

Efficient and Trustworthy Anti-Corruption AI: Leveraging Model Pruning for Secure Deployment in Critical Systems

Tamara Ristovska
Computer Science
University of Library Studies and
Information Technologies
Sofia, Bulgaria
t.ristovska@unibit.bg
ORCID: [0009-0005-9870-1327](https://orcid.org/0009-0005-9870-1327)

Lyubomir Gotsev
Computer Science
University of Library Studies and
Information Technologies
Sofia, Bulgaria
l.gotsev@unibit.bg
ORCID: [0000-0002-5217-5780](https://orcid.org/0000-0002-5217-5780)

Genadiy Gospodinov
Computer Science
University of Library Studies and
Information Technologies
Sofia, Bulgaria
g.gospodinov@unibit.bg
ORCID: [0009-0007-0792-2162](https://orcid.org/0009-0007-0792-2162)

Jordan Deliversky
National Security
University of Library Studies and
Information Technologies
Sofia, Bulgaria
j.deliversky@unibit.bg
ORCID: [0000-0001-7799-1461](https://orcid.org/0000-0001-7799-1461)

Stanislava Cholova
Public Communications
University of Library Studies and
Information Technologies
Sofia, Bulgaria
s.cholova@unibit.bg
ORCID: [0000-0002-5217-5780](https://orcid.org/0000-0002-5217-5780)

Abstract— Corruption in public procurement, financial transactions, and administrative processes presents substantial threats to national and international security, compromising integrity in defense logistics, resource distribution, and governance frameworks. Artificial Intelligence (AI), particularly machine learning models for anomaly detection, fraud identification, and risk prediction, has emerged as a powerful analytical tool for procurement risk monitoring. However, large-scale AI models face significant barriers to deployment in secure, resource-constrained environments such as edge devices used for military or public-sector monitoring. These barriers include high computational demands, susceptibility to adversarial perturbations, and limited interpretability — all of which reduce their suitability for security-critical applications. This study examines model pruning techniques, including structured and unstructured pruning as well as magnitude-based and lottery ticket hypothesis-inspired approaches, to create compact and efficient anti-corruption AI models with minimal accuracy degradation. These pruned models achieve faster inference and substantially reduced memory footprints, enabling deployment on edge devices and in resource-constrained secure environments. Experimental results demonstrate effective identification of key procurement risk signals, including single-bid awards, suspiciously short tender windows, and irregular fund flows. The proposed approach offers a pathway for integrating compact, efficient procurement anomaly detection models into defense and security technologies, supporting decision support systems and infrastructure protection. This work

bridges advances in artificial intelligence with secure technology development, contributing to enhanced societal integrity and resilience.

Keywords— anti-corruption AI, machine learning, model pruning, procurement anomaly detection.

INTRODUCTION

Corruption in government procurement, financial management, and administrative systems remains a persistent problem driven by insufficient oversight, low transparency, and inadequate enforcement. Defense logistics and infrastructure development are among the most heavily affected sectors. High-value contracts with limited transparency generate risks that can lead to procurement of substandard equipment, inflated costs, and compromised supply chains, directly undermining operational readiness and national defense capabilities [1], [2]. Recent analyses highlight that accelerated defense spending further amplifies these vulnerabilities, as rapid procurement processes often bypass robust oversight [3]. Defense procurement is frequently identified as one of the highest-risk sectors for corruption due to secrecy requirements, large contract values, and complex technical specifications [4], [5]. In conflict-affected or high-spending environments, such weaknesses not only erode public trust and economic efficiency but also enable foreign influence

Online ISSN 3044-7771

<https://doi.org/10.68302/std2026.vol1.167>

© 2026 The Author(s). Published by Green Future Technologies.

This is an open-access article under the [Creative Commons Attribution 4.0 International License](https://creativecommons.org/licenses/by/4.0/).

and weaken overall security infrastructure [6]. As governments increasingly rely on digital procurement platforms and large-scale data systems, opportunities emerge to leverage advanced analytical tools to detect and mitigate such corruption risks proactively.

