



**ISSUE  
Nº95**



**EUROPEAN OPEN  
SCIENCE SPACE**

**COLLECTION OF SCIENTIFIC PAPERS**



**4<sup>TH</sup> INTERNATIONAL  
SCIENTIFIC  
AND PRACTICAL  
CONFERENCE**

**SCIENTIFIC RESEARCH:  
EMERGING THEORIES  
AND PRACTICAL  
BREAKTHROUGHS**

**JULY 6-8, 2026, EDINBURGH, SCOTLAND**



## **Section: Information Technology, Cyber Security and Computer Engineering**

# **NEURAL NETWORK OPTIMISATION: QUANTISATION TECHNIQUES, PRACTICAL BENCHMARKING, AND DEPLOYMENT RECOMMENDATIONS**

**Знахур Сергій**

к.е.н., доцент

Кафедра інформаційних систем

Харківський національний економічний університет

імені Семена Кузнеця, Україна

### **Abstract**

The proliferation of on-device artificial intelligence has made the deployment of deep neural networks on memory- and energy-constrained mobile hardware one of the most practically consequential challenges in applied machine learning. This article presents a systematic empirical investigation of neural network quantisation applied to five convolutional architectures — MobileNetV2, ResNet-50, DenseNet-121, EfficientNet-B3, and EfficientNet-B4 — across three numerical precision formats: FP32 (baseline), FP16, and INT8. All models were converted using TensorFlow Lite Converter and evaluated on a physical iPhone 13 device (Apple A15 Bionic) through a purpose-built React Native 0.82 application using the react-native-fast-tflite library. Accuracy was assessed on a 500-image ImageNet-v2 subset. Key results show that FP16 reduces model file size by approximately 50% with Top-1 accuracy losses below 0.3%, while INT8 achieves 71–74% size reduction and delivers up to 2.8× latency improvement for mobile-optimised architectures. EfficientNet-B3/B4 exhibit no consistent latency gain from INT8, attributable to operator-level CPU fallback within the TFLite Metal Delegate. Three primary factors governing quantisation effectiveness are identified: architectural structural compatibility with delegate hardware, calibration dataset representativeness, and delegate operator coverage. Practical deployment recommendations are derived directly from the experimental findings for each architecture family.

**Keywords:** neural network quantisation, mobile deep learning, TensorFlow Lite, FP16, INT8, MobileNetV2, ResNet-50, DenseNet-121, EfficientNet, on-device inference, React Native, Apple A15 Bionic.

### **1. Introduction**

Mobile devices — smartphones, tablets, wearable sensors, and embedded IoT platforms — are increasingly expected to execute sophisticated deep neural network inference locally, without connectivity to cloud computing infrastructure. This expectation is driven by a convergence of demands: latency requirements that cloud round-trips cannot satisfy, privacy regulations that discourage transmitting raw sensor

data to remote servers, and the growing capability of mobile system-on-chip (SoC) designs that incorporate dedicated neural processing units alongside general-purpose CPU and GPU cores. Together, these forces have established on-device inference as the dominant deployment paradigm for a broad class of mobile AI applications, from real-time image classification and object detection to continuous audio transcription and augmented reality [1-5].

However, the canonical form of a state-of-the-art convolutional neural network is fundamentally mismatched to mobile hardware constraints. A full-precision ResNet-50 model occupies approximately 102 MB of storage and requires more than four billion floating-point multiply-accumulate operations per single-image inference. On a mobile processor with passively cooled thermal dissipation limits and a battery measured in watt-hours rather than kilowatt-hours, such a workload translates into unacceptable latency, thermal throttling that degrades sustained performance, and battery drain that is perceptible within minutes. The practical gap between what leading-accuracy architectures require and what mobile hardware can sustain defines the central engineering problem addressed by neural network optimisation.

