

Multi-Response Optimization of AHSS Spot-Welding Parameters: Desirability, Trade-offs and Sensitivity in a Synthetic Case Study

Seshabattar Phaneendra¹, Dr. Dheeraj Nimawat²

¹Research Scholar, Department of Mechanical Engineering, JS University, Shikohabad, UP

²Supervisor, Department of Mechanical Engineering, JS University, Shikohabad, UP

ABSTRACT

An optimum welding schedule is conditional on the response model, the admissible search region and the decision preferences used to compare candidates. This paper examines those dependencies using the explicitly synthetic DP780-labelled and DP980-labelled response surfaces contained in an AHSS spot-welding thesis. Tensile-shear and cross-tension loads are represented by fitted quadratic models in welding current, weld time and electrode force. A discrete search evaluates 28,701 points inside the convex hull of a three-factor Box–Behnken design before applying assumed width and indentation constraints. Equal-weight geometric desirability selects 9.20 kA, 310 ms and 4.00 kN for the DP780-labelled scenario, with fitted loads of 8.866 and 4.106 kN. For the DP980-labelled scenario, it selects 9.10 kA, 295 ms and 4.00 kN, with fitted loads of 9.611 and 4.245 kN. These are conditional numerical selections, not qualified process settings. Preference-weight changes alter the selected combinations, especially force in the first scenario. Baseline comparisons, interior alternatives, local perturbations and an additional worst-neighbourhood calculation expose differences between nominal performance and decision stability. The analysis retains the source model's adverse cross-tension lack-of-fit result and does not treat additional random draws as physical confirmation. Its contribution is a reproducible decision audit that makes preferences and feasibility assumptions visible. Experimental qualification remains necessary before any schedule can be recommended for actual AHSS joining.

Keywords: multi-response optimization; resistance spot welding; desirability; sensitivity; robust selection; synthetic surrogate; AHSS.

INTRODUCTION

A process optimization result is often reported as a short list of preferred settings. Such a list is useful only if the reader also knows what was optimized, which candidates were allowed, how competing responses were valued and what uncertainty remains. These questions are especially important when an optimization algorithm is applied to a fitted surface rather than directly to a fully qualified physical process. The numerical maximum belongs to the mathematical problem that was specified; it does not automatically belong to the welding system that motivated the problem.

Resistance spot welding provides a practical setting for studying this distinction. A joint may perform differently under different loading configurations, and selecting a schedule solely on one peak-load response can conceal an unfavourable balance elsewhere. Kimchi and Phillips (2023) provide the automotive resistance-spot-welding context. The present paper concentrates on the decision layer that follows modelling, rather than deriving a physical welding law or reporting new mechanical tests.

The source thesis proposes DP780 and DP980 sheet-steel investigations and includes a revised chapter containing assumed quadratic response scenarios. Its numerical example is explicitly described as uncalibrated and synthetic. That provenance is retained here. The scenario labels do not establish actual material behaviour, and the load equations do not derive from measured thermal cycles, nugget geometry or fracture tests. The calculation offers a controlled way to investigate decision rules while the proposed laboratory programme remains unexecuted.

The question addressed is how a reported optimum changes when the analyst makes reasonable but consequential choices. Those choices include the lower and target values used to scale responses, the relative weight assigned to tensile-shear and cross-tension performance, the baseline used to calculate percentage improvement and the neighbourhood over which stability is examined. Rather than hide these choices inside software defaults, the paper states them as part of the problem definition.

The manuscript extends the thesis presentation through a decision-focused audit and a deterministic worst-neighbourhood comparison. It does not claim a novel optimization algorithm. Mathematical optimization is a mature subject, with relevant foundations in Antoniou and Lu (2021) and Stein (2024a, 2024b). The contribution here is the traceable connection between a small response-surface model, an explicit feasible set and the language used to communicate a conditional recommendation.

This paper shares the synthetic dataset with a companion manuscript on model identifiability and diagnostics. A third manuscript proposes experimental validation. The present contribution is therefore not an independent replication of the modelling results. Any submission must disclose the shared data and thesis basis. The separation into manuscripts should be reconsidered if a target journal judges that one integrated methods paper would communicate the contribution more effectively.

2. Decision problem and data provenance

2.1 Responses, factors and scenario labels

The candidate schedule contains welding current I , weld time t and electrode force F . Their planning ranges are 6–10 kA, 150–350 ms and 2.5–4.5 kN. The responses are tensile-shear load, abbreviated TS, and cross-tension load, abbreviated CT, both expressed in kilonewtons. Higher fitted values are treated as preferable within the assumed example. This preference is a modelling decision and does not replace product-specific acceptance criteria.

Both response surfaces use coded coordinates $x_1 = (I - 8)/2$, $x_2 = (t - 250)/100$ and $x_3 = F - 3.5$. Each surface contains an intercept, three linear terms, three squared terms and three pairwise interactions. The coefficients were obtained by fitting seventeen synthetic row means from a Box–Behnken layout. The generator and noise assumptions are disclosed in the source thesis and the companion modelling paper, while the fitted coefficients are repeated here to make the optimization calculation independently inspectable.

Table 1. Fitted coefficients from synthetic row means; use coded coordinates.