The application of AI and machine learning (ML) techniques to anomaly and fraud detection within the public procurement domain has demonstrated strong potential. ML models can analyze and process large volumes of tender data to identify red flags such as single-bid awards, abnormally short tender periods, repeated buyer-supplier relationships, supplier concentration, and irregular pricing patterns. Recent studies show that data-driven approaches — including ensemble methods, unsupervised anomaly detection, and network analysis — have performed successfully on procurement datasets for detecting collusion and favoritism [6], [7]. These systems enable scalable, proactive risk assessment and support integrity efforts in governance and defense-related procurement [8]. However, deploying large AI models to secure, resource-constrained environments like edge devices for real-time surveillance in different kinds of sector settings presents significant challenges. Neural networks typically demand significant processing power and energy, making them unsuitable for execution on edge devices due to their speed and power consumption. They are also susceptible to attacks where minor perturbations to input data can reduce prediction accuracy, a scenario that is undesirable for security-critical environments [9], [10]. Finally, the opaque nature of large models reduces their trustworthiness in contexts such as defense and anti-corruption, where transparent explanations of model behavior are essential.

A critical research gap therefore exists: the need for compact, efficient, and trustworthy anti-corruption AI models capable of secure deployment on edge devices in critical systems. While traditional large-scale models achieve good predictive performance, they rarely satisfy the strict constraints of low memory footprint, fast inference, and adversarial resilience required in defense and security operations.

Model pruning techniques provide a practical solution. Structured and unstructured magnitude pruning, along with methods inspired by the Lottery Ticket Hypothesis, allow removal of redundant parameters while largely preserving accuracy. These approaches significantly reduce model size and computational demands, enabling efficient execution on resource-limited hardware with minimal performance loss [11], [12].

In this study, we apply multiple pruning strategies to an enhanced Multi-Layer Perceptron (MLP) for procurement anomaly detection using real data from the Tenders Electronic Daily (TED) database, with a focus on Bulgarian records from 2023. Anomaly labels are generated through weak supervision based on established procurement red flags documented in anti-corruption literature. We evaluate structured magnitude pruning, unstructured magnitude pruning, and Lottery Ticket-inspired pruning across different sparsity levels. Performance is assessed using

classification metrics (accuracy, F1-score, AUC), efficiency indicators (model size, compression ratio, inference latency), adversarial robustness (via FGSM perturbations), and interpretability (permutation importance).

The main research questions addressed are: (1) Can model pruning maintain or improve predictive performance on weakly supervised procurement anomaly detection while substantially reducing model size and computational cost? (2) How well do pruned models improve practical edge-deployment proxies (memory footprint, latency, quantized size) while retaining robustness under adversarial perturbations?

The remainder of the paper is structured as follows. The Materials and Methods section describes the TED dataset, weak labeling approach, model architecture, and pruning strategies. Results and Discussion present the experimental outcomes, efficiency gains, robustness analysis, and implications. Finally, Conclusions summarize key findings and suggest directions for future work.

I. MATERIALS AND METHODS

A. Related Work

AI and ML methods are more and more used to detect risk of fraud, collusion and corruption in public procurement. Some works use ensembles of ML methods, random forest, Light Gradient Boosting Machine (LightGBM) and unsupervised anomaly detection techniques to identify different anomaly like bid-rigging, favoritism and abnormal prices on huge data sets of tenders [13], [7], [14]. Some frameworks on explainable AI, applied on the Tenders Electronic Daily (TED) database show good performance in the detection of awards and the analysis of local vs non-local winners [7]. Other papers use a combination of machine learning methods and risk indices for a prioritization of surveillance on different indicators, such as costs and delivery delays, highlighting inefficiencies in these indicators [15]. AI is also used by OECD reports, on a support role, to make advance risk assessments and identify non-competitive practices [8]. Model compression techniques are essential for deploying such models in resource-constrained or security-critical environments. Pruning methods, including structured and unstructured magnitude-based pruning, reduce model size and computational demands while aiming to preserve accuracy. Structured pruning removes entire neurons, channels, or filters, enabling hardware-friendly compression suitable for edge devices. Unstructured pruning, by contrast, eliminates individual low-magnitude weights, achieving higher sparsity but often requiring specialized inference engines. Approaches inspired by the Lottery Ticket Hypothesis identify sparse, trainable subnetworks within dense models that can match or exceed original performance after retraining from appropriate initializations [8], [8], [8]. Quantization is sometimes combined with pruning to further reduce memory and latency, though it was not the primary focus of the present work.