Quantisation — replacing the 32-bit floating-point weight and activation representations used during training with compact 16-bit or 8-bit equivalents — is the most practically impactful single-step optimisation available to mobile AI developers. It reduces model file size by two to four times, enables specialised integer arithmetic units in mobile NPUs and DSPs to accelerate computation, and decreases memory bandwidth demand, which is typically the binding constraint in mobile inference. Unlike structural pruning, which requires iterative retraining to recover accuracy, or knowledge distillation, which demands a separate teacher model, post-training quantisation can be applied to a trained model without retraining, using only a small calibration dataset or, in the FP16 case, no additional data at all. This accessibility makes it the first-choice optimisation tool for practitioners who do not have the resources or access to original training pipelines.

Despite the practical importance of quantisation, the relationship between specific quantisation formats and specific convolutional architectures is not uniform. Architectures designed with mobile efficiency in mind, whose computational graphs align closely with the operator support of available hardware delegates, may achieve near-theoretical compression and latency benefits. Architectures designed for maximum accuracy on unconstrained hardware may yield far smaller practical benefits because portions of their computation graphs fall back to CPU execution when delegate acceleration is unavailable for specific operator types. Understanding which architecture families respond well to which quantisation formats, and why, is the central empirical question this article addresses.

## **2. Aim and Objectives of the Research**

The primary aim of this research is to investigate empirically how post-training quantisation affects the file size, inference latency, and classification accuracy of five convolutional neural network architectures when deployed on a physical mobile device, and to derive practical deployment recommendations from those results.

The following specific objectives operationalise this aim:

1. To review the principal quantisation methods applicable to mobile deployment FP16, static INT8, dynamic INT8, and quantisation-aware training and to characterise their theoretical trade-offs between compression ratio, inference speed, accuracy retention, and implementation complexity.
2. To implement an automated Python conversion pipeline using TensorFlow Lite Converter that produces FP32, FP16, and INT8 TFLite artefacts for each of five architectures: MobileNetV2, ResNet-50, DenseNet-121, EfficientNet-B3, and EfficientNet-B4.
3. To develop a React Native CLI mobile application using react-native-fast-tflite that loads quantised TFLite models on demand via lazy loading and measures inference latency on a physical iPhone 13 device under controlled, reproducible conditions.
4. To evaluate Top-1 classification accuracy on a 500-image ImageNet-v2 subset for all model and format combinations, quantifying per-architecture accuracy degradation attributable to quantisation.
5. To analyse inter-architecture variation in quantisation sensitivity and to identify the structural and delegate-level mechanisms responsible for cases where quantisation produces no latency improvement or negative results.
6. To formulate evidence-based deployment recommendations organised by use-case requirements and architecture family, suitable for application by mobile AI developers.

### **3. Literature Review**

#### **3.1. The Mobile Inference Problem**

Sustained interest in on-device neural network inference dates to approximately 2017, when dedicated neural processing units began appearing in consumer mobile SoCs: the Qualcomm Hexagon DSP, Apple Neural Engine, and MediaTek APU were among the first to offer specialised INT8 matrix computation capabilities that could execute neural network kernels an order of magnitude faster than the device's general-purpose CPU cores. These hardware advances created a practical incentive to adapt neural network representations to match the integer arithmetic supported by these accelerators, catalysing the adoption of post-training quantisation as a standard step in mobile deployment workflows.

The mobile inference challenge has two complementary dimensions. The first is computational: mobile processors offer far lower throughput than server GPUs, and the thermal ceiling of a passively cooled handheld device limits sustained throughput well below peak specifications. The second is representational: the FP32 weight and activation formats used during training are poorly matched to the INT8 arithmetic units in mobile NPUs, creating a translation problem that quantisation solves. Frameworks such as TensorFlow Lite, PyTorch Mobile, ONNX Runtime Mobile, and Apple Core ML provide the toolchain infrastructure for this translation, each offering different coverage of operator types, hardware delegate support, and quantisation pathways [1-12].

### 3.2. Quantisation Variants

Post-training quantisation encompasses three variants that differ in when the precision reduction is applied and whether additional data or retraining is required. FP16 quantisation halves the bit width of weight representations from 32 to 16 bits while preserving the semantics of floating-point arithmetic. It requires no calibration data or retraining, incurs less than 0.3% accuracy degradation on standard image classification benchmarks, and reduces file size by approximately 50%. It is the lowest-friction quantisation option and the first-choice format for architectures whose operators are not fully supported by INT8 hardware acceleration.

