

Lightweight, Uncertainty-Aware Conformalized Visual Odometry

Alex C. Stutts, Danilo Erricolo, Theja Tulabandhula, and Amit Ranjan Trivedi

Abstract—Data-driven visual odometry (VO) is a critical subroutine for autonomous edge robotics, and recent progress in the field has produced highly accurate point predictions in complex environments. However, emerging autonomous edge robotics devices like insect-scale drones and surgical robots lack a computationally efficient framework to estimate VO’s predictive uncertainties. Meanwhile, as edge robotics continue to proliferate into mission-critical application spaces, awareness of the model’s predictive uncertainties has become crucial for risk-aware decision-making. This paper addresses this challenge by presenting a novel, lightweight, and statistically robust framework that leverages conformal inference (CI) to extract VO’s uncertainty bands. Our approach represents the uncertainties using flexible, adaptable, and adjustable prediction intervals that, on average, guarantee the inclusion of the ground truth across all degrees of freedom (DOF) of pose estimation. We discuss the architectures of generative deep neural networks for estimating multivariate uncertainty bands along with point (mean) prediction. We also present techniques to improve the uncertainty estimation accuracy, such as leveraging Monte Carlo dropout (MC-dropout) for data augmentation. Finally, we propose a novel training loss function that combines interval scoring and calibration loss with traditional training metrics—mean-squared error and KL-divergence—to improve uncertainty-aware learning. Our simulation results demonstrate that the presented framework consistently captures true uncertainty in pose estimations across different datasets, estimation models, and applied noise types, indicating its wide applicability.

I. INTRODUCTION

As machine learning (ML) becomes more prevalent in privacy, safety, and mission-critical robotics, the ability to quantitatively and visually assess predictive uncertainties of ML models is becoming essential for risk-aware control and decision-making. Two types of uncertainty can result from data-driven learning: *epistemic* and *aleatoric*. Epistemic uncertainty arises from the variance in the training data and can often be reduced with more data. On the other hand, aleatoric uncertainty arises from random distortions in the data, such as blurriness, occlusions, overexposure, *etc.*, which additional training data cannot mitigate.

Although mitigating epistemic uncertainty is challenging, detecting, explaining, and handling aleatoric uncertainty is even more difficult. Therefore, for mission-critical risk-aware robotics, it is desired that predictive models under aleatoric uncertainty provide confidence measures and express input/model-dependent predictive uncertainties, along with the point (mean) prediction. Additionally, the uncertainty estimates must be extracted with minimal additional

Acknowledgement: This work was supported in part by COGNISENSE, one of seven centers in JUMP 2.0, a Semiconductor Research Corporation (SRC) program sponsored by DARPA. Authors are with the University of Illinois Chicago (UIC), Chicago, IL, Email: astutt2@uic.edu

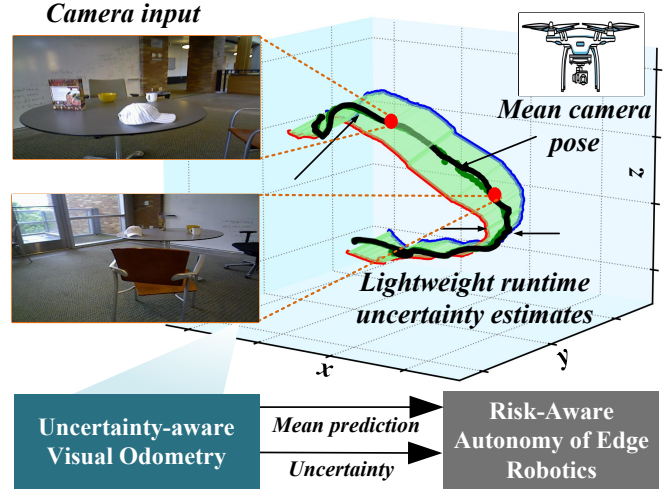


Fig. 1: **Lightweight uncertainty-aware visual odometry (VO):** In this paper, we present a lightweight framework for uncertainty-aware visual odometry for edge robotics where computing resources are limited while the predictions need to be made in real time. Our framework exploits *conformal inference* and presents four novel methodologies to extract point prediction and upper/lower bounds under the designated uncertainty coverage.

computing cost since many edge robotic systems have limited computing and storage capacity due to cost, footprint, legacy design, battery power, and other factors.

Meanwhile, given their computational intensity, traditional methods for uncertainty-aware predictions, such as Bayesian ML models, are unsuited for edge robotics. For example, in Bayesian ML models, a posterior distribution of weights is learned from the training data using Bayesian principles. Weights are then sampled from the posterior, and the network output is statistically computed and weighed against each sample’s probability. Generating uncertainty-aware predictions, thus, involves sampling many model parameters from the posterior and evaluating predictions at each sample, which can be impractical under typical time and resource constraints for various edge robotics applications.

A more computationally-efficient alternative to Bayesian inference is Variational inference (VI) [1]. VI considers a family of approximate densities F , from which a member $q^{(w)}$ is learned by minimizing the Kullback-Leibler (KL) divergence to the posterior. However, selecting a flexible family that closely matches the true posterior in practice can be challenging. Additionally, VI struggles to approximate complex posteriors with multiple modes or sharp peaks.

Recently, conformal inference (CI) of ML models, also known as conformal prediction, was developed to overcome the above limitations of traditional uncertainty-aware prediction frameworks [2]–[7]. Unlike classical statistical

inference, which relies on the data distribution to capture prediction uncertainty and can be sensitive to model misspecification, CI is a distribution-free approach that guarantees the statistical validity of uncertainty intervals given a finite training sample. CI can also assess the degree of conformity of each new observation to the available data and uses this information to construct an uncertainty interval calibrated to a user-specified coverage rate. Moreover, CI can be combined with any underlying model with an inherent notion of uncertainty to provide uncertainty quantification that is both statistically valid and model-agnostic. One of the challenges in making CI practical is that the uncertainty estimates are often conservative, and our work overcomes this and other limitations, providing a compelling case for their use in robotics applications with stringent resource usage requirements.

