

Understanding OPF Redispatch Actions through LLM-Assisted Counterfactual Analysis

Ferinar Moaidi^{*†}, Ricardo J. Bessa, *Fellow, IEEE*^{*}

^{*} Institute for Systems and Computer Engineering, Technology and Science (INESC TEC), Porto, Portugal
Email: {ferinar.moaidi, ricardo.j.bessa}@inesctec.pt

[†] Faculty of Engineering of the University of Porto (FEUP), Porto, Portugal

Abstract—Modern power systems increasingly rely on AC optimal power flow (AC-OPF) to support generation redispatch and relieve network congestion. However, the resulting dispatch decisions often remain difficult to interpret for operators, grid users, and regulators. This paper introduces an interpretable and adaptive framework that combines physical operational analysis (POA) with a multi-agent large language model architecture to generate concise, operator-oriented explanations of AC-OPF redispatch decisions. A reflective multi-agent loop composed of an Actor, Critic, Distiller, and Refiner iteratively improves the explanation process by extracting compact operational rules and integrating them into the Actor prompt through a sliding-window refinement mechanism. Evaluation on the IEEE 39-bus system shows that the proposed agent delivers physically consistent, causally coherent, and highly readable redispatch explanations, while generalizing well to unseen operating snapshots. Using multi-criteria metrics reflecting both accuracy and usefulness, it achieved approximately 96% performance with low computational overhead.

Index Terms—Optimal Power Flow, Large Language Models, Counterfactual Analysis, Power System Operation, Sensitivity Analysis, Interpretability.

I. INTRODUCTION

In modern power systems, it is essential to maintain secure and economical operation under increasing renewable energy integration and uncertain load patterns. This requires continuous adjustment of generator outputs, a process commonly referred to as *redispatch*. The primary objectives of redispatch are to relieve congestion and restore voltage limits under varying operating conditions. Despite the widespread use of AC-OPF, its solutions remain largely opaque to human operators. Optimization solvers output dispatch setpoints and dual

variables. However, they offer limited insight into *why* certain generators are adjusted or *how* specific constraints drive the solution. As a result, redispatch decisions derived from AC-OPF are often perceived as “black boxes”, providing minimal interpretability in operational contexts where human understanding and justification are essential.

Recent advances in Large Language Models (LLMs) have opened a new frontier for interpretable optimization. LLMs are increasingly used to bridge the gap between formal mathematical solvers and human operators by providing natural language interfaces, scenario exploration, and verbal reasoning over model internals. For instance, Zhang et al. (2025) introduced *OptiChat* [1], a dialogue system driven by LLMs that enables users to interrogate Pyomo-based optimization models. OptiChat supports “what-if” and “why-not” analyses, extracting variable structures, constraints, and sensitivities to produce counterfactual explanations for model outputs. Similarly, Chen et al. (2024) [2] proposed an LLM interface for diagnosing infeasible optimization problems, allowing operators to explore relaxation options and trade-offs in natural language. Their model demonstrated promising descriptive capabilities but limited quantitative grounding. LLMs are increasingly employed as autonomous reasoning agents capable of iterative improvement through feedback. Recent developments have introduced frameworks in which an LLM not only generates textual outputs but also learns from verbal feedback, reflection, and memory consolidation to enhance reasoning quality over time. This paradigm, referred to as *verbal reinforcement learning*, provides a natural and interpretable extension to reinforcement learning principles in language-based reasoning systems. Architectures such as Reflexion [3], Language Agent Tree Search (LATS) [4], and Self-Evaluation Memory [5] represent a significant step beyond traditional explainable optimization. Unlike prior LLM approaches that produce static textual interpretations, these models dynamically refine their explanations based on physical feedback.