Term	780 TS	780 CT	980 TS	980 CT
Intercept	8.4888	3.7924	9.1588	3.9989
x_1	0.7897	0.3581	0.9920	0.3373
x_2	0.4882	0.3569	0.4073	0.4338
x_3	0.1626	0.2970	0.3537	0.2955
x_1^2	-0.7085	-0.4231	-0.8430	-0.4813
x_2^2	-0.5230	-0.3560	-0.5360	-0.4864
x_3^2	-0.3547	-0.2818	-0.4067	-0.3146
x_1x_2	0.1770	0.1039	0.1313	0.0892
x_1x_3	0.0758	0.0992	0.0819	0.0323
x_2x_3	-0.0827	0.0672	-0.1945	0.0411

The grade labels are convenient identifiers, not material certificates. The DP980-labelled intercept is larger because the synthetic generator was constructed that way. No inference about actual joint superiority follows from that fact. Likewise, the preferred 1.2 mm sheet thickness supplies thesis context but is not an input to a calibrated mechanical model. The optimizer cannot infer chemistry, coating or electrode geometry from a nominal steel designation.

2.2 Why nominal performance is insufficient

A nominal optimum can be attractive for several different reasons. It may occupy a broad plateau where nearby settings perform similarly. It may instead be a sharp peak that loses value under small perturbations. It may be close to a feasibility boundary or close to a region where the statistical model has little support. These cases should not be treated as equivalent merely because their printed objective values have similar decimal precision.

The decision therefore needs at least three views: performance at the selected point, performance under declared changes in the decision rule and performance under declared changes in the implemented settings. None of these views constitutes physical validation. Together, however, they make the computational recommendation more honest by showing what is stable and what is contingent within the model.

2.3 Feasibility is not qualification

The word feasible is used here in a strictly mathematical sense. A candidate is feasible when it lies in the specified statistical region and satisfies two assumed geometric inequalities. These conditions do not establish absence of expulsion, cracking, excessive electrode wear or unacceptable fracture mode. The distinction is essential because a numerical constraint often looks authoritative when written in engineering units, even if its threshold was selected only for illustration. Dwivedi (2022) discusses joining mechanisms and performance, which helps explain why real qualification extends beyond one optimized load value. In this paper, those physical considerations remain requirements for later testing. They are not silently represented by the two simple geometry equations. The optimizer should be understood as selecting candidates for further investigation, not issuing a welding procedure specification.

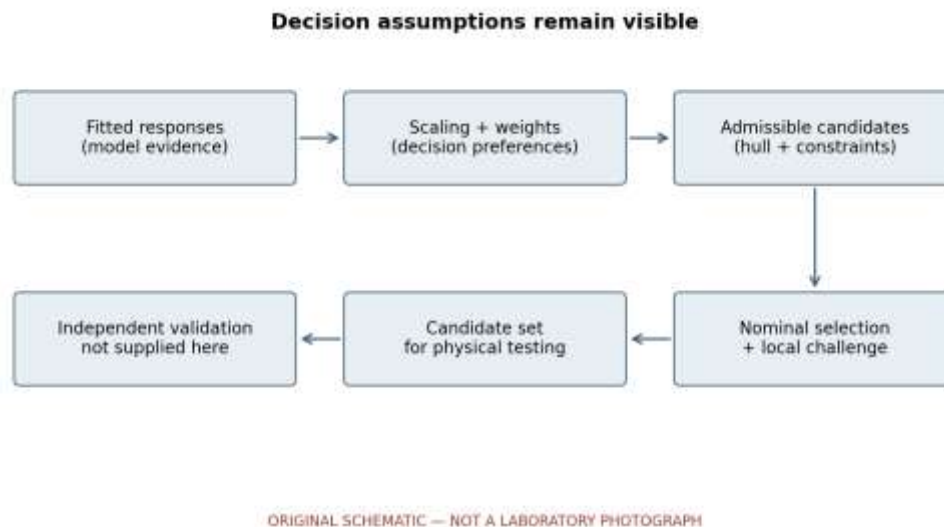


Figure 1. Decision workflow separating response evidence, preferences, constraints and validation. Original explanatory schematic.

3. Objective function and admissible domain

3.1 Individual desirability functions

Each predicted response is converted to a dimensionless score between zero and one. For TS, the score is the fitted response minus an assumed lower value, divided by the difference between an assumed target and that lower value. The result is clipped to the unit interval. CT is transformed in the same way using its own limits. A response at or below the lower value receives zero; a response at or above the target receives one.

These transformations make it possible to combine responses with different numerical scales. They also introduce preferences. A one-kilonewton increase does not have the same desirability effect when the scaling interval is narrow as

when it is wide. The lower and target values are therefore not innocuous formatting constants. They are part of the substantive decision and must be justified in a physical application.

Table 2. Assumed desirability bounds; these are not standards or product limits.

Scenario	TS lower	TS target	CT lower	CT target
DP780	7.50	10.00	3.00	5.00
DP980	8.00	11.00	3.20	5.30

For the DP780-labelled scenario, TS is scaled between 7.5 and 10.0 kN, while CT is scaled between 3.0 and 5.0 kN. For the DP980-labelled scenario, the corresponding intervals are 8.0–11.0 kN and 3.2–5.3 kN. These are assumed teaching values inherited from the thesis calculation. They are not taken from a standard, a vehicle manufacturer's specification or experimental acceptance testing.

The clipping rule creates plateaus and discontinuities in derivatives at the scaling limits. As a result, the combined objective is not simply another unconstrained quadratic, even though the underlying response surfaces are quadratic. An analyst who differentiates the load equations alone has not necessarily solved the stated desirability problem. The response transformations and all feasibility conditions must remain part of the calculation.