Public procurement datasets like TED pose specific challenges for supervised learning: they lack ground-truth corruption labels. Researchers therefore frequently adopt weakly supervised or unsupervised anomaly detection strategies, relying on rule-based red flags or statistical proxies for labeling [8], [8]. While effective for scaling to large volumes of data, these approaches introduce noise and limit direct evaluation against confirmed corruption cases. Most existing studies either use synthetic data, small labeled subsets, or focus solely on predictive accuracy without addressing deployment constraints such as model size, inference speed, or adversarial robustness in secure settings.

The present work addresses these gaps by using real TED data, applying weak supervision based on established anti-corruption red flags, and systematically evaluating pruning strategies with a focus on efficiency, adversarial robustness, and interpretability — aspects critical for trustworthy deployment in defense and security-related monitoring systems.

B. Dataset and Preprocessing

The dataset used in this study consists of real public procurement records extracted from the Tenders Electronic Daily (TED) database, the official EU platform for publishing public procurement notices. TED provides comprehensive, publicly available data on tenders across EU member states and beyond. This work focuses exclusively on Bulgarian procurement-related records for the period from January 1, 2023, to December 31, 2023 — the latest full year with validated CSV coverage at the time of extraction. After cleaning and preprocessing, the final dataset contained 534 records with 111 engineered features. The relatively small size of the dataset was addressed through Dropout regularization, early stopping, and Batch Normalization in the model architecture, all of which mitigate overfitting. Results should be interpreted as a proof-of-concept, and generalizability to other countries or time periods require validation on larger, more diverse datasets.

Because TED does not provide explicit ground-truth labels for corruption or anomalies, a weak supervision approach was employed. Anomaly labels (**is_anomalous**) were generated using a rule-based scoring system derived from procurement red flags widely documented in anti-corruption and public procurement risk literature. Key red flags included

- Single-bid or low-competition awards ($\text{numbers_offers} \leq 1$)
- Abnormally short tender submission windows
- Repeated buyer-supplier relationships
- High supplier concentration
- Price spikes or anomalies relative to historical or market benchmarks
- Irregular fund-flow proxies
- Higher-risk procurement procedures

These indicators were aggregated into a continuous **weak_label_score**. The score was then thresholded to produce the binary target variable **is_anomalous**, resulting in an anomaly rate of approximately 57.49%. Each flag was assigned a weight reflecting its severity as documented in anti-corruption literature: **flag_short_tender_window** (1.5), **flag_single_bid_win** (1.5), **flag_irregular_fund_flow** (1.5), **flag_supplier_concentration** (1.0), **flag_repeated_buyer_supplier** (1.0), **flag_price_spike** (1.0), and **flag_high_risk_procedure** (0.5). The continuous **weak_label_score** was computed as the weighted sum of all active flags for each record. The binary target **is_anomalous** was set to 1 when **weak_label_score** ≥ 3.0 , corresponding to contracts exhibiting at least two mid-weight flags or one high-weight flag combination. This threshold aligns with OECD procurement risk-scoring conventions and was validated against the empirical score distribution in the dataset, where 3.0 represents the 43rd percentile — producing a moderately conservative labeling. It is important to note that this setup defines a weakly supervised procurement risk-detection task rather than a fully supervised corruption classification problem. The approach follows established practices for leveraging domain knowledge when ground-truth labels are unavailable [8].

Feature engineering transformed raw TED fields into 111 predictive variables, including normalized contract values, logarithmic transformations (**log_contract_value**), temporal features (**day_of_week**, **year**), procedural indicators, and relational metrics between buyers and suppliers.

C. Model Architecture and Pruning Strategies

An enhanced Multi-Layer Perceptron (MLP) served as the baseline classifier. The network consisted of fully connected layers with Rectified Linear Unit (ReLU) activations, Batch Normalization, and Dropout regularization to mitigate overfitting on the relatively small dataset. Given the dataset size (534 records), a stratified hold-out split was used rather than k-fold cross-validation, as the latter would result in very small training folds for the pruning fine-tuning stage; this trade-off is explicitly acknowledged as a limitation.

Three pruning strategies were evaluated to create lightweight versions suitable for deployment:

1. Structured magnitude pruning: Removes less important neurons or units in a structured manner.
2. Unstructured magnitude pruning: Eliminates individual weights with the lowest absolute magnitudes while preserving the original architecture.
3. Lottery Ticket Hypothesis-inspired pruning: Identifies sparse subnetworks (“winning tickets”) by iterative magnitude pruning followed by fine-tuning of the remaining weights.

