interpreting a chart, and answering a written question about it.

Neural-Symbolic AI

An approach that combines neural networks with symbolic reasoning, rules, logic, or structured knowledge. It aims to combine pattern learning with explicit reasoning.

Post-Training

Techniques applied after initial pretraining, such as fine-tuning, instruction tuning, alignment training, safety tuning, or distillation.

Pruning

Removing parts of a model that appear less important in order to improve efficiency. Pruning must be tested carefully to avoid performance loss.

Quantization

A technique that reduces the precision of model numbers to make models smaller, faster, or cheaper to run. It can sometimes reduce quality.

Reasoning Model

A model designed or trained to perform better on multi-step problem solving, planning, coding, mathematics, or structured analysis. Reasoning models may take longer and cost more per response.

Retraining

Training a model again, often with updated data, corrected labels, new objectives, or improved methods. Retraining may be needed when drift or errors appear.

Scaling Laws

Empirical relationships showing how model performance changes with more data, compute, and parameters. Scaling laws help guide decisions about model size and training resources.

Sparse Model

A model in which only part of the network or representation is active for a given input. Sparsity can improve efficiency and interpretability in certain contexts.

World Model

A representation of how an environment works that helps an AI system predict consequences and plan actions. The term is used in robotics, reinforcement learning, and discussions of advanced AI systems.

Alphabetical Index of Terms

A Ablation

Study — Machine Learning and Model Development
 Access Control — Security, Privacy, and Misuse
 Accountability — AI Safety, Ethics, and Governance
 Accuracy — Evaluation, Quality, and Performance
 Activation Function — Machine Learning and Model Development
 Active Learning — Advanced and Emerging Concepts
 Adapter — Advanced and Emerging Concepts
 Adversarial Example — Security, Privacy, and Misuse
 Adversarial Machine Learning — Security, Privacy, and Misuse
 Agent — Language Models and Generative AI
 Agentic Workflow — Advanced and Emerging Concepts
 AI Governance — AI Safety, Ethics, and Governance
 AI-Assisted Creation — Creative, Media, and Business Workflow Terms
 Algorithm — Core AI Concepts
 Alignment — AI Safety, Ethics, and Governance
 Annotation — Data, Datasets, and Representation
 API — Deployment, Operations, and Enterprise AI
 Application Layer — Deployment, Operations, and Enterprise AI
 Artificial General Intelligence (AGI) — Advanced and Emerging Concepts
 Artificial Intelligence (AI) — Core AI Concepts
 Attention Mechanism — Language Models and Generative AI
 Audit Log — Deployment, Operations, and Enterprise AI
 Autoencoder — Language Models and Generative AI
 Automation — Core AI Concepts
 Autonomous Agent — Advanced and Emerging Concepts
 Autonomy — Core AI Concepts
 Autoregressive Model — Language Models and Generative AI

B Backpropagation

— Machine Learning and Model Development
 Batch Size — Machine Learning and Model Development
 Benchmark — Machine Learning and Model Development
 Benchmark Leakage — Evaluation, Quality, and Performance
 Bias — AI Safety, Ethics, and Governance
 Black Box — AI Safety, Ethics, and Governance
 Brand Voice — Creative, Media, and Business Workflow Terms

C Caching

— Deployment, Operations, and Enterprise AI
 Calibration — Evaluation, Quality, and Performance
 Catastrophic Forgetting — Advanced and Emerging Concepts
 Causal Inference — Advanced and Emerging Concepts
 Chain-of-Thought Prompting — Language Models and Generative AI
 Class Imbalance — Data, Datasets, and Representation
 Classification — Machine Learning and Model Development
 Closed Model — Deployment, Operations, and Enterprise AI
 Cloud AI — Deployment, Operations, and Enterprise AI
 Clustering — Machine Learning and Model Development

Co-Creation — Creative, Media, and Business Workflow Terms
 Cognitive Computing — Core AI Concepts
 Commercial Use — Creative, Media, and Business Workflow Terms
 Completion — Language Models and Generative AI
 Computational Intelligence — Core AI Concepts
 Concept Drift — Data, Datasets, and Representation
 Confabulation — Language Models and Generative AI
 Confusion Matrix — Evaluation, Quality, and Performance
 Consent — AI Safety, Ethics, and Governance
 Content Credentials — AI Safety, Ethics, and Governance
 Content Provenance — AI Safety, Ethics, and Governance
 Context Poisoning — Security, Privacy, and Misuse
 Context Window — Data, Datasets, and Representation
 Continual Learning — Advanced and Emerging Concepts
 Corpus — Data, Datasets, and Representation
 Counterfactual Explanation — Advanced and Emerging Concepts
 Creative Direction — Creative, Media, and Business Workflow Terms
 Cross-Validation — Machine Learning and Model Development
 Curation — Creative, Media, and Business Workflow Terms

