

Quantum-Inspired Risk Portfolio Optimization via Exponential Tilting

Fabio Vitale[†], Francesca Cibrario^{*}, Davide Veronelli^{*}, Chiara Vercellino[†], Valeria Zaffaroni^{*},
Giacomo Vitali[‡], Gregorio Antonio Polito^{*}, Emanuele Dri[‡],
Carla Monferrato[†], Olivier Terzo[‡], Davide Corbelleto^{*}, Marco Bianchetti^{*}

^{*}Intesa Sanpaolo, Torino, Italy

francesca.cibrario@intesaspaolo.com, davide.veronelli@intesaspaolo.com,
valeria.zaffaroni@intesaspaolo.com, gregorio.polito@intesaspaolo.com,
davide.corbelleto@intesaspaolo.com, marco.bianchetti@intesaspaolo.com,

[†]Intesa Sanpaolo Innovation Center, Torino, Italy

fabio.vitale@intesaspaolo.com, carla.monferrato@intesaspaolo.com,

[‡]Fondazione LINKS, Torino, Italy

chiara.vercellino@linksfoundation.com, giacomo.vitali@linksfoundation.com,
emanuele.dri@linksfoundation.com, olivier.terzo@linksfoundation.com

Abstract—Financial portfolio optimization is often framed through mean–variance objectives, yet real profit-and-loss (P&L) scenario distributions are typically discrete and non-Gaussian. To explore realistic P&L distributions, we propose an objective function based on exponential tilting and certainty equivalent under constant absolute risk aversion. We then derive a second-order approximation, recovering a quadratic optimization problem solvable through mixed integer quadratic programming (MIQP) and quadratic unconstrained binary optimization (QUBO). Constraints in the QUBO formulation are incorporated via Lagrangian relaxation and slack variable approaches, enabling a direct comparison between classical constrained MIQP and quantum-inspired methods. In order to study how exponential tilting affects mean, standard deviation, and tail-risk measures such as Value at Risk and Conditional Value at Risk, we design an experimental protocol based on scanning a grid of tilt parameters and solving the associated optimization problems. The results on a realistic trading portfolio show how the tilts can be tuned to optimize different risk-return profiles while remaining competitive with those found using other state-of-the-art approaches. Moreover, performance analysis suggests that, even using classical solvers, the QUBO formulation with Lagrangian relaxation scales better than the MIQP approach, providing the best constraint encoding strategy.

Index Terms—quantum-inspired optimization, QUBO, portfolio optimization, exponential tilting, Esscher transform, mean-variance

I. INTRODUCTION

Portfolio optimization is the process of selecting the best asset allocation among available choices in order to balance expected return and risk. In a scenario-based setting, portfolio performance is evaluated over a finite collection of possible market states, each producing a profit-and-loss (P&L) outcome, i.e., a change in the economic value of the portfolio. In this framework, a standard approach in quantitative finance is the *mean–variance* criterion, introduced by Markowitz [17], which selects portfolios by maximizing the average P&L across scenarios (the mean) while penalizing the variability of P&L across scenarios (the variance), thus favoring portfolios that achieve higher returns with limited dispersion of outcomes. This approach is frequently applied to relatively simple investment portfolios (e.g., N stocks with continuous quantities) and is suitable when the P&L distribution is Gaussian, since in that

case the portfolio risk is fully characterized by the variance. However, heterogeneous trading portfolios of commercial or investment banks typically encompass millions of financial instruments (such as stocks, securities, loans, derivatives, funds, etc.), featuring any level of complexity (plain-vanilla, exotics, etc.), depending on a wide range of underlying risk factors (interest rates, equity, etc.). The resulting P&L distributions can be markedly non-Gaussian, implying that risk cannot be fully captured by variance alone. Practitioners may therefore wish to look at the P&L distribution from different perspectives, for example by focusing on adverse or favorable outcomes. Moreover, in this realistic setting, the portfolio optimization problem consists in defining a strategy to build a new portfolio $\mathcal{T}$ improving an already owned *initial portfolio* $\mathcal{P}$, adding an appropriate selection of liquid financial instruments.

The objective of this work is to deal with this configuration and go beyond the mean–variance criterion while preserving a quadratic objective function. This enables the use of mixed-integer quadratic programming to solve small- and medium-scale portfolio optimization problems with current classical computing resources, thereby paving the way for the application of quantum computing to larger instances.

Our approach starts from the initial portfolio $\mathcal{P}$, identifies a suitable portfolio $\mathcal{B}$, hereafter called *base-point* portfolio, and uses a softmax map applied to its P&L distribution to reweight scenarios and obtain a new non-uniform probability measure. This construction is equivalent to the *Esscher transform* or exponential tilting [11, 23]. The portfolio $\mathcal{B}$ is “suitable” in the sense that it should be easily computable and reasonably in a neighborhood of promising solutions. A tilt parameter controls how P&L scenarios are reweighted: positive tilts concentrate weight on scenarios corresponding to high P&L values of $\mathcal{B}$, and conversely for negative tilts.

In order to find an optimal portfolio $\mathcal{T}$, we maximize a single monetary value, the *certainty equivalent* (CE) [13], associated with a P&L distribution under the weighted probability measure. To define the CE, we use the constant absolute risk aversion (CARA) exponential utility function [13, 10], which provides a tractable way of penalizing risk while maintaining a clear monetary interpretation. Finally, we replace the exact

CE with its second-order approximation, which can be used as a tilted mean–variance quadratic objective function.

To validate the proposed theoretical framework, we adopt the setup introduced in [3], where the initial trading portfolio is optimized by selecting a specific set (and related discretized quantities) of *unique eligible instruments* (UEIs) subject to predefined constraints. The base-point portfolio required by our approach is taken to be the solution of a constrained mean–variance relaxed problem using quadratic programming (QP), with continuous decision variables representing the quantity of each UEI added to the initial portfolio. For a grid of tilt values, we construct the corresponding objective functions and solve the constrained discretized problem using MIQP to evaluate the impact of tilting. Finally, for a subset of selected tilts, we demonstrate that the proposed approach can, in principle, be addressed using quantum computing techniques such as quantum annealing or variational quantum algorithms. We formulate a QUBO model by embedding constraints into penalty terms, using both slack variables and a Lagrangian relaxation approach. The QUBO formulation is validated on classical hardware using the solver GUROBI [14] and the quantum-inspired Simulated Bifurcation method [1], demonstrating the superior performance of the Lagrangian relaxation-based constraint encoding strategy. Notably, a preliminary scaling analysis, conducted by increasing the size of the optimization problem, demonstrates that on classical hardware the unconstrained approach can provide feasible results, requiring fewer hardware resources than the constrained version, while the quantum-inspired methodology demonstrates better memory management.

Main contributions.

- **Exponential tilting as a weighting method.** We introduce a one-parameter family of Esscher/softmax scenario weights around a fixed base-point portfolio of interest. Negative tilts emphasize downside base-point scenarios, while positive tilts emphasize upside ones.
- **Quadratic objective function with clear financial interpretation.** Using a CARA certainty equivalent, approximated to the second-order, we obtain a quadratic objective function with a clear risk–return meaning: it rewards expected gains and penalizes residual variability under the tilted scenario measure, while retaining non-Gaussian distributional information through the tilted weights.
- **QUBO formulation and constraints encoding.** The quadratic objective function is suitable for a QUBO formulation and, therefore, for quantum optimization. We embed financially relevant constraints into the quadratic objective function through penalties and Lagrangian relaxation, avoiding the introduction of additional binary variables.
- **Experimental validation and scalability.** Finally, experimental results validate the practical effectiveness of the proposed formulation, showing that the QUBO formulation with Lagrangian relaxation achieves competitive performance while preserving computational tractability on classical hardware.

II. RELATED WORK

A. Risk-aware portfolio optimization

Recent work in risk-aware trading portfolio optimization is formulated as the selection of a *tradable* strategy that improves

a given objective function of an initial portfolio while satisfying risk and business constraints. A representative framework is Risk-aware Trading Portfolio Optimization (RATPO) and its particle-swarm solver RATS [3].

In RATPO, the starting point is an initial portfolio (e.g., related to a single trading desk or the whole trading book) at some fixed date, and a user-specified universe of eligible optimization instruments, i.e., the liquid market instruments quoted by exchanges, brokers, or market makers, typically used by traders to hedge Delta, Vega, and Gamma risks. RATS introduces the notion of unique eligible instruments: each UEI is an eligible optimization instrument identified by a tuple of discrete parameters (e.g., underlying, maturity, strike, option type) with unitary notional value. An *eligible optimization strategy* (EOS) is then a combination of UEIs weighted with their corresponding discrete notional amounts, subject to constraints that encode feasibility and risk limits. Operationally, the decision variables specify which UEI-notional value combinations are selected, i.e., elements of a discrete Cartesian product of instrument-parameters and notional value domains [3], rather than continuous portfolio weights in a frictionless mean–variance model. The objective function includes market risk measures commonly used for risk management and capital purposes, such as Value at Risk (VaR) and Conditional Value at Risk (CVaR), calculated from empirical discrete profit-and-loss (P&L) distributions obtained by full portfolio revaluation in historical risk-factor scenarios. VaR is a market risk metric that estimates the maximum potential loss of a portfolio over a defined period within a specific confidence level. CVaR, a.k.a. Expected Shortfall (ES), measures the average loss exceeding the VaR threshold. It is often used to provide a more robust measure of extreme loss than VaR alone [25, 26]. The constraints addressed in RATPO are the following.

- **EOS composition:** prescribes the EOS composition in terms of maximum number and types of UEIs (e.g., one Futures or single stock and two call/put options) for each relevant underlying risk factor, encoding the traders’ optimization strategy.
- **UEI quantities:** prescribes discrete notional values from a fixed discrete set for each UEI, encoding the market tradable sizes.
- **Greeks constraints:** prescribes the maximum allowed variation of Delta, Gamma and Vega sensitivities, encoding the rigidity of the portfolio risk profile decided by traders.

Our work is based on constraints, discrete instrument choices, and datasets (see Section VI-A for more details) introduced in RATPO. It complements classical mean–variance and VaR/CVaR-based portfolio optimization [26, 30, 28, 25], building a novel quadratic objective function without imposing convexity or differentiability, as for particle-swarm optimization [27, 3]. In contrast with meta-heuristic search over structured trade sets, our approach keeps a single MIQP/QUBO solve by building a quadratic surrogate anchored at a fixed base-point, while still allowing tail emphasis through tilts and scenario reweighting. To build our objective function, we use an explicit Esscher/softmax reweighting of scenarios based on base-point P&L scores. Esscher transforms are classical in quantitative finance, for example, in option pricing and in Lévy-based asset models [11, 23]. Related ideas of scenario reweighting also appear in entropy-pooling approaches, where

posterior scenario probabilities are obtained by relative-entropy minimization under stress views [20].

B. Quantum portfolio optimization

Combinatorial optimization consists of finding the optimum solution of a problem defined by an objective function from a finite set of candidates. A brute-force strategy that evaluates all possible candidate solutions becomes computationally intractable as the dimension of the problem increases. Classical heuristics can be used to approach large scale combinatorial optimization problems providing high-quality solutions within reasonable computational time, without guaranteeing solution optimality. In this context, quantum optimization has emerged as an alternative computational paradigm that aims to tackle combinatorial problems through quantum computing.

Quantum computation can be realized through two principal paradigms: gate-based and analog quantum computing. QUBO models, which can be directly mapped to Ising Hamiltonians, provide a common mathematical framework compatible with both approaches. Gate-based quantum computers allow for the solution of QUBO problems through hybrid quantum-classical algorithms, such as the quantum approximate optimization algorithm (QAOA) [6], in which a parameterized quantum circuit is iteratively updated through a classical loop minimizing the expectation value of the Hamiltonian. Alternatively, analog quantum approaches, such as quantum annealing [7], map the Hamiltonian into the system's energy and exploit adiabatic evolution to drive the system toward its ground state, i.e. the optimum of the optimization problem.

A key challenge in applying these paradigms to industrial combinatorial optimization problems, which typically involve constraints, continuous or non-binary variables, and non-quadratic objective functions, lies in translating such problems into QUBO form.

A comprehensive review of quantum approaches applied to portfolio optimization is provided in [22]. From early studies based on mean-variance portfolio optimization formulated as QUBO problems [18], subsequent research has progressed in two main directions: (i) QUBO reformulations of more realistic portfolio models and (ii) hybrid quantum-classical approaches to extend the perimeter of addressable models relaxing QUBO requirements.

