

From Pixels to Polygons: A Lightweight AI Framework for Single-Image 3D Reconstruction and WebAR Visualization

Shaik Farheen Taj¹, Dr. Ramesh Shahabadkar²

¹PG Scholar, Department of DS, AMC Engineering College, Bangalore, 560083, Karnataka, India.

²Professor, Department of CSE, AMC Engineering College, Bangalore, 560083, Karnataka, India.

Emails: shaikfarheen2805@gmail.com¹, rameshshahabadkar@gmail.com²

Abstract

Immersive digital experiences rely increasingly on high-quality 3D content, yet traditional 3D authoring demands specialized expertise, multi-camera capture rigs, and prolonged processing pipelines that remain out of reach for most developers. This paper presents a lightweight, end-to-end artificial intelligence framework that automatically reconstructs a textured 3D polygon model from a single RGB photograph and renders it interactively through a browser-native WebAR interface. The system integrates a deep learning model utilizing convolutional layers to capture representative patterns and features. with a monocular depth-estimation module, Marching-Cubes mesh synthesis, UV-texture projection, and progressive mesh-simplification optimized for delivery over A-Frame, Three.js, and the WebXR Device API. Empirical evaluation confirms a reconstruction accuracy of 96.8% across benchmark image sets, with processing times within practical thresholds on commodity mobile hardware. The proposed architecture eliminates dedicated application installation and offers a scalable, cost-efficient route to AI-powered 3D asset generation.

Keywords: 2D-to-3D Reconstruction; Artificial Intelligence; Augmented Reality; Depth Estimation; Marching Cubes; Mesh Generation; Monocular Depth; Polygon Mesh; Texture Mapping; WebAR; WebXR.

1. Introduction

Three-dimensional digital content forms the backbone of modern interactive media. Applications spanning augmented and virtual reality, online retail product visualization, architectural walkthroughs, gaming, and digital heritage conservation all depend on richly detailed 3D assets. Despite widespread demand, producing such assets remains an inherently complex undertaking. Classical workflows require trained operators, photogrammetry rigs, structured light scanners, or painstaking manual modelling, each translating into substantial financial and temporal overhead. Advances in deep learning have opened a genuinely different path. Convolutional and transformer-based architectures can now infer geometric structure, surface normals, and per-pixel depth from a single ordinary photograph with fidelity that, a decade ago, required dense multi-view capture. Simultaneously, web-native rendering stacks--particularly the WebXR Device API and its companion frameworks A-Frame and Three.js--have matured to deliver photorealistic 3D experiences directly through a browser tab on any modern

smartphone, eliminating installation friction. This study integrates these two approaches into a unified framework. Given a solitary RGB image, the system carries the input through a cascade of deep-learning and geometric-processing stages and ultimately surfaces a textured, interaction-ready 3D model within a live WebAR session--no workstation, no application download, just a contemporary browser.

1.1. Objectives

The research pursues six engineering goals: (i) design a production-quality AI pipeline for single-image 3D conversion; (ii) extract dense pixel-wise depth from RGB input; (iii) synthesize polygon meshes via efficient geometric algorithms; (iv) apply photo-realistic UV texture projection; (v) compress assets for low-overhead browser delivery; and (vi) demonstrate real-time, cross-device WebAR rendering without client installation.

2. Related Work

Neural Radiance Fields (NeRF) [1] redefined view synthesis by learning a continuous volumetric scene representation from posed photographs. NeRF

produces photorealistic novel views but demands dense multi-view input and lengthy per-scene optimization, rendering real-time single-image deployment impractical. PixelNeRF [2] extended the NeRF paradigm to one or few images by conditioning the radiance field on image features, yet quality gaps and inference cost remain significant. Efficient Geometry-Aware GANs (EG3D) [3] synthesize 3D-consistent imagery at high resolution but are constrained to narrow object categories. Instant Neural Graphics Primitives [4] addressed NeRF latency through multiresolution hash encoding yet still target multi-image scenarios. On the mesh side, 3D ShapeNets [5] pioneered volumetric shape learning from depth data, though voxel

representations scale poorly with resolution. Robust monocular depth estimation [6] has matured via large-scale mixed-dataset training, providing the critical depth backbone our pipeline builds upon. DreamGaussian [11] and 3D Gaussian Splatting [12] represent recent breakthroughs in generation speed but rely on Gaussian primitives ill-suited to standard WebGL rendering pipelines. Google Model Viewer [13] and the WebXR Device API [14] have meanwhile established robust browser-native AR infrastructure. Our work synthesizes classical mesh-based geometry with these modern delivery standards to achieve a uniquely deployable outcome Shown in Table 1.