In particular, leveraging these advantages of CI for edge robotics, in this work, we investigate the conformalization of visual odometry (VO), an essential task for autonomous navigation that estimates the position and orientation (i.e., pose) of a camera mounted on a moving vehicle relative to the environment as the camera moves. The resulting motion estimates from VO are used for typical autonomy objectives such as three-dimensional (3D) reconstruction, mapping, localization, *etc.* Especially since VO relies solely on cameras, which are passive, low power, and have a small footprint, VO-based motion estimates are suitable for ego-motion tracking under stringent footprint and battery power constraints. Fig. 1 overviews the proposed techniques.

Exploring CI for lightweight uncertainty estimates in VO, our work makes the following key contributions:

- We present *four frameworks* for extracting VO-uncertainty bands by univariate and multivariate conformalization. The presented frameworks vary in computing workload, uncertainty accuracy, interval adaptiveness, data dependency, *etc.*, aiming to provide an entire spectrum of solutions for extracting VO's uncertainty estimates under varying computing resources, training data, and timing constraints for edge robotics.
- We present a novel loss function for joint training of point (mean) prediction and heteroskedastic uncertainty bounds. The loss function combines an interval score function and combined calibration loss function with mean squared error (MSE) loss and Kullback-Leibler (KL) divergence. Our novel loss function improves the accuracy of uncertainty coverage and enables computationally efficient predictive architecture for the simultaneous extraction of point predictions and uncertainties.
- We also discuss data augmentation for uncertainty-aware VO by providing additional training set by Monte Carlo Dropout (MC-Dropout), thus enabling the distillation of uncertainty estimates from sophisticated, computationally expensive frameworks to conformal quantile bounds.

II. RELATED WORKS

Under VO, the camera's ego motion is estimated based on the changes in visual information captured by the camera.

The traditional methods utilized techniques that involved identifying and tracking distinctive visual features, such as corners or edges, for VO [8]. Direct methods [9] estimated the camera's motion by analyzing the changes in the intensity values of the pixels in the images captured by the camera. Hybrid methods for VO combined feature-based and direct methods; for example, feature-based techniques were used for initial camera motion and then refined using direct methods [10]. Structure from motion (SfM) is popularly used for simultaneously estimating the 3D structure of the environment and camera's motion [11].

A recent trend for VO is to use deep neural networks (DNN) to directly learn the relation between the camera's visual field and its ego-motion from data. PoseNet [12] demonstrated the feasibility of training convolutional neural networks for end-to-end tracking of the ego-motion of monocular cameras. PoseLSTM [13] utilized long short-term memory (LSTM) to improve ego-motion accuracy compared to frame-based methods such as PoseNet. DeepVO [14] demonstrated the combined strength of recurrent convolutional neural networks (RCNNs) in learning features and modeling sequential relationships between consecutive camera frames. UnDeepVO [15] proposed an unsupervised deep learning approach capable of absolute scale recovery.

While many prior studies have focused on improving the accuracy of VO, only a few have addressed the challenge of extracting its predictive uncertainties, especially under time and computing resource constraints; for instance, [16] integrated deep learning-based depth predictions with particle filtering to achieve uncertainty-aware visual localization. However, particle filtering-based uncertainty quantification can be computationally expensive for large or high-dimensional state spaces. Kalman filters are a more computationally efficient alternative to particle filters, but their usage is limited to linear systems with Gaussian noise assumptions. By integrating Monte-Carlo Dropout (MC-dropout) [17] with deep learning-based predictors, such as PoseNet, [18] showed uncertainty-aware VO estimations. However, the method requires performing a sufficient number of dropout iterations and evaluating predictions for each, making it challenging to implement under time and computing resource constraints. Besides, MC-dropout, as a basis for uncertainty, lacks statistical rigor and ultimately leads to conservative prediction intervals that are not particularly adaptable to the data. D3VO [19] also utilized Bayesian techniques for uncertainty quantification, demonstrating state-of-the-art accuracy but with a large computational expense. With similar computational overhead, MDN-VO [20] combined a RCNN and mixture density network for uncertainty estimations based on maximizing the likelihood of pose estimations across a sequence of images. In contrast, the authors in [21] presented a lightweight deterministic uncertainty estimator as a small neural network that could be applied to VO or any data-driven deep neural network model. The method uncovers spatial and semantic model uncertainty with significantly less computation but lacks statistical guarantees.

Unlike the above, our methods build on conformal in-

Table I: Qualitative comparison of proposed techniques on VO-uncertainty band extraction

Method	CQR (Sec. III-A)	CSP (Sec. III-B)	MCQR w/ MC-dropout (Sec. IV-A)	CJP (Sec. IV-B)
Uncertainty Bands	Univariate	Univariate	Multivariate	Multivariate
Disjoint Bands	No	Yes	No	No
Computational Complexity	Low	Low	High	Medium
Interval Calibration	Fixed	Fixed	Fixed	Tunable
Interval Adaptiveness	Conservative	Semi-tunable	Flexible	Tunable
Uncertainty Accuracy	Low	Medium	High	High
Training Time	Short	Medium	Long	Medium
Data Dependency	Low	High	High	Low

ference, which can provide statistically valid uncertainty intervals while not requiring heavy computational budgets. Combining CI with an underlying model with an inherent notion of uncertainty (e.g., conditional quantile regression) guarantees true value coverage within the uncertainty intervals even when the model’s predictions are poor. However, CI alone is not highly adaptable to heteroskedastic data, as prediction sets can appear fixed and weakly dependent on the model’s predictors. For example, conformal classification can produce rather conservative prediction sets given their sole reliance on softmax scores, as mentioned in [5]. For regression problems, conformalized quantile regression (CQR) was introduced in [22]. CQR produces adaptive and distribution-free prediction intervals with a guaranteed marginal coverage rate, such as 90%. While typical full conformal prediction assumes that samples are drawn exchangeably (i.e., conditionally, i.i.d, on the joint probability distribution function between input and labels), CQR departs from this assumption to achieve a finite sample coverage guarantee through a technique called split conformal prediction. Various forms of CI, including conformal classification and CQR, were initially developed for one-dimensional (1D) data. They have since been extended to multivariate cases, making them particularly valuable for our application to VO.