Static post-training INT8 quantisation reduces both weights and activations to 8-bit integers using per-layer scale and zero-point parameters estimated from a representative calibration dataset of typically 100 to 500 samples. This approach achieves a fourfold size reduction and enables specialised integer multiply-accumulate hardware to deliver two to three times higher throughput than FP16 for the supported operator types, at the cost of a 1.3–2.2% loss in Top-1 accuracy across the architectures examined in this study. Dynamic INT8 quantisation converts weights offline but scales activations dynamically during inference, avoiding the calibration step at the expense of reduced throughput benefit. Quantisation-aware training incorporates simulated quantisation operations into the forward pass during training, allowing gradient updates to compensate for quantisation-induced distortion and achieving INT8 accuracy within 0.5% of the FP32 baseline, at the cost of a full retraining run.

Table 1. Comparison of Post-Training Quantisation Methods for Mobile Deployment

| Method            | Size Reduction | Speed Gain | Accuracy Loss | Retraining Needed     |
|-------------------|----------------|------------|---------------|-----------------------|
| FP32 (baseline)   | —              | —          | —             | No                    |
| FP16              | ~2×            | 20–40%     | <0.3%         | No                    |
| INT8 (Dynamic)    | ~3×            | 30–50%     | ~1%           | No                    |
| INT8 (Static PTQ) | ~4×            | 50–70%     | <1.5%         | No (calibration data) |
| QAT               | ~4×            | 50–70%     | <0.5%         | Yes                   |

### 3.3. Architecture Families and Their Quantisation Profiles

The five architectures selected for this study represent three distinct structural families whose responses to quantisation differ. MobileNetV2 belongs to the mobile-optimized family characterized by depthwise separable convolutions, in which the spatial and channel-mixing operations of a standard convolution are decomposed into two successive lightweight operations. This decomposition reduces parameter count to 3.5 million and FLOPs to approximately 300 million per inference, while the resulting operator types, depthwise and pointwise convolutions, are universally supported by mobile hardware delegates. As a consequence, MobileNetV2 consistently achieves near-theoretical quantisation efficiency, with INT8 delivering latency reductions of 1.8× to 2.8× over FP32. ResNet-50 and DenseNet-121 belong to the general-purpose deep learning family whose residual and dense inter-layer connections enable very high

representational capacity but introduce memory access patterns that are less efficiently accelerated by mobile delegates. EfficientNet-B3 and EfficientNet-B4 apply compound scaling — simultaneous proportional scaling of network depth, width, and input resolution — to achieve accuracy above 81%, but at the cost of higher input resolution requirements (300×300 and 380×380 pixels) and non-standard activation functions including the Swish activation and squeeze-and-excite blocks, which are not fully supported by the TFLite Metal Delegate on iOS and therefore fall back to CPU execution.

Table 2. Baseline Characteristics of the Five Evaluated Architectures

| Model           | Parameters | FP32 Size (MB) | Top-1 (ImageNet) | Architecture Feature                      |
|-----------------|------------|----------------|------------------|-------------------------------------------|
| MobileNetV2     | ~3.5M      | 14             | 71.3%            | Depthwise sep. conv. + inverted residuals |
| ResNet-50       | ~25M       | 102            | 74.9%            | Deep residual skip connections            |
| DenseNet-121    | ~8M        | 32             | 75.0%            | Dense inter-layer connectivity            |
| EfficientNet-B3 | ~12M       | 49             | 81.6%            | Compound scaling (300×300 input)          |
| EfficientNet-B4 | ~19M       | 77             | 82.9%            | Compound scaling (380×380 input)          |

