

Proposals

- Proposal: ImportExport module
- The design of MKLDNN integration
- Logging MXNet Data for Visualization in TensorBoard
- Improved Exception Handling in MXNet - Phase 2
- Thread-safety in MXNet
- MXNet C++ package improvements
- Parallel Inference in MXNet
- Dynamic shape
- Optimize dynamic neural network models with control flow operators
- KVStore for Distributed Model Parallel FC(nsoftmax).
- End2End Fixpoint Finetuning Plug-In
- KVStore for Partial FC(nsoftmax) for Large Scale Data Training.
- Unified integration with external backend libraries
- Website Update Proposals
 - Website Redesign Information Architecture & Wireframe
 - Website Global Search Feature
 - MXNet Website 2.0 - Slow Site Speed in China
 - Merge beta.mxnet.io To MXNet Official Website
 - Merge numpy.mxnet.io to MXNet Official Website
- Sparse Matrix Dot Product Parallel Algorithm Design
- Automated Flaky Test Detector
 - Flaky Test Bot Integration Plan
- SVRG Optimization in MXNet Python Module
- Extend MXNet Distributed Training by MPI AllReduce
- Single machine All Reduce Topology-aware Communication
- Horovod-MXNet Integration
- Runtime Integration with TensorRT
- MXNet Java Inference API
- MXNet nGraph integration using subgraph backend interface
- Upstreaming of Mxnet HIP port
- Data IO
 - Image Transforms and RecordIO file Creation
- MXNet with Intel MKL-DNN - Performance Benchmarking
- Benchmarking MXNet with different OpenMP implementations
- Gluon API for Scala/ Clojure/ Java packages
- MXNet end to end models
- Extending MXNet Model save/load APIs
- [WIP] LR Finder
 - [WIP] Differential Learning Rate
- New Clojure NDArray/Symbol APIs
- Archived Proposals
 - Apache MXNet Design Proposal Template
 - Fused RNN Operators for CPU
 - Gluon - Audio
 - Gluon Sparse Support
 - Improved Exception Handling in MXNet
 - Label Prediction by the @mxnet-label-bot
 - Machine Learning Based GitHub Bot
 - Make cmake default build system
 - Model Backwards Compatibility Check
 - PR Best Practices
- MXNet Graph Optimization and Quantization based on subgraph and MKL-DNN
- Enable Operator Level Parallelism under Subgraph
- Higher Order Gradient Calculation
- MXNet Memory Profiling Enhancements
- Gluon Fit API - Tech Design
 - Callback Design for Fit Loop
- Conversion from FP32 to Mixed Precision Models
- MXNet Operator Benchmarks
- MXVM: Operator graph 2.0
- GPU Pointwise fusion
- Fixing Duplication in Operator Profiling
- Custom Operator Profiling Enhancement
- Design of CNML/CNRT Integration
- Make MXNet deploy it's own distribution
- Bring your own Accelerator
- Using TensorInspector to help debug operators
- Pull Request review bot
- Dynamic CustomOp Support
- BytePS-MXNet Integration
- Large Tensor Support
- MXNet FFI for Operator Imperative Invocation