Preserving the QUBO formulation enables hardware-agnostic implementations by encoding non-quadratic or constrained portfolio optimization problems into the quadratic binary objective while remaining faithful to the original financial intent. Linear inequality constraints are commonly handled by introducing slack variables [12] or by unbalanced penalization [21]. The former approach improves constraint enforcement, but increases the number of binary variables, thus expanding the effective search space and the associated quantum resource requirements. The latter, while variable efficient, requires careful coefficient tuning to avoid infeasible low energy states. Higher-order objective functions and nonlinear constraints can be reduced to quadratic form by introducing auxiliary variables [16]. Following this direction, several approaches have been proposed beyond the standard mean-variance QUBO approach. In [24] portfolio variance minimization is combined with a linear penalty enforcing a minimum return threshold instead of directly minimizing it. Binary variables indicate whether a specific asset is selected. In [2] and [19], binary variables instead encode discretized asset

allocation quantities. The former proposes QUBO formulations based on the Herfindahl-Hirschman Index as a risk measure and return on capital as a performance metric. The latter instead proposes a QUBO objective derived from the Sharpe ratio.

The second direction of research takes advantage of the flexibility of hybrid quantum-classical algorithms to relax the requirement of quadratic unconstrained objective functions. These methods allow the inclusion of higher-order or nonlinear objectives and explicit constraints management [9], but often require deeper quantum circuits or non-native interactions, which may limit their applicability on current noisy intermediate-scale quantum (NISQ) hardware. In [29], higher-order portfolio optimization is addressed using gate-based QAOA, extending the Hamiltonian to incorporate skewness and kurtosis. Additionally, Value-at-Risk (VaR) objectives have been incorporated into the optimization process by combining quantum annealing with an external classical optimization loop [31].

This work proposes a QUBO preserving reformulation inspired by the portfolio optimization framework described in [3], designed to remain compatible with both gate-based and analog quantum paradigms. By limiting the introduction of auxiliary decision variables, the approach aims to facilitate future implementation on larger-scale quantum hardware.

III. SCENARIO-BASED FRAMEWORK AND TILTING

This section introduces the scenario-based framework and the exponential tilting used throughout the paper. We describe (i) the problem setting, introducing P&Ls and the decision variable, (ii) the relevant metrics used in the optimization process (tilted P&L mean and variance under Esscher/softmax weights), (iii) the Markowitz optimization setting and the choice of the tilting base-point $\bar{\mathbf{x}}$.

A. Problem setting

Let N_{UEI} be the number of possible UEIs, each described by its specific tuple of discrete parameters that, with the corresponding notional values, make up the EOS. For each UEI, we choose discrete notional values from its corresponding domain with cardinality equal to N_{notional} (see Section II-A). Thus, we may enumerate all possible pairs UEI-notional value using an index $u = 1, \dots, d$ with $d := N_{\text{UEI}} \cdot N_{\text{notional}}$.

Let $s = 1, \dots, S$ index the risk factors scenarios and the corresponding P&L scenarios. A P&L scenario represents the change of the economic value of the portfolio according to the risk factors scenario. Then, for each u and s , let $X_{s,u}$ denote the *unit-notional* P&L contribution of UEI indexed by u in scenario s , and let θ_u denote its notional value. The matrix $\mathbf{G} \in \mathbb{R}^{S \times d}$ collects the notional value-scaled scenario contributions $G_{s,u} := \theta_u X_{s,u}$. The P&L distribution of any portfolio $\mathcal{H}$ is indicated by the vector $\mathbf{z}(\mathcal{H}) \in \mathbb{R}^S$ whose s -th component is $z_s(\mathcal{H})$.

In our setting, an EOS can be identified by a decision vector $\mathbf{x} \in \{0, 1\}^d$, whose u -th component takes the value 1 if the pair UEI-notional value indexed by u is to be added to the initial portfolio $\mathcal{P}$ and 0 otherwise. Every EOS $\mathbf{x}$ defines a new portfolio $\mathcal{T}(\mathbf{x})$, obtained by adding the EOS to $\mathcal{P}$, whose P&L distribution is

$$\mathbf{z}(\mathcal{T}(\mathbf{x})) = \mathbf{z}(\mathcal{P}) + \mathbf{G}\mathbf{x}.$$

B. Weighted metrics and exponential tilting

A probability measure can be identified with a real vector $\mathbf{w} = (w_1, \dots, w_S)^\top$, whose components are non-negative, sum to 1 and represent the probability of each scenario. Given two vectors $\mathbf{f}, \mathbf{g} \in \mathbb{R}^S$, representing the values of functions defined on the scenario set, we can compute their expected value and covariance under $\mathbf{w}$

$$\mathbb{E}_{\mathbf{w}}[\mathbf{f}] = \sum_{s=1}^S w_s f_s = \mathbf{f}^\top \mathbf{w}, \quad (1)$$

$$\begin{aligned} \text{Cov}_{\mathbf{w}}(\mathbf{f}, \mathbf{g}) &= \sum_{s=1}^S w_s f_s g_s - \sum_{s, s'=1}^S w_s w_{s'} f_s g_{s'} \\ &= \mathbf{g}^\top (\text{diag}(\mathbf{w}) - \mathbf{w} \mathbf{w}^\top) \mathbf{f}. \end{aligned} \quad (2)$$

Observe that Eq. (1) and (2) can be extended to matrices $\mathbf{M} \in \mathbb{R}^{S \times d}$, so that

$$\begin{aligned} \mathbb{E}_{\mathbf{w}}[\mathbf{M}] &= \mathbf{M}^\top \mathbf{w} \in \mathbb{R}^d, \\ \text{Cov}_{\mathbf{w}}(\mathbf{f}, \mathbf{M}) &= \mathbf{M}^\top (\text{diag}(\mathbf{w}) - \mathbf{w} \mathbf{w}^\top) \mathbf{f} \in \mathbb{R}^d, \\ \text{Cov}_{\mathbf{w}}(\mathbf{M}, \mathbf{M}) &= \mathbf{M}^\top (\text{diag}(\mathbf{w}) - \mathbf{w} \mathbf{w}^\top) \mathbf{M} \in \mathbb{R}^{d \times d}. \end{aligned}$$

Assume that the base-point portfolio $\mathcal{B}$ is $\mathcal{T}(\bar{\mathbf{x}})$ for a given $\bar{\mathbf{x}}$ (defined in Section III-C). The exponential tilting is defined by reweighting scenarios through an exponential score. For a scalar tilt parameter $t \in \mathbb{R}$, we define the exponentially tilted (Esscher) weights vector $\pi(t|\mathcal{B})$, having as s entry

$$\pi_s(t|\mathcal{B}) := \frac{w_s \exp(t z_s(\mathcal{B}))}{\sum_{s'=1}^S w_{s'} \exp(t z_{s'}(\mathcal{B}))}. \quad (3)$$

In a finite scenario set, Eq. (3) is a softmax reweighting of the base-point P&L $\mathbf{z}(\mathcal{B})$: as $|t|$ increases, the weights become more concentrated in extreme scenarios, while $t = 0$ recovers the base weights $\mathbf{w}$. For $t > 0$ the reweighting emphasizes positive P&L scenarios, whereas for $t < 0$ it emphasizes negative P&L scenarios.

We define the tilted mean vector $\bar{\mathbf{b}}_{t,\mathcal{B}} \in \mathbb{R}^d$ and the covariance matrix of scenario contributions $\bar{\mathbf{Q}}_{t,\mathcal{B}} \in \mathbb{R}^{d \times d}$ as the exponentially weighted sample mean and covariance

$$\bar{\mathbf{b}}_{t,\mathcal{B}} := \mathbb{E}_{\pi(t|\mathcal{B})}[\mathbf{G}], \quad \bar{\mathbf{Q}}_{t,\mathcal{B}} := \text{Cov}_{\pi(t|\mathcal{B})}(\mathbf{G}, \mathbf{G}).$$

Additional details are provided in Appendix A.

We can define the P&L *difference* with respect to the base-point

$$\mathbf{y}(\mathbf{x}) := \mathbf{z}(\mathcal{T}(\mathbf{x})) - \mathbf{z}(\mathcal{T}(\bar{\mathbf{x}})) = \mathbf{G}(\mathbf{x} - \bar{\mathbf{x}}). \quad (4)$$

We work with $\mathbf{y}(\mathbf{x})$ for three reasons: (i) the static term $\mathbf{z}(\mathcal{P})$ cancels out, so the tilted moments depend only on the decision-dependent part; (ii) under CARA preferences, translation invariance makes the comparison between portfolios naturally depending on P&L difference relative to a fixed base-point, so the absolute level $\mathbf{z}(\mathcal{T}(\bar{\mathbf{x}}))$ does not affect the solution of the exact optimization problem (see Section IV for details on CARA); (iii) with $\bar{\mathbf{x}}$ chosen as a candidate in the neighborhood of good solutions, deviations from $\bar{\mathbf{x}}$ under the tilted surrogate reflect the effect of exponential reweighting rather than correction of an inefficient starting portfolio $\mathcal{P}$.

Under $\pi(t|\mathcal{B})$, the mean and variance of $\mathbf{y}(\mathbf{x})$ are

$$\mathbb{E}_{\pi(t|\mathcal{B})}[\mathbf{y}(\mathbf{x})] = \bar{\mathbf{b}}_{t,\bar{\mathbf{x}}}^\top (\mathbf{x} - \bar{\mathbf{x}}), \quad (5)$$

$$\mathbb{V}_{\pi(t|\mathcal{B})}[\mathbf{y}(\mathbf{x})] = (\mathbf{x} - \bar{\mathbf{x}})^\top \bar{\mathbf{Q}}_{t,\bar{\mathbf{x}}} (\mathbf{x} - \bar{\mathbf{x}}). \quad (6)$$

These two quantities are used as coefficients of the quadratic objective function introduced in Section IV. The tilted weights depend on the entire set of base-point P&Ls $\mathbf{z}_s(\mathcal{B})$ through exponential reweighting, so varying t changes the tilted mean and covariance. In this sense, exponential tilting allows the resulting quadratic coefficients to retain information from the base-point distribution beyond the Gaussian case, although such features are not modeled explicitly.

C. Base-point portfolio as relaxed Markowitz solution

The Markowitz problem consists in finding $\mathbf{x}_M$ such that

$$\mathbf{x}_M = \arg \min_{\mathbf{x} \in \{0,1\}^d} \{\text{obj}_M(\mathcal{T}(\mathbf{x}))\}, \quad (7)$$

where

$$\text{obj}_M(\mathcal{T}(\mathbf{x})) := \frac{1}{2} \mathbb{V}_{\mathbf{m}}[\mathbf{z}(\mathcal{T}(\mathbf{x}))] - \mathbb{E}_{\mathbf{m}}[\mathbf{z}(\mathcal{T}(\mathbf{x}))]; \quad (8)$$

with $\mathbf{m}$ we denote the uniform probability measure, whose entries are all equal to $1/S$. Defining

$$\begin{aligned} \mu_{\mathbf{G}} &:= \mathbb{E}_{\mathbf{m}}[\mathbf{G}], \\ \Sigma_{\mathbf{G}\mathbf{G}} &:= \text{Cov}_{\mathbf{m}}(\mathbf{G}, \mathbf{G}), \\ \Sigma_{\mathbf{G}\mathcal{P}} &:= \text{Cov}_{\mathbf{m}}(\mathbf{z}(\mathcal{P}), \mathbf{G}), \end{aligned}$$

we can derive

$$\begin{aligned} \mathbb{E}_{\mathbf{m}}[\mathbf{z}(\mathcal{T}(\mathbf{x}))] &= \mathbb{E}_{\mathbf{m}}[\mathbf{z}(\mathcal{P})] + \mu_{\mathbf{G}}^\top \mathbf{x}, \\ \mathbb{V}_{\mathbf{m}}[\mathbf{z}(\mathcal{T}(\mathbf{x}))] &= \mathbb{V}_{\mathbf{m}}[\mathbf{z}(\mathcal{P})] + \mathbf{x}^\top \Sigma_{\mathbf{G}\mathbf{G}} \mathbf{x} + 2 \Sigma_{\mathbf{G}\mathcal{P}}^\top \mathbf{x}. \end{aligned}$$

We choose $\bar{\mathbf{x}}$ to be the continuous solution, found through quadratic programming, of the relaxed Markowitz problem

$$\bar{\mathbf{x}} = \arg \min_{\mathbf{x} \in [0,1]^d} \{\text{obj}_M(\mathcal{T}(\mathbf{x}))\}. \quad (9)$$

The rationale for this choice and alternative base-point selections are discussed in Appendix B.

IV. QUADRATIC OBJECTIVE AS A SECOND-ORDER CARA CERTAINTY EQUIVALENT

This section introduces the tilted quadratic objective function proposed in this work: a second-order approximation of a CARA certainty equivalent. The key idea is to optimize the metrics related to the P&L difference relative to the chosen base-point portfolio $\mathcal{B}$, while measuring both expected gain and residual dispersion under the Esscher-tilted weights.

We proceed as follows: (i) we introduce the certainty equivalent as a cash-valued scalar criterion that maps P&L distributions to comparable scores; (ii) we motivate the specific choice of CARA (exponential) preferences; (iii) we show that the exact CARA certainty equivalent is non-quadratic and derive the second-order (mean–variance) surrogate.