3.2 Combined desirability and preference weights

The combined score is $D = dTS^w \times dCT^{(1-w)}$, where w is the TS preference weight. Equal weighting uses $w = 0.5$ and reduces to the square root of the product of the two individual scores. The geometric construction means that a zero value for either response makes the combined score zero when both weights are positive. This feature encodes a non-compensatory boundary: an unacceptable response cannot be rescued by a very favourable other response.

Equal numerical weights do not mean equal physical importance in every sense. Because each response is first scaled by a different interval, its effect on the combined objective depends on both the weight and the normalization. A statement that the responses were weighted equally should therefore be accompanied by the scaling table. Otherwise, readers cannot reconstruct the actual trade-off being imposed.

Three TS weights are examined: 0.25, 0.50 and 0.75. These are sensitivity scenarios rather than empirically estimated stakeholder preferences. Their purpose is to reveal whether the selected schedule is fragile to a reasonable change in priorities. The preferred weight for an industrial application would need to be agreed with the parties responsible for structural performance and manufacturing constraints.

3.3 Statistical support region

Candidates satisfy $|x_i| \leq 1$ for each factor and $|x_1| + |x_2| + |x_3| \leq 2$. This region is the convex hull of the edge-midpoint Box–Behnken settings. It excludes the unsupported cube corners where all three factors are simultaneously extreme. Using this boundary does not guarantee accurate interpolation, but it prevents the optimizer from exploiting regions with no design support merely because the polynomial happens to predict high values there.

The distinction between range and support is particularly important in multi-factor studies. A current within its stated range and a force within its stated range do not automatically make their combination equally supported. Statistical support belongs to combinations, not individual coordinates considered separately. Reporting only minimum and maximum values for each factor can therefore obscure a consequential restriction on the search.

3.4 Assumed geometric constraints

Two deterministic quantities are computed from the coded factors. The assumed width is $wN = 5.1 + 0.45x_1 + 0.30x_2 - 0.12x_3 - 0.12x_1^2 - 0.08x_2^2$ mm. The assumed indentation is $hN = 0.11 + 0.025x_1 + 0.018x_2 - 0.012x_3 + 0.008x_1^2$ mm. Candidates must satisfy $wN \geq 4.5$ mm and $hN \leq 0.20$ mm. The same functions are used for both labelled scenarios.

These functions are neither fitted measurements nor finite-element predictions. Their role is to illustrate how constraints interact with a load-based objective. Using identical geometry equations across the two scenarios is an additional simplification, not evidence that the two steel grades develop identical nuggets. The geometric variable names identify an analogy; they do not establish an observation.

Any future application should replace these functions with independently supported quantities and include uncertainty in the constraints. A candidate predicted to exceed a minimum width by a very small margin may be inappropriate if measurement and model uncertainty are larger than that margin. Deterministic feasibility can therefore be an initial screening rule, but it should not be confused with a specified probability of meeting a physical requirement.

4. Computational procedure

4.1 Enumerated grid search

The coded grid uses forty-one levels for current, forty-one for time and twenty-one for force. In physical units, the increments are 0.1 kA, 5 ms and 0.1 kN. The complete rectangular product contains 35,301 points. Restriction to the declared hull leaves 28,701 candidates before the two geometric constraints are applied. The finite set is small enough to evaluate directly without a stochastic search algorithm.

Every retained point receives fitted TS and CT responses, individual desirabilities and a combined score. Points failing the assumed geometry rules receive zero score. The maximum is selected using a fixed candidate order so that ties are reproducible. The calculation identifies the best point on this grid, not a proven continuous global optimum. A finer grid could select a slightly different schedule even when the response equations and preferences remain unchanged.

Direct enumeration is attractive here because it makes the search auditable. Evolutionary or other heuristic methods can be appropriate for larger problems, as discussed more generally by Feng et al. (2021), but their use would add unnecessary algorithmic variability to this small example. The numerical complexity of an optimizer should match the problem rather than serve as a substitute for a clear objective function.

4.2 Baseline and interior comparisons

The baseline is fixed at 7 kA, 200 ms and 3.5 kN. It is an assumed comparison setting, not an existing qualified production schedule. Improvement percentages are calculated from fitted means relative to that baseline. Because the denominator is selected by the analyst, the percentages are conditional and should always be accompanied by absolute response values.

An interior alternative is constructed by halving every coded coordinate of the equal-weight optimum. This moves the candidate toward the design centre. The point is not asserted to be robust simply because it is more central. Its purpose is to provide a clearly defined comparator with less extreme coordinates, allowing the reader to see the performance cost of a conservative geometric move without inventing a production advantage.

4.3 Local perturbations

The nominal optimum is challenged one factor at a time by ± 0.2 kA, ± 10 ms and ± 0.1 kN. These perturbations correspond to ± 0.1 in the relevant coded coordinate. The resulting response values are deterministic predictions from the fitted model. They are not additional welds, and the perturbation sizes are not measured controller tolerances. The test asks how the polynomial changes locally under a declared numerical disturbance.

One-factor perturbations have a useful diagnostic role because they separate local directional effects. They do not establish tolerance to simultaneous disturbances. A schedule may be insensitive to each factor considered separately but still experience an important response change when several factors move together through interaction terms. That limitation motivates the additional neighbourhood calculation described next.

4.4 Additional worst-neighbourhood analysis

For each candidate, an additional computational audit considers the twenty-seven combinations of coded offsets -0.1 , 0 and $+0.1$ in all three factors. A candidate is retained for this audit only when every challenged point remains inside the declared hull and satisfies the assumed geometry constraints. The candidate's robustness score is the minimum equal-weight desirability over those twenty-seven points. The largest such minimum identifies a worst-neighbourhood selection on the same nominal grid.