This work is part of the AI4REALNET (*AI for REAL-world NET-work operation*) project, which received funding from the European Union’s Horizon Europe Research and Innovation programme under the Grant Agreement No 101119527. This project is funded by the European Union and SERI. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union. Neither the European Union nor the granting authority can be held responsible for them. Ferinar Moaidi was supported by FCT (*Fundação para a Ciência e a Tecnologia*) within the Ph.D. Grant 2023.04869.BD.

A. Original Contributions

The proposed framework introduces a novel integration of POA and LLM reasoning to enhance the interpretability of power system redispatch decisions by transforming “black-box” outputs into transparent, physically grounded explanations. Its main contributions are: 1) interpretable redispatch via counterfactual sensitivities, with an interpretation layer that converts AC-OPF results into counterfactual insights by leveraging sensitivity matrices, enabling a localized understanding of how small redispatch actions affect line flows and slack variable relaxations, and allowing interpretable what-if/why-not reasoning within the feasible space; 2) a structured reasoning loop composed of an *Actor*, *Critic*, *Distiller*, and *Refiner*, which interprets optimization outputs while preserving consistency with power-system physics through explanation hypotheses, physical plausibility checks, compact operational rule extraction, and natural-language prompt refinement; and 3) a verbal reasoning loop, building on Reflexion and LATS, in which linguistic critique guides subsequent reasoning episodes, with short-term memory capturing transient reasoning contexts and long-term memory impacting the actor prompt, thereby enhancing interpretability and consistency over multiple redispatch analyzes.

B. Structure

The remainder of this paper is organized as follows. Section II presents the AC-OPF redispatch formulation and numerical solution approach. Section III describes the physical operational analysis methods applied to the AC-OPF outputs. Section IV introduces the proposed LLM-based adaptive interpretation approach. Section V reports validation on the IEEE 39-bus system. Section VI concludes the paper and outlines future work.

II. REDISPATCH WITH AC-OPF

The operational problem is formulated as a two-stage AC-OPF with an initial *dispatch* stage and a subsequent *redispatch* stage. The dispatch stage emulates the market-based operating point, while redispatch adjusts generation closer to real time to relieve congestion and voltage-limit violations under updated conditions.

In the dispatch stage, generator outputs and voltages are optimized subject to AC power-flow equations and operational limits. The objective includes conventional generation costs, renewable curtailment, and penalized slack variables to ensure feasibility. AC power balance is enforced at each bus, and generators operate within their physical limits, while renewable units are modeled as available but curtailable.

In the redispatch stage, updated load and renewable conditions are applied, and conventional generators are adjusted around the dispatch solution:

$$P_{g,i}^{\text{RD}} = P_{g,i}^{\text{DA}} + \Delta P_i^+ - \Delta P_i^-. \quad (1)$$

The redispatch objective minimizes deviations from the dispatch schedule:

$$\begin{aligned} \min J^{\text{RD}} = & \sum_{i \in \mathcal{G}^{\text{conv}}} \left(c_i^{\text{up}} \Delta P_i^+ + c_i^{\text{down}} \Delta P_i^- \right) \\ & + \sum_{i \in \mathcal{R}} c_i^{\text{res}} C_i^{\text{res,RD}} + M \sum_{i \in \mathcal{B}} \Phi_i^{\text{RD}}. \end{aligned} \quad (2)$$

where $P_{g,i}^{\text{RD}}$ and $P_{g,i}^{\text{DA}}$ denote the redispatch and dispatch active power of generator i , respectively, ΔP_i^+ and ΔP_i^- are the upward and downward redispatch adjustments, c_i^{up} and c_i^{down} represent the corresponding redispatch costs, $C_i^{\text{res,RD}}$ denotes renewable curtailment, and Φ_i^{RD} are slack variables penalized by a large constant M .

III. PHYSICAL OPERATIONAL ANALYSIS

The Physical Operational Analysis (POA) module does not constitute a separate optimization method. Instead, it refers to the information extracted from the AC-OPF redispatch solution to interpret the physical drivers of the result. In this work, POA is limited to the active constraints and associated dual variables of the OPF solution, together with PTDF and LODF sensitivity indicators. Its purpose is to explain why a given redispatch pattern emerges by identifying the constraints associated with the solution, quantifying their influence through dual variables, and using sensitivity factors to construct counterfactual operating insights.

