

Constraint Locality in Learned Syndrome Decoding for 3D Occupancy Grids

Brian Guarraci
Galatheus
brian@galatheus.com

May 2026

Abstract

We study learned reconstruction for underdetermined binary syndrome systems arising from 32^3 occupancy grids. A syndrome $\mathbf{H}\mathbf{x} = \mathbf{s} \pmod{2}$ specifies an exponentially large coset, so reconstruction depends on a data-dependent learned proposal conditioned on coarse occupancy, low-frequency spectral side information, and syndrome-derived features. The central question is not whether a syndrome can uniquely reconstruct a grid, but which constraint structures help a learned proposal refine its prediction.

The main empirical finding is a constraint-locality effect. Hard algebraic projection with fixed high-weight checks can degrade a plausible learned prediction, while belief propagation on local overlapping surface checks can improve it. We frame the result as a design rule for learned syndrome-based occupancy coding: constraints should be local, overlapping, and targeted to uncertain regions of the proposal. We do not claim a complete production codec or state-of-the-art compression system; the contribution is the experimental isolation of when algebraic constraints help or hurt learned reconstruction.

1 Introduction

Binary syndrome decoding—recovering $\mathbf{x} \in \{0,1\}^N$ from $\mathbf{s} = \mathbf{H}\mathbf{x} \pmod{2}$ given a parity check matrix $\mathbf{H} \in \{0,1\}^{M \times N}$ —is a fundamental problem in coding theory [1]. When $M < N$, the system is underdetermined: the syndrome constrains $\mathbf{x}$ to a coset of size $2^{N-\text{rank}(\mathbf{H})}$, but does not uniquely determine it. The decoder must select from this exponentially large set using additional information.

In classical channel coding, the additional information is the channel observation (soft bit likelihoods). In compressed sensing [4], it is a sparsity prior. In this work, we study a setting where the additional information is a **learned spatial proposal**: a neural network trained on data from the target domain (3D voxel grids of real objects) that predicts likely solutions from spectral side-information and syndrome-derived features.

This setting arises naturally in 3D volumetric compression, where the encoder computes a syndrome of a binary voxel grid and transmits it alongside spectral coefficients. In this paper we keep the scope narrow: we evaluate 3D occupancy-grid reconstruction and use compression as the motivating rate-distortion setting. We do not rely on cross-domain analogies for the core result.

This is the compression member of a four-paper program on learned recovery over constraint systems. Its role is deliberately local: establish when parity checks help or hurt a learned spatial proposal. Companion papers use that constraint-locality lesson as framing for surface-code QEC, prove the axial-parity gauge theorem, and test a structure-aware neural decoder for the X-cube fracton code; none of those claims is needed for the empirical result reported here.

The work is closest in spirit to syndrome source coding with decoder side information [13, 12] and neural decoders for linear codes [5, 6]. We do not derive a new rate-distortion theorem and

we do not claim a state-of-the-art neural compression system in the lineage of end-to-end learned image codecs [7, 8]. Instead, we isolate a design question that appears once a learned decoder is already strong: which parity-check structures improve the learned proposal, and which structures damage it?

Our contributions:

1. **Diagnosis**: we evaluate how a fixed high-weight parity system interacts with a learned spatial proposal on 32^3 ShapeNet occupancy grids (Section 3).
2. **A constraint locality principle**: we compare hard algebraic projection against BP refinement with local overlapping surface checks and identify the structural conditions under which syndrome information improves the proposal (Section 4). We support the observation with a simple check-weight analysis (Theorems 1–2).
3. **Content-adaptive constraints**: we show that $\mathbf{H}$ can be derived from shared side-information (low-frequency spectral coefficients) available to both encoder and decoder, requiring no additional constraint-location signaling bits (Section 6).
4. **Density-stratified evidence**: we show that local BP helps in every occupancy quartile and helps most on the densest validation objects, while hard projection hurts in every quartile.

2 Problem Setting

Let $\mathbf{x} \in \{0, 1\}^N$ be a binary vector with spatial structure (a d^3 voxel grid, $N = d^3$). The encoder computes:

- A syndrome $\mathbf{s} = \mathbf{H}\mathbf{x} \bmod 2 \in \{0, 1\}^M$ from a parity check matrix $\mathbf{H}$.
- Spectral side-information $\mathbf{c}$ (e.g., low-frequency transform coefficients of $\mathbf{x}$).