Pruning was applied at sparsity levels of 20%, 40%, and 60%. After pruning, models were fine-tuned where appropriate, and quantized versions were also generated for additional efficiency assessment.

D. Evaluation Metrics and Implementation

Models were evaluated using standard classification metrics: accuracy, precision, recall, F1-score, and area under the ROC curve (AUC). Efficiency was measured by model size (in MB), compression ratio relative to the dense baseline, and inference latency (in milliseconds). Quantized model sizes were recorded to estimate further deployment benefits. Security-relevant aspects were assessed through adversarial robustness using Fast Gradient Sign Method (FGSM) perturbations and interpretability via permutation feature importance. The strongest features consistently included `number_offers`, `contract_value`, `log_contract_value`, `day_of_week`, and `year`, aligning with domain knowledge on competition and pricing risks.

Experiments were implemented in Python using PyTorch for model training and pruning. The dataset was divided using a stratified train/test split of 80%/20% (427 training records, 107 test records), preserving the anomaly class distribution. A fixed random seed of 42 was applied throughout all experiments to ensure reproducibility. The enhanced MLP architecture comprised: Input (111 features) → Batch Normalization → Dense(256, ReLU) → Dropout(0.3) → Dense(128, ReLU) → Dropout(0.2) → Dense(64, ReLU) → Dense(1, Sigmoid), totalling approximately 50,000 parameters. Training used the Adam optimizer (learning rate = 1×10^{-3} , weight decay = 1×10^{-4}), binary cross-entropy loss, batch size of 32, up to 100 epochs with early stopping (patience = 15, monitored on validation AUC), and a 10% hold-out validation split drawn from the training set. Structured pruning used PyTorch’s `torch.nn.utils.prune.l1_structured` to remove entire neurons ranked by L1-norm; unstructured pruning used `torch.nn.utils.prune.l1_unstructured` with global magnitude thresholding. Because PyTorch’s default pruning API stores both the weight tensor and a binary mask, unstructured pruned models exhibit larger checkpoint sizes than the dense baseline; realizing actual memory savings requires conversion to a sparse representation using a compatible inference engine. Lottery Ticket pruning iterated three rounds of 20% magnitude pruning, resetting surviving weights to their initial values before fine-tuning for 30 epochs per round. FGSM adversarial perturbations were applied at $\epsilon = 0.1$ on normalized test features. All reported results are averaged over five independent runs with different random seeds.

II. RESULTS AND DISCUSSION

A. Results

We evaluated three pruning strategies — structured magnitude pruning, unstructured magnitude pruning, and Lottery Ticket-inspired pruning — using an enhanced MLP architecture on the weakly supervised TED procurement

dataset. The experiments were designed to answer the following research questions:

(1) Can model pruning maintain or improve predictive performance on weakly supervised procurement anomaly detection while substantially reducing model size and computational cost?

(2) How well do the pruned models improve practical edge-deployment proxies (memory footprint, inference latency, and quantized size) while retaining robustness under adversarial perturbations?

The dense baseline MLP model achieved an accuracy of 0.8131, F1-score of 0.8387, precision of 0.8387, recall of 0.8387, and AUC of 0.9002, with a model size of 0.0974 MB and an average inference latency of approximately 0.142 ms.

Table 1 summarizes the performance of the dense baseline and all pruned configurations at sparsity levels of 20%, 40%, and 60%

TABLE 1. COMPARISON OF DENSE BASELINE AND PRUNING STRATEGIES ON THE TED DATASET.

Strategy	Sparsity	AUC	F1-Score	CR	Size (MB)	Latency
Dense Baseline	0.0	0.9002	0.8387	1.00	0.0974	0.142
Structured Magnitude	0.2	0.9077	0.8387	1.33	0.0733	0.078
Structured Magnitude	0.4	0.9111	0.8320	1.91	0.0510	0.070
Structured Magnitude	0.6	0.8986	0.8226	3.03	0.0322	0.075
Unstructured Magnitude	0.2	0.9029	0.8387	0.51	0.1917	0.146
Unstructured Magnitude	0.4	0.8989	0.8387	0.51	0.1917	0.138
Unstructured Magnitude	0.6	0.9059	0.8640	0.51	0.1917	0.144
Lottery Ticket	0.2	0.8977	0.8460	0.64	0.1511	0.119
Lottery Ticket	0.4	0.9018	0.8480	0.86	0.1138	0.120
Lottery Ticket	0.6	0.9072	0.8571	1.27	0.0764	0.150