The extracted information is organized by first identifying the active constraints of the OPF solution. For a generic inequality constraint $g_i(x) \leq b_i$, the constraint is considered active at the optimum x^* when $|g_i(x^*) - b_i| < \varepsilon$ and $\lambda_i^* > 0$, where λ_i^* is the associated Lagrange multiplier. The collection of such constraints forms the active set, i.e., the limits that are binding, or nearly binding, at the optimum and therefore shape the feasible redispatch pattern. The dual variables provide a quantitative measure of the influence of each active constraint on the optimal objective, with marginal sensitivity

$$\frac{\partial f^*}{\partial b_i} = -\lambda_i^*. \quad (3)$$

Large dual magnitudes therefore indicate operational limits that strongly influence the optimal redispatch.

To interpret how redispatch actions affect network constraints, PTDFs and LODFs are used as first-order screening indicators under the standard lossless DC network approximation [6], [7]. For incremental net

injections $\Delta \mathbf{p}$ satisfying $\mathbf{1}^\top \Delta \mathbf{p} = 0$, the approximate branch-flow change is

$$\Delta \mathbf{f} \approx \mathbf{H} \Delta \mathbf{p}, \quad (4)$$

where $\mathbf{H}$ is the PTDF matrix. Thus, PTDFs identify relieving versus aggravating redispatch directions and estimate the order of magnitude of redispatch required to relieve congested branches. Beyond redispatch directionality, POA also evaluates approximate $N-1$ security implications using LODFs [8], [9]. For an outage of line k , the post-contingency flow on monitored line ℓ is approximated by

$$f_\ell^{\text{post}} \approx f_\ell^{\text{pre}} + \text{LODF}_{\ell,k} f_k^{\text{pre}}, \quad (5)$$

which provides a rapid screening estimate of outage-induced rerouting. POA combines active constraints, duals, PTDF, and LODF to produce interpretable explanations and counterfactual redispatch insights.

IV. MULTI-AGENT LLM FRAMEWORK

A multi-agent LLM framework (Figure 1) was developed to generate physically grounded, operator-oriented explanations for AC-OPF redispatch outcomes. The framework does not replace the optimization model or solve an additional OPF instance; instead, it transforms the numerical evidence produced by the OPF and the POA module into a concise explanation resembling the reasoning of a transmission system operator. The explanation is explicitly constrained by physical evidence from the snapshot, including pre- and post-OPF operating points, dual variables, sensitivities, and LODF-based security indicators.

Let $\mathcal{S}_t$ denote the system snapshot at time index t :

$$\mathcal{S}_t = (\mathbf{x}_t^{\text{pre}}, \mathbf{x}_t^{\text{opf}}, \boldsymbol{\lambda}_t, \boldsymbol{\mu}_t, \boldsymbol{\Gamma}_t, \mathbf{L}_t), \quad (6)$$

where $\mathbf{x}_t^{\text{pre}}$ and $\mathbf{x}_t^{\text{opf}}$ are the pre- and post-redispatch states, $\boldsymbol{\lambda}_t$ and $\boldsymbol{\mu}_t$ are the dual variables, $\boldsymbol{\Gamma}_t$ contains sensitivity information, and $\mathbf{L}_t$ contains LODF values for qualitative contingency screening. The central component is the Actor agent A_θ , which maps the snapshot and prompt p_t to an operator-facing explanation:

$$e_t = A_\theta(\mathcal{S}_t, p_t). \quad (7)$$

The Actor is instructed to rely only on the provided snapshot, dual variables, sensitivities, and LODF values, and not to solve OPF or power flow. Its role is to identify the dominant active constraint, explain why redispatch occurred, construct at least one explicit counterfactual, and compare the chosen action against plausible alternatives.