The decoder observes $(\mathbf{s}, \mathbf{c})$ and must recover $\mathbf{x}$.

With the corrected fixed constraint set, $M = 2,977$ and $N = 32,768$; the syndrome eliminates only 2,977 bits of entropy, leaving $2^{29,791}$ valid solutions. The spectral coefficients provide a soft spatial proposal $P_\theta(x_i = 1 | \mathbf{c}, \mathbf{s})$ through the learned decoder, but the decoder’s fidelity is bandwidth-dependent: the post-projection L2 result is IoU 0.6947 ± 0.0072 , while the L4 learned proposal before projection is IoU 0.9370 ± 0.0028 . The question is: **how should the decoder combine the algebraic constraints with the spatial proposal?**

2.1 Three Decoding Operations

(A) Hard algebraic projection. Given a learned proposal $\tilde{\mathbf{x}}$, compute its residual $\mathbf{r} = \mathbf{s} - \mathbf{H}\tilde{\mathbf{x}}$ over $\text{GF}(2)$ and apply a right-inverse correction $\hat{\mathbf{x}} = \tilde{\mathbf{x}} + \mathbf{L}\mathbf{r}$, where $\mathbf{H}\mathbf{L} = \mathbf{I}_M$. This satisfies the syndrome exactly but chooses a correction by algebraic convenience rather than spatial coherence. The canonical solution $\mathbf{L}\mathbf{s}$ is the special case $\tilde{\mathbf{x}} = 0$.

(B) Learned proposal decoding. Train a neural network $f_\theta(\mathbf{x}_0(\mathbf{s}), \mathbf{c}, \mathbf{s}) \rightarrow \hat{\mathbf{x}}$ on domain data, where $\mathbf{x}_0(\mathbf{s})$ is a canonical syndrome solution. The network learns what spatially plausible solutions look like, but its output is not guaranteed to satisfy the syndrome exactly.

(C) Proposal + constraint refinement. Use the learned proposal as an initial estimate, then refine via belief propagation on the Tanner graph [2, 3] of $\mathbf{H}$ with the proposal’s confidence as variable beliefs and $\mathbf{s}$ as check values.

We show that (A) fails in this setting, (B) is strong, and (C) is better than (B) *only when $\mathbf{H}$ has the right structure*.

3 Diagnosis: Why Hard Algebraic Projection Fails

The clean rsfix evaluation of operation (A) uses 32^3 ShapeNet voxel grids [10] with the corrected axis + plane + global parity checks ($M = 2,977$). After learned proposal decoding (operation B,

mean IoU 0.9370 ± 0.0028 at spectral level 4), we project onto the constraint surface via the right inverse.

Table 1: Clean hard algebraic projection diagnosis (200 ShapeNet validation samples, spectral level 4, mean $\pm$ SEM over seeds 42/43/44). The algebraic correction is usually wrong even though the corrected constraint set has only 1.6% violated checks before projection.

Metric	Value
Constraint violations before projection	47.6 ± 4.5 / 2,977 (1.6%)
Bits flipped by projection	52.5 ± 5.2 / 32,768 (0.16%)
IoU before projection	0.9370 ± 0.0028
IoU after projection	0.8782 ± 0.0085 (-0.0588 ± 0.0058)
Mean $ \text{logit} $ of flipped bits	8.665 ± 0.020
Mean $ \text{logit} $ of kept bits	14.151 ± 0.016
Fraction of helpful flips	5.3% $\pm$ 0.5%

The projection flips only 52.5 bits on average (0.16% of the grid), but 94.7% of those flips move *away* from the ground truth. The right inverse finds the minimum-Hamming-weight correction to satisfy the syndrome, but this algebraically minimal correction has no relationship to the spatially correct one. The projection operates at the scale of full rows (32 voxels), while errors are localized at object surfaces (2–8 voxels). This scale mismatch is the root cause.

4 The Constraint Locality Principle

We compare five post-processing strategies applied to the same learned proposal output (Table 2). Hard projection uses a right inverse; BP refinement uses the proposal’s soft output as variable beliefs.

Table 2: Clean constraint design ablation. Same learned proposal (IoU 0.9370 ± 0.0028), different post-processing. 200 ShapeNet validation samples, spectral level 4, mean $\pm$ SEM over seeds 42/43/44. Extra bits are additional local syndrome values beyond the fixed global-syndrome payload; this is a constraint-design ablation, not a rate-matched codec comparison.