A. Certainty equivalent as a cash-valued score

The optimization variable $\mathbf{x} \in \{0,1\}^d$ induces a distribution of the P&L difference across scenarios, $\mathbf{y}(\mathbf{x})$, under the probability measure $\pi(t|\mathcal{B})$ (hereafter also denoted respectively by $\mathbf{y}$ and π). To decide whether one portfolio is preferable to another, we therefore need a scalar criterion that maps each such P&L distribution to a single cash-valued score (e.g., in euros), so that different portfolios can be compared by one number.

A standard approach to obtain such a scalar score is through expected-utility. Rather than optimizing expected utility directly, we use its CE: the deterministic cash amount

that delivers the same expected utility as the risky P&L, providing a single monetary score that summarizes the entire distribution of P&L difference across scenarios [13].

In our context, given a utility function $U_\gamma : \mathbb{R} \rightarrow \mathbb{R}$, i.e., a numerical representation of the decision-maker's preferences over monetary outcomes under uncertainty, and a probability measure $\mathbf{w}$, the certainty equivalent $\text{CE}_\gamma^\mathbf{w}(\mathbf{y})$ is defined as the sure cash amount c that satisfies

$$U_\gamma(c) = \mathbb{E}_\mathbf{w}[U_\gamma(\mathbf{y})], \quad (10)$$

where U_γ is applied to each component of the vector $\mathbf{y}$. Each y_s is viewed as a possible outcome of a discrete random quantity on the finite scenario set $\{1, \dots, S\}$ under $\mathbf{w}$. Maximizing $\text{CE}_\gamma^\mathbf{w}(\mathbf{y})$ is equivalent to selecting the portfolio with the largest risk-adjusted cash equivalent value under the chosen probability measure $\mathbf{w}$, i.e., the highest guaranteed cash amount that the investor would accept instead of the risky P&L difference. In particular, when the utility function is concave, it penalizes unfavorable outcomes more than it rewards favorable ones, even when their expected P&L is similar; therefore, portfolios with higher downside risk have lower CE.

We use a CE-based objective function because it is (i) *monetary*: one score in, e.g., euro; (ii) *risk-aware*: it penalizes unfavorable outcomes through the curvature of U_γ ; (iii) *compatible with scenario reweighting*: changing $\mathbf{w}$ (here via Esscher tilting) changes the expectation operator, without altering the optimization pipeline.

B. Rationale for CARA utility

In our setting, three operational requirements narrow the choice of the utility function.

- 1) **Cash-valued output.** The criterion must produce a score in the same monetary unit as the P&L (e.g., EUR), to maintain direct comparability between different portfolios.
- 2) **Translation invariance.** Adding an amount c to the P&L should shift the evaluation by exactly c : $\text{CE}_\gamma^\mathbf{w}(\mathbf{y} + c\mathbf{1}) = \text{CE}_\gamma^\mathbf{w}(\mathbf{y}) + c$, where $\mathbf{1}$ is the vector whose each component is one. This ensures that the optimal solution to our optimization problem does not depend on the fixed base-point portfolio $\mathcal{B}$ defined by $\bar{\mathbf{x}}$.
- 3) **Well-defined for negative outcomes.** The P&L difference can be negative (the new portfolio may underperform the base-point in some scenarios), so the utility must also be defined for negative values.

A formulation that satisfies all three requirements simultaneously is provided by the CARA family

$$U_\gamma(y) := -\exp(-\gamma y), \quad y \in \mathbb{R}, \quad \gamma > 0, \quad (11)$$

where the parameter γ is the absolute risk-aversion coefficient (units: e.g., $1/\text{EUR}$). The utility function in Eq. (11) is well-defined for negative outcomes, the closed-form CE (derived in Section IV-C) provides a cash-valued output in monetary units and the translation invariance follows from the exponential structure.

C. From exact CARA CE to the quadratic objective function

Substituting the CARA utility (11) into the CE definition (10) and solving for c yields the closed form

$$\text{CE}_\gamma^\mathbf{w}(\mathbf{y}) = -\frac{1}{\gamma} \log \mathbb{E}_\mathbf{w}[\exp(-\gamma \mathbf{y})], \quad (12)$$

where the exponential function is applied component-wise.

To use the CE based on CARA utility as an objective function suitable for QUBO, we develop a quadratic approximation of Eq. (12). For a fixed tilted measure $\pi(t|\mathcal{B})$, we consider $\text{CE}_\gamma^{\pi(t|\mathcal{B})}(\mathbf{y})$ as a function of the risk-aversion parameter γ and use the usual local expansion for small risk aversion around $\gamma = 0$, i.e., around the risk-neutral limit in which the certainty equivalent reduces to the expected P&L; see, e.g., [13, 10]. Indeed,

$$\mathbb{E}_\pi[\exp(-\gamma \mathbf{y})] = 1 - \gamma \mathbb{E}_\pi[\mathbf{y}] + \frac{\gamma^2}{2} \mathbb{E}_\pi[\mathbf{y}^2] + \mathcal{O}(\gamma^3)$$

where $\mathbf{y}^2$ is the vector whose components are the square of the components of $\mathbf{y}$. Applying $\log(1+v) = v - \frac{v^2}{2} + \mathcal{O}(v^3)$ with

$$v = -\gamma \mathbb{E}_\pi[\mathbf{y}] + \frac{\gamma^2}{2} \mathbb{E}_\pi[\mathbf{y}^2] + \mathcal{O}(\gamma^3),$$

gives $v^2 = \gamma^2 (\mathbb{E}_\pi[\mathbf{y}])^2 + \mathcal{O}(\gamma^3)$, and therefore

$$\begin{aligned} \log \mathbb{E}_\pi[\exp(-\gamma \mathbf{y})] &= \\ &= -\gamma \mathbb{E}_\pi[\mathbf{y}] + \frac{\gamma^2}{2} \mathbb{E}_\pi[\mathbf{y}^2] - \frac{\gamma^2}{2} (\mathbb{E}_\pi[\mathbf{y}])^2 + \mathcal{O}(\gamma^3) = \\ &= -\gamma \mathbb{E}_\pi[\mathbf{y}] + \frac{\gamma^2}{2} \text{Var}_\pi(\mathbf{y}) + \mathcal{O}(\gamma^3). \end{aligned} \quad (13)$$

Substituting (13) into (12) and omitting $\mathcal{O}(\gamma^2)$, we obtain

$$\text{CE}_\gamma^{\pi(t|\mathcal{B})}(\mathbf{y}(\mathbf{x})) \approx \mathbb{E}_{\pi(t|\mathcal{B})}[\mathbf{y}(\mathbf{x})] - \frac{\gamma}{2} \text{Var}_{\pi(t|\mathcal{B})}[\mathbf{y}(\mathbf{x})]. \quad (14)$$

By maximizing this quantity, configurations with a higher tilted mean and lower tilted variance are rewarded, in coherence with the Markowitz approach. Higher-order distributional features (skewness, kurtosis, ...) are not explicitly modeled in the quadratic surrogate; their influence enters only indirectly through the tilted weights $\pi(t|\mathcal{B})$ used to compute the two moments in Eq. (14).

The expansion is locally accurate when γ is small relative to the scale of $\mathbf{y}$ components, as in the case of our experimental settings. Accordingly, we use the resulting tilted mean-variance form as a quadratic surrogate of the exact CARA certainty equivalent under the tilted measure. The resulting objective is a standard and self-consistent quadratic criterion for any $\gamma > 0$: the second-order expansion motivates the functional form and clarifies its economic content (reward versus penalty structure), but the quadratic surrogate does not rely on γ being small in order to define a well-posed optimization problem.

It is important to note that the framework involves two parameters with distinct roles:

- the tilt t defines the measure $\pi(t|\mathcal{B})$, i.e., which scenarios are emphasized (upside for $t > 0$, downside for $t < 0$);
- the risk-aversion coefficient γ controls the tilted mean-variance trade-off once the measure is fixed: larger γ emphasizes the second term of Eq. (14).

In our experiments, we couple the two parameters via $\gamma = |t|$. Substituting the tilted mean of Eq. (5) and variance of Eq. (6) into Eq. (14) yields the objective function to be minimized

$$\begin{aligned} \text{obj}_t(\mathcal{T}(\mathbf{x})) &:= -\bar{\mathbf{b}}_{t,\mathcal{B}}^\top (\mathbf{x} - \bar{\mathbf{x}}) \\ &\quad + \frac{|t|}{2} (\mathbf{x} - \bar{\mathbf{x}})^\top \bar{\mathbf{Q}}_{t,\mathcal{B}} (\mathbf{x} - \bar{\mathbf{x}}), \end{aligned} \quad (15)$$

where the tilted mean vector $\bar{\mathbf{b}}_{t,\mathcal{B}}$ and the covariance matrix $\bar{\mathbf{Q}}_{t,\mathcal{B}}$ are the weighted sample averages defined in Section III-B.

Under this coupling, more severe tilting (larger $|t|$) is accompanied by a stronger dispersion emphasis, reflecting the view that a practitioner choosing a more extreme scenario reweighting may pay more attention to risk. Moreover, the coupling $\gamma = |t|$ does not imply symmetry between positive and negative tilts, as the measure $\pi(t|\mathcal{B})$ depends on the sign of t through the exponential reweighting of scenarios. Consequently, the tilted moments $\bar{\mathbf{b}}_{t,\mathcal{B}}$ and $\bar{\mathbf{Q}}_{t,\mathcal{B}}$ generally differ for t and $-t$, and the resulting optimization problems need not produce symmetric solutions.

In our experiments, the tilt grid is restricted to moderate values of $|t|$, to remain in the same qualitative regime as the second-order surrogate in Eq. (14).

V. QUBO FORMULATION

In Section V-A we present a QUBO formulation of the tilted quadratic objective function defined in Eq. (15), and in the subsequent Sections (V-B, V-C, V-D) we then show how to incorporate as penalty terms the business constraints defined in Section II-A.

A. QUBO objective function

We begin by deriving a QUBO representation of the tilted quadratic objective function defined in Eq. (15), which serves as the foundation for the optimization framework.

A QUBO problem consists of finding $\mathbf{x}_{\text{opt}}$ such that

$$\mathbf{x}_{\text{opt}} = \arg \min_{\mathbf{x} \in \{0,1\}^d} \{ \mathbf{x}^\top \mathbf{Q}_{\text{obj}} \mathbf{x} \}. \quad (16)$$

where $\mathbf{Q}_{\text{obj}} \in \mathbb{R}^{d \times d}$ is a symmetric matrix.

In our context, $\mathbf{x} \in \{0,1\}^d$ is a binary decision vector whose component x_u takes the value 1 if a specific UEI and its notional value is included in the EOS and 0 otherwise (as discussed in Section III-A). Up to constant terms, the optimization problem based on the objective function $\text{obj}_t(\mathcal{T}(\mathbf{x}))$ in Eq. (15) can be written, by expanding and collecting linear and quadratic terms, as

$$\mathbf{x}_t = \arg \min_{\mathbf{x} \in \{0,1\}^d} \left\{ \mathbf{x}^\top \left(\frac{|t|}{2} \bar{\mathbf{Q}}_{t,\mathcal{B}} - \text{diag}(\mathbf{b}) \right) \mathbf{x} \right\}, \quad (17)$$

that satisfies the formulation of Eq. (16), where

$$\mathbf{b} = \bar{\mathbf{b}}_{t,\mathcal{B}} + |t| \bar{\mathbf{Q}}_{t,\mathcal{B}} \bar{\mathbf{x}}. \quad (18)$$

In this formulation, minimizing the objective function can be interpreted as selecting specific UEIs, together with a notional value, to be added to the initial portfolio $\mathcal{P}$ to obtain an optimized final portfolio configuration $\mathcal{T}$.

B. EOS composition constraint

Let I denote the set of indices corresponding to instruments of a given type i and associated with the same underlying asset. If n_i is the maximum allowed number of instruments of type i , the constraint can be written as

$$\sum_{u \in I} x_u = n_i.$$

It is worth noting that this equality implicitly imposes an inequality constraint. In fact, the x_u in the above expression may correspond to a notional value equal to zero, which is equivalent to not selecting an instrument of the given type.

This constraint can be mapped into the QUBO framework through the quadratic penalty term

$$\left(\sum_{u \in I} x_u - n_i \right)^2, \quad (19)$$

which is added to the objective function for each underlying asset and type of instrument i , scaled by an appropriate penalty coefficient.

C. UEI quantities constraints

Since the choice of a notional value for a UEI corresponds to set to 1 only a single binary variable among those in $\mathbf{x}$ associated with that UEI, we must ensure that

$$\sum_{u \in U} x_u \leq 1 \quad \forall U,$$

where U denotes the set of indices corresponding to different notional values of the same UEI. The total number of such sets U is N_{UEI} . The inequality is fundamental because enforcing equality would encourage the addition of the UEI into the EOS, thereby creating an inconsistency with the EOS composition constraint. With this formulation, the discretization of the notional values is naturally enforced.

