

EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS

EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS IS AN EMERGING AREA OF RESEARCH AND PRACTICAL INTEREST WITHIN THE FIELD OF ARTIFICIAL INTELLIGENCE AND NATURAL LANGUAGE PROCESSING. AS LARGE LANGUAGE MODELS (LLMs) LIKE GPT AND OTHERS CONTINUE TO GROW IN COMPLEXITY AND CAPABILITY, UNDERSTANDING THE NATURE OF THE DATA THEY WERE TRAINED ON BECOMES INCREASINGLY IMPORTANT. EXTRACTING TRAINING DATA FROM THESE MODELS INVOLVES TECHNIQUES THAT ATTEMPT TO RECOVER OR INFER ORIGINAL DATA POINTS OR PATTERNS USED DURING THE MODEL'S TRAINING PHASE. THIS ARTICLE EXPLORES THE METHODOLOGIES, CHALLENGES, ETHICAL IMPLICATIONS, AND APPLICATIONS RELATED TO EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS. IT ALSO DISCUSSES SECURITY CONCERNS AND THE BALANCE BETWEEN TRANSPARENCY AND PRIVACY. THE SUBSEQUENT SECTIONS PROVIDE A COMPREHENSIVE OVERVIEW OF THE TECHNICAL PROCESSES, RISKS, AND ONGOING RESEARCH EFFORTS IN THIS DOMAIN.

- UNDERSTANDING LARGE LANGUAGE MODELS AND THEIR TRAINING DATA
- TECHNIQUES FOR EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS
- CHALLENGES AND LIMITATIONS OF DATA EXTRACTION
- ETHICAL AND PRIVACY CONSIDERATIONS
- APPLICATIONS AND IMPLICATIONS OF TRAINING DATA EXTRACTION
- SECURITY RISKS AND MITIGATION STRATEGIES

UNDERSTANDING LARGE LANGUAGE MODELS AND THEIR TRAINING DATA

LARGE LANGUAGE MODELS ARE ADVANCED AI SYSTEMS TRAINED ON MASSIVE CORPORA OF TEXT DATA TO GENERATE HUMAN-LIKE LANGUAGE AND PERFORM DIVERSE NATURAL LANGUAGE TASKS. THESE MODELS LEARN STATISTICAL PATTERNS AND SEMANTIC RELATIONSHIPS FROM THEIR TRAINING DATASETS, WHICH OFTEN INCLUDE BOOKS, ARTICLES, WEBSITES, AND OTHER TEXT SOURCES. THE TRAINING DATA FOR LLMs IS TYPICALLY VAST, DIVERSE, AND AGGREGATED FROM MULTIPLE DOMAINS, MAKING IT CHALLENGING TO PINPOINT INDIVIDUAL ENTRIES OR PROPRIETARY CONTENT. UNDERSTANDING THE NATURE AND SCOPE OF TRAINING DATA IS ESSENTIAL FOR APPRECIATING THE COMPLEXITIES INVOLVED IN EXTRACTING SUCH INFORMATION FROM THE MODELS THEMSELVES.

COMPOSITION OF TRAINING DATA

THE TRAINING DATASETS USED FOR LARGE LANGUAGE MODELS GENERALLY CONSIST OF TEXT COLLECTED FROM PUBLICLY AVAILABLE SOURCES SUCH AS WEB PAGES, DIGITAL LIBRARIES, AND LICENSED DATASETS. THIS DATA IS PREPROCESSED AND CLEANED TO IMPROVE QUALITY AND RELEVANCE BEFORE MODEL TRAINING. THE DIVERSITY AND VOLUME OF TRAINING DATA CONTRIBUTE TO THE MODEL'S ABILITY TO GENERALIZE AND GENERATE COHERENT, CONTEXTUALLY APPROPRIATE RESPONSES ACROSS VARIOUS TOPICS.