Post-processing	Weight	Overlap	Extra bits	Δ IoU
None (proposal only)	—	—	0	baseline
Hard projection (global)	32–1024	low	0	-0.0588 ± 0.0058
<i>BP refinement with surface constraints:</i>				
Surface patches (2^3)	8	yes	4,205	$+0.0363 \pm 0.0026$
Surface gradient runs	4	yes	12,682	$+0.0603 \pm 0.0031$
Surface neighbor pairs	2	yes	12,806	$+0.0623 \pm 0.0028$

Numbers from `reproduce_results.py` Sections 5–6, clean rsfix artifacts in `data/clean-v2-may02-rsfix/`.

The learned decoder gives a smooth nominal rate–distortion diagnostic as more low-frequency spectral side-information is made available (Table 3). These numbers are post-hard-projection reconstructions before surface-BP refinement; the L4 row matches the hard-projection setting in Table 1. The table measures the bandwidth available to the learned proposal plus the cost of the fixed syndrome, under an 8-bit-per-retained-coefficient accounting model.

Because the ShapeNet validation subset is occupancy-sparse, we also stratify the L4 results by occupied-voxel density (Table 4). The strata are rank quartiles of the same 200 validation samples, so each row has equal support. IoU ignores true negatives, so these are not background-accuracy effects: the table measures overlap on occupied structure.

Table 3: Clean nominal rate–distortion results on 200 ShapeNet validation samples, mean $\pm$ SEM over seeds 42/43/44. Bits count the 2,977-bit syndrome, 8 coarse bits, and 8 bits per retained low-frequency spectral coefficient; no entropy-coded coefficient bitstream is implemented.

Spectral level	Bits	Mean IoU
L2	3,201	0.6947 ± 0.0072
L3	3,985	0.7941 ± 0.0081
L4	5,729	0.8781 ± 0.0085
L5	8,817	0.9380 ± 0.0072
L6	13,633	0.9729 ± 0.0036
L7	20,561	0.9908 ± 0.0017
L8	29,985	0.9970 ± 0.0006

Table 4: Clean L4 results stratified by occupancy rank quartile. Each stratum contains 50 validation objects per seed; entries are mean $\pm$ SEM over seeds 42/43/44. Local BP helps in every quartile and helps most in the densest, hardest objects.

Stratum	Occupancy	Raw IoU	Hard Δ	Patch Δ	Neighbor Δ
Q1	0.62–1.68%	0.946 ± 0.003	-0.059 ± 0.007	$+0.035 \pm 0.002$	$+0.053 \pm 0.003$
Q2	1.68–1.92%	0.947 ± 0.003	-0.053 ± 0.004	$+0.034 \pm 0.002$	$+0.052 \pm 0.003$
Q3	1.93–2.27%	0.944 ± 0.003	-0.058 ± 0.006	$+0.035 \pm 0.003$	$+0.056 \pm 0.003$
Q4	2.28–5.12%	0.911 ± 0.002	-0.065 ± 0.006	$+0.041 \pm 0.003$	$+0.088 \pm 0.002$

Generated by `stratify_compression_results.py`; aggregate JSON:
`data/clean-v2-may02-strata/strata_summary.json`.

The Q4 monotonic increase aligns with the locality argument in Section 5. Denser objects expose more occupied surface, which in turn places more local checks within the BP support. The single-error informativeness $P_1(w, p)$ is a per-check quantity, but reconstruction quality scales with the number of locally informative checks that cover uncertain voxels. Q4 has the largest fraction of such voxels by construction, so a fixed local check family yields the largest ΔIoU there. The same effect is invisible to hard projection because its high-weight checks have negligible P_1 regardless of which stratum they cover.

All three successful BP variants share three properties:

1. **Locality**: few variables per check (≤ 8), so corrections are spatially contained.
2. **Overlap**: adjacent checks share variables, enabling BP to propagate information between patches.
3. **Targeting**: checks placed where the proposal is uncertain (object boundaries), not uniformly over the grid.

Hard projection fails because it operates at the wrong spatial scale (full rows of 32 voxels) and uses minimum Hamming weight rather than spatial coherence as its correction objective.

5 Why Constraint Weight Controls Local Informativeness

The empirical results in Table 2 show that syndrome-guided correction can either help or hurt a learned proposal. A simple check-weight calculation explains why high-weight checks are often uninformative for local correction.