The counterfactual logic is guided by first-order physical indicators. If branch ℓ^* is the dominant limiting

element, the effect of a balanced redispatch between buses i and j is approximated by

$$\Delta f_{\ell^*} \approx (\text{PTDF}_{\ell^*,i} - \text{PTDF}_{\ell^*,j}) \Delta P, \quad (8)$$

so PTDFs identify relieving versus aggravating redispatch directions. Dual variables are interpreted according to physical role: power-balance duals act as nodal marginal prices, while duals of thermal, generator, reactive, or voltage limits indicate active operational constraints. The Actor therefore combines primal evidence (e.g., overloaded line or saturated generator) with dual evidence to produce a numerically grounded explanation rather than a generic summary. LODFs are used only for qualitative $N-1$ screening.

The generated explanation is then evaluated by an AI Critic,

$$\mathbf{v}_t = C^{\text{AI}}(\mathcal{S}_t, e_t), \quad (9)$$

which checks whether the explanation identifies the driving constraint, includes a counterfactual, provides approximate quantitative reasoning, and remains operationally readable. The Critic feedback is compressed by a Rule Distiller,

$$r_t = D(e_t, \mathbf{v}_t), \quad (10)$$

which extracts reusable operational rules. These rules are incorporated by the Prompt Refiner over a sliding window of recent snapshots,

$$p_{t+1} = R(p_t, W_t), \quad (11)$$

so that learning occurs at the prompt level rather than by updating model parameters. The overall closed-loop process is

$$\mathcal{S}_t \xrightarrow{A_\theta} e_t \xrightarrow{C^{\text{AI}}} \mathbf{v}_t \xrightarrow{D} r_t \xrightarrow{R} p_{t+1}. \quad (12)$$

The complete procedure is summarized in Algorithm 1. The framework therefore formalizes explanation generation as a physically constrained inference-and-refinement process in which the Actor produces counterfactual explanations from AC-OPF and POA outputs, the Critic verifies technical coherence, the Distiller extracts transferable rules, and the Refiner updates the Actor prompt so that future explanations become progressively more consistent, physically grounded, and operationally useful.

V. EXPERIMENTAL RESULTS

A. Case-study description

The proposed interpretable AC-OPF redispatch framework was evaluated on the IEEE 39-bus benchmark system. All AC-OPF instances were solved using IPOPT [10] with the MA57 sparse symmetric linear solver.

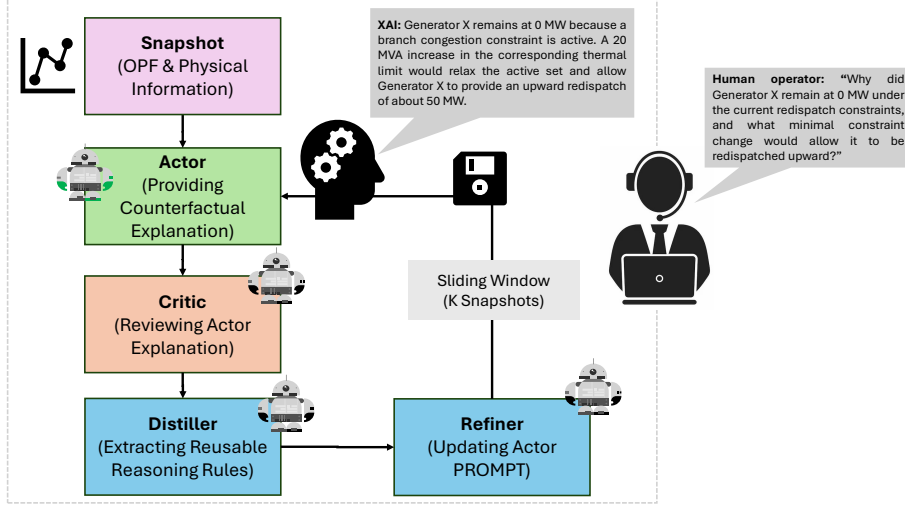


Fig. 1: Multi-agent LLM framework for AC-OPF redispatch explanation in control rooms