III. EXTRACTING UNIVARIATE BANDS OF UNCERTAINTIES FOR VISUAL ODOMETRY (VO)

This section introduces two methods for extracting predictive uncertainty bands in VO: univariate conformalized quantile regression and conformalized set prediction. The methods presented here are lightweight but produce rectangular uncertainty bands, which can be too conservative. The first method uses quantile regression on each coordinate of the position and orientation output vectors of VO, which are combined to produce an overall quantile region. The second method maps the regression problem in VO to a classification problem and performs a set prediction, which can produce disjoint uncertainty bands. The presented techniques are compared in Table I and discussed below:

A. VO-Uncertainty by Univariate CQR

We employ conformalized quantile regression (CQR) [22] to extract univariate uncertainty bands for each output coordinate in VO. To test our approach, we use PoseNet [12] with a ResNet34-based feature extractor for conformalization. However, the presented techniques are generalizable to other

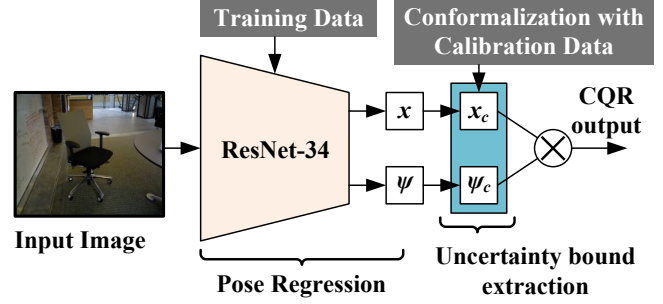


Fig. 2: VO uncertainty estimation using univariate CQR: Features of the input image and extracted and regressed to six-dimensional pose. Subsequently, each pose dimension’s predictions are conformalized using the held-out calibration data. Finally, multivariate uncertainty regions are predicted by multiplying 1D uncertainty bands.

feature-based and direct predictors for VO. The PoseNet model outputs translational elements x , y , and z , as well as rotational elements p , q , and r of angle w in radians. To reduce the number of response variables from seven to six, we convert the orientation quaternion (w, p, q, r) to Euler angles (roll ϕ , pitch θ , yaw ψ). After pre-training, we apply CQR to obtain prediction intervals along each one-dimensional variable, which reveals the predictive uncertainty. Fig. 2 depicts the process for this method.

The main goal of CQR is to produce prediction intervals where the probability of true realizations falling within these intervals is guaranteed to be equivalent to or greater than some chosen coverage rate p , i.e., $\mathbb{P}\{Y \in C(X) \mid X = x\} \geq p$. This property is known as *marginal coverage*, and is calculated over all training and test set samples. To compute $C(X)$, we use Random Forest-based quantile regression, which is less prone to over-fitting, requires fewer data and computing, and results in lower and upper bound percentiles that capture heteroskedasticity by expanding and contracting based on the underlying uncertainty in the prediction. The quantiles are then conformalized (corrected) by a calibration set that is split a priori from the training set. The following equations represent the three steps of the procedure:

$$\{Q_l, Q_h\} \leftarrow QR(\{(X_i, Y_i)\}), i \in I_{train} \quad (1)$$

$$E_i \leftarrow \max\{Q_l(X_i) - Y_i, Y_i - Q_h(X_i)\}, i \in I_{cal} \quad (2)$$

$$C_{cal}(X) = [Q_l(X) - Q_{cal}(E, I_{cal}), Q_h(X) + Q_{cal}(E, I_{cal})], \text{ where} \quad (3a)$$

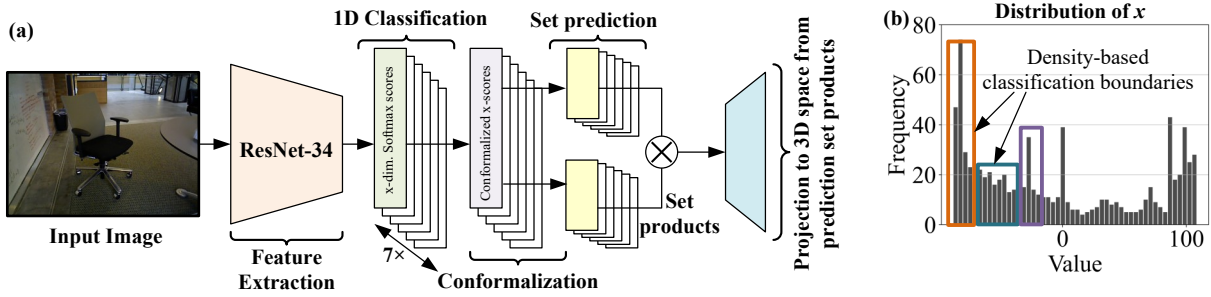


Fig. 3: **Uncertainty-aware VO by conformal set prediction (CSP)**: (a) The architecture of CSP for VO. The output of the feature extraction unit is fed to parallel 1D softmax classifiers. Subsequently, classifier outputs are conformalized for the predictive set generation along each dimension. Next, the sets are multiplied and projected to 3D space to generate uncertainty-aware rectangular predictive regions. (b) Our approach also leverages density-aware classification boundaries along each dimension. Based on the training set, frequently visited locations are packed into classes with a narrow gap between upper and lower boundaries for high precision determination when projecting to 3D space.

$$Q_{cal}(E, I_{cal}) := (1 - \alpha)(1 + 1/|I_{cal}|) - \text{th} \quad (3b)$$

quantile of $\{E_i : i \in I_{cal}\}$

Here, given a certain miscoverage rate α , the quantile regression algorithm (QR) fits two conditional quantile functions Q_l and Q_h on the training set I_{train} , which contains samples X and labels Y . The effectiveness of the initial prediction interval $[Q_l(X), Q_h(X)]$ in covering Y_i is then measured using conformal scores E_i , which are evaluated on the held out calibration set I_{cal} . If Y_i is outside the boundaries, E_i measures the distance from the nearest boundary. If Y_i falls within the desired boundaries, E_i measures the larger of the two distances, accounting for both undercoverage and overcoverage. Finally, the calibrated prediction interval C_{cal} is formed using $Q_{cal}(E, I_{cal})$, which is the empirical $(1 - \alpha)(1 + 1/|I_{cal}|)$ -th quantile of E_i , $i \in I_{cal}$. Using a calibration set ensures that the resulting intervals have the desired miscoverage rate.