D Data

— Data, Datasets, and Representation
 Data Augmentation — Data, Datasets, and Representation
 Data Cleaning — Data, Datasets, and Representation
 Data Curation — Data, Datasets, and Representation
 Data Drift — Data, Datasets, and Representation
 Data Poisoning — Security, Privacy, and Misuse
 Data Provenance — Data, Datasets, and Representation
 Data Retention — Security, Privacy, and Misuse
 Dataset — Data, Datasets, and Representation
 De-Identification — Data, Datasets, and Representation
 Decision Boundary — Machine Learning and Model Development
 Decision Support System — Core AI Concepts
 Decoder — Language Models and Generative AI
 Deep Learning — Machine Learning and Model Development
 Deepfake — Security, Privacy, and Misuse
 Dense Model — Advanced and Emerging Concepts
 Derivative Work — Creative, Media, and Business Workflow Terms
 Differential Privacy — Security, Privacy, and Misuse
 Diffusion Model — Language Models and Generative AI
 Disclosure — AI Safety, Ethics, and Governance
 Disparate Impact — AI Safety, Ethics, and Governance
 Distillation — Machine Learning and Model Development

E Edge

AI — Deployment, Operations, and Enterprise AI
 Editorial Review — Creative, Media, and Business Workflow Terms
 Embedding — Data, Datasets, and Representation
 Embodied AI — Advanced and Emerging Concepts
 Emergent Behavior — Advanced and Emerging Concepts
 Encoder — Language Models and Generative AI
 Endpoint — Deployment, Operations, and Enterprise AI

Epoch — Machine Learning and Model Development
 Evals — Evaluation, Quality, and Performance
 Evaluation — Evaluation, Quality, and Performance
 Expert System — Core AI Concepts
 Explainability — AI Safety, Ethics, and Governance

F_{F1}

Score — Evaluation, Quality, and Performance
 Fairness — AI Safety, Ethics, and Governance
 Fallback — Deployment, Operations, and Enterprise AI
 False Negative — Evaluation, Quality, and Performance
 False Positive — Evaluation, Quality, and Performance
 Feature — Machine Learning and Model Development
 Feature Engineering — Machine Learning and Model Development
 Federated Learning — Data, Datasets, and Representation
 Few-Shot Prompting — Language Models and Generative AI
 Fine-Tuning — Machine Learning and Model Development
 Foundation Model — Language Models and Generative AI
 Function Calling — Language Models and Generative AI

G Generalization

— Machine Learning and Model Development
 Generative Adversarial Network (GAN) — Language Models and Generative AI
 Generative AI — Language Models and Generative AI
 Generative Fill — Creative, Media, and Business Workflow Terms
 Gradient Descent — Machine Learning and Model Development
 Ground Truth — Data, Datasets, and Representation
 Grounding — Language Models and Generative AI
 Guardrails — AI Safety, Ethics, and Governance

H Hallucination

— Language Models and Generative AI
 Harmful Bias — AI Safety, Ethics, and Governance
 Human Escalation — Deployment, Operations, and Enterprise AI
 Human Evaluation — Evaluation, Quality, and Performance
 Human-in-the-Loop — Core AI Concepts
 Hyperparameter — Machine Learning and Model Development

I Image-to-Image

Generation — Creative, Media, and Business Workflow Terms
 Impact Assessment — AI Safety, Ethics, and Governance
 Incident Disclosure — AI Safety, Ethics, and Governance
 Inference — Core AI Concepts
 Inpainting — Creative, Media, and Business Workflow Terms
 Instruction Tuning — Machine Learning and Model Development
 Intelligence Amplification — Core AI Concepts
 Interpretability — AI Safety, Ethics, and Governance
 Iteration — Creative, Media, and Business Workflow Terms