Algorithm 1: Multi-Agent LLM Prompt Refinement

Input: Snapshots $\{\mathcal{S}_t\}_{t=1}^T$, initial prompt p_0 , memory $\mathcal{M}$, interval K

Output: Final prompt p_{final} , updated memory $\mathcal{M}$

```

1  $p \leftarrow p_0$ ;
2 for  $t \leftarrow 1$  to  $T$  do
3    $e_t \leftarrow A_\theta(\mathcal{S}_t, p)$ ;
4    $\mathbf{v}_t \leftarrow C^{\text{AI}}(\mathcal{S}_t, e_t)$ ;
5    $r_t \leftarrow D(e_t, \mathbf{v}_t)$ ;
6    $\mathcal{M} \leftarrow \mathcal{M} \cup \{r_t\}$ ;
7   if  $t \geq K$  and  $t \bmod K = 0$  then
8      $W_t \leftarrow \{r_{t-K+1}, \dots, r_t\}$ ;
9      $p \leftarrow R(p, W_t)$ ;
10  $p_{\text{final}} \leftarrow p$ ;
11 return  $p_{\text{final}}, \mathcal{M}$ ;

```

A total of 500 operating snapshots were generated for training by introducing uncertainty in load and renewable generation. Load uncertainty was derived from scaled secondary substation measurements from the Iowa Distribution Test System [11]. Renewable uncertainty was based on wind power measurements from wind plants in France, used after replacing 40% of the system generators with wind units. Wind curtailment was allowed as a redispatch action, up to 30% of the available wind power. The day-ahead dispatch used the point forecast proposed in [12] based on NWP data generated at 0h00, while the redispatch stage used the very-short-term forecast model in [12]. An additional 100 previously unseen operating snapshots were reserved for testing.

The LLM agents used task-specific temperature settings: 0.7 for the Actor, 0.2 for both the Critic and Distiller, 0.4 for the Refiner, and 0.1 for the Evaluator. All experiments were conducted using GPT-5.4.

B. Illustrative Example

This subsection presents one example dialogue generated by the proposed LLM agent to illustrate how it interprets an OPF snapshot using information from both the OPF solution and the POA. The case is informative because the overloads are distributed across related areas, yet the final redispatch is still driven mainly by a single dominant thermal constraint, namely line 1–2. Before redispatch, five thermal overloads were reported: line 1–2 at 143.02%, line 2–3 at 134.77%, line 5–10 at 121.88%, line 14–15 at 119.26%, and line 15–18 at 107.46%. These were not isolated violations, since line 1–2 and line 2–3 belong to the same corridor, while line 14–15 and line 15–18 form a related path around bus 15. After redispatch, all were reduced below 100%, with line 1–2 ending at 93.61%. This line remained the main limiting element, as it had the highest pre-redispatch overload and was the only line with a nonzero post-redispatch thermal dual (−6.761713). Note that the nonzero shadow price arises because, in AC-OPF with soft constraints, a constraint can remain locally binding within numerical tolerance and still influence the optimal solution; the dual reflects marginal cost sensitivity rather than strict violation. The realized redispatch reduced Conventional Generator 8 at bus 38 by 20 MW, curtailed wind by 74.404 MW at bus 31 and 8.344 MW at bus 29, and increased output from Generator 6 at bus 36 by 40 MW and Generator 2 at bus 32 by 62 MW, with no load curtailment. Note that the reported redispatch

quantities are taken from the AC-OPF solution; PTDfFs are used only for explanatory first-order screening, so exact lossless DC power balance is not expected because AC losses are represented.

a) *Question 1. Why not do more downward redispatch on Generator 8 at bus 38 instead of 1 MW more wind curtailment at bus 31?:*

The closed marginal swap is +1 MW injection from reducing wind curtailment at bus 31 and -1 MW withdrawal from extra downward redispatch of Generator 8 at bus 38.