B. VO-Uncertainty by Conformalized Set Prediction (CSP)

This section describes our method, conformalized set prediction (CSP), for converting each of the multiple univariate pose regression problems in VO to pose classification using one-hot encoding and generating conformalized prediction sets by discretizing the flying space. We use the same feature extractor as discussed in the previous section. However, instead of relying on quantile regression to determine the model's uncertainties, we leverage the model's normalized exponential scores (i.e., softmax scores) for conformal set prediction. A notable advantage of this approach is the ability to predict disjoint regions of uncertainty, unlike the previous one that only predicts contiguous regions of uncertainty.

The approach begins by discretizing the drone's flying space into K sets along each dimension x , y , and z . To achieve this, a non-uniform space discretization is followed, wherein the histograms of the training set trajectories along each dimension are divided into K quantiles to determine class boundaries. The result is a one-hot encoding matrix based on K , and a lightweight neural network classifier is built with a programmable output layer of size K . For each dimension of the pose, a network is trained to output reliable softmax scores that can be conformalized. Fig. 3 illustrates the softmax classifiers for each dimension sharing their feature extraction head for parameter sharing and

computational efficiency. The product of the predicted set of classes along each dimension produces the net uncertainty regions. Notably, this method can represent frequently visited locations with higher precision, represented by a class with a narrower gap between upper and lower boundaries. Conversely, less frequently visited locations are represented with lower precision.

Under the above classification treatment of VO, to perform conformal set prediction, a calibration set is first separated from the training set to compute conformal scores. It is essential to ensure the statistical guarantee of conformal prediction by ensuring that the average probability of correct classifications within the prediction sets $C(X) \subset \{1, \dots, K\}$ is almost exactly $1 - \alpha$, where α denotes an arbitrary miscoverage rate (e.g., 10%). This marginal coverage property can be expressed mathematically as $1 - \alpha \leq \mathbb{P}\{Y \in C(X) \mid X = x\} \leq 1 - \alpha + \frac{1}{n+1}$, where n is the size of the calibration set. Conformal scores are then obtained by subtracting the softmax of the correct class for each input from one. Finally, $\hat{q}$ is computed as the $\text{ceil}(\frac{(n+1)(1-\alpha)}{n})$ empirical quantile of these conformal scores, and the conformalized prediction set $C(X) \mid X = x$ is formed by incorporating classes with softmax scores greater than $(1 - \hat{q})$.

Importantly, since the predicted set along each dimension may include classes that are not proximal, the above method has the ability to produce disjoint uncertainty regions. In contrast, the previous method in Sec. III-A (and multivariate methods presented in the later discussion) can only output contiguous uncertainty regions. The class labels in the above conformal set prediction can also be determined by discretizing the entire 3D space; however, the necessary classes for matching precision to dimension-wise discretization would grow exponentially. We also found that the above conformalization procedure may sometimes be too strict, especially when the conditional softmax outputs are not representative. However, it can be improved by using the softmax outputs of all classes in gathering the conformal scores as in [5].

IV. EXTRACTING MULTIVARIATE BANDS OF PREDICTIVE UNCERTAINTIES IN VISUAL ODOMETRY (VO)

This section presents two more approaches for extracting predictive uncertainty bands in VO: multivariate conformal-

ized quantile regression with MC-dropout, and conformalized joint prediction. The first approach demonstrates the novel usage of MC-dropout as a data augmentation technique in improving the performance of a conditional variational autoencoder (CVAE). The second approach jointly trains predictive uncertainty bands and pose estimation using a new loss function, thus reducing computational resources for applications where uncertainty estimates are always required.

A. VO-Uncertainty by Multivariate CQR (MCQR)

To generate more informative VO-uncertainty bands that consider the correlation among pose dimensions, we adopt the multivariate conformalized quantile regression algorithm in [23] and enhance it with MC-dropout (MCQR w/ MC-dropout). This method entails utilizing a CVAE to acquire a proper latent representation of the sample distribution $Y | X$ and applying an extension of directional quantile regression (DQR) to create quantile regions. The architecture of the CVAE, shown in Fig. 4(a), is a variant of a variational autoencoder (VAE) that conditions the generative model on labels during both the encoding and decoding steps. This additional information improves the model's ability to produce targeted outputs. Instead of directly fitting quantile functions to the data with simple regression, a separate neural network is developed to learn the best conditions for sampling the latent representation $Z | X$ of the complete multivariate response variable based on $Y | X$. At the same time, a quantile region is generated in this dimension-reduced latent space, which has a more informative, flexible, and arbitrary shape when propagated through the decoder.

Compared to the procedure for extracting univariate bands in Sec. III-A, extracting multivariate bands of uncertainty here using a VAE predictor requires more training data due to the model's complexity and higher number of parameters. Therefore, to improve the accuracy of VAE-based uncertainty bands, we utilize MC-dropout to generate additional training data. Training VAEs with the additional data enables uncertainty distillation from the MC-dropout procedure into lightweight conformalized quantile bands. However, the underlying pose estimation model's performance heavily influences this approach's effectiveness. Poor predictions of mean estimates can hinder learning with additional data.

B. VO-Uncertainty by Conformalized Joint Prediction (CJP)

This section discusses the joint training of multivariate VO-uncertainty bands and the mean pose, i.e., conformalized joint prediction (CJP). The approach reduces the computational costs of uncertainty-aware VO by sharing parameters and network layers for mean prediction and uncertainty bands, which is particularly useful for platforms where uncertainty estimation must always be ON.