MODEL ARCHITECTURE AND DATA ENCODING

LLMs USE SOPHISTICATED NEURAL NETWORK ARCHITECTURES, SUCH AS TRANSFORMERS, WHICH ENCODE TRAINING DATA INTO DISTRIBUTED REPRESENTATIONS OR EMBEDDINGS. THESE EMBEDDINGS CAPTURE SEMANTIC AND SYNTACTIC INFORMATION BUT DO NOT STORE TRAINING DATA VERBATIM. INSTEAD, THE MODEL INTERNALIZES PATTERNS AND ASSOCIATIONS, MAKING DIRECT

RETRIEVAL OF ORIGINAL TRAINING INPUTS INHERENTLY DIFFICULT BUT NOT IMPOSSIBLE UNDER CERTAIN CONDITIONS.

TECHNIQUES FOR EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS

EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS INVOLVES METHODS THAT ATTEMPT TO RECONSTRUCT OR APPROXIMATE THE ORIGINAL DATA USED DURING TRAINING. VARIOUS TECHNIQUES HAVE BEEN DEVELOPED TO PROBE MODELS FOR MEMORIZED INFORMATION, ASSESS DATA LEAKAGE RISKS, AND ANALYZE MODEL BEHAVIOR IN DETAIL.

MEMBERSHIP INFERENCE ATTACKS

MEMBERSHIP INFERENCE ATTACKS ARE DESIGNED TO DETERMINE WHETHER A SPECIFIC DATA POINT WAS PART OF A MODEL'S TRAINING DATASET. THESE ATTACKS EXPLOIT THE MODEL'S RESPONSES TO QUERIES AND ANALYZE CONFIDENCE SCORES OR OUTPUT DISTRIBUTIONS TO INFER MEMBERSHIP STATUS. THIS TECHNIQUE CAN REVEAL SENSITIVE INFORMATION AND IS A CRITICAL CONCERN FOR PRIVACY IN MACHINE LEARNING.

MODEL INVERSION AND RECONSTRUCTION

MODEL INVERSION ATTACKS AIM TO RECONSTRUCT INPUT DATA BY LEVERAGING THE MODEL'S OUTPUTS AND LEARNED REPRESENTATIONS. THIS METHOD ATTEMPTS TO REVERSE-ENGINEER TRAINING EXAMPLES BY QUERYING THE MODEL WITH CRAFTED INPUTS AND INTERPRETING THE RESPONSES TO RECREATE ORIGINAL OR SIMILAR DATA POINTS. ALTHOUGH CHALLENGING, THESE ATTACKS DEMONSTRATE THE POTENTIAL FOR TRAINING DATA EXTRACTION UNDER CERTAIN CIRCUMSTANCES.

PROMPT ENGINEERING AND OUTPUT SAMPLING

CAREFULLY DESIGNED PROMPTS AND ITERATIVE SAMPLING STRATEGIES CAN COAX A LANGUAGE MODEL TO GENERATE TEXT THAT CLOSELY RESEMBLES ITS TRAINING DATA. BY EXPLOITING MODEL TENDENCIES TO MEMORIZE AND REPRODUCE RARE OR UNIQUE SEQUENCES, RESEARCHERS CAN EXTRACT FRAGMENTS OF TRAINING CONTENT. THIS APPROACH REQUIRES EXTENSIVE EXPERIMENTATION AND MAY NOT RELIABLY RECOVER COMPLETE DATA POINTS BUT CAN UNCOVER SENSITIVE OR PROPRIETARY INFORMATION INADVERTENTLY MEMORIZED BY THE MODEL.

CHALLENGES AND LIMITATIONS OF DATA EXTRACTION

EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS IS FRAUGHT WITH TECHNICAL AND PRACTICAL CHALLENGES. THE COMPLEXITY OF NEURAL ARCHITECTURES, THE SHEER SIZE OF DATASETS, AND THE PROBABILISTIC NATURE OF MODEL OUTPUTS LIMIT THE PRECISION AND SCOPE OF DATA RECOVERY EFFORTS. UNDERSTANDING THESE CHALLENGES HELPS CONTEXTUALIZE THE CAPABILITIES AND CONSTRAINTS OF CURRENT EXTRACTION TECHNIQUES.