On line 1-2, the PTDfF effect is $\text{PTDF}(31) - \text{PTDF}(38) = -0.0125 - 0.3183 = -0.3308$, so this swap would relieve line 1-2 by about 0.3308 MW. On line 2-3, the effect is $0.0164 - 0.2706 = -0.2542$ MW, so it would also relieve line 2-3.

Economically, $\Delta C = -30.000 + 32.190 = +2.190$ €/MWh, so this alternative is more expensive. It was also blocked by a local operating limit: Generator 8 was already providing 20 MW of downward redispatch and had a binding RampDown dual of -55.811888, which indicates that further downward movement was constrained.

This counterfactual would become feasible if Generator 8 had at least 1 MW of additional downward ramp margin. It would become economically preferable if the Gen 8 downward redispatch price fell by at least 2.190 €/MWh, from 32.190 to 30 €/MWh or lower.

b) *Question 2. Why not replace 1 MW of wind curtailment at bus 31 with 1 MW of downward redispatch on Generator 7 at bus 37?:*

The closed swap is +1 MW injection at bus 31 and -1 MW withdrawal at bus 37.

On line 1-2, the PTDfF effect is $\text{PTDF}(31) - \text{PTDF}(37) = -0.0125 - 0.2787 = -0.2912$ MW, so this would relieve the dominant line. On line 2-3, the effect is $0.0164 - 0.4821 = -0.4657$ MW, so it would also relieve the same corridor.

However, on line 14-15, the effect is $0.0400 - (-0.3555) = +0.3955$ MW, so it would worsen that already overloaded related line. Economically, $\Delta C = -30.000 + 21.616 = -8.838$ €/MWh, so it would look cheaper. Even so, it is not viable because it shifts congestion onto line 14-15.

For 1 MW of this swap to become feasible, line 14-15 would need at least 0.3955 MW of additional thermal margin. The realized redispatch is therefore the more plausible local solution because it reduces injections at buses 29, 31, and 38, and each of those buses has a positive PTDfF on line 1-2 and line 2-3. In that way, the chosen package consistently unloads the dominant corridor without requiring emergency load curtailment.

c) *Question 3. Is there any major N-1 cascading risk?:*

A clearer screened risk is the outage of line 14-15 onto monitored line 2-17. The pre-outage flow on line 2-17 was $0.3511 \times 500 = 175.6$ MVA, while the flow on the outaged line 14-15 was $1.1926 \times 540 = 644.0$ MVA.

With LODF = 0.7823, the screened post-outage flow becomes $175.6 + 0.7823 \times 644.0 = 679.4$ MVA, which is 135.9% of the 500 MVA limit. This gives first-order evidence that redispatch was also needed to reduce screened outage risk, not only to remove the base-case thermal overloads.

C. Quantitative Evaluation

The generated AC-OPF redispatch explanations were evaluated on previously unseen operating snapshots using a dedicated evaluator agent. The evaluator does not re-solve OPF or assess stylistic plausibility in isolation; instead, it checks whether each explanation correctly justifies the redispatch decision using only the information available in the corresponding snapshot. In particular, it verifies physical consistency, counterfactual validity, numerical grounding, and operator usefulness.

For each unseen snapshot S_n , $n = 1, \dots, N_{\text{test}}$, the Actor generates an explanation e_n . The evaluator then assesses the pair (S_n, e_n) through 20 binary verification questions:

$$q_{n,m} \in \{0, 1\}, \quad m = 1, \dots, 20, \quad (13)$$

where $q_{n,m} = 1$ indicates that the corresponding claim is correct and supported by the snapshot, and $q_{n,m} = 0$ indicates that it is incorrect, unsupported, or unverifiable. This makes the evaluation deliberately strict.