Fig. 4(b) shows the proposed VAE architecture for simultaneously extracting pose predictions and upper/lower uncertainty bounds based on true conditional quantiles. In the network, we adopt the parametric rectified linear unit (PReLU) as an activation function to introduce additional learnable parameters, better handle negative values and avoid

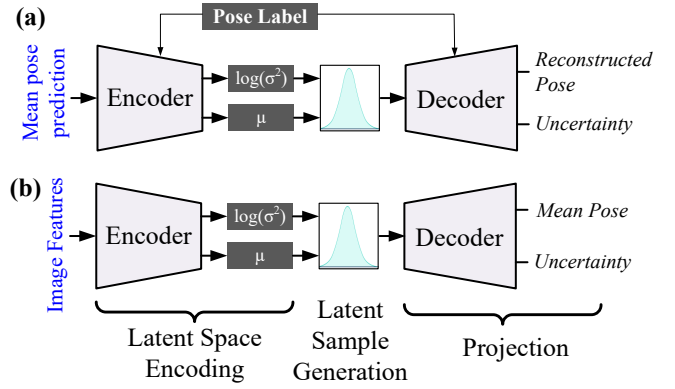


Fig. 4: **VAE architecture for joint training of point (mean) prediction and uncertainty bounds:** (a) Predictions from a pose estimation network (such as PoseNet) are fed to a conditional VAE for pose reconstruction and uncertainty estimation by first encoding to the latent space and then generatively passing through a decoder. (b) Comparatively, the network here jointly learns point prediction and uncertainty estimates directly from the features extracted from camera images. The novel loss function in Eq. (6) is used for the integrated training. The simultaneous extraction of mean and uncertainty bounds provides better computational efficiency for applications where uncertainty awareness must always be ON.

the dying ReLU problem (values become zero for any input). In Fig. 4, the network first extracts features from the input image using a pre-trained convolutional neural network (e.g., MobileNet [24]), followed by encoding to a 3D latent space with a spherical Gaussian probability distribution. Next, a sample is taken from the latent space and propagated through the decoder to generate the multivariate pose and uncertainty-bound estimations at the desired coverage rate.

The pinball loss function, a.k.a. quantile loss, used for conformalized quantile regression in [22], [23] is inadequate for such integrated training of the mean prediction and uncertainty bands. Thus, we propose a new training loss function combining the traditional mean-squared error and Kullback-Leibler (KL) divergence losses of a VAE, MSE_{loss} and KL_{loss} , with two additional loss components $INTSCORE_{loss}$ and $COMCAL_{loss}$. This total loss function, $\mathcal{L}_{Total}$, capitalizes on two unique findings reported in [25].

The first finding in the proposed loss function is the unconventional use of the interval score function (i.e., Winkler score [26], [27]) to train an uncertainty quantification model to consistently output centered prediction intervals within a specified quantile range. Notably, the interval score function is typically used to evaluate prediction intervals, not optimize them. This approach is advantageous for generating informative prediction intervals over pose trajectories. We choose two percentiles (e.g., $\alpha_l = 5\%$ and $\alpha_h = 95\%$) as uncertainty bounds, compute the interval score loss shown below for each percentile, and then take their expected value:

$$INTSCORE_{loss} = (Q_h - Q_l) + \frac{2}{\alpha} (Q_l - y) \mathbb{I}\{y < Q_l\} + \frac{2}{\alpha} (y - Q_h) \mathbb{I}\{y > Q_h\} \quad (4)$$

Here, y represents the pose label, Q_l and Q_h are each dimension's lower and upper quantile estimates, and $\mathbb{I}$ is the indicator function. Including the interval score in $\mathcal{L}_{Total}$ is instrumental in linking pose reconstruction optimization via

MSE_{loss} to quantile optimization over each dimension while honoring the dimensions' covariance and correlation.

The second finding is combined calibration loss, denoted as $COMCAL_{loss}$. $COMCAL_{loss}$ establishes an explicit and controllable balance between prediction interval calibration and sharpness rather than relying on the loosely implicit balance in basic pinball loss. $COMCAL_{loss}$ comprises two objectives: one that minimizes the difference between p_{avg}^{cov} and the chosen marginal coverage rate p (CAL_{obj}), and another that minimizes the distance between Q_l and Q_h ($SHARP_{obj}$). p_{avg}^{cov} is the estimated average probability that the pose label values lie within $[Q_l, Q_h]$. The two objectives may conflict with each other since one aims to minimize over-coverage and under-coverage, while the other aims to minimize the length of the prediction interval. This necessitates including a hyper-parameter λ to strike the appropriate balance. Thus, $COMCAL_{loss}$ is defined as follows:

$$CAL_{obj} =$$

$$\begin{aligned} & \mathbb{I}\{p_{avg}^{cov} < p\} \times \frac{1}{N} \sum_{i=1}^N [(y_i - Q_{l,h}(x_i)) \mathbb{I}\{y_i > Q_{l,h}\}] + \\ & \mathbb{I}\{p_{avg}^{cov} > p\} \times \frac{1}{N} \sum_{i=1}^N [(Q_{l,h}(x_i) - y_i) \mathbb{I}\{y_i < Q_{l,h}\}] \end{aligned} \quad (5a)$$

$$SHARP_{obj} = \frac{1}{N} \sum_{i=1}^N \begin{cases} Q_l(x_i) - Q_h(x_i), & p \leq 0.5 \\ Q_h(x_i) - Q_l(x_i), & p > 0.5 \end{cases} \quad (5b)$$

$$COMCAL_{loss} = (1 - \lambda) \times CAL_{obj} + \lambda \times SHARP_{obj} \quad (5c)$$

where x_i and y_i are individual images and labels of a training batch and $Q_{l,h}$ is meant to denote that CAL_{obj} is computed separately for both Q_l and Q_h and then summed.

The complete loss function combining MSE loss, KL divergence, interval score, and calibration loss in our approach is thus given as:

$$\begin{aligned} \mathcal{L}_{Total} = & MSE_{loss}(y, \hat{y}) + KL_{loss}(\mu, \log(\sigma^2)) \\ & + INTSCORE_{loss}(y, Q_l, Q_h, \{\alpha_l, \alpha_h\}) \\ & + COMCAL_{loss}(y, p_{avg}^{cov}, Q_l, Q_h) \end{aligned} \quad (6)$$

where $\hat{y}$ is the reconstructed pose. μ and $\log(\sigma^2)$ are the hyper-parameters of the VAE's latent space. We estimate p_{avg}^{cov} for each pose dimension with every training batch to enable flexible calibration during training.

Traditional CQR involves splitting the training set into two subsets and conformalizing the quantiles afterward (i.e., split conformal prediction). This is done to avoid full conformal prediction, which requires many more calibration steps than just one, sacrificing speed for statistical efficiency. However, cross-conformal prediction provides a viable alternative for conformalization by striking a reasonable balance between computational and statistical efficiency, utilizing multiple calibration steps across the entire training set [28]. Our approach is similar, representing a unique form of cross-conformal prediction. We assume that p_{avg}^{cov} , which is averaged over randomized training batches, is sufficiently close to an average over the entire training set.