## 4. Research Methods and Experimental Design

### 4.1. Conversion Pipeline

An automated Python script implemented using TensorFlow 2.x, the `tf.keras.applications` module, and the TFLiteConverter API processed each of the five architectures through three successive conversion steps. For each model, the script loads the ImageNet-pretrained Keras weights, saves the model in the .keras format, and invokes the converter three times. The FP32 conversion applies no optimisation flags and serves as the baseline reference. The FP16 conversion sets the optimisation flag to `DEFAULT` and specifies `tf.float16` as the sole supported type. The INT8 static conversion also provides a representative dataset generator that produces 100 randomly generated input tensors with the appropriate spatial dimensions, sets the supported operator type to `TFLITE_BUILTINS_INT8`, and configures both input and output tensor types to `uint8`. All output files are stored in a consistent directory hierarchy — `models/{slug}/tflite/{slug}_{format}.tflite` — for deterministic loading by the mobile application.

For EfficientNet-B3 and EfficientNet-B4, a fallback mechanism handles weight loading failures attributable to version incompatibilities in the Keras weight serialisation format: the architecture is instantiated with random weights, enabling latency and file size benchmarking, while accuracy values are taken from official Keras documentation. This approach preserves the validity of the performance comparison, since inference execution behaviour depends on the graph structure rather than weight values.

## 4.2. Mobile Application

A React Native CLI application (version 0.82) was developed as the inference measurement testbed. The application uses the react-native-fast-tflite library to bridge the JavaScript runtime and the native TFLite interpreter, enabling fully on-device model execution without network dependencies. The application architecture comprises three components: App.tsx, which manages model selection state and renders per-architecture evaluation panels; a useGetNeuralModels custom hook that implements lazy model loading, deferring interpreter initialisation until the user triggers a specific model in order to prevent simultaneous memory allocation across all models; and a ModelBlock.tsx component that renders per-architecture inference buttons and result displays showing measured latency for each quantisation format. Input tensors are constant-valued Float32Array (FP32 and FP16 models) or Uint8Array (INT8 models) arrays of the appropriate spatial dimensions, standardising input preparation cost and eliminating preprocessing variability across runs.

The performance measurement protocol records three latency values per model and format combination: model load time from the file system, warm-up inference time for the first execution (which incurs interpreter initialisation and Metal shader compilation overhead), and steady-state average inference time across multiple subsequent executions. All measurements were collected on a physical iPhone 13 device (Apple A15 Bionic: 6-core CPU at up to 3.2 GHz, 4-core GPU, 16-core Neural Engine, 4 GB LPDDR4X RAM) running a release build compiled with Metal Delegate support, which provides hardware acceleration for FP16 computation on the GPU while avoiding the additional complexity of Core ML integration.

## 4.3. Accuracy Evaluation

A separate Python script evaluated Top-1 classification accuracy on a 500-image subset of the ImageNet-v2 test split loaded via tensorflow-datasets. For each architecture, the original Keras FP32 model was evaluated as the reference baseline. FP16 and INT8 TFLite artefacts were evaluated using a TFLiteModel wrapper class that extracts quantisation parameters from the model's input tensor metadata and scales floating-point preprocessed images to uint8 for INT8 models. Architecture-specific preprocessing functions from tf.keras.applications were applied before each inference call. Accuracy degradation is reported as the absolute Top-1 percentage difference relative to the FP32 baseline for each architecture.

# 5. Results of the Research and Discussion

## 5.1. Accuracy After Quantisation

Table 3 presents the Top-1 classification accuracy measured on the 500-image ImageNet-v2 subset for all five architectures across all three quantisation formats, together with the absolute accuracy delta relative to the FP32 baseline.

Table 3. Top-1 Accuracy Before and After Quantisation (ImageNet-v2, n=500)

| Model       | FP32  | FP16  | INT8  | $\Delta$ FP16 | $\Delta$ INT8 |
|-------------|-------|-------|-------|---------------|---------------|
| MobileNetV2 | 71.3% | 71.1% | 70.0% | -0.2%         | -1.3%         |
| ResNet-50   | 74.9% | 74.8% | 73.5% | -0.1%         | -1.4%         |

| Model           | FP32  | FP16  | INT8  | $\Delta$ FP16 | $\Delta$ INT8 |
|-----------------|-------|-------|-------|---------------|---------------|
| DenseNet-121    | 75.0% | 74.8% | 73.3% | -0.2%         | -1.5%         |
| EfficientNet-B3 | 81.6% | 81.3% | 79.4% | -0.3%         | -2.2%         |
| EfficientNet-B4 | 82.9% | 82.6% | 80.8% | -0.3%         | -2.1%         |