This constraint can be formulated as a QUBO penalty term as follows

$$\sum_{u=1}^{N_{\text{notional}}-1} x_u \sum_{u'=u+1}^{N_{\text{notional}}} x_{u'} \quad \forall U. \quad (20)$$

D. Greeks constraints

The Greeks measure the sensitivity of the portfolio value to small changes in underlying risk factors. The Greeks of our interest are Delta (Δ), Vega ($\mathcal{V}$), and Gamma (Γ). In practice, they encode the risk profile of the initial portfolio $\mathcal{P}$ decided by traders, according to their market view. The Greeks constraints prescribe their maximum allowed variation, to control the change in the risk profile of the optimized portfolio $\mathcal{T}$ with respect to $\mathcal{P}$.

Let $\mathcal{G}_{\mathcal{P}}$ and $\mathcal{G}_{\mathcal{T}}$ denote the value of one of the Greeks of the portfolios $\mathcal{P}$ and $\mathcal{T}$, respectively, $\text{tol}_{\mathcal{G}}$ the maximum allowed percentage variation of $\mathcal{G}_{\mathcal{P}}$, and g_u the Greek value associated with the u -th pair UEI-notional value. The constraint is

$$|\mathcal{G}_{\mathcal{T}} - \mathcal{G}_{\mathcal{P}}| = \left| \sum_u g_u x_u \right| < C_{\mathcal{G}} \quad (21)$$

where $C_{\mathcal{G}} := \text{tol}_{\mathcal{G}} |\mathcal{G}_{\mathcal{P}}|$.

We consider two strategies to embed this constraint in the formulation: slack variables and Lagrangian relaxation. The first approach adopts the **slack variable** technique [12] to reformulate the inequality constraint as a quadratic penalty term

$$\left(\sum_u g_u x_u - C_{\mathcal{G}} + V \right)^2, \quad (22)$$

where $V = \sum_{j=0}^m 2^j v_j$ is a slack variable, discretized in binary form, whose resolution controls the desired precision of the Greek constraint. The binary slack variable components v_j must be added to the decision variables of the QUBO problem, extending the solution space by $m = \lceil \log_2(2C_{\mathcal{G}} + 1) \rceil$. Besides the increased dimensionality of the QUBO problems, note that we have three Greek constraints for each problem instance, and therefore three sets of binary slack variables. Expressing the

constraints in this way introduces additional penalty parameters whose appropriate values are not trivial to determine to reach optimality.

For this reason, we introduce an alternative method to incorporate these constraints into the QUBO objective function, based on **Lagrangian relaxation** [8]. This approach relies on treating the constraint as a *complicating constraint* and incorporating it into the objective function via a multiplier. To do so, we need to express the constraint as $l(\mathbf{x}) \leq 0$ for some function l that must be quadratic to be embedded in the objective function. The absolute value in Eq. (21) is addressed by squaring the constraint, yielding

$$l(\mathbf{x}) := \left(\frac{\sum_u g_u x_u}{C_G} \right)^2 - 1,$$

where we divide by C_G for numerical stability. Now, we define the final objective function as

$$\text{obj}_\lambda(\mathcal{T}(\mathbf{x})) = \mathbf{x}^T \mathbf{Q}_{\text{pen}} \mathbf{x} + \lambda l(\mathbf{x}) \quad (23)$$

for some positive multiplier λ that maximizes

$$f(\lambda) := \min_{\mathbf{x} \in \{0,1\}^d} \text{obj}_\lambda(\mathcal{T}(\mathbf{x})). \quad (24)$$

Here, the matrix $\mathbf{Q}_{\text{pen}}$ incorporates both the objective component $\mathbf{Q}_{\text{obj}}$ and the penalty terms associated with the constraints discussed in Sections V-B and V-C. To find an adequate value for λ , we adopt an iterative gradient-based method [8]: at the i -th iteration, given some $\lambda = \lambda^{(i)}$, we use a QUBO solver to find $\mathbf{x}^{(i)} \in \{0,1\}^d$ that minimizes $\text{obj}_{\lambda^{(i)}}(\mathcal{T}(\mathbf{x}))$. The function l evaluated in $\mathbf{x}^{(i)}$ represents a valid supergradient of f in $\lambda^{(i)}$, since it can be shown that

$$f(\lambda) - f(\lambda^{(i)}) \leq l(\mathbf{x}^{(i)})(\lambda - \lambda^{(i)}),$$

for every $\lambda > 0$. So, the next multiplier is

$$\lambda^{(i+1)} = \max(0, \lambda^{(i)} + \mu^{(i)} l(\mathbf{x}^{(i)}))$$

where $\mu^{(i)}$ is a properly chosen iteratively decreasing factor (see Section VI).

The same approach can be extended to manage all Greeks, by introducing a Lagrangian multiplier (λ_Δ , λ_Γ and λ_ν) and a constraint function (l_Δ , l_Γ and l_ν) for each Greek. Hence, the complete Lagrangian dual problem is

$$\max_{\lambda_\Delta, \lambda_\Gamma, \lambda_\nu \geq 0} \min_{\mathbf{x} \in \{0,1\}^d} \left(\mathbf{x}^T \mathbf{Q}_{\text{pen}} \mathbf{x} + \sum_{k \in \{\Delta, \Gamma, \nu\}} \lambda_k l_k(\mathbf{x}) \right).$$

Note that Lagrangian multipliers are optimized simultaneously.

VI. EXPERIMENTS

In our experimental analysis, we first focused on validating the proposed tilting framework, assessing the impact that different tilt values have on the optimization process. We adopted the constrained formulation of the optimization problem and a reduced dataset to provide greater guarantees of obtaining a feasible, optimal solution using MIQP.

The results were analyzed to identify a subset of representative tilts used to test the QUBO formulation. Each constrained tilted model was translated into its equivalent unconstrained representation and executed using MIQP and the Simulated Bifurcation quantum-inspired method.

To assess scalability, we further investigated the computational tractability of larger problem instances under MIQP.

Using the full dataset and the Markowitz formulation, we evaluated the execution times of MIQP solvers (for both constrained and unconstrained versions of the problem) and of the Simulated Bifurcation heuristic.

A. Dataset and problem setting

For our tests, we adopted the same dataset and settings as used in [3]. The dataset consists of 8 stocks and 5 stocks indexes. For each underlying, the unique eligible instruments are identified by a choice of instrument type (call, put, future, or stock), strike, and maturity. See [3, Table 3] for details. To replicate the setting, we build $\mathcal{T}$ by selecting at most two options and one stock/future per underlying asset. For sensitivities, we restrict our experiments to a maximum fixed percentage variation $\text{tol}_\Delta = \text{tol}_\Gamma = \text{tol}_\nu = 100\%$.

The number N_{notional} of discretized notional values for each underlying is 21, corresponding to 10 positive amounts, 10 negative and 0, all equally spaced within the same interval centered at 0. In other words, notional values are discretized using a step that is a tenth of the maximal allowed notional value. Most of our experiments were performed replicating the small-sized application setting of [3] that limits the dataset to only one underlying: the Euro Stoxx 50 Index (.STOXX50E). Hereafter, the subset of all unique eligible instruments associated to this underlying is named *reduced* dataset. The combinations of instrument type, strike, and maturity for .STOXX50E are 42, yielding a dimension $d = 42 \times 21 = 882$. However, some final experiment (see Section VI-D) has been performed using the complete dataset, hereafter named *full* dataset, for a total of 556 UEIs and a global dimension $d = 21 \times 556 = 11\,676$.

Actually, for numerical reasons, the dataset was pre-processed by dividing all entries in $\mathbf{G}$ by the standard deviation $\sigma_{\mathbf{G}}$, which is 19 051.49 for the reduced dataset and 18 473.17 for the full dataset. Adopting the proposed framework using a normalized dataset is formally equivalent to using the original data and scaling tilts by the same factor $\sigma_{\mathbf{G}}$. The reader should be aware that the tilt values reported in this paper refer to the normalized dataset.

Finally, in accordance with the Markowitz formulation (see Section III-C), the base probability measure $\mathbf{w}$ applied to scenarios before tilting is the uniform distribution $\mathbf{m}$.

B. Tilting framework validation

We considered all tilts on a grid with a step size of 0.025, ranging from -0.1 to 0.2 , plus some more tilt in some sub-interval that required further investigation. For each tilt in this set, we employed GUROBI to solve the problem defined in Eq. (17), subject to the constraints described in Section V handled explicitly and not as penalties. As already mentioned in Section III-C, the base-point $\bar{\mathbf{x}}$ required by the model is chosen to be the continuous solution of the mean-variance problem. In order to find such a solution, we adopted CVXOPT through QPSOLVERS [4], which has proven to be very fast and essentially negligible in terms of time during the execution of the entire process.

Each solution $\mathbf{x}_t$ (one for each tilt t) was used to build a new portfolio $\mathcal{T} = \mathcal{T}(\mathbf{x}_t)$. The P&L distribution of each portfolio was then primarily evaluated in terms of its mean, standard deviation, VaR, and CVaR at 99% and 97.5% levels. Figure 1 shows the trends of the mean (in blue), standard deviation (in orange), VaR (in green) and CVaR at 99% (in red) of the solutions as the tilt varies. The metrics values are

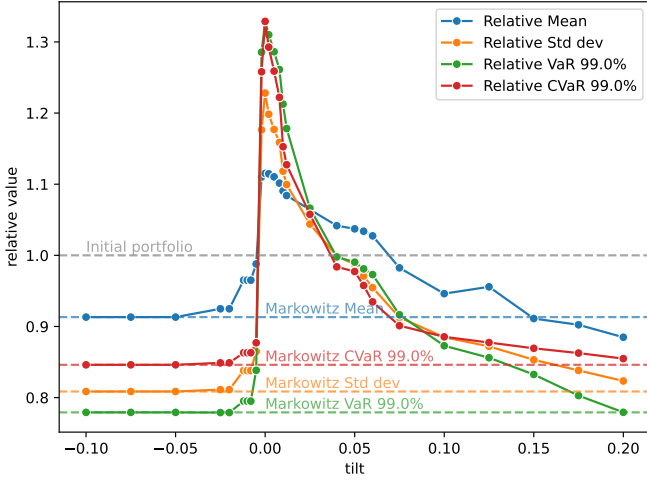


Fig. 1: Plot of relative (with respect to the initial portfolio) mean, standard deviation, VaR and CVaR (at 99% level), in dependence of tilts values. For the sake of clarity, VaR and CVaR at 97.5% level are not presented, since they behave very similarly to VaR and CVaR at 99%.

normalized by dividing by the corresponding value of the initial portfolio $\mathcal{P}$. Therefore, a normalized mean greater than one indicates an improved performance relative to $\mathcal{P}$. For VaR and CVaR, the interpretation requires care. In our convention, these risk metrics are negative, and an improvement corresponds to values closer to zero (smaller absolute magnitude). Hence, if both the VaRs of $\mathcal{P}$ and $\mathcal{T}$ are $\text{VaR}_{\mathcal{P}} < 0$ and $\text{VaR}_{\mathcal{T}} < 0$, the normalized ratio $\text{VaR}_{\mathcal{T}}/\text{VaR}_{\mathcal{P}}$ is positive and an improved (less negative) $\text{VaR}_{\mathcal{T}}$ generates a ratio below one (and analogously for CVaR). In this figure, the initial portfolio metrics are highlighted by the gray dashed line, and all have values equal to 1. Instead, colored dashed lines correspond to the same metrics evaluated on the solution of the Markowitz problem (see Eq. (7)).

All metrics shown exhibit similar behavior, flattening toward the edges of the grid and displaying a cusp-like maximum at zero. This latter feature was expected *a priori* and is due to the factor $|t|$ that multiplies the quadratic component: when t is zero, only the linear part, i.e. the mean, is optimized, as the quadratic risk term does not contribute.

Starting from zero and moving toward negative tilts, we observe a steep slope. This is because the quadratic component gradually gains importance, and negative tilts emphasize adverse scenarios; consequently, the optimization process focuses on minimizing losses in these scenarios rather than maximizing returns. Conversely, as the weight of the quadratic component continues to increase for positive tilts, these emphasize favorable scenarios where achieving higher returns is more advantageous than simply reducing risk. Therefore, the slope is less steep on this side.

The flattening observed at the edges is expected, since extreme tilts effectively select a very limited number of scenarios through exponential tilting. This becomes particularly intuitive when considering the limiting behavior: for very large tilts, the weighting scheme concentrates on a single scenario, and the solution to the corresponding optimization problems therefore becomes invariant. For negative tilts, flattening occurs rapidly and the solution converges to the Markowitz optimum.

This behavior was unexpected and is likely related to the specific characteristics of the dataset used in the experiments. In practice, solving the Markowitz problem on this dataset appears equivalent to solving it on an effectively single-scenario dataset, namely, the one emphasized by exponential tilting.

Moreover, it emerged experimentally that with this particular reduced dataset, the quadratic part in the Markowitz objective function is much more significant than the linear part and, since we are constrained to choose just a limited number of UEs to replicate the setting of [3], solving the Markowitz problem is indeed equivalent to minimize the variance. This explains why, varying the tilts, we cannot obtain any solution reducing the standard deviation with respect to the Markowitz solution.