As will be demonstrated later in Sec. V, this approach yields superior outcomes compared to the other methods with a small increase in model complexity and size. The results underscore the method's accuracy, adaptability, and flexibility in producing conformalized prediction intervals for multivariate pose data. The method maintains its strength even when swapping the underlying model from a larger, more accurate one, such as ResNet34, to one better suited for lightweight edge applications like MobileNetV2. Furthermore, its advantages are consistent across multiple datasets, implying its robustness to numerous applications.

V. RESULTS AND DISCUSSIONS

This section presents a chronological review of the results from the four frameworks proposed for conformalized visual odometry (CVO) and compared in Table I. The aim of these methods is to develop reliable prediction intervals that can capture true uncertainty in data-driven pose estimation with guaranteed confidence. While each approach is based on PoseNet [12], they can be applied to other pose estimation models. In highlighting the differences between the methods, we use the data from scene 2 of RGB-D Scenes Dataset v.2 [29] as visuals of its pose trajectory are easy to comprehend. All results illustrate uncertainty bands that exhibit at least 90% guaranteed marginal coverage through CI.

In Sec. III-A, we introduced univariate conformalized quantile regression (CQR) to generate uncertainty bands. However, as shown in Fig. 5, the resulting prediction interval is overly cautious and box-shaped despite its ability to adjust to the data's heteroscedasticity. This is because the method employs quantile regression on the distinct variables of the multivariate pose response and then merges them without considering their covariance and correlation. As a result, this method lacks sufficient information and is not particularly useful in most cases, despite its low computational complexity. This drawback is consistent across all datasets.

In Sec. III-B, we presented conformalized set prediction (CSP) to generate uncertainty bands. As illustrated in Fig. 5, this approach adapts well to the data, even with non-continuous uncertainty boundaries. The method's adaptiveness is deeply dependent on the model's softmax scores and the proper discretization of pose values, which can be adjusted by changing the number of classes K . Like the first method, this approach has difficulty handling multivariate data, which limits its ability to estimate true uncertainty accurately. Although the average interval length is significantly shorter, the conservativeness of the rectangular uncertainty region has only scaled linearly, limiting its usefulness.

In Sec. IV-A, we introduced the third method, which utilizes multivariate conformalized quantile regression with MC-dropout (MCQR w/ MC-dropout) to generate uncertainty bands. This approach showcases the novel use of MC-dropout as a data augmentation technique, in addition to regularization and uncertainty estimation, to enhance learning. This method is the first step towards addressing the limitations of the previous methods by considering the relationships between pose dimensions. As shown in Fig. 5,

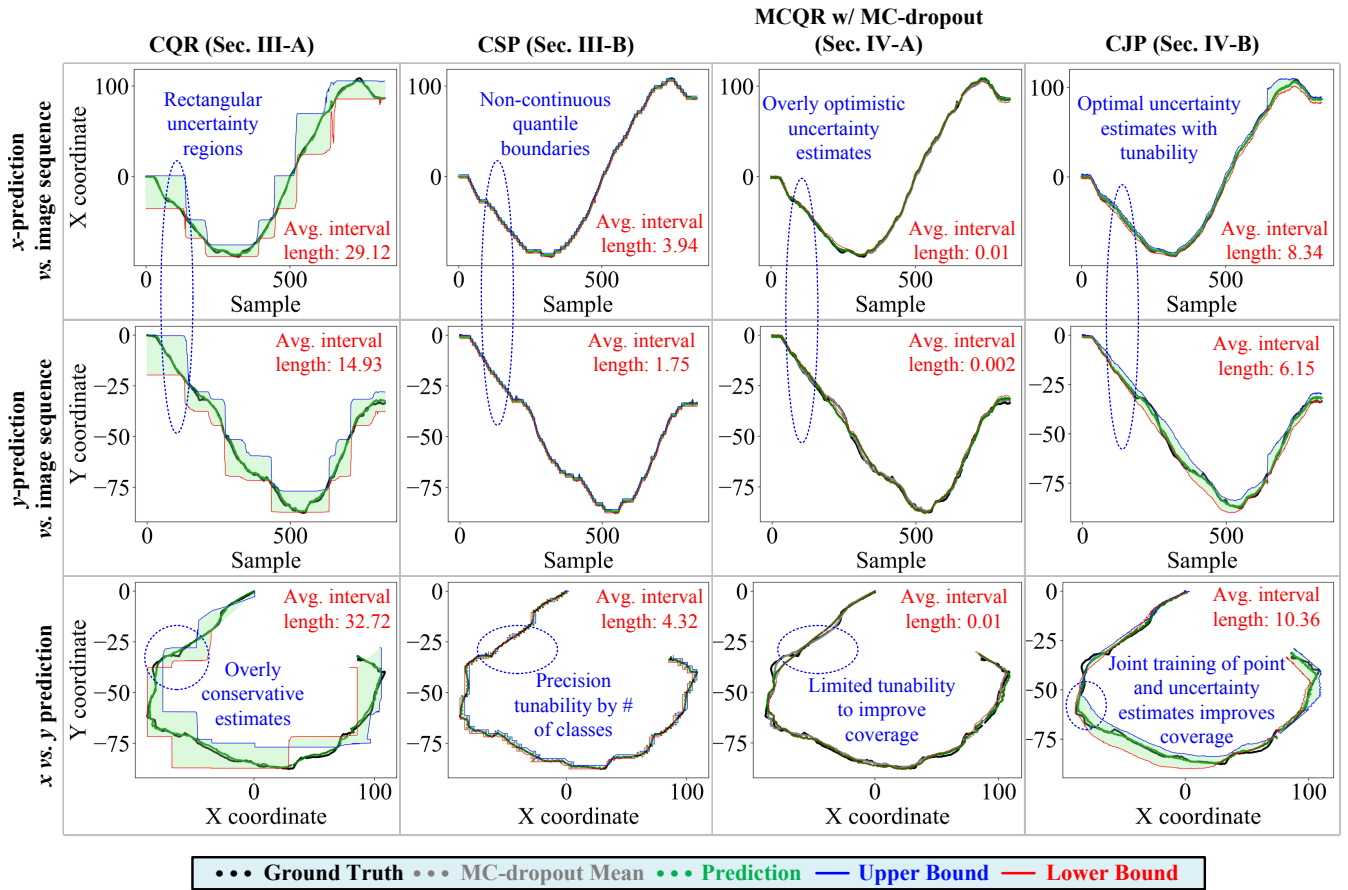


Fig. 5: Comparison of uncertainty-aware prediction intervals over 1D and 2D pose estimations from the four proposed conformalized visual odometry (CVO) methods in Secs. III and IV: Differences in average interval length, contiguity, precision, accuracy, coverage, and adaptiveness are highlighted across the methods using Scene 2 of RGB-D Scenes Dataset v.2.