Theorem 1 (Correction Utility). *Let $\hat{\mathbf{x}}$ be an estimate of $\mathbf{x}^* \in \text{GF}(q)^n$ where each symbol is independently correct with probability $p \in (1/2, 1)$. For a linear check of weight w (i.e.,*

$\sum_{i=1}^w h_i x_i = s$), the probability that the check contains exactly one error is

$$P_1(w, p) = w(1-p)p^{w-1}. \quad (1)$$

In this single-error regime, the check residual determines the correction value once the erroneous symbol is identified by the prior or by overlapping checks. This function is maximized at $w^* = -1/\ln p$ and satisfies

$$P_1(w, p) \leq w \exp(-(w-1)\ln(1/p)), \quad (2)$$

which decays exponentially for $w \gg w^*$.

Proof. Each of the w symbols in the check is independently wrong with probability $1-p$. The event “exactly one wrong” has probability $\binom{w}{1}(1-p)^1 p^{w-1}$, giving (1). Setting $\partial P_1/\partial w = 0$ (treating w as continuous) yields $(1-p)p^{w-1}[1+w\ln p] = 0$, so $w^* = -1/\ln p$. The bound follows from $p^{w-1} = \exp((w-1)\ln p)$. $\square$

Example. For a prior with $p = 0.88$: $w^* \approx 8$, $P_1(4) = 0.33$, $P_1(64) = 5 \times 10^{-4}$. With $m = 32$ checks, local checks ($w=4$) yield ~ 10.6 single-error-informative checks; global checks ($w=64$) yield ~ 0.016 .

Theorem 2 (Expected utility of residual-only correction). *Consider a weight- w check over $\text{GF}(2)$ on a codeword $\mathbf{x}^*$, and an estimate $\hat{\mathbf{x}} = \mathbf{x}^* + \mathbf{e}$ whose entries within the check support are independently in error with probability $1-p$. Define the uniform-residual correction policy: whenever the check residual $r \neq 0$, pick a support position i uniformly at random and flip $\hat{x}_i$. Conditional on $r \neq 0$ (equivalently, the support contains an odd number of errors), the expected change in error count is*

$$\mathbb{E}[\Delta_{\text{err}} \mid r \neq 0] = \frac{1}{Z} \sum_{\substack{k=1 \\ k \text{ odd}}}^w \binom{w}{k} (1-p)^k p^{w-k} \frac{w-2k}{w}, \quad (3)$$

with $Z = \sum_{k \text{ odd}} \binom{w}{k} (1-p)^k p^{w-k}$. For $p > 1/2$ this expectation is strictly positive whenever $w \geq 3$, with leading order $\mathbb{E}[\Delta_{\text{err}} \mid r \neq 0] = (w-2)/w + O((1-p)^2)$ as $p \rightarrow 1$.

Proof. Conditional on exactly k errors in the support, the random correction position lies on an error symbol with probability k/w (the flip removes it: $\Delta_{\text{err}} = -1$) and on a correct symbol with probability $(w-k)/w$ (the flip introduces a new error: $\Delta_{\text{err}} = +1$). The conditional expectation is therefore $(w-2k)/w$. Averaging against the distribution on k given $r \neq 0$ yields (3). As $p \rightarrow 1$, the conditional mass concentrates at $k = 1$, where $(w-2)/w > 0$ for $w \geq 3$. $\square$

Corollary 3 (A prior is necessary). *No correction policy that selects a support position using only the residual value can achieve negative expected Δ_{err} for $w \geq 3$ and $p > 1/2$.*

Proof. The residual is invariant under permutations of the support, so any policy depending only on the residual is symmetric across positions and equivalent in expectation to the uniform-residual policy. Theorem 2 then applies. $\square$

Corollary 4 (Why high-weight checks are low-utility). *For an independent prior with per-symbol accuracy p , the probability that a check is in the unambiguous single-error regime is $P_1(w, p) = w(1-p)p^{w-1}$. This probability decays exponentially once $w \gg -1/\ln p$. High-weight checks therefore spend most of their violations in a regime where the residual alone cannot identify the correction; the syndrome must be resolved by an external prior or by overlapping checks, not by the single check.*

Theorem 2 replaces an earlier deterministic claim that every multi-error correction is harmful, which is false because multi-error cancellations can sometimes reduce the error count. The probabilistic version gives the right operational content: without an external information source — the learned proposal, or overlapping local checks — the residual alone cannot lower the expected error count for $w \geq 3$ in the high-prior-accuracy regime ($p \rightarrow 1$). The same conclusion holds qualitatively over $\text{GF}(q)$ with $q > 2$ for the analogous policy that pairs a uniform position choice with the unique residual-cancelling correction at that position. The practical point is that large checks are rarely locally informative, while small overlapping checks create many single-error or near-single-error neighborhoods where the learned proposal and the syndrome can cooperate.