DATA MEMORIZATION VS. GENERALIZATION

LARGE LANGUAGE MODELS ARE DESIGNED TO GENERALIZE RATHER THAN MEMORIZE SPECIFIC TRAINING EXAMPLES. WHILE SOME MEMORIZATION OCCURS, ESPECIALLY FOR RARE OR UNIQUE DATA POINTS, MOST INFORMATION IS ENCODED AS GENERALIZED PATTERNS. THIS DISTINCTION MAKES IT DIFFICULT TO EXTRACT EXACT TRAINING DATA, AS MODELS TYPICALLY GENERATE

NOVEL CONTENT INSPIRED BY, BUT NOT IDENTICAL TO, THE INPUT DATA.

SCALE AND COMPLEXITY

THE VAST SCALE OF TRAINING DATASETS—OFTEN COMPRISING BILLIONS OF TOKENS—AND THE HIGH DIMENSIONALITY OF MODEL PARAMETERS CONTRIBUTE TO THE DIFFICULTY OF EXTRACTING MEANINGFUL DATA. SEARCHING FOR SPECIFIC TRAINING ITEMS WITHIN THIS ENORMOUS SPACE REQUIRES SIGNIFICANT COMPUTATIONAL RESOURCES AND SOPHISTICATED ALGORITHMS, WHICH MAY STILL YIELD LIMITED RESULTS.

MODEL UPDATES AND FINE-TUNING

MODELS FREQUENTLY UNDERGO UPDATES, RETRAINING, OR FINE-TUNING ON NEW DATA, WHICH CAN OVERWRITE OR OBSCURE PREVIOUSLY LEARNED INFORMATION. THIS DYNAMIC NATURE OF LARGE LANGUAGE MODELS COMPLICATES EXTRACTION ATTEMPTS, AS DATA REPRESENTATIONS EVOLVE OVER TIME AND MAY NO LONGER CORRESPOND TO ORIGINAL TRAINING INPUTS.

ETHICAL AND PRIVACY CONSIDERATIONS

THE PRACTICE OF EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS RAISES IMPORTANT ETHICAL AND PRIVACY ISSUES. SINCE TRAINING DATASETS OFTEN CONTAIN PERSONAL, SENSITIVE, OR COPYRIGHTED INFORMATION, UNAUTHORIZED EXTRACTION CAN LEAD TO DATA BREACHES, PRIVACY VIOLATIONS, AND INTELLECTUAL PROPERTY CONCERNS.

DATA PRIVACY RISKS

EXTRACTED DATA MIGHT EXPOSE PERSONALLY IDENTIFIABLE INFORMATION (PII) OR CONFIDENTIAL CONTENT INADVERTENTLY INCLUDED IN TRAINING SETS. THIS RISK NECESSITATES STRICT ADHERENCE TO DATA PROTECTION REGULATIONS AND RESPONSIBLE AI DEVELOPMENT PRACTICES TO PREVENT MISUSE OR HARM TO INDIVIDUALS WHOSE DATA MAY BE INVOLVED.

INTELLECTUAL PROPERTY AND COPYRIGHT

TRAINING DATA CAN INCLUDE COPYRIGHTED WORKS, CREATING LEGAL CHALLENGES WHEN FRAGMENTS OF SUCH CONTENT ARE EXTRACTED OR REPRODUCED BY THE MODEL. THE BOUNDARIES BETWEEN FAIR USE, DATA OWNERSHIP, AND AI-GENERATED OUTPUTS ARE COMPLEX AND CONTINUALLY EVOLVING, EMPHASIZING THE NEED FOR TRANSPARENT POLICIES REGARDING TRAINING DATA EXTRACTION.

RESPONSIBLE DISCLOSURE AND USAGE

RESEARCHERS AND ORGANIZATIONS INVOLVED IN EXTRACTING TRAINING DATA MUST FOLLOW ETHICAL GUIDELINES, ENSURING THAT ANY FINDINGS ARE DISCLOSED RESPONSIBLY AND THAT DATA IS NOT EXPLOITED FOR MALICIOUS PURPOSES. TRANSPARENCY IN METHODOLOGY AND COLLABORATION WITH STAKEHOLDERS HELP MITIGATE ETHICAL CONCERNS.