This is a deterministic bounded-disturbance exercise. It does not assign probabilities to the offsets and does not estimate production reliability. The twenty-seven points sample the disturbance box at its corners, edge centres, face centres and centre; they do not by themselves prove that every interior disturbance has been checked for a general nonlinear objective. The result is therefore named a discrete worst-neighbourhood score rather than an exact continuous robust optimum.

The additional calculation is distinguished from the original thesis result. It uses the same fitted coefficients but a different selection rule, providing a new analytical comparison without pretending that new empirical data have been collected. The

code evaluates support and geometry for every challenged point, preventing a nominal candidate from receiving a favourable robust score through predictions outside the declared admissible region.

5. Numerical results

5.1 Equal-weight selections

The DP780-labelled equal-weight grid selection is 9.20 kA, 310 ms and 4.00 kN. Its fitted TS and CT loads are 8.866 and 4.106 kN. The DP980-labelled selection is 9.10 kA, 295 ms and 4.00 kN, with fitted loads of 9.611 and 4.245 kN. The common selected force is a property of this numerical realization and objective, not evidence for a universal optimum force for AHSS.

Table 3. Equal-weight grid selections and fitted synthetic responses.

Scena rio	I (k A)	t (ms)	F (kN)	TS (kN)	CT (kN)
DP780	9.2 0	310. 00	4.00	8.866	4.106
DP980	9.1 0	295. 00	4.00	9.611	4.245

At the first selected point, the assumed geometry equations produce a width of approximately 5.418 mm and indentation of 0.133 mm. For the second, they give approximately 5.370 mm and 0.128 mm. These values demonstrate the arithmetic of the constraints. They must not be inserted into an experimental results table as measured nugget diameter or surface indentation.

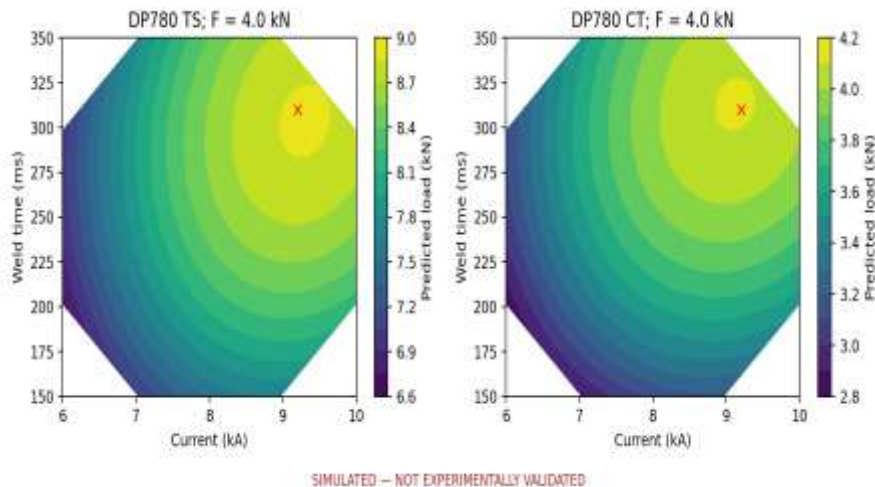


Figure 2. DP780-labelled fitted TS and CT surfaces at 4.0 kN force. The marked point is the equal-weight grid selection; blank regions lie outside the design hull. Synthetic predictions only.

The surface plots retain the two responses separately. Combining them into a single colour map would be convenient, but it would hide whether a favourable objective value arose from balanced gains or from compensation between differently scaled responses. The selected point is marked on each response panel, and regions outside the hull remain blank. No contour is presented as a measured thermal boundary or weld lobe.

5.2 Baseline-dependent improvement

Relative to the assumed baseline, the first scenario's fitted TS and CT responses increase by 16.88% and 25.73%. The corresponding increases for the second scenario are 17.96% and 25.08%. These percentages are correctly calculated for the stated model and baseline. They are not demonstrated improvements over a factory process, and they cannot be generalized to another starting schedule.

Table 4. Baseline and candidate comparison; fitted synthetic means only.

Scenario	Schedule	TS (kN)	CT (kN)	TS change %	CT change %
DP780	Baseline	7.586	3.266	0.00	0.00
DP780	Grid selection	8.866	4.106	16.88	25.73
DP780	Interior	8.795	4.015	15.94	22.94
DP980	Baseline	8.147	3.394	0.00	0.00
DP980	Grid selection	9.611	4.245	17.96	25.08
DP980	Interior	9.498	4.193	16.58	23.54

The absolute load values are more portable than the percentage language because they retain the response scale, but even they remain synthetic. A reader should not infer that the baseline is deliberately poor or that the selected schedule is physically attainable. The comparison serves to demonstrate how optimization reporting depends on reference conditions. Choosing the design centre instead would produce different improvement percentages while leaving the fitted surfaces unchanged.

5.3 Preference sensitivity

When the TS weight is 0.25, the DP780-labelled optimum is 9.20 kA, 310 ms and 4.10 kN. At equal weight, force decreases to 4.00 kN. At a TS weight of 0.75, the selection becomes 9.20 kA, 305 ms and 3.80 kN. The change reveals a trade-off in the assumed surfaces: the schedule preferred when CT receives more emphasis is not identical to the schedule preferred when TS dominates.