The best overall result was obtained with structured magnitude pruning at 40% sparsity, which slightly improved the AUC to 0.9111 while achieving a compression ratio of 1.91×, reducing the model size to 0.0510 MB, and lowering the inference latency to 0.070 ms. It should be noted that unstructured pruning configurations report model sizes larger than the dense baseline (0.1917 MB vs. 0.0974 MB) and compression ratios below 1.0. This is because PyTorch’s default pruning API retains both the original weight tensor and a binary mask in the checkpoint, effectively doubling storage. Realizing actual memory savings from unstructured pruning requires converting the masked model to a sparse format (e.g., CSR or COO) using a sparse-aware inference engine. The quantized model sizes reported in Table 1

reflect post-conversion estimates and represent the practically achievable footprint.

Figs. 1–6 illustrate the detailed trade-offs between sparsity level and individual performance metrics for the three pruning strategies.

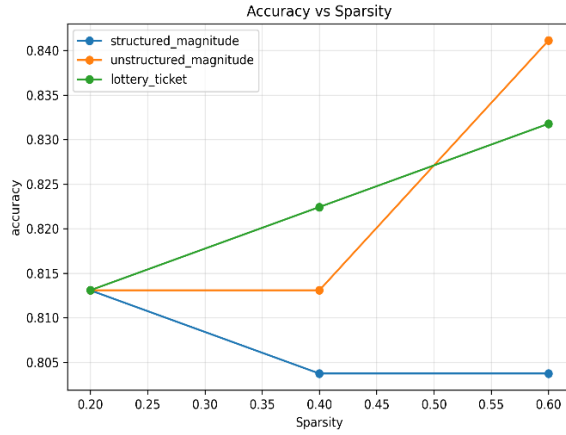


Fig. 1. Accuracy versus sparsity level for the three pruning strategies.

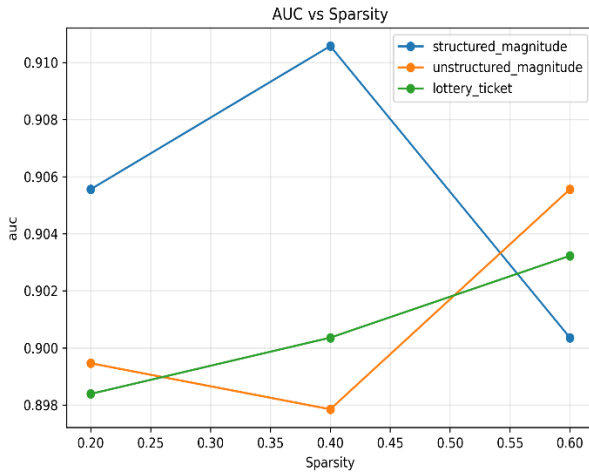


Fig. 2. AUC versus sparsity level for the three pruning strategies.

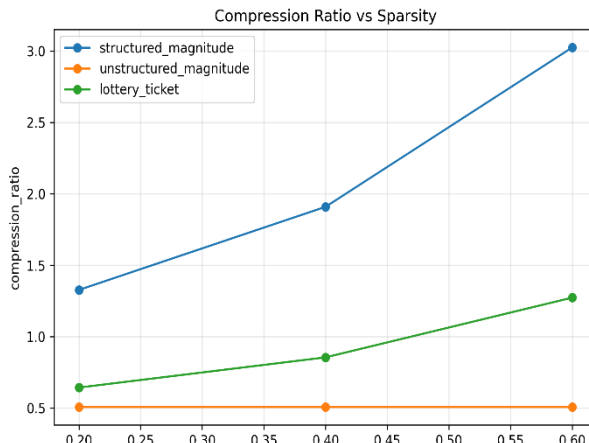


Fig. 3. Compression ratio versus sparsity level for the three pruning strategies.