APPLICATIONS AND IMPLICATIONS OF TRAINING DATA EXTRACTION

DESPITE THE CHALLENGES AND RISKS, EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS HAS VALUABLE APPLICATIONS ACROSS MULTIPLE DOMAINS. UNDERSTANDING THESE APPLICATIONS PROVIDES INSIGHT INTO WHY THIS AREA REMAINS AN ACTIVE RESEARCH FOCUS.

MODEL AUDITING AND TRANSPARENCY

EXTRACTION TECHNIQUES ENABLE AUDITORS AND DEVELOPERS TO VERIFY WHAT TYPES OF DATA MODELS HAVE LEARNED, ASSESS BIAS, AND IDENTIFY POTENTIAL DATA LEAKAGE. THIS TRANSPARENCY HELPS IMPROVE MODEL ROBUSTNESS AND ACCOUNTABILITY, FOSTERING TRUST IN AI SYSTEMS.

DATA PROVENANCE AND COMPLIANCE

ORGANIZATIONS CAN USE TRAINING DATA EXTRACTION METHODS TO VERIFY COMPLIANCE WITH DATA USAGE POLICIES, ENSURING THAT MODELS DO NOT RETAIN OR REPRODUCE UNAUTHORIZED OR SENSITIVE INFORMATION. THIS IS CRITICAL FOR MEETING REGULATORY REQUIREMENTS AND MAINTAINING ETHICAL STANDARDS.

SECURITY TESTING AND VULNERABILITY ASSESSMENT

BY SIMULATING DATA EXTRACTION ATTACKS, SECURITY TEAMS EVALUATE MODEL VULNERABILITIES AND DEVELOP MITIGATION STRATEGIES. THESE ASSESSMENTS HELP SAFEGUARD AI SYSTEMS AGAINST PRIVACY BREACHES AND ADVERSARIAL EXPLOITATION.

SECURITY RISKS AND MITIGATION STRATEGIES

THE POTENTIAL FOR EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS INTRODUCES SECURITY RISKS THAT REQUIRE PROACTIVE COUNTERMEASURES. ADDRESSING THESE RISKS INVOLVES BOTH TECHNICAL SOLUTIONS AND POLICY MEASURES.

TECHNIQUES TO REDUCE DATA MEMORIZATION

APPROACHES SUCH AS DIFFERENTIAL PRIVACY, REGULARIZATION, AND DATA SANITIZATION CAN LIMIT THE EXTENT TO WHICH MODELS MEMORIZE SPECIFIC TRAINING EXAMPLES. IMPLEMENTING THESE TECHNIQUES DURING TRAINING REDUCES THE LIKELIHOOD OF DATA LEAKAGE THROUGH EXTRACTION.

ACCESS CONTROL AND USAGE MONITORING

RESTRICTING MODEL ACCESS AND MONITORING USAGE PATTERNS HELP DETECT AND PREVENT UNAUTHORIZED EXTRACTION ATTEMPTS. RATE LIMITING, QUERY AUDITING, AND ANOMALY DETECTION ARE PRACTICAL MEASURES TO SAFEGUARD SENSITIVE INFORMATION.

ONGOING RESEARCH AND BEST PRACTICES

CONTINUOUS RESEARCH INTO MODEL INTERPRETABILITY, PRIVACY-PRESERVING MACHINE LEARNING, AND SECURE AI DEPLOYMENT IS ESSENTIAL FOR DEVELOPING ROBUST DEFENSES AGAINST TRAINING DATA EXTRACTION. ADOPTING INDUSTRY BEST PRACTICES AND UPDATING FRAMEWORKS AS NEW THREATS EMERGE ENSURES RESPONSIBLE AI DEVELOPMENT.

SUMMARY OF KEY POINTS

- EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS INVOLVES TECHNIQUES SUCH AS MEMBERSHIP INFERENCE, MODEL INVERSION, AND PROMPT ENGINEERING.
- LLMs ENCODE TRAINING DATA IN COMPLEX, DISTRIBUTED REPRESENTATIONS, MAKING EXACT DATA RECOVERY DIFFICULT BUT NOT IMPOSSIBLE.
- CHALLENGES INCLUDE MODEL GENERALIZATION, DATASET SCALE, AND EVOLVING MODEL PARAMETERS.
- ETHICAL CONCERNS CENTER ON PRIVACY VIOLATIONS, INTELLECTUAL PROPERTY RIGHTS, AND RESPONSIBLE USE OF EXTRACTED DATA.
- APPLICATIONS INCLUDE MODEL AUDITING, COMPLIANCE VERIFICATION, AND SECURITY TESTING.
- MITIGATION STRATEGIES ENCOMPASS PRIVACY-PRESERVING TRAINING METHODS, ACCESS CONTROL, AND CONTINUOUS RESEARCH.