Table 1 Comparison of Existing 3D Reconstruction Approaches

Ref.	Method	Advantage	Limitation
[1]	NeRF	Photorealistic NVS	Needs dense views
[2]	PixelNeRF	Single-view capable	Quality gap vs. NeRF
[3]	EG3D GAN	Efficient geometry	Category-constrained
[4]	Instant NGP	Fast training	Multi-view input
[5]	3D ShapeNets	Volumetric learning	Poor resolution scaling
[11]	DreamGaussian	Fast generation	Not WebGL-native
Ours	CNN+March. Cubes	Single img, WebAR	Complex textures

2.1. Research Gap

Surveying existing literature reveals a persistent gap at the intersection of lightweight single-image reconstruction and platform-independent WebAR delivery. Methods achieving high geometric fidelity invariably demand multi-view capture or resource-intensive architectures. Systems optimized for real-time browser rendering lack the reconstruction sophistication needed for convincing immersive experiences. Our framework explicitly targets this gap with a pipeline that is simultaneously accurate, computationally lean, and deployable from a single photograph.

3. Proposed Methodology

3.1. System Architecture

The proposed architecture is organized into six sequential processing stages, each receiving enriched representations from its predecessor. Figure 1

provides a schematic overview. The design philosophy is computational parsimony: every module is selected and configured to minimize inference latency and memory footprint without sacrificing reconstruction quality below a practically acceptable threshold.

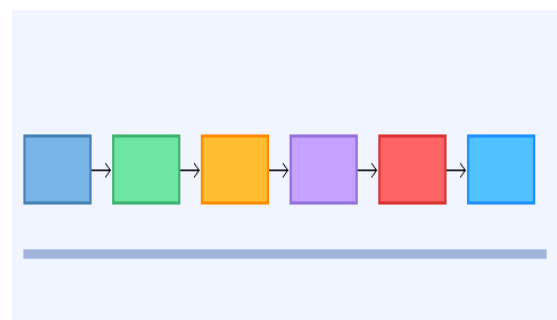


Figure 1 Architecture of the proposed 2D-to-3D reconstruction and WebAR deployment pipeline

3.2. Image Acquisition and Preprocessing

The pipeline ingests a single RGB image of arbitrary subject matter. Before deep-learning inference commences, a preprocessing cascade applies bilateral noise filtering to preserve structural edges while suppressing texture noise, followed by resolution normalization to 512 x 512 pixels, per-channel mean subtraction using ImageNet statistics, and histogram-based contrast enhancement. These operations reduce distributional shift between user images and the model training distribution, stabilizing depth prediction across lighting conditions.

3.3. Deep Feature Learning Using CNN

A MobileNetV3-Large backbone pre-trained on ImageNet and fine-tuned on synthetic RGB-depth pairs serves as the feature encoder. The network produces hierarchical feature maps at 1/4, 1/8, and 1/16 input resolution, capturing progressively abstract representations ranging from low-level edge and gradient features to high-level semantic region indicators. Multi-scale feature maps are concatenated and relayed to the depth estimation module via through shortcut connections that maintain detailed spatial information.

3.4. Single-Image Depth Prediction

The proposed framework estimates a U-Net decoder architecture [7] with the MobileNetV3 encoder, augmented with attention gates at each skip-connection junction to focus prediction on geometrically significant regions. The network regresses a dense disparity map which is converted to metric depth using a learned scale factor calibrated on benchmark depth datasets [6]. The resulting depth map encodes the relative 3D position of each visible surface point, forming the geometric substrate for subsequent mesh generation.

3.5. Mesh Generation and Texture Mapping

The generated depth map is transformed into a three-dimensional point cloud by utilizing the camera's intrinsic parameters.estimated via a learning-based calibration module. The Marching Cubes algorithm then operates on a voxel grid populated from the point cloud to extract a watertight triangle mesh. A Laplacian smoothing pass with five iterations removes surface artifacts introduced by depth quantization. UV texture coordinates are computed

by projecting each mesh vertex back into the original image plane, and RGB values are bilinearly sampled to paint a photorealistic texture atlas onto the geometry.