In summary, the figure shows that varying the tilts produces solutions targeting different P&L objectives. This can also be observed in Figures 2 and 3, where solutions with tilts -0.02 and 0.055 exhibit markedly different, but interesting, features. Indeed, we identified two neighborhoods around these values, namely $[0.04, 0.06]$ and $[-0.025, -0.005]$, that preserve this behavior and are therefore worth further investigation.

Let us start from the interval $[0.04, 0.06]$. The solutions for these tilts differ from the Markowitz one, and in particular they perform worse with respect to risk-related metrics (such as standard deviation, VaR and CVaR), whereas the Markowitz solution delivers a smaller return. However, they perform better than the initial portfolio with respect to every considered metric. This fact can be observed in Table I and Figure 1, where the blue points of the interval exceed 1, while all other points lie below 1.

The solutions found for tilts in the interval $[-0.025, -0.005]$, on the other hand, are intermediate between the initial portfolio, which is more return-oriented, and the mean-variance solution, which is more conservative. However, they are characterized by a higher mean-to-standard deviation ratio. Indeed, these solutions also exhibit good mean-to-VaR ratios, as shown in Table I. The table also reports the performances obtained through the RATPO algorithm, which slightly outperforms others with respect to the mean-to-VaR ratio at the confidence level 99%, this is expected since both metrics contribute to the construction of the RATPO fractional objective function. However, it performs worse according to all other metrics than the solutions in the interval.

C. Validating QUBO formulation

Once we have validated the theoretical tilting framework, we solved again the problem in the QUBO formulation for some of the most interesting tilts: -0.005 , -0.02 , 0.04 and 0.055 .

We compared the two strategies for handling the Greeks constraints presented in Section V-D: a slack variable encoding and a Lagrangian relaxation approach, in the latter case we used an update step $\mu^{(i)} = 1/\sqrt{i}$ satisfying the standard convergence requirements of the supergradient method [8].

We tested both methods in combination with two different QUBO solvers: GUROBI's wrapper for QUBO formulations [15], hereafter named GUROBI_Q, and Simulated Bifurcation [1]. Furthermore, as a baseline for optimality, we relied on the constrained formulation of the MIQP problem, discussed in Section VI-B.

The experiments were conducted on two Ubuntu 22.04.3 LTS machines equipped with AMD EPYC 7763 CPUs. The

Instance	Mean	Std dev	VaR 99.0%	CVaR 99.0%	VaR 97.5%	CVaR 97.5%	Mean/Std	Mean/VaR99%	Mean/VaR97.5%
Initial	33 918	302 757	-735 750	-775 588	-714 031	-755 669	0.1120	-0.0461	-0.0475
RATPO	32 634	255 585	-582 060	-677 105	-581 533	-629 582	0.1277	-0.0561	-0.0561
Markowitz	30 974	244 804	-573 327	-656 230	-535 283	-614 779	0.1265	-0.0540	-0.0579
t=-0.1	30 974	244 804	-573 327	-656 230	-535 283	-614 779	0.1265	-0.0540	-0.0579
t=-0.075	30 974	244 804	-573 327	-656 230	-535 283	-614 779	0.1265	-0.0540	-0.0579
t=-0.05	30 974	244 804	-573 327	-656 230	-535 283	-614 779	0.1265	-0.0540	-0.0579
t=-0.025	31 371	245 604	-573 145	-658 376	-541 731	-615 761	0.1277	-0.0547	-0.0579
t=-0.02	31 371	245 604	-573 145	-658 376	-541 731	-615 761	0.1277	-0.0547	-0.0579
t=-0.012	32 734	253 705	-584 977	-669 412	-575 860	-627 194	0.1290	-0.0560	-0.0568
t=-0.01	32 734	253 705	-584 977	-669 412	-575 860	-627 194	0.1290	-0.0560	-0.0568
t=0.0	37 825	371 832	-977 771	-1 030 644	-808 508	-1 004 207	0.1017	-0.0387	-0.0468
t=0.025	36 100	316 066	-784 533	-820 297	-744 161	-802 415	0.1142	-0.0460	-0.0485
t=0.04	35 333	302 494	-734 122	-763 045	-713 917	-748 584	0.1168	-0.0481	-0.0495
t=0.05	35 182	298 851	-728 705	-757 940	-696 159	-743 323	0.1177	-0.0483	-0.0505
t=0.055	35 068	293 951	-721 733	-742 723	-665 803	-732 228	0.1193	-0.0486	-0.0527
t=0.06	34 845	288 977	-715 806	-724 739	-638 419	-720 273	0.1206	-0.0487	-0.0546
t=0.075	33 322	276 354	-674 291	-698 916	-597 866	-686 603	0.1206	-0.0494	-0.0557
t=0.1	32 137	268 027	-643 121	-686 775	-581 849	-664 948	0.1199	-0.0500	-0.0552
t=0.125	32 417	264 024	-629 867	-680 532	-578 225	-655 199	0.1228	-0.0515	-0.0561
t=0.15	30 905	258 297	-612 509	-674 232	-576 940	-643 370	0.1196	-0.0505	-0.0536
t=0.175	30 607	253 780	-590 695	-669 025	-576 633	-629 860	0.1206	-0.0518	-0.0531
t=0.2	30 012	249 288	-573 406	-662 992	-572 541	-618 199	0.1204	-0.0523	-0.0524

TABLE I: Metrics of a selection of solutions found by the proposed framework compared to mean–variance solution, RATPO solution and the initial portfolio. Here, all columns, but the last three that are dimensionless, are in Euros and represent actual P&L figures. Solutions for tilts in the interval $[0.04, 0.06]$ outperform the initial portfolio with respect to all displayed metrics. Interval $[-0.025, -0.005]$ instead presents solutions having high mean-standard deviation and mean-VaR ratios. Highlighted solutions are represented also on Figures 2 and 3.

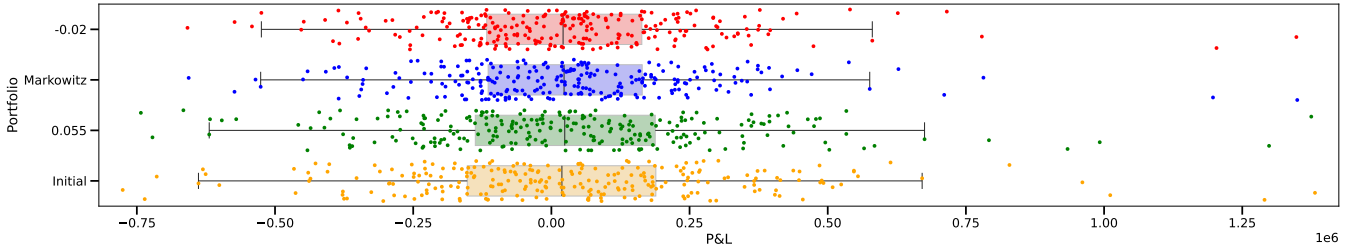


Fig. 2: Box plots comparing the P&L distributions of the initial portfolio (orange), the Markowitz solution (blue), and the solutions for tilts 0.055 (green) and -0.02 (red), illustrating how different tilts capture different distribution features.

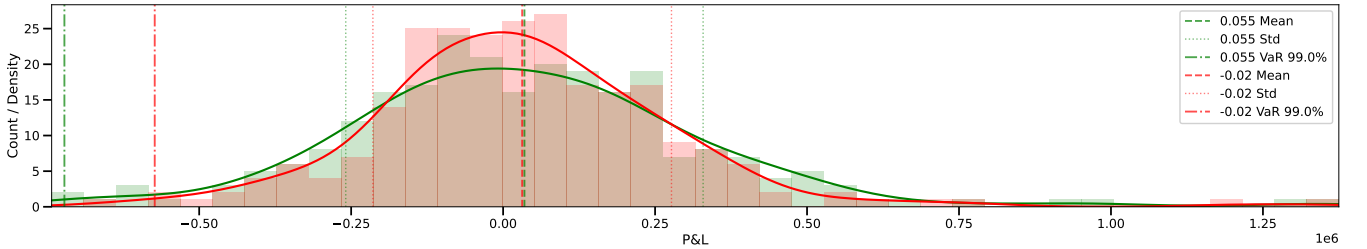


Fig. 3: Histograms comparing the P&L distribution of the solution for tilts 0.055 (in green) and -0.02 (in red). Vertical lines report the mean, the VaR at 99.0% and the mean $\pm$ the standard deviation. It should be noted that while the former shows a higher mean (i. e. expected return), the latter reduces risk-related metrics.

GUROBI solver was run on a CPU-only machine (Machine A) with 15 CPU cores and 57 GB RAM, while the Simulated Bifurcation solver was executed on a GPU-enabled machine (Machine B) with 7 CPU cores, 28 GB RAM, and an NVIDIA A30 GPU.

The slack variable formulation proved to be ineffective in the tested instances. Despite exploring multiple values for the penalty coefficients associated with the Greeks constraints and attempting to scale them with respect to the objective magnitude, the QUBO solvers were unable to obtain satisfactory solutions within a time limit of 1 hour. In particular,

GUROBI_Q consistently returned suboptimal solutions, while Simulated Bifurcation failed to find feasible solutions. A detailed report of the results obtained with this methodology can be found in Appendix C.1.

On the other hand, Lagrangian relaxation proved to be quite efficient, obtaining the optimal solutions for all tilts except $t = -0.02$, when solving with GUROBI_Q. The solutions were found in a few minutes, the longest run being for 0.055 which took less than 15 minutes (see Table II). The Simulated Bifurcation solver proved to be less reliable and exhibited convergence behavior that was difficult to control. In general, it

Instance	GUROBI_Q		GUROBI MIQP		Simulated Bifurcation	
	time	optimality gap	time	optimum	time	optimality gap
Markowitz (reduced)	9.6 min	0.00%	< 10s	-43.5591	>1.5 h	unfeasible solution
$t = -0.02$	20 s	1.77%	< 10s	-0.7100	30 min	0.00%
$t = -0.005$	11.3 min	0.00%	< 10s	-0.1368	>1.5 h	unfeasible solution
$t = 0.04$	4.5 min	0.00%	< 10s	-0.0710	17 min	0.00%
$t = 0.055$	6.6 min	0.00%	< 10s	-0.0847	36 min	1.50%
Markowitz (full)	2.8 min	0.85%	15 min	-117.47*	>1.5 h	unfeasible solution

TABLE II: Comparison of execution times of the constrained optimization problem (GUROBI MIQP) and the unconstrained versions using Lagrangian relaxation (GUROBI_Q and Simulated Bifurcation). The optimum column displays the optimal value if available, otherwise the best known bound*, found with the GUROBI MIQP solver. The optimality gap, for GUROBI_Q and Simulated Bifurcation, is defined as the relative difference between the best objective value found and the optimum of GUROBI MIQP. Times reported for GUROBI_Q/Simulated Bifurcation, correspond to the execution time required to obtain the solution used in the optimality gap evaluation; however, the solvers were allowed to run for longer ($\approx 30/100$ min).

underperformed GUROBI_Q in most cases, with the exception of $t = -0.02$, for which it identified an optimal solution that GUROBI_Q did not obtain. In our experiments, we considered up to 200 Lagrangian dual iterations; each QUBO subproblem was solved with a time limit of 10 seconds using GUROBI_Q and 30 seconds using the Simulated Bifurcation solver.

Overall, the best QUBO configuration was Lagrangian relaxation combined with GUROBI_Q, which recovered all optimal solutions except for one tilted instance in the reduced dataset. See Appendix C.2 for further details.

D. Scaling-up

All tests that have been discussed so far were conducted using a reduced dataset consisting of a single underlying (see Section VI-A). To assess the impact of increased problem dimensionality on the runtime of our model, and to evaluate the effective need for quantum resources for bigger instances, we attempted to solve the mean-variance problem formulated by Eq. (7) using the QUBO formulation on the full dataset. This setting increased the dimension of the solution vector $\mathbf{x}$ from 882 to 11 676.

One of the main reasons for choosing GUROBI was its high efficiency when using the reduced dataset, consistently finding constrained solutions within a few seconds.

However, even GUROBI struggles to find a solution to constrained instances of problems using the complete dataset. In particular, on Machine B (32 GB of RAM), the MIQP solver ran out of memory without producing any solution, while GUROBI_Q was able to find a suboptimal solution. On Machine A (57 GB of RAM), GUROBI MIQP solver again ran out of memory before the optimal solution to the constrained problem was found, after approximately 15 minutes, returning a high-quality solution. It is important to note that the pre-solving phase requires about 10 minutes, thus imposing a lower bound on the total runtime. Alternatively, GUROBI_Q with Lagrangian relaxation is capable of providing a good quality solution, although not as good as that found by GUROBI in the constrained formulation, in 2.8 minutes and could be run longer to refine such a solution without encountering out-of-memory issues.

This evidence suggests that, as the problem dimension grows, the QUBO approach may scale better than the constrained one. The constrained formulation becomes increasingly limited by GUROBI memory usage and pre-solve time, whereas the QUBO route remains practically viable.