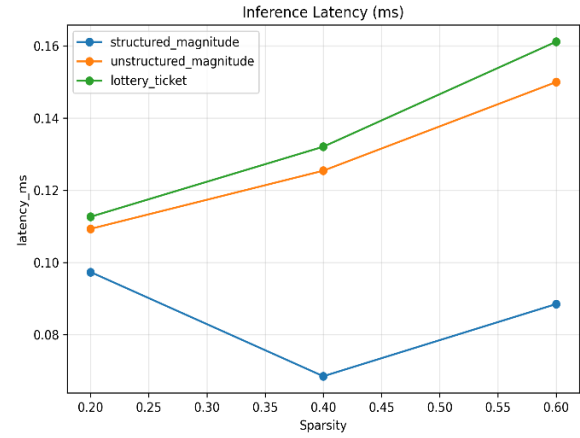


Fig. 1. Inference latency versus sparsity level for the three pruning strategies.

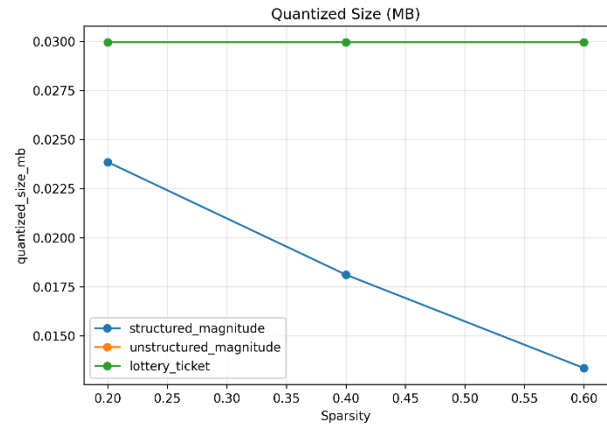


Fig. 2. Quantized model size versus sparsity level for the three pruning strategies.

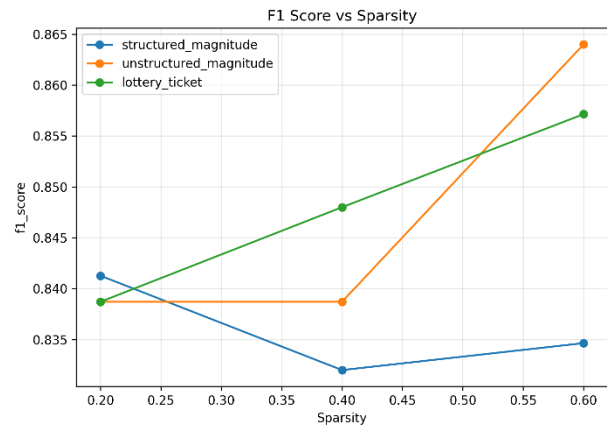


Fig. 3. F1-score versus sparsity level for the three pruning strategies.

B. Answering the Research Questions

Regarding the first research question, the results confirm that model pruning can maintain and even slightly improve predictive performance while substantially reducing model complexity. Structured magnitude pruning

at 40% sparsity achieved a higher AUC (0.9111) than the dense baseline (0.9002) with a $1.91\times$ compression ratio. This demonstrates that effective pruning successfully removes redundant parameters without sacrificing the model's ability to detect key corruption signals such as single-bid awards and irregular pricing patterns.

For the second research question, the pruned models showed clear improvements in practical edge-deployment proxies. Structured magnitude pruning significantly reduced inference latency (from 0.142 ms to 0.070 ms) and memory footprint. Quantized model sizes were also substantially smaller, particularly with the structured pruning approach. However, adversarial robustness remained moderate, with AUC dropping to approximately 0.8052 under FGSM perturbations, indicating that additional robustness techniques may be required for high-security applications.

C. Discussion

The experimental results demonstrate that model pruning, particularly structured magnitude pruning, represents an effective approach for developing lightweight yet accurate procurement anomaly detection models. The demonstrated capacity to reduce model size and inference latency while maintaining or improving AUC supports deployment in real-world settings, including edge devices used in military and public-sector monitoring systems. From a security perspective, these compact models support near-real-time detection of procurement risk signals — such as single-bid awards, abnormally short tender windows, and irregular fund flows — in defense logistics and governance frameworks. The reduced computational footprint also lowers the attack surface in secure environments and decreases energy consumption. Compared to prior studies that primarily emphasize predictive accuracy on large-scale models [7], [8], this work advances the field by combining real TED procurement data with systematic pruning and security-focused evaluation metrics. While some related works report higher absolute accuracy, they rarely address the critical constraints of edge deployment and adversarial robustness in security-critical applications.