this method has the tightest uncertainty estimates, but they represent the opposite extreme of the conservatism observed in the univariate methods. However, this is inconsistent across all datasets; the method performs much better in some other datasets, indicating that it is highly data-dependent. This inconsistency in uncertainty accuracy is particularly noticeable in smaller datasets, as is the case here. Furthermore, this method has the highest computational complexity, requiring training and optimizing multiple networks. Although it can potentially produce flexible and highly accurate uncertainty intervals, its inconsistency across datasets and computational inefficiency limit its wide applicability.

The final method, described in Sec. IV-B, demonstrates a novel combination of multivariate conformalized quantile regression and the mean prediction to produce uncertainty bounds using a new loss function that remarkably achieves multiple objectives simultaneously. We deem this approach conformalized joint prediction (CJP). Fig. 5 demonstrates its significant improvements over the other methods. Firstly, the prediction interval optimally expands and contracts in response to the data, being neither too dispersed nor sharp in each plot. The increase in interval length (and hence uncertainty) at the middle and end points in the trajectory corresponds to worse lighting conditions in the scene's image sequence. Secondly, with the tunable calibration parameter

$\lambda \in \{0, 1\}$, we can optimize the network for sharpness and/or coverage. Here, λ was set to 0.5 to showcase the balanced standard. Additionally, Fig. 6 shows results of this method when switching the feature extraction model from ResNet34 to MobileNetV2, effectively reducing its learning capacity by over 80%. It handles the same image sequence with comparable accuracy even when introducing different forms of noise, such as gaussian, salt and pepper, and speckle. Notably, the uncertainty increases with this transition and even more with added noise. The speckle noise exasperates light variation the most and thus has the largest average interval length. An increase in the average uncertainty intervals indicates that the framework can capture the sensor's uncertainties along with the model's predictive uncertainties. The optimality of this method in extracting true pose estimation uncertainty is consistent across several datasets and conditions, indicating that it can fit many applications.

VI. CONCLUSION

This paper introduced and compared four novel frameworks for capturing aleatoric and epistemic uncertainty in data-driven VO through conformal inference while considering computational resource limitations. The presented frameworks trade off essential metrics of uncertainty-aware inference, such as computational workload, adaptiveness of

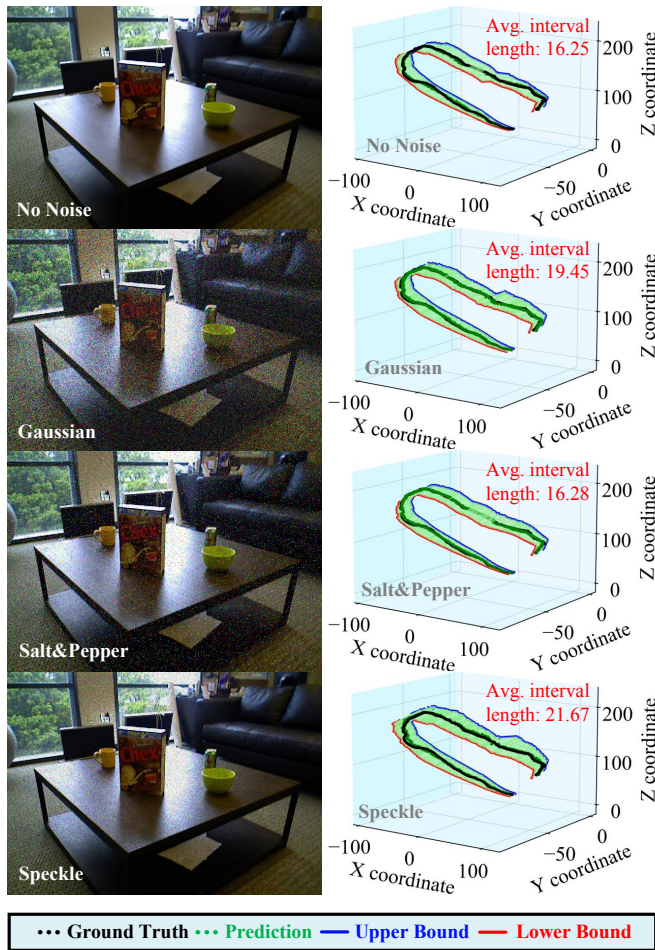


Fig. 6: Comparison of uncertainty-aware prediction intervals over 3D pose estimations from CJP under different noise conditions: After switching from the ResNet34 feature extractor to MobileNetV2, gaussian, salt and pepper, and speckle noise were applied to Scene 2 of RGB-D Scenes Dataset v.2 to further stress test the method and demonstrate its resilience. These conditions are meant to simulate real-world disturbance.

uncertainty interval, training time, *etc.*, aiming to provide a spectrum of solutions against widely varying robustness and resource constraints on edge robotics. We discussed architectures of generative deep neural networks for estimating multivariate uncertainty bands along with point (mean) prediction, data augmentation techniques to improve the accuracy of uncertainty estimation, and a novel training loss function that combines interval scoring and calibration loss with traditional training metrics such as KL-divergence to improve uncertainty-aware learning.