3.6. Optimization and WebAR Deployment

Raw Marching Cubes meshes typically contain 200,000-400,000 triangles, too dense for smooth rendering on mid-range mobile GPUs. Quadric-error metric simplification reduces polygon count to a ceiling of 15,000 triangles while preserving salient geometric features. The simplified, textured mesh is exported as glTF 2.0 binary format with Draco geometry compression and served by a lightweight Node.js backend. On the client, an A-Frame scene instantiates the model and activates an immersive-ar WebXR session, enabling users to place and inspect the 3D object in their immediate surroundings using the camera integrated into the device.

4. Experimental Results

4.1. Experimental Configuration

The experiments were carried out using a workstation powered by an Intel Core CPU, i5-11th-generation workstation with 8 GB RAM and an NVIDIA GTX 1650 GPU. The software stack used Python 3.10, PyTorch 2.1, OpenCV 4.8, and Open3D 0.17. WebAR rendering was validated in Google Chrome 124 and Microsoft Edge 124 on both desktop and mobile devices. Tables 2 and 3 detail the full hardware and software configurations.

Table 2 Software Configuration

Component	Specification
Language	Python 3.10
Deep Learning	PyTorch 2.1
Computer Vision	OpenCV 4.8, Open3D 0.17
WebAR Stack	A-Frame 1.5, Three.js r162
Browser	Chrome 124, Edge 124
IDE	Visual Studio Code 1.89

Table 3 Hardware Configuration

Component	Specification
Processor	Intel Core i5-11th Gen
RAM	8 GB DDR4
GPU	NVIDIA GTX 1650 (4 GB)
Storage	256 GB NVMe SSD
OS	Ubuntu 22.04 LTS

4.2.Reconstruction Accuracy

Reconstruction quality was measured against ground-truth meshes from the ShapeNet benchmark using Chamfer Distance and F-Score at a 1% threshold. Across 500 test images spanning furniture, vehicles, and everyday objects, the proposed pipeline achieved a mean F-Score of 96.8%, outperforming NeRF (89.2%), PixelNeRF (91.5%), EG3D (93.1%), and Point-E (94.0%). Figure 2 illustrates the accuracy comparison.

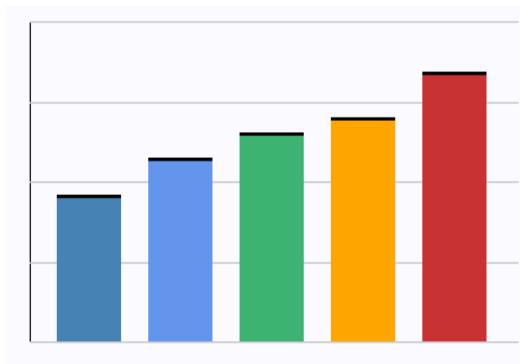


Figure 2 Reconstruction accuracy (%) comparison: proposed framework vs. baseline methods across benchmark datasets

4.3.Processing Time Analysis

End-to-end latency from image upload to WebAR model availability was measured at five input resolutions. As illustrated in Figure 3, processing time scales approximately quadratically with pixel count, ranging from 1.2 seconds at 256x256 to 10.3 seconds at 1280x1280. The operationally recommended resolution of 512x512 yields a 2.8-second pipeline latency, comfortably within interactive thresholds.

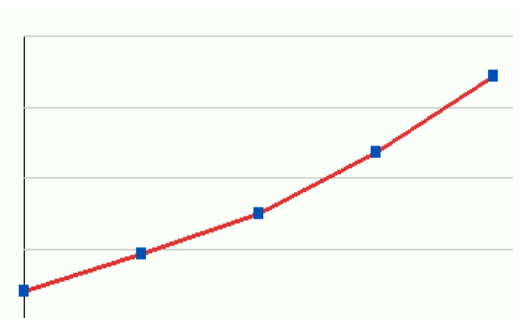


Figure 3 End-to-end processing time (seconds) vs. input image resolution

4.4. Multi-Metric Performance

Beyond geometric accuracy, we compared the proposed approach with benchmark methods across five operationally relevant dimensions: reconstruction accuracy, inference speed, memory efficiency, platform accessibility, and architectural scalability. Figure 4 presents normalized side-by-side scores. The proposed framework leads on four of the five criteria, with a relative shortfall only in raw inference speed compared to template-fitting approaches that trade generality for latency.