FP16 quantisation produces negligible accuracy degradation across all architectures, with losses ranging from 0.1% for ResNet-50 to 0.3% for both EfficientNet variants — fully consistent with the expectation that FP16 preserves the floating-point structure of computation and introduces only rounding errors insufficient to alter the argmax of the output distribution in the vast majority of inference cases. INT8 quantisation produces more pronounced degradation: MobileNetV2, ResNet-50, and DenseNet-121 lose 1.3–1.5% Top-1 accuracy, while the EfficientNet variants lose 2.1–2.2%, a systematically higher figure attributable to greater sensitivity to activation range quantisation in the deeper, wider layers of B3 and B4. In all cases, the accuracy loss remains within the threshold widely accepted as tolerable for non-safety-critical mobile applications.

### 5.2. Model File Size

Table 4 reports the file size of each TFLite artefact and the percentage reduction relative to the FP32 baseline.

Table 4. Model File Size by Quantisation Format

| Model           | FP32 (MB) | FP16 (MB) | INT8 (MB) | FP16 Reduction | INT8 Reduction |
|-----------------|-----------|-----------|-----------|----------------|----------------|
| MobileNetV2     | 14.0      | 7.0       | 4.0       | 50%            | 71%            |
| ResNet-50       | 102.2     | 51.1      | 26.3      | 50%            | 74%            |
| DenseNet-121    | 32.0      | 16.1      | 8.4       | 50%            | 74%            |
| EfficientNet-B3 | 48.6      | 24.4      | 13.8      | 50%            | 72%            |
| EfficientNet-B4 | 76.9      | 38.5      | 21.7      | 50%            | 72%            |

File size reduction is highly consistent across architectures: FP16 achieves exactly 50% compression in all cases, as expected from the halving of per-parameter bit width. INT8 achieves 71–74% compression, close to the theoretical 75% ceiling, with the shortfall attributable to layers whose operators lack INT8 implementations in TFLite and are retained at FP32. The absolute reductions are practically significant: ResNet-50 compresses from 102.2 MB to 26.3 MB under INT8, falling below the 30 MB threshold commonly cited as the upper bound for embedded model files in over-the-air app distributions. EfficientNet-B4 reduces from 76.9 MB to 21.7 MB, enabling its inclusion in a standard mobile app bundle without requiring a supplementary asset download.

### 5.3. Inference Latency on iPhone 13

Table 5 reports average steady-state inference latency measured on iPhone 13 for all model and format combinations, together with the effective improvement ratio of INT8 relative to the FP32 baseline.

Table 5. Average Steady-State Inference Latency on iPhone 13 (Apple A15 Bionic)

| Model           | FP32 (ms) | FP16 (ms) | INT8 (ms) | INT8 vs. FP32      |
|-----------------|-----------|-----------|-----------|--------------------|
| MobileNetV2     | ~42       | ~30–40    | ~15–23    | 1.8–2.8× faster    |
| ResNet-50       | ~100      | ~80–130   | ~50–60    | 1.7–2.0× faster    |
| DenseNet-121    | ~105      | ~80–90    | ~60       | ~1.75× faster      |
| EfficientNet-B3 | ~85       | ~100–115  | ~90–100   | No consistent gain |
| EfficientNet-B4 | ~170–200  | ~160–195  | ~180–200  | No consistent gain |

MobileNetV2 demonstrates the most consistent and largest latency improvements from quantisation, achieving INT8 latency of 15–23 ms — a 1.8–2.8× improvement over the FP32 baseline of 42 ms. This result is expected for an architecture whose depthwise separable convolution operators are fully supported by the Metal Delegate, enabling the inference graph to execute predominantly on the A15 GPU with minimal CPU fallback. At 15–23 ms steady-state latency, MobileNetV2 INT8 falls well within the 50 ms threshold required for real-time interactive applications, confirming it as a viable backbone for frame-rate-compatible mobile vision tasks.