For the DP980-labelled scenario, the corresponding selections are 9.00 kA, 300 ms and 4.00 kN; 9.10 kA, 295 ms and 4.00 kN; and 9.20 kA, 295 ms and 3.90 kN. Again, these differences describe preference dependence within the synthetic model. Their practical significance would depend on machine resolution, repeatability and the uncertainty of the physical responses, none of which is established here.

Table 5. Preference sensitivity; all settings selected from synthetic fitted models.

TS weight	DP780: kA / ms / kN	DP980: kA / ms / kN
0.25	9.20 / 310.00 / 4.10	9.00 / 300.00 / 4.00
0.5	9.20 / 310.00 / 4.00	9.10 / 295.00 / 4.00
0.75	9.20 / 305.00 / 3.80	9.20 / 295.00 / 3.90

The results discourage treating one reported coordinate triple as an immutable material property. Even without changing the data, an analyst can obtain a different preferred schedule by changing the value judgement represented by the objective. A transparent paper should state whether the weights were agreed in advance, derived from a specification or selected after inspecting results. The present weights are openly declared sensitivity cases.

5.4 Local and neighbourhood stability

Table 6 records fitted response changes under the one-factor perturbations. At a combined-objective optimum, it is not necessary for both individual responses to have zero slope in every direction. A small move can improve one response while reducing the other, leaving the combined objective lower. This is a defining feature of a trade-off solution and should not be mistaken for an inconsistency in the optimization.

Table 6. Local one-factor challenges; fitted synthetic loads in kN.

Change	780 TS	780 CT	980 TS	980 CT
Nominal	8.866	4.106	9.611	4.245
Current -0.2 kA	8.851	4.106	9.586	4.254
Current +0.2 kA	8.868	4.098	9.619	4.226
Time -10 ms	8.869	4.100	9.615	4.233
Time +10 ms	8.854	4.105	9.595	4.247
Force -0.1 kN	8.883	4.092	9.616	4.240
Force +0.1 kN	8.843	4.115	9.597	4.243

The discrete worst-neighbourhood analysis produces the comparison in Table 7. It distinguishes nominal desirability from the least favourable score among the twenty-seven declared challenges. A candidate may sacrifice a small amount of nominal performance to improve that minimum, or it may remain unchanged if the nominal optimum already balances the sampled disturbances well. The numerical table, rather than an assumed preference for central points, determines which situation occurs.

Table 7. Additional discrete worst-neighbourhood analysis; fixed synthetic models and 27 offsets.

Scen ario	Robust kA / ms / kN	Nomin al D	Nomina l worst	Robust D	Robust worst
DP78 0	9.20 / 305.00 / 4.00	0.5499	0.5402	0.5498	0.5416
DP98 0	9.10 / 295.00 / 4.00	0.5168	0.5089	0.5168	0.5089

For the DP780-labelled scenario, the discrete worst-neighbourhood selection is 9.20 kA, 305 ms and 4.00 kN. Its minimum challenged score is 0.5416, compared with 0.5402 at the nominal selection, while nominal desirability decreases slightly from 0.549870 to 0.549817. For the DP980-labelled scenario, the selected schedule is unchanged and the minimum challenged score is 0.5089. The small first-scenario improvement is a property of the fixed synthetic surfaces and sampled disturbances; its magnitude does not establish a practically significant improvement in welding.

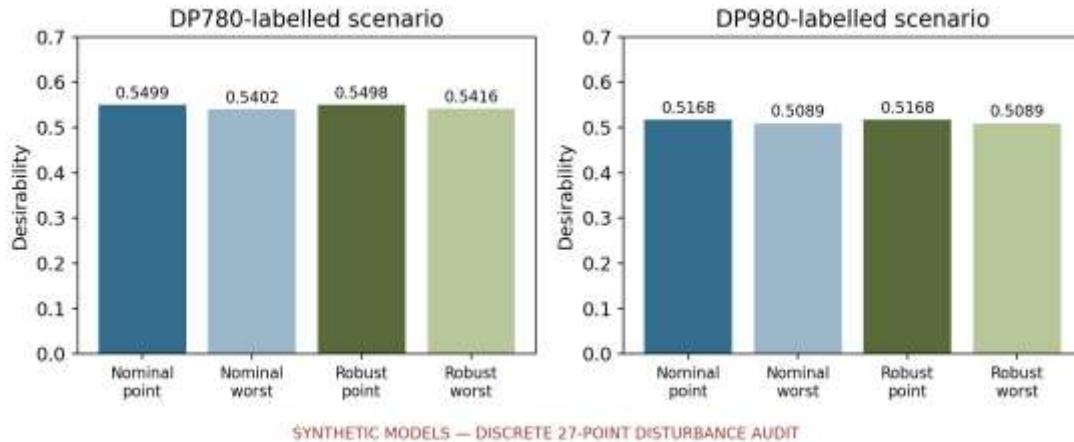


Figure 3. Nominal and minimum challenged desirability at nominal and discrete worst-neighbourhood selections. Twenty-seven deterministic disturbances are evaluated; no production reliability is implied.

This robustness audit is deliberately modest. It does not include electrode deterioration, batch chemistry, coating changes, correlated timing errors or a shift in the fitted coefficients. Those are different uncertainty sources. The result should therefore be read as local numerical stability with fixed coefficients, not as a reliability claim about actual production.

6. Model uncertainty and confirmation limits

6.1 Diagnostic concerns inherited from modelling

The optimization layer inherits the limitations of the surfaces on which it operates. In the source calculation, the DP780-labelled CT model has a lack-of-fit probability of 0.0304 at the row-mean analysis scale. The present manuscript does not resolve that diagnostic by changing the objective function. Its optimum remains an illustration conditional on a model that has not been cleared for physical prediction.