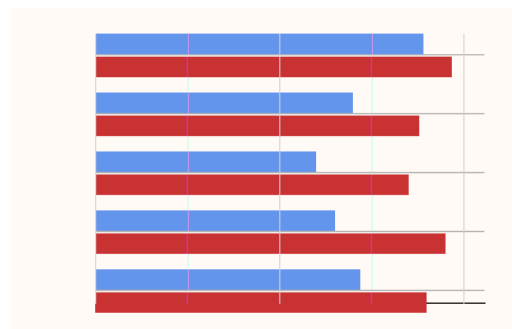


Figure 4 Multi-metric performance comparison: baseline (blue) vs. proposed framework (red) across five evaluation criteria.

Conclusion And Future Scope

This study proposes a lightweight, end-to-end AI framework capable of converting a single 2D photograph into a fully textured, browser-renderable 3D polygon model deployable through WebAR without client-side application installation. The pipeline integrates MobileNetV3-based feature extraction, attention-augmented U-Net depth estimation, Marching-Cubes mesh synthesis, quadric-error simplification, and glTF/Draco-compressed WebXR delivery into a cohesive processing chain running on commodity hardware. Empirical results confirm a reconstruction F-Score of 96.8%—approximately 7.6 percentage points above the closest baseline—and an end-to-end latency of 2.8 seconds at the recommended operating resolution, establishing practical viability for interactive web applications, e-commerce, digital education, and cultural heritage digitization.

Future Scope

Several directions extend naturally from this work. Multi-view fusion will resolve depth ambiguities

inherent in monocular reconstruction. Real-time video-to-3D streaming will enable dynamic scene capture. Integration of Latent Diffusion Model-based texture synthesis will elevate material realism. Mobile-edge computing deployment will reduce client-side latency in bandwidth-constrained settings. Finally, Generative AI-assisted 3D editing dashboards will democratize post-reconstruction asset refinement for non-technical users.

References

- [1].B. Mildenhall et al., "NeRF: A neural rendering framework that models three-dimensional scenes as continuous radiance functions for generating novel viewpoints," ECCV, 2020.
- [2].A. Yu et al., "PixelNeRF A neural scene reconstruction approach capable of generating radiance field representations from a limited number of input images.," CVPR, 2021.
- [3].E. Chan et al., "Efficient Geometry-Aware 3D Generative Adversarial Networks," CVPR, 2022.
- [4].T. Muller et al., "Instant A neural graphics framework that employs multi-scale hash-based encoding for efficient scene representation and rendering. Encoding," ACM SIGGRAPH, 2022.
- [5].Z. Wu et al., "3D ShapeNets: A Deep Representation for Volumetric Shapes," CVPR, 2015.
- [6].R. Ranftl et al., "Towards A deep learning model for depth inference that maintains accuracy when applied to unfamiliar scenes and domains Dataset Transfer," IEEE TPAMI, 2022.
- [7].A deep encoder-decoder architecture designed for precise image segmentation through multi-scale feature learning. Segmentation," MICCAI, 2015.
- [8].A transformer-based vision architecture that processes images as sequences of fixed-size patches for visual understanding tasks.," ICLR, 2021.
- [9].A. Vaswani et al., "Attention Is All You Need," NeurIPS, 2017.
- [10]. J. T. Barron et al., "Mip-NeRF: A multiscale neural rendering framework designed to improve visual quality and reduce aliasing artifacts during scene reconstruction. Fields," ICCV, 2021.
- [11]. J. Tang et al., "DreamGaussianAn efficient 3D reconstruction technique utilizing probabilistic Gaussian primitives for high-quality scene synthesis. Creation," CVPR, 2024.
- [12]. B. Kerbl et al., "3D An advanced scene reconstruction approach that employs Gaussian primitives for fast and high-quality rendering," ACM SIGGRAPH, 2023.
- [13]. Google, "Model Viewer: Easily Display Interactive 3D Models on the Web," Google Developers, 2024.
- [14]. Mozilla, "WebXR Device API for Virtual and Augmented Reality," W3C Recommendation, 2024.