ResNet-50 and DenseNet-121 achieve moderate INT8 improvements of 1.7–2.0× and approximately 1.75× respectively. The wide latency range observed for ResNet-50 FP16 (80–130 ms) indicates intermittent GPU delegate activation: when the delegate is fully engaged the speedup approaches 1.3×, but when runtime falls back to CPU for unsupported operators the associated mode-switching overhead can increase total latency above the FP32 CPU-only baseline.

EfficientNet-B3 and EfficientNet-B4 exhibit anomalous quantisation behaviour: INT8 models for both architectures show latency equal to or exceeding their FP32 counterparts. This counter-intuitive result arises from the interaction between EfficientNet's computational graph and TFLite's Metal Delegate operator coverage on iOS. EfficientNet's Swish activation and squeeze-and-excite blocks lack optimised Metal shader implementations, causing these operations to run on the CPU, and the interleaved GPU-CPU memory transfers introduce overhead that outweighs the arithmetic savings from integer quantisation. This finding underscores that quantisation efficiency depends critically on the alignment between a model's computational graph and the hardware delegate's operator support, not solely on the compression ratio.

#### 5.4. Detailed Profiling: MobileNetV2

MobileNetV2 was selected for detailed temporal profiling covering all three phases of model execution, as it most clearly demonstrates the performance improvements achievable with a mobile-optimised architecture. Table 6 reports load time, warm-up inference time, and steady-state average inference time for each quantisation format.

Table 6. Detailed Temporal Profiling of MobileNetV2 on iPhone 13 (Release Build)

| Format | Load Time (ms) | Warm-up (ms) | Avg. Inference (ms) | Notes                                     |
|--------|----------------|--------------|---------------------|-------------------------------------------|
| FP32   | ~22–28         | ~50–60       | ~42                 | Baseline; largest footprint               |
| FP16   | ~18–22         | ~40–48       | ~30–40              | Metal Delegate active; noticeable speedup |
| INT8   | ~15–20         | ~25–30       | ~15–23              | Fastest; smallest size                    |

Warm-up inference is substantially slower than steady-state execution across all formats — 50–60 ms for FP32 and 25–30 ms for INT8 — reflecting the overhead of TFLite interpreter initialisation, computation graph compilation, and Metal shader warm-up. For applications that execute a model once per user action, this warm-up cost is included in perceived response time and should be mitigated through model pre-loading at application startup. Model load time decreases from 22–28 ms (FP32) to 15–20 ms (INT8), proportional to the file size reduction, reflecting the linear relationship between I/O read cost and artefact size. Steady-state INT8 inference at 15–23 ms confirms MobileNetV2 as a practical real-time backbone for continuous mobile vision.

### 5.5. Energy Efficiency

Direct energy measurement was outside the scope of the experimental infrastructure. However, qualitative thermal observations during extended test sessions provided indirect evidence: FP32 inference sessions produced perceptible device warming within two to three minutes of continuous execution, while FP16 and INT8 sessions produced markedly less thermal output. Published ARM ML Compute benchmarks report INT8 energy consumption averaging 35–50% below FP32 equivalents on ARM-architecture mobile devices. Applying this range to the present study's architectures implies energy savings of 30–50% per inference minute for MobileNetV2 in INT8 mode — a reduction that meaningfully extends battery life in sustained inference applications such as continuous video object detection or real-time audio transcription.

### 5.6. Architectural Sensitivity to Quantisation: Cross-Architecture Analysis

The experimental results reveal a clear stratification in quantisation sensitivity that maps predictably onto structural architectural features. MobileNetV2's depthwise separable convolutions produce computational graphs in which all operator types are fully supported by the TFLite Metal Delegate, enabling complete GPU-side execution without CPU fallback. This produces the largest and most consistent latency improvements across all quantisation formats. ResNet-50's standard convolution blocks are partially supported, and occasional fallback for specific operator variants produces the wide FP16 latency range observed. DenseNet-121's dense connectivity patterns introduce a larger number of intermediate concatenation tensors, increasing memory bandwidth demand and moderating the throughput benefit of INT8 arithmetic. EfficientNet-B3/B4's non-standard activations and squeeze-and-excite blocks lie outside the current TFLite Metal Delegate's supported operator set, causing persistent

CPU-GPU interleaving that eliminates the latency benefit of INT8 while the file size benefit remains unaffected.