The 20 questions are grouped into eight categories: $\bar{\sigma}^{(1)}$ redispatch trigger identification, $\bar{\sigma}^{(2)}$ operational logic, $\bar{\sigma}^{(3)}$ PTDfF directional reasoning, $\bar{\sigma}^{(4)}$ counterfactual reasoning, $\bar{\sigma}^{(5)}$ numerical grounding and report discipline, $\bar{\sigma}^{(6)}$ economic and dual interpretation, $\bar{\sigma}^{(7)}$ N-1 security screening, and $\bar{\sigma}^{(8)}$ language and operator usefulness.

For each snapshot, P_n denotes the percentage score of the generated explanation, and the overall performance is summarized by the mean percentage score $\bar{P}$ over all test snapshots. Similarly, $\sigma_n^{(c)}$ denotes the score of category c for snapshot n , and $\bar{\sigma}^{(c)}$ denotes the corresponding average over the full test set. The aggregate metrics $(\bar{P}, \bar{\sigma}^{(c)})$ therefore provide a structured basis for comparing prompt variants and measuring explanation faithfulness to the available OPF and POA data.

The effectiveness of the reflective prompt-refinement framework was assessed using this evaluator. Table I compares the initial Actor prompt with the evolved prompt obtained after iterative refinement. The evolved

TABLE I: Comparison of the initial and evolved prompts using the category-wise mean scores $\bar{\sigma}^{(c)}$, $c = 1, \dots, 8$, and the overall percentage score $\bar{P}$.

Metric	Initial	Evolved	Improvement (%)
$\bar{\sigma}^{(1)}$	0.8950	0.9825	9.7765
$\bar{\sigma}^{(2)}$	0.8000	0.9930	24.1250
$\bar{\sigma}^{(3)}$	0.7900	0.9800	24.0506
$\bar{\sigma}^{(4)}$	0.8340	0.9650	15.7074
$\bar{\sigma}^{(5)}$	0.7270	0.8230	13.2050
$\bar{\sigma}^{(6)}$	0.7450	0.9900	32.8859
$\bar{\sigma}^{(7)}$	0.7200	0.9500	31.9444
$\bar{\sigma}^{(8)}$	0.8900	0.9900	11.2360
$\bar{P}$	80.0125	95.9188	19.8797

$\bar{\sigma}^{(1)}$: Redispatch trigger identification; $\bar{\sigma}^{(2)}$: Operational logic; $\bar{\sigma}^{(3)}$: PTDF directional reasoning; $\bar{\sigma}^{(4)}$: Counterfactual reasoning; $\bar{\sigma}^{(5)}$: Numerical grounding and report discipline; $\bar{\sigma}^{(6)}$: Economic and dual interpretation; $\bar{\sigma}^{(7)}$: Security screening; $\bar{\sigma}^{(8)}$: Language and operator usefulness.

prompt improves all eight categories, showing that the refinement process strengthens not only textual quality, but also physical grounding, numerical reliability, and operator relevance.

The overall score increased from 80.01% to 95.92%, corresponding to a relative improvement of 19.88%. The largest gains were obtained in economic and dual interpretation (+32.89%), security screening (+31.94%), operational logic (+24.13%), and PTDF directional reasoning (+24.05%), indicating that the evolved prompt more consistently explains network sensitivities, contingency effects, and economic drivers. Counterfactual reasoning, numerical grounding, and language quality also improved, while redispatch trigger identification remained strong and increased further from an already high baseline. Overall, the results show that sliding-window reflective refinement leads to broader, more balanced improvements in explanation quality and supports more physically grounded, operator-oriented redispatch justification.