J Jailbreak

— Security, Privacy, and Misuse
 JSON Schema — Language Models and Generative AI

K Knowledge

Graph — Data, Datasets, and Representation

KnowledgeRepresentation — Core AI Concepts

L Label

— Data, Datasets, and Representation

Label Leakage — Data, Datasets, and Representation

Large Language Model (LLM) — Language Models and Generative AI

Latency — Evaluation, Quality, and Performance

Latent Diffusion — Language Models and Generative AI

Latent Space — Data, Datasets, and Representation

Licensing — Creative, Media, and Business Workflow Terms

LLMOps — Deployment, Operations, and Enterprise AI

Loss Function — Machine Learning and Model Development

Low-Rank Adaptation (LoRA) — Advanced and Emerging Concepts

M Machine

Learning (ML) — Machine Learning and Model Development

MalwareAssistance — Security, Privacy, and Misuse

Membership Inference — Security, Privacy, and Misuse

Memorization — Security, Privacy, and Misuse

Metadata — Data, Datasets, and Representation

Metric — Evaluation, Quality, and Performance

Misuse — Security, Privacy, and Misuse

Mixture of Experts (MoE) — Advanced and Emerging Concepts

MLOps — Deployment, Operations, and Enterprise AI

Model — Core AI Concepts

Model Card — Evaluation, Quality, and Performance

Model Collapse — Advanced and Emerging Concepts

Model Deployment — Deployment, Operations, and Enterprise AI

Model Extraction — Security, Privacy, and Misuse

Model Inversion — Security, Privacy, and Misuse

Model Monitoring — Deployment, Operations, and Enterprise AI

Multimodal AI — Language Models and Generative AI

Multimodal Reasoning — Advanced and Emerging Concepts

N Narrow

AI — Core AI Concepts

Negative Prompt — Creative, Media, and Business Workflow Terms

Neural Network — Machine Learning and Model Development

Neural-Symbolic AI — Advanced and Emerging Concepts

Noncommercial Use — Creative, Media, and Business Workflow Terms

O Observability

— Evaluation, Quality, and Performance

On-Premises Deployment — Deployment, Operations, and Enterprise AI

Ontology — Data, Datasets, and Representation

Open-Weight Model — Deployment, Operations, and Enterprise AI

Optimizer — Machine Learning and Model Development

Orchestration — Deployment, Operations, and EnterpriseAI

Outpainting — Creative, Media, and Business WorkflowTerms

Overfitting — Machine Learning and Model Development

P Parameter

Parameter — Machine Learning and Model Development

Personally Identifiable Information (PII) — Security, Privacy, and Misuse

Pipeline — Deployment, Operations, and Enterprise AI

Positional Encoding — Language Models and Generative AI

Post-Training — Advanced and Emerging Concepts

Precision — Evaluation, Quality, and Performance

Privacy-Enhancing Technology — Security, Privacy, and Misuse

Probabilistic System — Core AI Concepts

Prompt — Language Models and Generative AI

Prompt Engineering — Language Models and Generative AI

Prompt Injection — Security, Privacy, and Misuse

Prompt Library — Creative, Media, and Business Workflow Terms

Pruning — Advanced and Emerging Concepts

Q Quantization

Quantization — Advanced and Emerging Concepts

R Rate

Rate Limit — Deployment, Operations, and Enterprise AI

Reasoning Model — Advanced and Emerging Concepts

Recall — Evaluation, Quality, and Performance

Red Teaming — AI Safety, Ethics, and Governance

Regression — Machine Learning and Model Development

Reinforcement Learning — Machine Learning and Model Development

Reliability — Evaluation, Quality, and Performance

Reproducibility — Evaluation, Quality, and Performance

Responsible AI — AI Safety, Ethics, and Governance

Retraining — Advanced and Emerging Concepts

Retrieval-Augmented Generation (RAG) — Language Models and Generative AI

Risk Management — AI Safety, Ethics, and Governance

Robustness — Evaluation, Quality, and Performance

Rollback — Deployment, Operations, and Enterprise AI

S Safety

Safety — AI Safety, Ethics, and Governance

Safety Filter — AI Safety, Ethics, and Governance

Sampling — Language Models and Generative AI

Scalability — Deployment, Operations, and Enterprise AI

Scaling Laws — Advanced and Emerging Concepts

Secure and Resilient AI — Security, Privacy, and Misuse

Secure Deployment — Security, Privacy, and Misuse

Seed — Creative, Media, and Business Workflow Terms

Self-Attention — Language Models and Generative AI

Self-Supervised Learning — Machine Learning and Model Development