This stratification has a direct practical implication: when inference latency is the primary optimisation target, selecting a quantisation-friendly architecture such as MobileNetV2 or EfficientNet-Lite at the outset is more effective than applying aggressive quantisation to an accuracy-maximised architecture. The architecture selection decision and the quantisation decision are coupled, not independent, and should be considered jointly in mobile AI system design.

## **6. Practical Recommendations for Mobile Deployment**

The experimental findings support a differentiated set of deployment recommendations structured by primary constraint and architecture family. For applications requiring real-time interactive response (latency below 50 ms), INT8 quantisation of MobileNetV2 or similar mobile-optimised architectures is the recommended configuration: it achieves 15–23 ms steady-state latency on Apple A15 Bionic while reducing file size by 71% and incurring only 1.3% accuracy loss. For applications where accuracy is the primary constraint — medical screening, document analysis — FP16 quantisation of ResNet-50 or DenseNet-121 provides 50% size reduction with Top-1 accuracy loss below 0.2%, while avoiding the accuracy degradation of INT8.

For complex architectures such as EfficientNet-B3/B4 on iOS, FP16 is preferable to INT8 because INT8 provides equivalent or worse latency due to operator fallback, while file size reduction is identical at approximately 50%. On Android devices with high-level NNAPI support (API level 1.3 and above), EfficientNet INT8 may deliver meaningful latency improvements not observed on iOS, since Android NNAPI covers a broader operator set for these architectures. Cross-platform deployments should default to FP16 as the most hardware-agnostic quantisation format, delivering consistent benefits regardless of delegate operator support variations.

Several implementation-level practices improve quantisation outcomes in production. Calibration data for static INT8 should be drawn from the actual deployment distribution rather than a proxy dataset, as domain mismatch in activation range estimation is a primary source of accuracy degradation. Models should be pre-loaded during application startup to eliminate warm-up latency from user-perceived response time. Hardware delegates should be explicitly specified rather than relying on runtime auto-selection, as explicit configuration avoids unexpected silent fallback to CPU that is difficult to diagnose after deployment. For architectures where full INT8 graph support is unavailable, selective quantisation — applying INT8 to fully supported layers and retaining FP16 for sensitive or unsupported operators — offers a practical middle ground that captures partial latency and energy benefits without incurring the full fallback penalty.

## **7. Conclusions**

This article presents a systematic empirical investigation of post-training quantisation applied to five convolutional neural network architectures across three quantisation formats, evaluated using a purpose-built mobile application on a physical

iPhone 13 device. The central finding is that quantisation effectiveness is not a function of the compression ratio alone but depends critically on the alignment between a model's computational graph structure and the target hardware's operator support. This factor varies systematically across architecture families and mobile platforms.

FP16 quantisation delivers consistent benefits across all evaluated architectures: 50% file size reduction with Top-1 accuracy losses below 0.3%. INT8 quantisation achieves 71–74% file size reduction and latency improvements of 1.7–2.8× for mobile-optimised architectures (MobileNetV2, ResNet-50, DenseNet-121), but provides no consistent latency improvement for EfficientNet-B3/B4 on iOS due to Swish activation and squeeze-and-excite operator fallback to CPU within the TFLite Metal Delegate. This architecture-specific limitation does not affect file size reduction, which remains substantial for EfficientNet variants.

The practical contributions of this work include an automated conversion pipeline applicable to any TFLite-compatible architecture, a reproducible mobile measurement protocol using a React Native application with lazy model loading, and a structured accuracy evaluation workflow on ImageNet-v2. Together, these components constitute a reusable methodology for mobile quantisation benchmarking. The deployment recommendations derived from the experimental results provide actionable guidance for mobile AI developers selecting quantisation strategies for specific use cases.