Notably, even if Simulated Bifurcation failed to find a feasible solution for the full dataset, it proved to be a memory-

efficient approach compared to GUROBI-based methods. This suggests that this QUBO formulation could further benefit from more reliable or better tunable quantum-inspired solvers, where improved controllability and robustness may enhance scalability and convergence on large-scale instances. In fact, in the implementation of the Simulated Bifurcation we used, the relationship between hyperparameter settings and solution quality is not fully transparent, making systematic tuning challenging and often leading to inconsistent performance across instances.

In conclusion, the scaling tests show that the proposed Lagrangian QUBO formulation offers a compact and scalable alternative to the explicitly constrained MIQP approach. Moreover, the experiments highlight the demanding increase in resource requirements of traditional methodologies as the problem size grows, raising the opportunity to efficiently exploit quantum hardware. Avoiding the slack variables overhead, our methodology allows preserving the original number of binary decision variables that is theoretically equal to the number of qubits required. In practice, the exact amount of quantum resources needed depends on the mapping on physical hardware and therefore on the strategy and technology adopted.

VII. CONCLUSION AND FUTURE WORK

We presented a framework for portfolio optimization that combines exponential tilting of empirical scenarios with quadratic surrogate models compatible with QUBO solvers. The method constructs tilted moments via Esscher reweighting of base-point P&L, yielding a tractable QUBO representation of a locally approximated CARA certainty-equivalent objective that enables the use of classical and quantum-inspired optimization tools. Varying the tilt parameter induces controlled changes in moments, allowing the optimization to explore different regimes without altering the underlying scenario set. Experiments on a realistic trading portfolio showed that varying the tilt parameter produces portfolios with meaningfully different risk-return profiles. For the QUBO formulation, Lagrangian relaxation proved effective in encoding sensitivity inequality constraints without expanding the binary search space. A preliminary scaling test (from 882 to 11 676 variables) suggested that the QUBO route remains practically viable in larger instances, while constrained MIQP encountered memory limitations.

Future work includes structured multi-tilt constructions, data-driven tilt selection targeting out-of-sample performance, and execution on quantum hardware to assess whether the observed

QUBO scalability translates into practical benefit relative to classical and quantum-inspired methods.

REFERENCES

- [1] Romain Ageron, Thomas Bouquet, and Lorenzo Pugliese. *Simulated Bifurcation (SB) algorithm for Python*. Version 2.0.0. Apr. 2025. URL: <https://github.com/bqth29/simulated-bifurcation-algorithm>.
- [2] Esteban Aguilera et al. “Multi-Objective Portfolio Optimization Using a Quantum Annealer”. In: *Mathematics* 12.9 (2024). ISSN: 2227-7390. DOI: [10.3390/math12091291](https://doi.org/10.3390/math12091291). URL: <https://www.mdpi.com/2227-7390/12/9/1291>.
- [3] Marco Bianchetti et al. *Risk-aware Trading Portfolio Optimization*. 2025. arXiv: [2503.04662](https://arxiv.org/abs/2503.04662) [q-fin.RM]. URL: <https://arxiv.org/abs/2503.04662>.
- [4] Stéphane Caron et al. *qpsolvers: Quadratic Programming Solvers in Python*. Version 4.8.2. 2025. URL: <https://github.com/qpsolvers/qpsolvers>.
- [5] Amir Dembo and Ofer Zeitouni. *Large Deviations Techniques and Applications*. 2nd ed. Vol. 38. Stochastic Modelling and Applied Probability. Corrected reprint of the 1998 second edition. Berlin, Heidelberg: Springer, 2010. DOI: [10.1007/978-3-642-03311-7](https://doi.org/10.1007/978-3-642-03311-7).
- [6] Edward Farhi, Jeffrey Goldstone, and Sam Gutmann. *A Quantum Approximate Optimization Algorithm*. 2014. arXiv: [1411.4028](https://arxiv.org/abs/1411.4028) [quant-ph]. URL: <https://arxiv.org/abs/1411.4028>.
- [7] A.B. Finnilla et al. “Quantum annealing: A new method for minimizing multidimensional functions”. In: *Chemical Physics Letters* 219.5 (1994), pp. 343–348. ISSN: 0009-2614. DOI: [https://doi.org/10.1016/0009-2614\(94\)00117-0](https://doi.org/10.1016/0009-2614(94)00117-0). URL: <https://www.sciencedirect.com/science/article/pii/0009261494001170>.
- [8] Marshall L Fisher. “The Lagrangian relaxation method for solving integer programming problems”. In: *Management science* 27.1 (1981), pp. 1–18.
- [9] Franz Georg Fuchs et al. “Constraint Preserving Mixers for the Quantum Approximate Optimization Algorithm”. In: *Algorithms* 15.6 (2022). ISSN: 1999-4893. DOI: [10.3390/a15060202](https://doi.org/10.3390/a15060202). URL: <https://www.mdpi.com/1999-4893/15/6/202>.
- [10] Hans U Gerber and Gérard Pafum. “Utility functions: from risk theory to finance”. In: *North American Actuarial Journal* 2.3 (1998), pp. 74–91.
- [11] Hans U. Gerber and Elias S. W. Shiu. *Option Pricing by Esscher Transforms*. Transactions of the Society of Actuaries 46, 99–191. 1994.
- [12] Fred Glover, Gary Kochenberger, and Yu Du. *A Tutorial on Formulating and Using QUBO Models*. 2019. arXiv: [1811.11538](https://arxiv.org/abs/1811.11538) [cs.DS]. URL: <https://arxiv.org/abs/1811.11538>.
- [13] Christian Gollier. *The economics of risk and time*. MIT press, 2001.
- [14] Gurobi Optimization, LLC. *Gurobi Optimizer Reference Manual*. 2026. URL: <https://www.gurobi.com>.
- [15] Gurobi Optimization, LLC. *Quadratic Unconstrained Binary Optimization (QUBO)*. <https://gurobi-optimods.readthedocs.io/en/stable/mods/qubo.html>. Accessed: 2026-04-17.
- [16] Avradip Mandal et al. *Compressed Quadraticization of Higher Order Binary Optimization Problems*. 2020. arXiv: [2001.00658](https://arxiv.org/abs/2001.00658) [quant-ph]. URL: <https://arxiv.org/abs/2001.00658>.
- [17] Harry Markowitz. “Portfolio Selection”. In: *The Journal of Finance* 7.1 (1952), pp. 77–91. ISSN: 00221082, 15406261. URL: <http://www.jstor.org/stable/2975974> (visited on 03/26/2026).
- [18] Mike Marzec. “Portfolio Optimization: Applications in Quantum Computing”. In: *SSRN Electronic Journal* (2013). SSRN working paper.
- [19] Mirko Mattesi et al. “Diversifying Investments and Maximizing Sharpe Ratio: A Novel Quadratic Unconstrained Binary Optimization Formulation”. In: *Quantum Reports* 6.2 (2024), pp. 244–262. ISSN: 2624-960X. DOI: [10.3390/quantum6020018](https://doi.org/10.3390/quantum6020018). URL: <https://www.mdpi.com/2624-960X/6/2/18>.
- [20] Attilio Meucci. “Fully Flexible Views: Theory and Practice”. In: *Risk* 21.10 (2008), pp. 97–102.
- [21] J A Montañez-Barrera et al. “Unbalanced penalization: a new approach to encode inequality constraints of combinatorial problems for quantum optimization algorithms”. In: *Quantum Science and Technology* 9.2 (Apr. 2024), p. 025022. ISSN: 2058-9565. DOI: [10.1088/2058-9565/ad35e4](https://doi.org/10.1088/2058-9565/ad35e4). URL: <http://dx.doi.org/10.1088/2058-9565/ad35e4>.
- [22] Abha Satyavan Naik et al. “From portfolio optimization to quantum blockchain and security: a systematic review of quantum computing in finance”. en. In: *Financial Innovation* 11.1 (Feb. 2025), p. 88. ISSN: 2199-4730. DOI: [10.1186/s40854-025-00751-6](https://doi.org/10.1186/s40854-025-00751-6). URL: <https://doi.org/10.1186/s40854-025-00751-6>.
- [23] Antonis Papapantoleon. “An introduction to Lévy processes with applications in finance”. In: *arXiv preprint arXiv:0804.0482* (2008).
- [24] Frank Phillipson and Harshit S. Bhatia. “Portfolio Optimisation Using the D-Wave Quantum Annealer”. In: *Computational Science – ICCS 2021*. Springer, 2021, pp. 45–59.
- [25] Amy V. Puelz. “Value-at-Risk Based Portfolio Optimization”. In: *Stochastic Optimization: Algorithms and Applications*. Ed. by Stanislav Uryasev and Panos M. Pardalos. Boston, MA: Springer US, 2001, pp. 279–302. ISBN: 978-1-4757-6594-6. DOI: [10.1007/978-1-4757-6594-6_13](https://doi.org/10.1007/978-1-4757-6594-6_13).
- [26] R. Tyrrell Rockafellar and Stanislav Uryasev. “Optimization of conditional value-at-risk”. In: *Journal of Risk* 2.3 (2000), pp. 21–41. DOI: [10.21314/JOR.2000.038](https://doi.org/10.21314/JOR.2000.038).
- [27] Yuhui Shi and Russell Eberhart. “A modified particle swarm optimizer”. In: *1998 IEEE International Conference on Evolutionary Computation Proceedings. IEEE World Congress on Computational Intelligence (Cat. No. 98TH8360)*. 1998, pp. 69–73. DOI: [10.1109/ICEC.1998.699146](https://doi.org/10.1109/ICEC.1998.699146).
- [28] Hamed Soleimani, Hamid Reza Golmakani, and Mohammad Hossein Salimi. “Markowitz-based portfolio selection with minimum transaction lots, cardinality constraints and regarding sector capitalization using genetic algorithm”. In: *Expert Systems with Applications* 36.3 (2009), pp. 5058–5063. ISSN: 0957-4174. DOI: [10.1016/j.eswa.2008.06.007](https://doi.org/10.1016/j.eswa.2008.06.007).

- [29] Valter Uotila, Julia Ripatti, and Bo Zhao. “Higher-Order Portfolio Optimization with Quantum Approximate Optimization Algorithm”. In: *2025 IEEE International Conference on Quantum Computing and Engineering (QCE)*. Albuquerque, NM, USA: IEEE, Aug. 2025, pp. 01–12. ISBN: 9798331557362. DOI: [10.1109 / QCE65121.2025.00244](https://doi.org/10.1109/QCE65121.2025.00244). URL: <https://ieeexplore.ieee.org/document/11249852/>.
- [30] David Wozabal, Ronald Hochreiter, and Georg Ch. Pflug. “A difference of convex formulation of value-at-risk constrained optimization”. In: *Optimization* 59.3 (2010), pp. 377–400. DOI: [10.1080/02331931003700731](https://doi.org/10.1080/02331931003700731).
- [31] Hanjing Xu et al. “Dynamic Asset Allocation with Expected Shortfall via Quantum Annealing”. en. In: *Entropy* 25.3 (Mar. 2023), p. 541. ISSN: 1099-4300. DOI: [10.3390/e25030541](https://doi.org/10.3390/e25030541). URL: <https://www.mdpi.com/1099-4300/25/3/541>.

APPENDIX A
SAMPLE AVERAGE APPROXIMATION AND ESSCHER TILTING

A.1. Sample Average Approximation (SAA) on a finite scenario set

Fix a base-point $\bar{\mathbf{x}}$ and a tilt t . We distinguish between two levels. At the *theoretical* level, scenarios are viewed as realizations of an underlying probability law $\mathbb{P}$. Under this law, the base-point P&L $z(\mathcal{B})$ is a function of scenarios and can be seen as a scalar random variable, and tilted moments are expectations under the Esscher-tilted measure $\pi(t|\mathcal{B})$ obtained from $\mathbb{P}$. We refer to these ideal quantities as *population* quantities.

At the *empirical* level, the law $\mathbb{P}$ is unknown: we observe only a finite scenario set. We write $z_s(\mathcal{B})$ for the realization of $z(\mathcal{B})$ in scenario s , and collect these S values in the vector $\mathbf{z}(\mathcal{B}) = (z_1(\mathcal{B}), \dots, z_S(\mathcal{B}))$. Thus, $\mathbf{z}(\mathcal{B})$ is the sample of realizations of the single random variable $z(\mathcal{B})$ across the scenario set; only this sample is observable. The sample average approximation (SAA) replaces each population expectation by a weighted average over these scenarios. Throughout, $\bar{\mathbf{x}}$ is held fixed (Appendix B.1), so $\mathbf{z}(\mathcal{B})$ is a fixed scenario vector and the only source of randomness is the scenario sampling.

At the theoretical level, the Esscher tilt reweights $\mathbb{P}$ by the base-point P&L. For any scenario quantity φ , the corresponding population tilted expectation is

$$\mathbb{E}_{\pi(t|\mathcal{B})}[\varphi] = \frac{\mathbb{E}_{\mathbb{P}}[\varphi e^{t z(\mathcal{B})}]}{\mathbb{E}_{\mathbb{P}}[e^{t z(\mathcal{B})}]}.$$

Since $\mathbb{P}$ is unknown, the SAA replaces the population expectations above by averages over the observed scenarios: the empirical mean serves as a surrogate for the unknown expectation under $\mathbb{P}$.