Nevertheless, the moderate drop in performance under adversarial perturbations highlights the need for complementary techniques such as adversarial training. Overall, this study provides a promising pathway for integrating efficient, trustworthy AI into secure systems for enhanced societal integrity and resilience in public procurement and defense sectors

III. CONCLUSIONS

This study investigated the effectiveness of model pruning techniques for developing lightweight procurement anomaly detection models suitable for deployment in secure and resource-constrained environments. Using the TED dataset focused on

Bulgarian public tenders in 2023, a baseline enhanced MLP was trained and compared against three pruning strategies. The experiments show that at 40% sparsity, structured magnitude pruning provided a good trade-off between predictive accuracy and model efficiency. It achieved an AUC of 0.9111 (vs 0.9002 for the dense baseline), while delivering a compression ratio of 1.91, reducing model size to 0.0510 MB, and decreasing inference time to 0.070 ms. This configuration is well suited for edge deployment in defense logistics, public sector monitoring, and infrastructure protection. Experiments demonstrated that model pruning can effectively reduce the key deployment barriers facing AI-based procurement risk monitoring systems — namely high computational and memory requirements and limited suitability for resource-constrained environments — while incurring only modest performance reductions with respect to key procurement risk signals such as single-bid awards, abnormally short tender windows, and price anomalies. Although the models showed moderate sensitivity to adversarial perturbations, the permutation importance analysis enhanced trustworthiness by identifying features aligned with established domain expertise in procurement risk assessment.

In conclusion, this work demonstrates that compact, efficient, and reasonably robust procurement anomaly detection models can be successfully developed through pruning techniques, contributing toward the broader goal of trustworthy AI for governance monitoring. The significant reduction in model size and latency opens promising opportunities for implementation on embedded and resource-limited hardware in real-world security-related applications. These results should be interpreted as a proof-of-concept demonstration, given the weakly supervised labeling setup and the single-country, single-year dataset scope; further validation on larger, multi-country datasets and confirmed ground-truth labels is required before operational deployment.

Future work will focus on expanding the dataset with multi-country TED records, combining pruning with adversarial training and quantization techniques to further enhance robustness and efficiency, and conducting practical deployment tests on tiny embedded devices and edge hardware. Integrating these lightweight models into operational decision support systems will be crucial for strengthening societal resilience against corruption in governance and defense sectors.

IV. ACKNOWLEDGMENTS

The present paper was prepared as a result of the scientific research activities conducted within scientific project “Research on the possibilities of preventing corruption activities” financed by the Bulgarian National Science Fund at the Ministry of Education and Science, “Competition for financial support of basic research projects – 2025”, Contract № KII-06-H95/7 by 09.12.2025