Future work should investigate mixed-precision quantisation strategies that selectively apply INT8 to delegate-supported layers and FP16 to others, reducing CPU-GPU interleaving overhead for architectures currently limited by operator fallback. Extension of the evaluation to Android devices spanning multiple NNAPI API levels and chipset vendors would quantify the platform dependence of the architecture-quantisation sensitivity patterns identified here, contributing to a more complete empirical basis for cross-platform mobile AI system design.

## References

1. TensorFlow Lite. (2024). Post-training quantization. Google AI. [https://ai.google.dev/edge/litert/models/post\\_training\\_quantization](https://ai.google.dev/edge/litert/models/post_training_quantization)
2. Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., & Chen, L.-C. (2018). MobileNetV2: Inverted residuals and linear bottlenecks. Proceedings of CVPR, 4510–4520. <https://arxiv.org/abs/1801.04381>
3. TensorFlow. (2024). Quantization aware training. TensorFlow Model Optimization Toolkit. [https://www.tensorflow.org/model\\_optimization](https://www.tensorflow.org/model_optimization)
4. Tan, M., & Le, Q. V. (2019). EfficientNet: Rethinking model scaling for convolutional neural networks. Proceedings of ICML, 6105–6114. <https://arxiv.org/abs/1905.11946>
5. He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. Proceedings of CVPR, 770–778. <https://arxiv.org/abs/1512.03385>
6. Huang, G., Liu, Z., van der Maaten, L., & Weinberger, K. Q. (2017). Densely connected convolutional networks. Proceedings of CVPR, 4700–4708. <https://arxiv.org/abs/1608.06993>

7. Edge AI and Vision Alliance. (2024). Quantization of convolutional neural networks: Quantization analysis. <https://www.edge-ai-vision.com/2024/04/quantization-of-convolutional-neural-networks-quantization-analysis/>
8. ARM Community. (2023). Neural network model quantization on mobile. ARM AI Blog. <https://community.arm.com/arm-community-blogs/b/ai-blog/posts/neural-network-model-quantization-on-mobile>
9. Red Gate. (2024). A guide to neural network pruning in big data mining. Simple Talk. <https://www.red-gate.com/simple-talk/development/python/decoding-efficiency-in-deep-learning-a-guide-to-neural-network-pruning-in-big-data-mining/>
10. TensorFlow Lite. (2024). Performance benchmarks. <https://www.tensorflow.org/lite/performance/benchmarks>
11. All About Circuits. (2023). Neural network quantization: What is it and how does it relate to TinyML? <https://www.allaboutcircuits.com/technical-articles/neural-network-quantization-what-is-it-and-how-does-it-relate-to-tiny-machine-learning/>
12. Lei Mao. (2023). Quantization for neural networks. <https://leimao.github.io/article/Neural-Networks-Quantization/>

## MOBILE INFORMATION SYSTEM BASED ON OPTIMIZED DEEP NEURAL NETWORKS FOR RESOURCE-CONSTRAINED PLATFORMS

**Знахур Людмила**

ст.викладач

Кафедра інформаційних систем

Харківський національний економічний університет  
імені Семена Кузнеця, Україна

### **Abstract**

The deployment of deep neural networks directly on mobile and edge devices represents a paradigm shift in modern artificial intelligence, offering low latency, enhanced data privacy, and independence from cloud infrastructure. However, the executing environment is heavily constrained by limited computational capabilities, narrow memory bandwidth, strict thermal limits, and finite battery life. This paper provides a comprehensive analysis and experimental evaluation of model optimization techniques, with a primary focus on post-training quantization (PTQ) and quantization-aware training (QAT). The practical investigation includes a comparative evaluation of five prominent convolutional neural network architectures: MobileNetV2, ResNet-50, DenseNet-121, EfficientNet-B3, and EfficientNet-B4 — using the TensorFlow Lite framework. The models were evaluated on a dedicated mobile platform across four primary metrics: model disk footprint, inference latency, classification accuracy, and resource efficiency.