Let $w_s > 0$, with $\sum_{s=1}^S w_s = 1$, be the base weights on the scenario set. The empirical (Esscher/softmax) weights are the ones defined in Eq. (3), and the empirical tilted expectation is the weighted average

$$\mathbb{E}_{\pi(t|\mathcal{B})}[\varphi] = \sum_{s=1}^S \pi_s(t|\mathcal{B}) \varphi_s.$$

Hence, tilted moments are computed by direct weighted averaging.

Assume the scenarios are independent and identically distributed draws from $\mathbb{P}$ with $w_s = 1/S$, and $\mathbb{E}_{\mathbb{P}}[e^{t z(\mathcal{B})}] < \infty$. By the strong law of large numbers applied to numerator and denominator, for any integrable φ the empirical weighted average converges to the population tilted expectation,

$$\sum_{s=1}^S \pi_s(t|\mathcal{B}) \varphi_s \xrightarrow[S \rightarrow \infty]{a.s.} \frac{\mathbb{E}_{\mathbb{P}}[\varphi e^{t z(\mathcal{B})}]}{\mathbb{E}_{\mathbb{P}}[e^{t z(\mathcal{B})}]}.$$

Thus the empirical tilted moments are consistent estimators of the population tilted moments at each fixed $(t, \bar{\mathbf{x}})$.

If base weights w_s are non-uniform, all formulas above remain valid.

A.2. Variational/KL viewpoint (score–entropy trade-off) – characterization of the Esscher weights.

The same weights arise as the unique optimizer of a strictly concave score–entropy trade-off: increasing the expected baseline score while staying close to the base measure. Concretely,

$$\pi(t|\mathcal{B}) = \arg \max_{p \in \Delta_S} \left\{ t \sum_{s=1}^S p_s z_s(\mathcal{B}) - \text{KL}(p||\mathbf{w}) \right\},$$

where Δ_S is the probability simplex and

$$\text{KL}(p||\mathbf{w}) = \sum_{s=1}^S p_s \log \left(\frac{p_s}{w_s} \right)$$

is the Kullback–Leibler divergence (relative entropy) of p with respect to the base measure $\mathbf{w}$. It is non-negative and vanishes if and only if $p = \mathbf{w}$. The KL divergence is not a metric – it is generally not symmetric and does not satisfy the triangle inequality – but, being non-negative and zero only at $p = \mathbf{w}$, it quantifies the discrepancy between the candidate measure p and the base measure $\mathbf{w}$. Hence, the optimization problem balances two competing effects: increasing the expected stressed score while remaining close to the original scenario measure. When $\mathbf{w}$ is uniform, $-\text{KL}(p||\mathbf{w})$ differs from the Shannon entropy $H(p)$ only by an additive constant, so the objective is equivalently a score–entropy trade-off. The associated value function is the log-sum-exp (a smooth maximum), explaining the increasing concentration of $\pi(t|\mathcal{B})$ as $|t|$ grows.

This variational characterization is a finite-scenario instance of the Donsker–Varadhan variational formula for relative entropy [5]. In the portfolio context, the same relative-entropy reweighting principle is closely related to the entropy-pooling construction of [20], where posterior scenario probabilities are obtained by minimal relative-entropy deformation of a base measure under views.

In conclusion, the softmax definition in Eq. (3) specifies how the tilted weights are computed; the variational characterization above clarifies what they represent. The tilted measure can be read as the solution of a constrained reweighting problem: among all measures on the scenario set, $\pi(t|\mathcal{B})$ is the one that attains a given level of expected base-point score with the smallest departure – in relative-entropy terms – from the base measure $\mathbf{w}$. This gives the tilt a precise meaning: t is the exchange rate between expected score and fidelity to $\mathbf{w}$, and the monotone concentration of the weights as $|t|$ grows reflects the score term progressively dominating the entropic penalty.

APPENDIX B

BASE-POINT SELECTION

B.1. Choosing the base-point $\bar{\mathbf{x}}$

The exponential tilting mechanism requires a base-point $\bar{\mathbf{x}}$, for which several choices are possible. In our experiments, we choose to set $\bar{\mathbf{x}}$ as the solution of the continuous relaxation of a standard mean–variance Markowitz problem computed under the base scenario weights, imposing the same constraints as in the downstream step. This problem is convex and yields a robust baseline from which to perform reweighting, as its solution typically lies in the neighborhood of high-quality portfolios. Other valid choices are, for example, the initial portfolio or a feasible point found using a greedy or multi-start heuristic.

In all cases, once $\bar{\mathbf{x}}$ is fixed, our pipeline constructs $\pi(t|\mathcal{B})$, then $\bar{\mathbf{b}}_{t,\bar{\mathbf{x}}}$, $\bar{\mathbf{Q}}_{t,\bar{\mathbf{x}}}$, and solves the resulting quadratic problem. It is worth noting that the procedure can be iterated by using the solution of the tilted problem as a new base-point, thereby further refining the portfolio obtained in the previous iteration.

B.2. Markowitz base-point properties

For completeness, let

$$f_{\text{abs}}(\mathbf{x}) := \frac{1}{2} \mathbf{x}^\top \Sigma_{GG} \mathbf{x} + \Sigma_{G\mathcal{P}}^\top \mathbf{x} - \boldsymbol{\mu}_G^\top \mathbf{x}$$

be the non-constant part of the Markowitz objective in Eq. (8).

The matrix Σ_{GG} is positive semidefinite (PSD). Indeed, recall from Eq. (2) that, under the uniform measure $\mathbf{m}$,

$$\Sigma_{GG} = \mathbf{G}^\top (\text{diag}(\mathbf{m}) - \mathbf{m}\mathbf{m}^\top) \mathbf{G},$$

where $\mathbf{G} \in \mathbb{R}^{S \times d}$. Let $\mathbf{C} := \text{diag}(\mathbf{m}) - \mathbf{m}\mathbf{m}^\top$. For any $\mathbf{u} \in \mathbb{R}^S$, writing $\bar{u} := \sum_{s=1}^S m_s u_s$,

$$\mathbf{u}^\top \mathbf{C} \mathbf{u} = \sum_{s=1}^S m_s u_s^2 - \left(\sum_{s=1}^S m_s u_s \right)^2 = \sum_{s=1}^S m_s (u_s - \bar{u})^2 \geq 0,$$

so $\mathbf{C}$ is PSD (this quantity is the variance of the scenario vector $\mathbf{u}$ under $\mathbf{m}$). Hence, for any $\mathbf{v} \in \mathbb{R}^d$, setting $\mathbf{u} = \mathbf{G}\mathbf{v}$,

$$\mathbf{v}^\top \Sigma_{GG} \mathbf{v} = (\mathbf{G}\mathbf{v})^\top \mathbf{C} (\mathbf{G}\mathbf{v}) \geq 0,$$

which proves that Σ_{GG} is PSD.

If the initial portfolio P&L is uncorrelated with the candidate instruments in the empirical scenario set, i.e., $\Sigma_{G\mathcal{P}} = \mathbf{0}$, then the objective reduces to the form

$$f_{\text{inc}}(\mathbf{x}) := \frac{1}{2} \mathbf{x}^\top \Sigma_{GG} \mathbf{x} - \boldsymbol{\mu}_G^\top \mathbf{x},$$

with the same PSD curvature matrix Σ_{GG} .

Consequently, both $f_{\text{abs}}(\mathbf{x})$ and $f_{\text{inc}}(\mathbf{x})$ are convex quadratic objectives on $\mathbb{R}^d$, and remain convex under standard convex constraints (e.g., box constraints $0 \leq x_u \leq 1$ and linear budget constraints).

APPENDIX C

INEQUALITY CONSTRAINTS: SLACK VARIABLES AND LAGRANGIAN RELAXATION

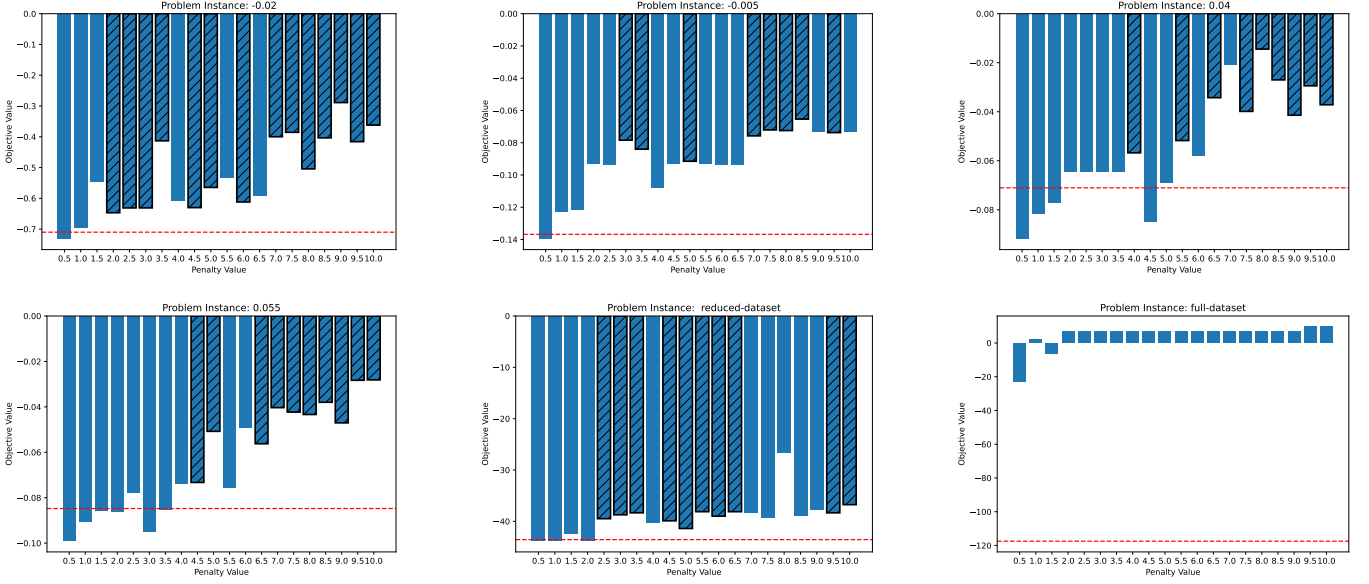
C.1. Slack variables experimental inefficiency

The slack variable approach to integrating the Greeks constraints into the QUBO formulation proved underperforming and difficult to generalize. In this method, the solution space is extended by introducing auxiliary binary variables to represent the feasible range of solutions satisfying the sensitivity constraints. The slack variable method is designed to yield a zero-valued penalty (or near-zero, depending on the discretization of the variable) whenever the solution complies with the Greeks constraints, and a large penalty when they are violated. This weight of the penalty is not fixed across instances of the optimization problem, but is rather calibrated relative to the other penalty terms and the scale of the objective function. Specifically, we weighted the penalty in Equation 22 with

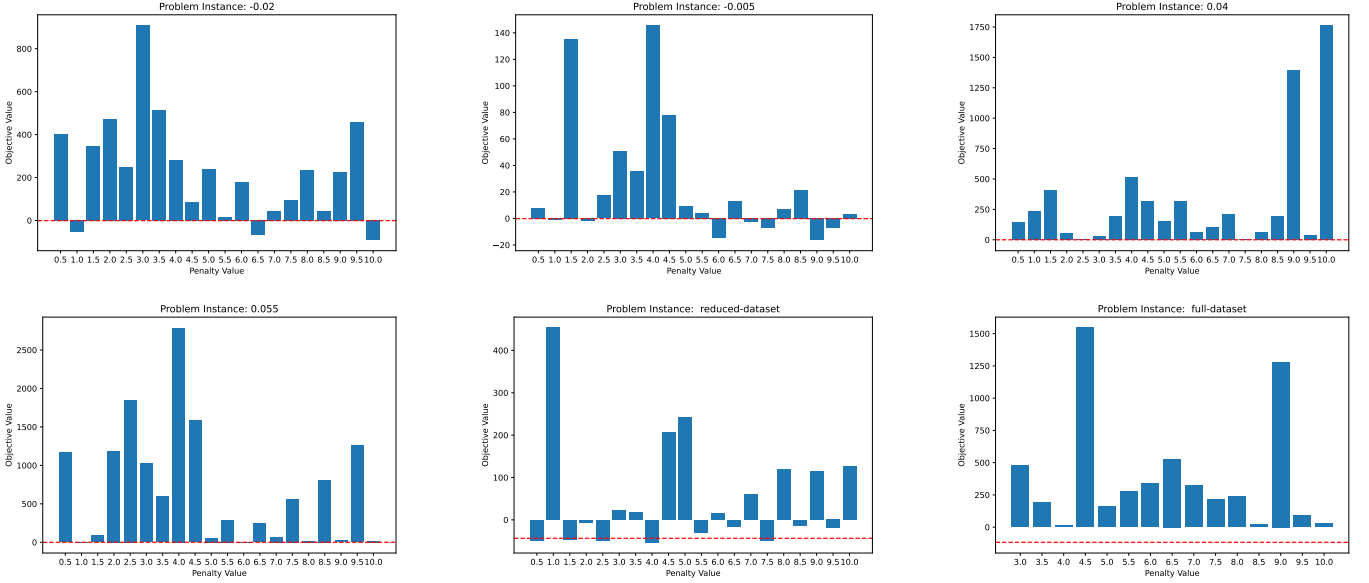
$$w_G = w_{\text{greek}} \frac{(w_{\text{EOS}} + w_{\text{UEI}})}{2C_G^2},$$

where w_{EOS} and w_{UEI} are the weights associated respectively with the EOS composition constraint and with the UEI quantities constraint. Instead, w_{greek} is a penalty value parameter that we tried to tune to allow each problem instance to reach feasible solutions concerning all the Greeks constraints.