Semantic Search — Language Models and Generative AI

Semi-Supervised Learning — Machine Learning and Model Development

Sensitive Data — Security, Privacy, and Misuse

Service Level Agreement (SLA) — Deployment, Operations, and Enterprise AI

Side-by-Side Evaluation — Evaluation, Quality, and Performance

Sparse Model — Advanced and Emerging Concepts
 Speech-to-Text — Creative, Media, and Business Workflow Terms
 Structured Output — Language Models and Generative AI
 Style Prompt — Creative, Media, and Business Workflow Terms
 Style Transfer — Creative, Media, and Business Workflow Terms
 Supervised Learning — Machine Learning and Model Development
 Symbolic AI — Core AI Concepts
 Synthetic Data — Data, Datasets, and Representation
 Synthetic Voice — Creative, Media, and Business Workflow Terms
 System Prompt — Language Models and Generative AI

T Task

— Core AI Concepts

Temperature — Language Models and Generative AI
 Test Data — Data, Datasets, and Representation
 Text-to-Image Generation — Creative, Media, and Business Workflow Terms
 Text-to-Speech — Creative, Media, and Business Workflow Terms
 Text-to-Video Generation — Creative, Media, and Business Workflow Terms
 Throughput — Evaluation, Quality, and Performance
 Token — Data, Datasets, and Representation
 Token Cost — Deployment, Operations, and Enterprise AI
 Tokenization — Data, Datasets, and Representation
 Tool Calling — Language Models and Generative AI
 Top-p Sampling — Language Models and Generative AI
 Training — Machine Learning and Model Development
 Training Data — Data, Datasets, and Representation
 Transfer Learning — Machine Learning and Model Development
 Transformer — Language Models and Generative AI
 Transparency — AI Safety, Ethics, and Governance
 Trustworthy AI — AI Safety, Ethics, and Governance

U Underfitting

— Machine Learning and Model Development

Unsupervised Learning — Machine Learning and Model Development
 Upscaling — Creative, Media, and Business Workflow Terms
 User Prompt — Language Models and Generative AI

V Validation

Data — Data, Datasets, and Representation
 Validity — Evaluation, Quality, and Performance
 Value Alignment — AI Safety, Ethics, and Governance
 Variational Autoencoder (VAE) — Language Models and Generative AI
 Vector — Data, Datasets, and Representation
 Vector Database — Data, Datasets, and Representation
 Vendor Lock-In — Deployment, Operations, and Enterprise AI
 Versioning — Deployment, Operations, and Enterprise AI
 Voice Cloning — Security, Privacy, and Misuse

W Watermarking

— AI Safety, Ethics, and Governance

Workflow — Core AI Concepts
 World Model — Advanced and Emerging Concepts

Z Zero-Shot

Source and Reference Note

This glossary synthesizes common terminology from AI, machine learning, generative AI, prompt engineering, responsible AI, security, privacy, creative workflow, and governance practice. It is an educational reference, not a legal standard, vendor manual, or substitute for official documentation.

The Haus Of Legends educational framing treats AI as a tool within a long history of creative technology. The glossary therefore emphasizes both technical literacy and human responsibility: concept, authorship, judgment, curation, verification, and ethical use.

- NIST AI Resource Center, AI RMF Glossary / The Language of Trustworthy AI: <https://airc.nist.gov/glossary/> NIST AI 100-1, Artificial Intelligence Risk Management Framework 1.0: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>
- NIST AI 600-1, Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>
- Google Machine Learning Glossary: <https://developers.google.com/machine-learning/glossary>
- OpenAI Prompt Engineering Guide: <https://developers.openai.com/api/docs/guides/prompt-engineering> OpenAI Function Calling Guide: <https://developers.openai.com/api/docs/guides/function-calling>
- OpenAI Structured Outputs Guide: <https://developers.openai.com/api/docs/guides/structured-outputs>
- The Haus Of Legends source material: Art and Technology: Project source file supplied by The Haus Of Legends

Editorial Use

This document may be used as a downloadable educational resource, website reference, classroom handout, course supplement, or companion document to beginner AI lessons. For publication, review current tool policies, legal requirements, and institutional standards before presenting the glossary as official guidance.