REFERENCES

- [1] D. M. Blei, A. Kucukelbir, and J. D. McAuliffe, "Variational Inference: A Review for Statisticians," *Journal of the American Statistical Association*, vol. 112, no. 518, pp. 859–877, 2017.
- [2] V. Vovk, A. Gammerman, and G. Shafer, *Algorithmic Learning in a Random World*. Berlin, Heidelberg: Springer-Verlag, 2005.
- [3] J. Lei, M. G'Sell, A. Rinaldo, R. J. Tibshirani, and L. Wasserman, "Distribution-Free Predictive Inference for Regression," 2018.
- [4] G. Shafer and V. Vovk, "A Tutorial on Conformal Prediction," *J. Mach. Learn. Res.*, vol. 9, p. 371–421, jun 2008.
- [5] A. N. Angelopoulos and S. Bates, "A Gentle Introduction to Conformal Prediction and Distribution-Free Uncertainty Quantification," 2021. [Online]. Available: <https://arxiv.org/abs/2107.07511>
- [6] B.-S. Einbinder, Y. Romano, M. Sesia, and Y. Zhou, "Training uncertainty-aware classifiers with conformalized deep learning," in *NeurIPS*, 2022.
- [7] S. Messoudi, S. Rousseau, and S. Destercke, "Deep conformal prediction for robust models," in *Information Processing and Management of Uncertainty in Knowledge-Based Systems*. Cham: Springer International Publishing, 2020, pp. 528–540.
- [8] D. Scaramuzza and F. Fraundorfer, "Visual odometry [tutorial]," *IEEE robotics & automation magazine*, vol. 18, no. 4, pp. 80–92, 2011.
- [9] H. Alismail, M. Kaess, B. Browning, and S. Lucey, "Direct visual odometry in low light using binary descriptors," *IEEE Robotics and Automation Letters*, vol. 2, no. 2, pp. 444–451, 2016.
- [10] N. Krombach, D. Droschel, and S. Behnke, "Combining feature-based and direct methods for semi-dense real-time stereo visual odometry," in *Intelligent Autonomous Systems 14: Proceedings of the 14th International Conference IAS-14 14*. Springer, 2017.
- [11] J. Gräter, T. Schwarze, and M. Lauer, "Robust scale estimation for monocular visual odometry using structure from motion and vanishing points," in *2015 IEEE Intelligent Vehicles Symposium (IV)*. IEEE, 2015, pp. 475–480.
- [12] A. Kendall, M. Grimes, and R. Cipolla, "PoseNet: A Convolutional Network for Real-Time 6-DOF Camera Relocalization," in *2015 IEEE International Conference on Computer Vision (ICCV)*, 2015.
- [13] F. Walch, C. Hazirbas, L. Leal-Taixé, T. Sattler, S. Hilsenbeck, and D. Cremers, "Image-Based Localization Using LSTMs for Structured Feature Correlation," in *2017 IEEE International Conference on Computer Vision (ICCV)*, 2017, pp. 627–637.
- [14] S. Wang, R. Clark, H. Wen, and N. Trigoni, "DeepVO: Towards end-to-end visual odometry with deep Recurrent Convolutional Neural Networks," in *2017 IEEE International Conference on Robotics and Automation (ICRA)*, 2017, pp. 2043–2050.
- [15] R. Li, S. Wang, Z. Long, and D. Gu, "UnDeepVO: Monocular Visual Odometry Through Unsupervised Deep Learning," in *2018 IEEE International Conference on Robotics and Automation (ICRA)*, 2018.
- [16] P. Shukla, A. C. Stutts, S. Ravi, T. Tulabandhula, A. R. Trivedi *et al.*, "Robust Monocular Localization of Drones by Adapting Domain Maps to Depth Prediction Inaccuracies," *arXiv preprint arXiv:2210.15559*.
- [17] Y. Gal and Z. Ghahramani, "Dropout as a bayesian approximation: Representing model uncertainty in deep learning," in *International Conference on Machine Learning*. PMLR, 2016, pp. 1050–1059.
- [18] G. Costante and M. Mancini, "Uncertainty Estimation for Data-Driven Visual Odometry," *IEEE Transactions on Robotics*, 2020.
- [19] N. Yang, L. von Stumberg, R. Wang, and D. Cremers, "D3VO: Deep Depth, Deep Pose and Deep Uncertainty for Monocular Visual Odometry," in *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020, pp. 1278–1289.
- [20] N. Kaygusuz, O. Mendez, and R. Bowden, "MDN-VO: Estimating Visual Odometry with Confidence," in *2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2021.
- [21] M. Lee, B. Mudassar, and S. Mukhopadhyay, "Lightweight Model Uncertainty Estimation for Deep Neural Object Detection," in *2022 International Joint Conference on Neural Networks (IJCNN)*, 2022.
- [22] Y. Romano, E. Patterson, and E. J. Candès, *Conformalized Quantile Regression*. Red Hook, NY, USA: Curran Associates Inc., 2019.
- [23] S. Feldman, S. Bates, and Y. Romano, "Calibrated multiple-output quantile regression with representation learning," *Journal of Machine Learning Research*, vol. 24, no. 24, pp. 1–48, 2023.
- [24] A. G. Howard, M. Zhu, B. Chen, D. Kalenichenko, W. Wang, T. Weyand, M. Andreetto, and H. Adam, "MobileNets: Efficient Convolutional Neural Networks for Mobile Vision Applications," 2017. [Online]. Available: <https://arxiv.org/abs/1704.04861>
- [25] Y. Chung, W. Neiswanger, I. Char, and J. Schneider, "Beyond pin-ball loss: Quantile methods for calibrated uncertainty quantification," *Advances in Neural Information Processing Systems*, vol. 34, 2021.
- [26] T. Gneiting and A. Raftery, "Strictly Proper Scoring Rules, Prediction, and Estimation," *Journal of the American Statistical Association*, vol. 102, pp. 359–378, 03 2007.
- [27] R. L. Winkler, "A Decision-Theoretic Approach to Interval Estimation," *Journal of the American Statistical Association*, 1972.
- [28] V. Vovk, "Cross-conformal predictors," *Annals of Mathematics and Artificial Intelligence*, vol. 74, 08 2012.
- [29] K. Lai, L. Bo, and D. Fox, "Unsupervised feature learning for 3D scene labeling," in *2014 IEEE International Conference on Robotics and Automation (ICRA)*, May 2014, pp. 3050–3057.