The other response models have favourable overall fit statistics, but that does not supply experimental validation. Their generators are quadratic, their noise is assumed and their domain is small. Beder (2022) provides the relevant linear-model background; the inferential point here is that an optimizer cannot upgrade the evidence supporting its inputs. A precise maximum of an uncertain model is still conditional on that uncertain model.

6.2 New random draws are not new experiments

The source thesis generates five additional synthetic values at selected schedules after model fitting. These draws are independent of the fitted sample conditional on the same generator, making them useful as a numerical holdout check. They do not test whether the generator describes real steel. Agreement shows consistency with the assumed data-generating mechanism, not confirmation of a weld qualification.

This distinction is important because the phrase confirmation experiment can carry more authority than the data justify. A true physical confirmation would require newly produced welds, traceable materials, calibrated testing and predefined success criteria. It would also need to record adverse outcomes rather than retain only specimens suitable for a preferred summary. A table of random draws cannot stand in for those activities.

6.3 Selection uncertainty

The present calculation treats fitted coefficients as fixed during the grid search and local disturbance audit. In reality, coefficient uncertainty can change the location of a preferred point. A future simulation extension could repeatedly generate or resample datasets, refit the surfaces and record the distribution of selected schedules. That would estimate selection stability under a specified statistical mechanism, but it would still not quantify unmodelled physical effects without supporting data.

Selection uncertainty is different from response uncertainty at one fixed candidate. A narrow predicted interval at a selected point does not automatically mean that the same point would be selected after a new dataset. Near-flat objectives can have many practically equivalent schedules whose ranking changes with small coefficient perturbations. Reporting a near-optimal set may therefore communicate more useful information than printing excessive decimal places for one maximizer.

6.4 Practical equivalence and machine resolution

The grid increments specify the numerical resolution of this search. They are not a statement about the welding controller's attainable or repeatable settings. If a real machine can only implement coarser changes, candidates should be re-evaluated on that implementable grid. Rounding the selected schedule after optimization without recalculating its responses and constraints can change the problem that was actually solved.

Conversely, searching a much finer grid than the measurement system can support may create an illusion of precision. A distinction between 9.10 and 9.11 kA is not automatically meaningful if the delivered current varies more than that difference. The correct response is to align numerical resolution, measurement capability and decision tolerance, not to assume that more digits imply better engineering.

7. Decision implications

7.1 Separate performance from acceptability

An objective ranks acceptable alternatives; it should not quietly replace the definition of acceptability. In the numerical example, lower desirability values act as simplified response thresholds, while geometry introduces additional inequalities. In a real application, some requirements may be mandatory regardless of weighted performance. The distinction should be established before optimization so that a high score cannot compensate for an unacceptable failure mode or a missing safety requirement.

Acceptance criteria also have provenance. They may come from a product specification, an agreed design requirement or a qualified procedure. If they are merely assumed for a teaching example, that must be stated. The current paper uses assumed bounds and therefore cannot claim standards compliance. It also does not reproduce or invent test-standard dimensions, acceptance limits or certification language.

7.2 Avoid universal factor rankings

The preference sensitivity shows why a universal statement such as current is the most important factor is too coarse for this problem. The relevant influence depends on response, location, coding range and objective. A factor can have a large effect on TS yet a different role in CT feasibility. It can also matter primarily through an interaction, making its apparent main effect an incomplete description.

A useful engineering summary would identify which setting changes improve a defined response locally, which constraints become active and which disturbances most reduce the agreed objective. These are actionable statements tied to a domain. They are more informative than a single rank order detached from units and acceptance conditions. The present tables provide the structure for such a summary without claiming that the synthetic rankings apply physically.

7.3 Candidate sets rather than a single recipe

The equal-weight, CT-emphasized, TS-emphasized, interior and worst-neighbourhood candidates form a small decision set. In a laboratory programme, these candidates could be compared using new welds under identical material and equipment conditions. A set-based approach makes the trade-off visible and reduces the temptation to treat the first numerical maximum as the only scientifically relevant schedule.

The choice among candidates should be based on observed performance and predefined priorities. An investigator should not favour the candidate that happens to deliver the largest isolated test result while ignoring scatter or failure morphology. If several schedules perform equivalently within uncertainty, the preferred operating point may depend on practical considerations such as controller resolution or distance from an observed defect boundary. Those considerations require measurement rather than assumption.

7.4 Transfer to another material system

The numerical framework can be reused, but the coefficients and thresholds cannot be transferred without evidence. A change in sheet thickness, coating, stack arrangement or electrode geometry creates a different process system. Reusing the same coordinate ranges might be a planning convenience, yet it does not establish that the same feasible region or optimum remains meaningful. The model must be recalibrated and the physical domain requalified.

Likewise, the two labelled scenarios do not establish a general grade effect. Comparing them is useful for inspecting how the decision procedure responds to different assumed surfaces. It is not equivalent to a controlled experiment comparing two steels. A real grade comparison would require a design that addresses procurement differences, batch effects and other potential confounders before attributing changes to nominal strength class.

8. Reproducibility and reporting requirements

The optimization can be reproduced from the fitted coefficient table, coding definitions, desirability limits, weight values, grid increments, hull rule and geometry functions. Candidate ordering should be retained to resolve exact ties. The nominal, interior and perturbation calculations require no new random draws. The worst-neighbourhood extension likewise uses deterministic evaluation of fixed surfaces.