FREQUENTLY ASKED QUESTIONS

WHAT DOES 'EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS' MEAN?

IT REFERS TO TECHNIQUES AND METHODS USED TO RETRIEVE OR INFER THE ORIGINAL TRAINING DATA OR PARTS OF IT THAT WERE USED TO TRAIN LARGE LANGUAGE MODELS, OFTEN BY PROBING OR QUERYING THE MODEL'S OUTPUTS.

WHY IS EXTRACTING TRAINING DATA FROM LARGE LANGUAGE MODELS A CONCERN?

EXTRACTING TRAINING DATA CAN RAISE PRIVACY AND SECURITY ISSUES, AS SENSITIVE OR PROPRIETARY INFORMATION MAY BE UNINTENTIONALLY MEMORIZED AND EXPOSED BY THE MODEL DURING INFERENCE.

WHAT METHODS ARE COMMONLY USED TO EXTRACT TRAINING DATA FROM LARGE LANGUAGE MODELS?

COMMON METHODS INCLUDE MEMBERSHIP INFERENCE ATTACKS, MODEL INVERSION ATTACKS, PROMPT ENGINEERING TO ELICIT MEMORIZED CONTENT, AND GRADIENT-BASED EXTRACTION TECHNIQUES.

CAN LARGE LANGUAGE MODELS UNINTENTIONALLY MEMORIZE AND REVEAL PERSONAL DATA?

YES, LARGE LANGUAGE MODELS CAN SOMETIMES MEMORIZE SPECIFIC DATA POINTS FROM THEIR TRAINING SETS, INCLUDING PERSONAL OR SENSITIVE DATA, AND INADVERTENTLY REVEAL THEM WHEN PROMPTED APPROPRIATELY.

How can we mitigate the risk of training data extraction from large language models?

Mitigation strategies include data anonymization, differential privacy during training, regularization techniques, limiting exposure to sensitive data, and careful prompt monitoring.

Is it possible to completely prevent training data extraction from large language models?

Completely preventing extraction is challenging because models inherently learn from data patterns, but employing privacy-preserving training methods can significantly reduce the risk.

What role does differential privacy play in protecting training data in language models?

Differential privacy adds noise to the training process, ensuring that the model's outputs do not reveal whether any single data point was part of the training set, thus protecting individual data privacy.

Are there ethical implications associated with extracting training data from language models?

Yes, extracting training data can violate data privacy, intellectual property rights, and user trust, raising ethical concerns about consent and responsible AI use.

How do researchers test the vulnerability of language models to training data extraction?

Researchers design and execute attacks such as membership inference, prompt-based extraction, and white-box analysis to evaluate how much and what type of training data can be recovered from a model.

Additional Resources

1. *Mining Minds: Techniques for Extracting Training Data from Large Language Models*

This book delves into the methodologies used to retrieve and analyze training data embedded within large language models. It covers various approaches, including prompt engineering, model inversion, and data reconstruction. Readers will gain insights into how these techniques reveal the underlying datasets that shape model behavior and performance.

2. *Inside the Black Box: Understanding Data Leakage in Large Language Models*

Focusing on the risks and mechanisms of data leakage, this book explores how training data can inadvertently be extracted from language models. It addresses privacy concerns, ethical implications, and strategies to mitigate unintentional data exposure. The book is essential for researchers and practitioners committed to responsible AI development.

3. *Reverse Engineering Language Models: Methods for Training Data Recovery*

This comprehensive guide presents technical methods for reverse engineering language models to recover portions of their training datasets. It discusses algorithmic strategies, experimental results, and the challenges faced during the extraction process. The text is suited for AI engineers and data scientists seeking to understand model transparency.