REFERENCES

- [1] Transparency International Defence & Security, "Governance First: Tackling Corruption Risks Amid Rising Defence Spending," Transparency International Defence & Security, Apr. 2026. [Online]. Available: <https://ti-defence.org/publications/governance-first-tackling-corruption-risks-amid-rising-defence-spending/>
- [2] Transparency International US and Transparency International Defence & Security, "Blissfully Blind: The New US Push for Defense Industrial Collaboration with Partner Countries and its Corruption Risks," Transparency International, Apr. 2024. [Online]. Available: <https://us.transparency.org/app/uploads/2024/04/Blissfully-Blind-US-Push-for-Defense-Industrial-Collaboration-WEB-New-1.pdf>
- [3] Transparency International Defence & Security, "Trojan Horse Tactics: Corruption Risk in Defence Spending," Transparency International Defence & Security, 2024. [Online]. Available: <https://ti-defence.org/wp-content/uploads/2024/04/Trojan-Horse-Tactics-corruption-risk-in-defence-spending.pdf>
- [4] Transparency International, "Corruption Perceptions Index 2025," Transparency International, Feb. 2026. [Online]. Available: <https://www.transparency.org/en/cpi/2025>
- [5] M. Resimić, "Harnessing artificial intelligence (AI) for anti-corruption," Transparency International Knowledge Hub, Oct. 2025. [Online]. Available: https://knowledgehub.transparencycdn.org/kproducts/Harnessing-artificial-intelligence-for-anti-corruption_01.10.2025.pdf
- [6] T. Zatonatska, O. Dluhopolskyi, O. Artiushenko, I. C. Lopes, A. Ignatyuk, and O. Liubkina, "Comparative analysis of ensemble machine learning models for risk-oriented monitoring of military procurement," *Journal of Risk and Financial Management*, vol. 19, no. 3, p. 170, 2026, doi: <https://doi.org/10.3390/jrfm19030170>.
- [7] C. Cernăzanu-Glăvan and A.-Ş. Bulzan, "Explainable AI and fuzzy linguistic interpretation for enhanced transparency in public procurement: Analyzing EU tender awards," *Mathematics*, vol. 13, no. 13, p. 2215, 2025, doi: <https://doi.org/10.3390/math13132215>.
- [8] OECD (2025), *Governing with Artificial Intelligence: The State of Play and Way Forward in Core Government Functions*, OECD Publishing, Paris, <https://doi.org/10.1787/795de142-en>.
- [9] G. Piras, M. Pintor, A. Demontis, B. Biggio, G. Giacinto, and F. Roli, "Adversarial pruning: A survey and benchmark of pruning methods for adversarial robustness," *Pattern Recognition*, vol. 168, p. 111788, 2025, doi: <https://doi.org/10.1016/j.patcog.2025.111788>.
- [10] A. Douzandeh Zenoozi, L. Erhan, A. Liotta, and L. Cavallaro, "A comparative study of neural network pruning strategies for industrial applications," *Frontiers in Computer Science*, vol. 7, p. 1563942, 2025, doi: <https://doi.org/10.3389/fcomp.2025.1563942>.
- [11] J. Frankle and M. Carbin, "The lottery ticket hypothesis: Finding sparse, trainable neural networks," *arXiv preprint arXiv:1803.03635*, 2018.
- [12] P. Austria *et al.*, "The effects of compounded reductions on adversarial robustness (pruning and robustness)," *OSTI*, 2025. Available: <https://www.osti.gov/servlets/purl/3002675>.
- [13] E. Schneider dos Santos, M. Machado dos Santos, M. Castro *et al.*, "Detection of fraud in public procurement using data-driven methods: a systematic mapping study," *EPJ Data Science*, vol. 14, p. 52, 2025, doi: <https://doi.org/10.1140/epjds/s13688-025-00569-3>
- [14] L. Potin, R. Figueiredo, V. Labatut, and C. Langeron, "Pattern mining for anomaly detection in graphs: Application to fraud in public procurement," in *Machine Learning and Knowledge Discovery in Databases: Applied Data Science and Demo Track (ECML PKDD 2023)*, G. D. F. Morales *et al.*, Eds. Lecture Notes in Computer Science, vol. 14174. Cham, Switzerland: Springer, 2023, pp. 1–? doi: https://doi.org/10.1007/978-3-031-43427-3_5.
- [15] U. U. Acikalin, M. K. Gorgun, M. Kutlu, and B. K. O. Tas, "How you describe procurement calls matters: Predicting outcome of public procurement using call descriptions," *Natural Language Engineering*, vol. 30, no. 6, pp. 1255–1276, 2024. doi:<https://doi.org/10.1017/S135132492300030X>
- [16] B. Liu *et al.*, "A Survey of Lottery Ticket Hypothesis," *arXiv preprint arXiv:2403.04861*, 2024. [Online]. Available: <https://arxiv.org/pdf/2403.04861>
- [17] R. Bluteau, R. Gras, and G. Peralta, "Genetic-based lottery ticket pruning for transformers in sentiment classification: Realized through lottery sample selection," *Algorithms*, vol. 18, no. 12, p. 798, 2025, doi: <https://doi.org/10.3390/a18120798>.
- [18] OECD, *Digital Transformation of Public Procurement*, OECD Publishing, Paris, 2025. Available: https://www.oecd.org/content/dam/oecd/en/publications/reports/2025/06/digital-transformation-of-public-procurement_90ace30d/79651651-en.pdf.
- [19] OECD, "AI in fighting corruption and promoting public integrity," in *Governing with Artificial Intelligence*, OECD Publishing, Sep. 2025. [Online]. Available: https://www.oecd.org/en/publications/governing-with-artificial-intelligence_795de142-en/full-report/ai-in-fighting-corruption-and-promoting-public-integrity_60f5c50a.html
- [20] Transparency International Defence & Security, "The Trust Gap: How Public Trust Links to Defence Spending," Transparency International Defence & Security, Mar. 2026. [Online]. Available: <https://ti-defence.org/the-trust-gap-how-public-trust-links-to-defence-spending/>