Figure 4a reports the results obtained across different penalty values for the 6 instances considered in the QUBO formulation, using GUROBI_Q. This solver exhibits stable behavior: as seen in the plots, increasing the penalty value tends to increase the likelihood of obtaining a feasible solution. However, none of the feasible solutions found matches the optimum of the constrained problem identified by GUROBI. The Simulated Bifurcation solver, by contrast, exhibits more unpredictable behavior. As reported in Figure 4b, it fails to find any feasible solution across all instances, despite being allotted up to 1 hour of runtime. These results were obtained with `mode='discrete'`, which is generally slower but more accurate than the ballistic variant, and `heated=True` that further improves solution quality at the cost of additional computation. We experimented also with the other solver configurations and, although this setting still did not yield feasible solutions, it was the only one to produce non-trivial, non-zero outputs.



(a) GUROBI_Q results.



(b) Simulated Bifurcation results.

Fig. 4: Results for the slack variable method using GUROBI_Q (a) and Simulated Bifurcation (b) solvers, varying penalty value w_{greek} . Red dashed lines indicate the optimal MIQP reference value; hatched bars denote feasible solutions.

C.2. Lagrangian relaxation experimental convergence

In principle, our Lagrangian relaxation-based methodology would provide a certificate of convergence to the optimum of the dual problem; however, neither GUROBI_Q nor Simulated Bifurcation guarantee optimality of the individual QUBO subproblems. We therefore adopted a practical stopping criterion: termination upon matching the optimum objective value of the constrained problem found by GUROBI, or after a maximum of 200 iterations, whichever occurred first.

When using GUROBI_Q, four out of six instances converged to the optimal solution within 27–68 iterations, as illustrated in Figure 5. The two remaining instances, $t = -0.02$ and the Markowitz approach applied to the full-dataset, did not reach the target objective value. Nevertheless, both cases yielded noteworthy results: for $t = -0.02$, a high-quality solution with only a 1.77% optimality gap was already obtained at the second iteration (Figure 6); for the full-dataset instance, the best solution was found at iteration 17 (Figure 7) and, although it fell short of the GUROBI MIQP optimum, it was produced in approximately 2.8 minutes, a time within which the MIQP solver had not yet completed its presolving phase, and which caused an out-of-memory error on Machine B.

We repeated the experiment using Simulated Bifurcation as the QUBO solver with `mode='discrete'` and `heated=True`. Even with this most promising configuration, identified testing the slack variables approach, the Lagrangian relaxation approach failed to find any feasible solution for all Markowitz instances and for the case $t = -0.005$. In this last, no optimal solution was identified within 200 iterations; nevertheless, a suboptimal solution within 1.5% of the optimum was obtained (Figure 8). For

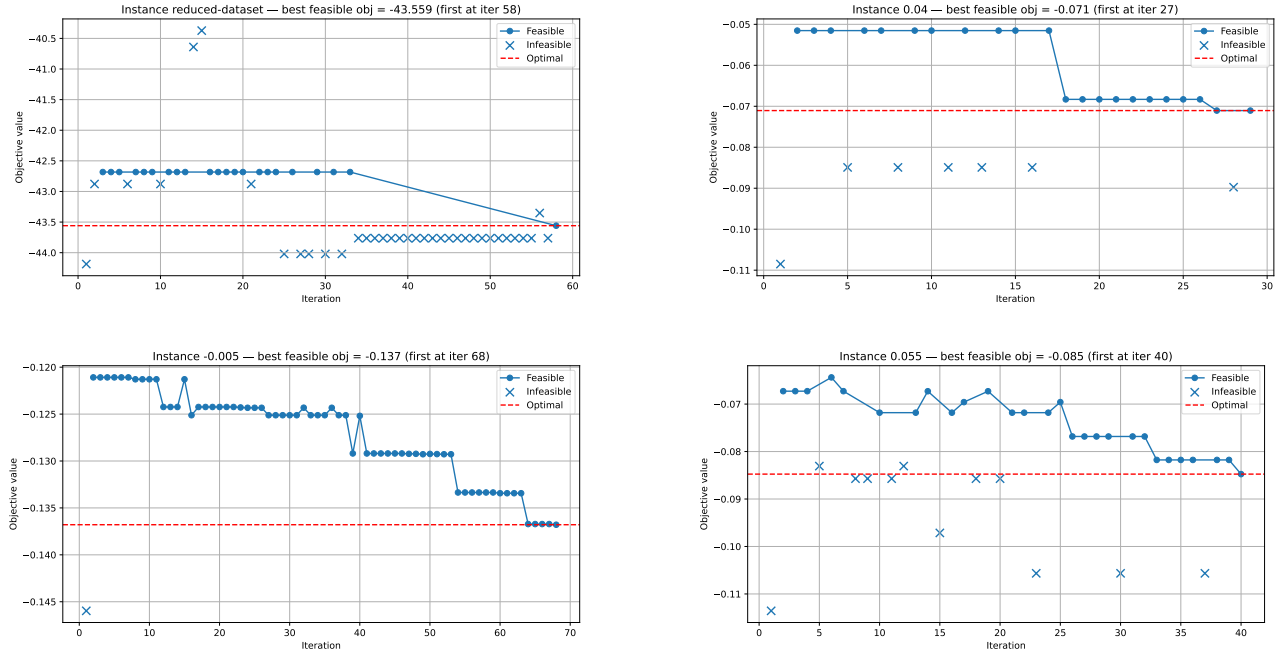


Fig. 5: Lagrangian relaxation convergence for the four instances in which GUROBI_Q reached the GUROBI MIQP optimal objective value (dashed line): reduced Markowitz dataset, $t = 0.04$, $t = -0.005$, and $t = 0.055$.

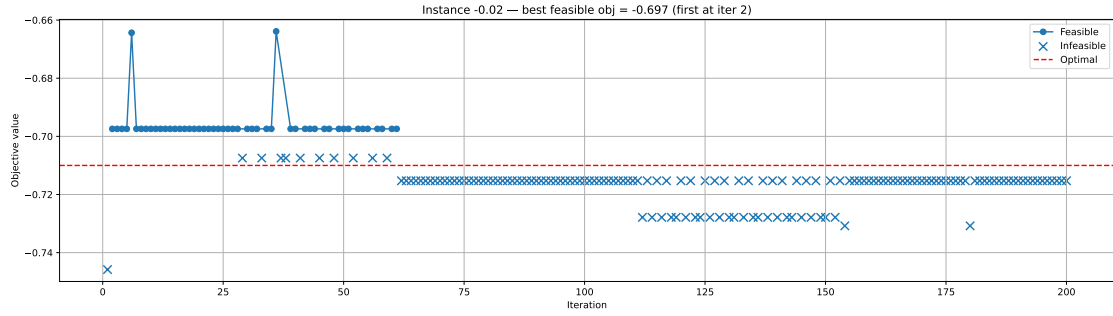


Fig. 6: Lagrangian relaxation convergence for instance $t = -0.02$, using GUROBI_Q as QUBO solver. The optimal objective value of the constrained problem solved using GUROBI (dashed line) is not reached; however, a solution with a 1.77% optimality gap is already obtained at the second iteration.

$t = 0.04$, an optimal solution was reached at iteration 34 (Figure 9b); while numerically distinct from the solution returned by GUROBI (both for the QUBO and constrained formulations), it represents the same financial portfolio. Notably, for $t = -0.02$, Simulated Bifurcation identified the optimal solution at iteration 60 (Figure 9a), which GUROBI_Q had failed to reach.

Overall, the behaviour of the Simulated Bifurcation solver is highly unpredictable: as visible across all convergence plots, the objective value exhibits pronounced oscillations across iterations and, compared to GUROBI_Q, a substantially higher proportion of iterates yield infeasible solutions.

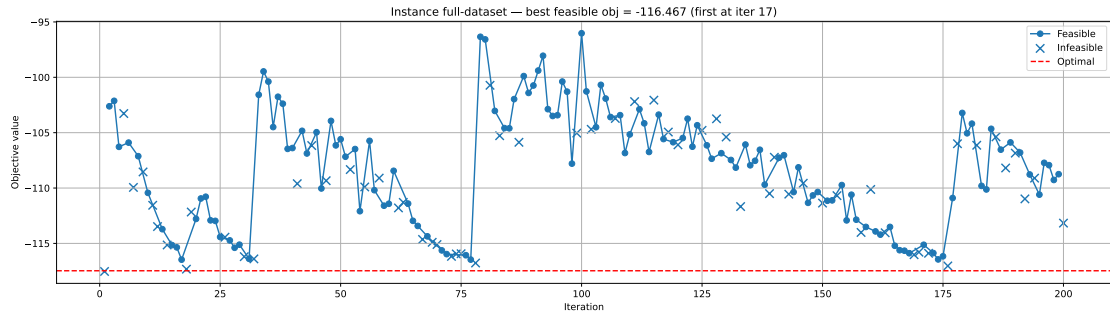


Fig. 7: Lagrangian relaxation convergence for the Markowitz approach applied to the full-dataset, using GUROBI_Q as QUBO solver. The optimal objective value of the constrained problem solved using GUROBI (dashed line) is not reached, however the best solution reached by GUROBI_Q was already found at iteration 17 in approximately 2.8 minutes.

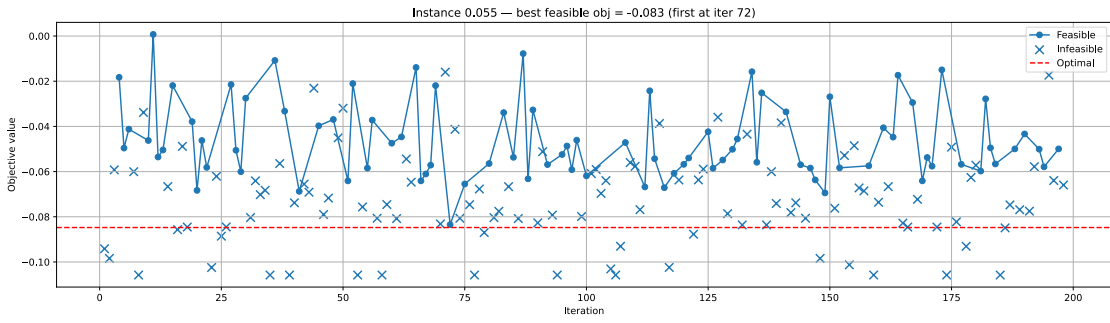
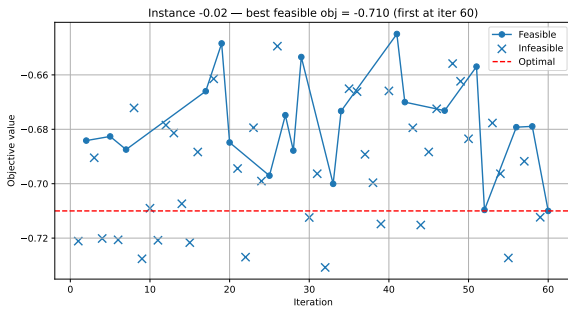
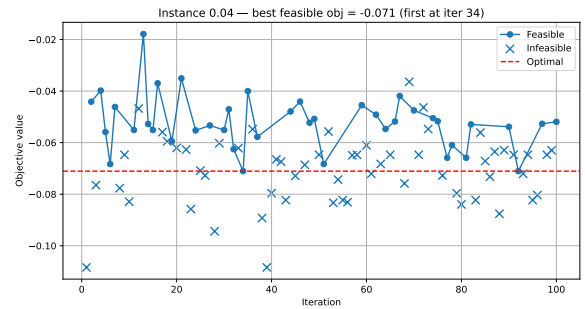


Fig. 8: Lagrangian relaxation convergence for instance $t = 0.055$ using Simulated Bifurcation. No optimal solution is found within 200 iterations; the best solution obtained is within 1.5% of the optimum in approximately 36 minutes.



(a) Instance $t = -0.02$: optimal solution found within 30 minutes



(b) Instance $t = 0.04$: optimal solution found within 17 minutes

Fig. 9: Lagrangian relaxation convergence using Simulated Bifurcation for instances $t = -0.02$ and $t = 0.04$, the two instances for which an optimal solution was found.