A reproduction should verify intermediate quantities before checking the final optimum. The unfiltered grid contains 35,301 points, and the hull-supported grid contains 28,701. The coded equal-weight selections are (0.60, 0.60, 0.50) and (0.55, 0.45, 0.50). Converting them back to physical units must reproduce the reported schedules. Checking these quantities can reveal errors in scaling or grid construction that a visually similar contour plot might conceal.

The reported improvements must use the same response predictions as the selection analysis. Mixing generator means with fitted means would change the comparison and should be labelled as a different calculation. Similarly, the interior point is defined in coded coordinates, not by halving the physical current, time and force. Confusing those operations would move the comparator outside the intended interpretation.

The manuscript preserves both numerical and semantic provenance. Every table identifies its quantities as synthetic, fitted or assumed. Figures are original computational graphics, not photographs of tested joints. The referenced books supply general conceptual context, whereas the thesis calculation supplies the synthetic coefficients. No reference is represented as the source of an invented experimental dataset.

9. Analytical interpretation of decision sensitivity

9.1 The local balance behind a compromise

Inside the region where both individual desirabilities are positive and below one, maximizing combined desirability is equivalent to maximizing its logarithm. The local objective becomes $w \log(dTS) + (1-w) \log(dCT)$. Its gradient is a weighted combination of response gradients divided by their distances above the respective lower limits. This expression makes the meaning of a compromise more transparent than the final coordinate triple alone.

For example, a response close to its lower limit can exert a strong influence on the logarithmic objective because a small absolute improvement represents a large proportional increase in its desirability. The same load increment has a different effect when the response is already far above the lower limit. Thus, the chosen lower bounds shape the local trade-off even before the explicit preference weight is changed. Calling the objective neutral would conceal that dependence.

At an unconstrained interior stationary point, the weighted gradient sum is zero. That does not require the TS gradient or the CT gradient to vanish separately. Their directional influences can oppose one another. At a constraint boundary, the feasible directions are restricted, and a zero unconstrained gradient is no longer the appropriate condition. These distinctions explain why inspecting individual response peaks is not enough to verify the combined constrained selection.

The grid search avoids solving these conditions analytically, but they remain useful checks. A purported optimum with a clear neighbouring improvement on the same admissible grid would indicate an implementation error. Conversely, a selected point that is not stationary in continuous coordinates may simply reflect grid resolution. The interpretation should follow the actual discrete problem rather than a continuous optimum that was never computed.

9.2 Dominance and preference

A candidate is dominated when another feasible candidate has at least as high a value for both responses and a strictly higher value for at least one. In the unsaturated positive desirability region, such a dominated point cannot improve the geometric objective for positive weights. This provides a useful conceptual separation: dominance removes clearly inferior alternatives, whereas preference weights choose among alternatives that trade one response against another.

Clipping complicates the language slightly. If a response exceeds its target, further increases receive no additional desirability, so two physically different response vectors may have the same objective contribution. Likewise, zero-score

regions can contain many candidates that are indistinguishable to the objective despite differing loads. Tie handling is therefore not merely a programming detail. It should be disclosed whenever the score transformation can create large equivalence classes.

An industrial decision may also include objectives not present here, such as cycle time, electrode consumption or energy. Adding them would change the dominance relation and possibly the preferred candidate set. They should not be invoked informally after a load-only optimization as though the selected point had already optimized them. A claim of improved productivity, for example, would require an appropriately measured and modelled productivity response.

9.3 Why moving toward the centre is not a proof of robustness

The interior comparator is geometrically conservative because it moves away from the coded extremes. Statistical and physical robustness are more complicated. A central point may have stronger design support but lower response margins. A less central point may lie on a broad objective plateau. Neither geometry alone nor nominal score alone determines how a candidate responds to disturbances.

The worst-neighbourhood calculation therefore evaluates the declared disturbances explicitly. Its feasibility requirement is also important: every challenged point must remain in the admissible region. If a candidate lies near the hull boundary, a favourable nominal score does not rescue it when the selected disturbance set leaves that region. This is a deliberate decision convention, not a discovery that all boundary settings are physically unsafe.

The disturbance set should be chosen to match the intended question. A symmetric coded box is convenient for this example, but actual machine errors may be asymmetric, correlated or constrained by controller logic. Replacing that box with a measured uncertainty set would be a meaningful extension. The resulting robust candidate could differ even when the nominal response model remains the same.

9.4 A worked interpretation of the first scenario

For the DP780-labelled equal-weight selection, the fitted TS response is about 1.3665 kN above its lower desirability bound of 7.5 kN. Dividing by the 2.5 kN scaling interval gives a TS desirability of approximately 0.5466. The fitted CT response is about 1.1063 kN above its lower bound of 3.0 kN; dividing by its 2.0 kN interval gives approximately 0.5532. Their geometric mean is about 0.5499.

This calculation explains why the reported objective is far below one even though the candidate is the best on the declared grid. The targets are not both attained by the fitted responses. A score of 0.55 is not a fifty-five-percent probability of success, nor does it mean that fifty-five percent of a product specification has been satisfied. It is simply the numerical result of the chosen transformations and weights.

The same distinction applies when scores are compared between the two labelled scenarios. Because their scaling limits differ, an identical desirability would not imply identical load values or identical structural performance. Cross-scenario comparisons should therefore retain the original response units and the scaling table. The dimensionless score is designed for within-problem decision making, not as a universal measure of weld quality.