4. *Data Forensics in AI: Extracting and Analyzing Training Sets from Neural Networks*

Combining principles of data forensics and artificial intelligence, this book explores how investigators can extract and analyze training data from neural network models. It covers forensic tools, case studies, and

LEGAL CONSIDERATIONS IN UNCOVERING DATA ORIGINS. THE BOOK IS VALUABLE FOR LEGAL EXPERTS, DATA SCIENTISTS, AND AI AUDITORS.

5. *ETHICS AND EXTRACTION: NAVIGATING TRAINING DATA RETRIEVAL FROM LANGUAGE MODELS*

THIS BOOK TACKLES THE ETHICAL LANDSCAPE SURROUNDING THE EXTRACTION OF TRAINING DATA FROM LARGE LANGUAGE MODELS. IT DISCUSSES PRIVACY RIGHTS, CONSENT, AND THE IMPACT OF DATA RECOVERY ON INTELLECTUAL PROPERTY. READERS ARE ENCOURAGED TO CONSIDER ETHICAL FRAMEWORKS AND BEST PRACTICES WHEN ENGAGING IN DATA EXTRACTION.

6. *TRAINING DATA RECONSTRUCTION: TECHNIQUES AND CHALLENGES IN LARGE LANGUAGE MODELS*

HIGHLIGHTING THE TECHNICAL CHALLENGES OF RECONSTRUCTING TRAINING DATA, THIS BOOK SURVEYS CURRENT TECHNIQUES SUCH AS MEMBERSHIP INFERENCE ATTACKS AND GRADIENT INVERSION. IT PROVIDES EXPERIMENTAL ANALYSES AND DISCUSSES THE IMPLICATIONS FOR MODEL SECURITY. THE BOOK SERVES AS A RESOURCE FOR AI RESEARCHERS FOCUSED ON MODEL VULNERABILITY.

7. *LANGUAGE MODELS UNVEILED: EXTRACTING KNOWLEDGE AND DATA FROM AI SYSTEMS*

THIS WORK EXPLORES HOW KNOWLEDGE AND TRAINING DATA CAN BE EXTRACTED FROM LANGUAGE MODELS BEYOND MERE TEXT GENERATION. IT INCLUDES DISCUSSIONS ON KNOWLEDGE DISTILLATION, PROBING, AND MODEL INTERPRETABILITY. THE BOOK IS USEFUL FOR THOSE INTERESTED IN UNDERSTANDING AND HARNESSING LATENT DATA WITHIN AI SYSTEMS.

8. *PRIVACY IN THE AGE OF AI: PROTECTING AND EXTRACTING TRAINING DATA FROM LANGUAGE MODELS*

ADDRESSING PRIVACY CONCERNS, THIS BOOK EXAMINES THE BALANCE BETWEEN PROTECTING SENSITIVE TRAINING DATA AND THE TECHNICAL PROCESSES INVOLVED IN DATA EXTRACTION. IT REVIEWS PRIVACY-PRESERVING TECHNIQUES LIKE DIFFERENTIAL PRIVACY AND FEDERATED LEARNING. THE TEXT IS AIMED AT DEVELOPERS AND POLICYMAKERS FOCUSED ON SAFEGUARDING DATA.

9. *FROM DATA TO MODEL AND BACK: CIRCULAR INSIGHTS INTO TRAINING DATA EXTRACTION*

THIS BOOK PROVIDES A HOLISTIC VIEW OF THE CYCLICAL RELATIONSHIP BETWEEN TRAINING DATA AND MODEL OUTPUTS, EMPHASIZING METHODS TO EXTRACT DATA THAT INFORM MODEL REFINEMENT. IT DISCUSSES FEEDBACK LOOPS, DATA AUGMENTATION, AND ITERATIVE EXTRACTION TECHNIQUES. THE BOOK IS TARGETED AT AI PRACTITIONERS INTERESTED IN CONTINUOUS LEARNING AND DATA TRANSPARENCY.

Extracting Training Data From Large Language Models

Extracting Training Data From Large Language Models

Related Articles

- [film soper irani](#)
- [eureka math lesson 8 answer key](#)
- [examples of manual filing systems](#)

[Back to Home](#)