VI. CONCLUSIONS

This paper proposed an interpretable multi-agent large language model framework for explaining AC-OPF redispatch decisions in a physically grounded and operator-oriented manner by combining AC-OPF outputs, physical operational analysis, and a reflective reasoning pipeline composed of an Actor, Critic, Distiller, and Refiner. The framework generates concise counterfactual explanations tied to system state, active constraints, dual variables, sensitivity information, and LODF-based security indicators, while iterative prompt refinement through sliding-window memory improves linguistic clarity, numerical defensibility, physical consistency, and operator relevance. Evaluation on the IEEE 39-bus system showed that the evolved prompt outperformed the initial prompt on unseen snapshots, increasing the overall score from

80.01% to 95.92%, with the largest gains in security reasoning, numerical consistency, and economic interpretation. Overall, the results highlight the potential of large language models as trustworthy explanation engines for power-system operation, with future work focusing on larger networks and human-in-the-loop evaluation.

CODE AVAILABILITY

The code will be released as open-source in the GitHub repository (<https://github.com/AI4REALNET/XLLM>) of the AI4REALNET project and will be made publicly available upon acceptance of the paper.

ACKNOWLEDGMENTS

The authors would like to thank Micael Filipe Simões of INESC TEC for his valuable guidance and support in implementing the AC-OPF formulation using IPOPT.

REFERENCES

- [1] H. Chen, G. E. Constante-Flores, K. S. I. Mantri, S. M. Kompalli, A. S. Ahluwalia, and C. Li, "Optichat: Bridging optimization models and practitioners with large language models," *arXiv preprint arXiv:2501.08406*, 2025.
- [2] H. Chen, G. E. Constante-Flores, and C. Li, "Diagnosing infeasible optimization problems using large language models," *INFOR: Information Systems and Operational Research*, vol. 62, no. 4, pp. 573–587, 2024.
- [3] N. Shinn, F. Cassano, S. Koyejo, and N. Goodman, "Reflexion: Language agents with verbal reinforcement learning," *arXiv preprint arXiv:2303.11366*, 2023.
- [4] S. Zhou, K. Chen, Y. Chen, and W. Wang, "Language agent tree search unifies reasoning, acting, and planning in language models," *arXiv preprint arXiv:2308.08155*, 2023.
- [5] Y. Xie, C. Zhang, and L. Gao, "Self-evaluation guided beam search for reasoning in large language models," *arXiv preprint arXiv:2305.00633*, 2023.
- [6] R. D. Zimmerman, C. E. Murillo-Sánchez, and R. J. Thomas, "MATPOWER: Steady-state operations, planning, and analysis tools for power systems research and education," *IEEE Transactions on Power Systems*, vol. 26, no. 1, pp. 12–19, 2011.
- [7] B. Stott, J. Jardim, and O. Alsac, "DC Power Flow Revisited," *IEEE Transactions on Power Systems*, vol. 24, no. 3, pp. 1290–1300, 2009.
- [8] T. Güler, G. Gross, and M. Liu, "Generalized line outage distribution factors," *IEEE Transactions on Power Systems*, vol. 22, no. 2, pp. 879–881, 2007.
- [9] J. Guo, Y. Fu, Z. Li, and M. Shahidehpour, "Direct calculation of line outage distribution factors," *IEEE Transactions on Power Systems*, vol. 24, no. 3, pp. 1633–1634, 2009.
- [10] A. Wächter and L. T. Biegler, "On the implementation of an interior-point filter line-search algorithm for large-scale nonlinear programming," *Mathematical Programming*, vol. 106, no. 1, pp. 25–57, 2006.
- [11] Z. Wang, "Iowa distribution test systems," December 2022. [Online]. Available: <http://wzy.ece.iastate.edu/Testsystem.html>
- [12] R. J. Bessa, F. Moaidi, J. Viana, and J. R. Andrade, "Uncertainty-aware procurement of flexibilities for electrical grid operational planning," *IEEE Transactions on Sustainable Energy*, vol. 15, no. 2, pp. 789–802, 2024.