9.5 Reporting a recommendation without overstating it

A defensible numerical recommendation has several clauses. It identifies the model and dataset, names the finite grid, states the objective and constraints, lists the selected physical settings and describes the validation still required. Omitting any one of those clauses can change the practical meaning. For instance, deleting the phrase on the declared grid implies a stronger continuous search result, while deleting the word synthetic implies a different evidence source.

The language should also distinguish the best candidate from an acceptable candidate. The highest score may still fail an external requirement not represented in the model. If the model predicts that no candidate meets a mandatory threshold, the correct conclusion may be that the domain or process needs further investigation. Optimization does not guarantee the existence of a suitable engineering solution merely because its software returns a maximum.

The present result is therefore a candidate-selection demonstration. It supports the next question to test, not the answer to a physical qualification question. The value of the exercise lies in exposing the assumptions that a later experiment must examine: load response shape, constraint accuracy, priority agreement and tolerance to realistic disturbances. Those assumptions are the bridge between numerical calculation and engineering evidence.

9.6 Separating computational cost from engineering value

The grid has only tens of thousands of candidates, and each candidate requires evaluation of short polynomial expressions. Computational speed is therefore unlikely to be the main limitation of this example. The expensive information is the physical evidence that would be needed to replace the assumed surfaces. An algorithm that saves a fraction of a second cannot compensate for missing calibration, poorly defined constraints or an unrepresentative material sample.

This observation changes how an optimization study should be evaluated. A comparison between algorithms should account for objective evaluations, stopping rules and reproducibility, but it should also ask whether algorithmic differences matter relative to response uncertainty and machine resolution. If two methods select practically equivalent schedules under the same uncertain model, a claim of engineering superiority based on a tiny score difference would be difficult to justify.

A useful next investment is often information gathering rather than a more elaborate search. Additional observations near a sensitive trade-off region can clarify response shape. Replication can improve precision at a candidate. A measured defect boundary can replace an assumed constraint. These actions answer different questions, so the experiment should target the uncertainty most likely to change the decision rather than distribute new tests without a defined purpose.

The present case also illustrates the importance of preserving failed candidate evaluations. A candidate rejected because a challenged point leaves the hull is different from one rejected because the assumed width is too small. Recording rejection reasons makes the decision auditable and reveals which restrictions are influential. A single final score obscures that information and can make a constraint-driven result look like a purely load-driven optimum.

In a physical follow-up, the optimization report should therefore include the selected candidates, the dominant rejection reasons and the measurements needed to resolve the remaining uncertainty. The report would then function as a research decision aid rather than as an unexplained recipe. This is the intended role of the current numerical example: to organize a transparent next experiment, not to claim that the engineering problem has already been solved.

CONCLUSIONS

Equal-weight desirability selects 9.20 kA, 310 ms and 4.00 kN for the DP780-labelled synthetic scenario and 9.10 kA, 295 ms and 4.00 kN for the DP980-labelled scenario on the declared grid. These are conditional mathematical selections. They are not proven best welding parameters, and their response values are not laboratory measurements.

Changing response priorities changes the selected schedules without changing the underlying data. Baseline-dependent percentages, local perturbations and the discrete worst-neighbourhood audit reveal additional dependencies that a single optimum table would hide. The analysis therefore supports reporting a transparent candidate set and its assumptions rather than an unconditional recipe.

The source model's diagnostic limitations remain active, including the adverse DP780-labelled CT lack-of-fit result. Physical recommendations require new material-specific welding and testing, with agreed acceptance rules and complete process records. Optimization is a decision tool applied to evidence; it is not a mechanism for converting an assumed surface into experimental proof.

Declarations and relationship to the thesis

This draft is derived from the supplied AHSS spot-welding thesis and shares its synthetic dataset with the companion response-surface modelling manuscript. The deterministic worst-neighbourhood calculation is an additional analysis, not a new physical experiment. Thesis author, institution, year and repository details must be completed before submission. Actual authors must supply affiliations, contributions, funding and conflict-of-interest statements. AI assistance was used in drafting and computational presentation and must be disclosed according to the target journal's policy. No laboratory photographs, metallographic observations or production qualification results are claimed. Separate publication of related manuscripts requires editorial assessment of overlap and distinct contribution.

REFERENCES

1. Antoniou, A., & Lu, W.-S. (2021). Practical optimization: Algorithms and engineering applications (2nd ed.). Springer. <https://doi.org/10.1007/978-1-0716-0843-2>
2. Beder, J. H. (2022). Linear models and design. Springer. <https://doi.org/10.1007/978-3-031-08176-7>
3. Dwivedi, D. K. (2022). Fundamentals of metal joining: Processes, mechanism and performance. Springer. <https://doi.org/10.1007/978-981-16-4819-9>
4. Feng, L., Hou, Y., & Zhu, Z. (2021). Optinformatics in evolutionary learning and optimization. Springer. <https://doi.org/10.1007/978-3-030-70920-4>
5. Kimchi, M., & Phillips, D. H. (2023). Resistance spot welding: Fundamentals and applications for the automotive industry (2nd ed.). Springer. <https://doi.org/10.1007/978-3-031-25783-4>
6. Stein, O. (2024a). Basic concepts of global optimization. Springer. <https://doi.org/10.1007/978-3-662-66240-3>
7. Stein, O. (2024b). Basic concepts of nonlinear optimization. Springer. <https://doi.org/10.1007/978-3-662-69741-2>
8. Source thesis (unpublished document supplied for this work): Optimization of Process Parameters for Spot Welding of Advanced High-Strength Steels. Author, institution and year were not supplied and must be completed before submission. The source thesis is separate from the published-book references above.