6 Content-Adaptive Constraints

The constraint matrix $\mathbf{H}$ need not be fixed. We derive it from the spectral side-information $\mathbf{c}$ (low-frequency transform coefficients) that both encoder and decoder share:

1. Reconstruct the low-frequency spectral field $\rightarrow$ soft spatial proposal.
2. Detect object surfaces via morphological operations on the thresholded proposal.
3. Place overlapping 2^3 patch checks at surface locations.

Since both sides perform identical surface detection from the same spectral side-information, the constraint locations require **zero additional signaling bits**. A codec that replaces fixed global checks with local surface checks would still transmit the local syndrome values themselves ($\sim 4,200$ bits at 32^3 for patch checks). Table 2 therefore treats these as extra diagnostic bits rather than folding them into the rate-distortion curve.

7 Learned Proposal and Evaluation Protocol

The learned proposal is a flat 3D residual network [9] (width 192, depth 12, 25M params) with no downsampling. Input: canonical base solution + spectral field + coarse occupancy map, conditioned on the syndrome via MLP embedding. Trained on ShapeNet [10] 32^3 voxels with random spectral bandwidth per sample.

We note that a small width-48 depth-8 ($\sim 1.68\text{M}$ param) model achieves IoU comparable to the width-192 model on this task — the learned proposal is heavily over-parameterized at 32^3 . The constraint system, not model capacity, is the binding factor for reconstruction quality at this resolution. We report this only as a qualitative scaling observation; the multi-seed results in Tables 1–4 use the width-192 model.

Tables 1–4 use the width-192 clean v2b models on 200 ShapeNet validation samples, reproduced via `reproduce_results.py`. Full provenance is stored in the local clean-v2 rsfix and strata data directories. Appendix A reports an auxiliary temporal diagnostic that uses no checkpoint-sensitive randomness and is identical across the three reproduce jobs.

8 Discussion

When do constraints help? Our experiments show that syndrome constraints improve a learned proposal only in a narrow regime: the successful variants all use local checks (few variables), overlapping structure (shared variables between adjacent checks), and surface targeting (placed where the proposal is uncertain). Hard algebraic projection with global checks actively degrades quality. This constrains the design space for syndrome-based spatial coding.

The learned proposal does most of the work. The clean learned proposal achieves IoU 0.9370 ± 0.0028 at spectral level 4 before exact hard projection. Surface BP adds $+0.0363$ to $+0.0623$ depending on the local constraint family. Post-hoc constraint enforcement is therefore

a refinement mechanism, not the main source of reconstruction quality. Constraint design should be evaluated jointly with proposal quality: strong learned proposals are the prerequisite, and local constraints provide targeted refinements.

Limitations. We tested voxel compression at 32^3 resolution only. The current experiments isolate reconstruction behavior but do not yet constitute a complete entropy-coded production codec. Rate claims must include all transmitted payloads: fixed syndrome bits, surface-refinement syndromes, coarse occupancy bits, spectral coefficients, and any variable-rate selection metadata. The current spectral side field is a controlled low-frequency diagnostic with nominal 8-bit-per-coefficient accounting, not a finished quantized coefficient bitstream. Larger resolutions and head-to-head comparison against established voxel and point-cloud codecs (e.g., Draco [11]) remain future work.

9 Conclusion

We studied underdetermined syndrome decoding for 3D occupancy grids with learned spatial proposals. The main result is not that fixed parity checks solve compression by themselves; the learned proposal does most of the reconstruction work. Syndrome constraints help only when they are local, overlapping, and targeted at uncertain regions of the proposal. This constraint-locality principle gives a concrete design rule for future learned syndrome-based occupancy codecs and a clear failure mode for hard algebraic projection with high-weight checks.

Code and exact reproduction commands are released with the paper family in the companion portfolio repository.

A Auxiliary Temporal Diagnostic

When multiple frames of the same scene are available, each frame provides an independent noisy estimate of the true solution. We include this diagnostic only to show that local BP outputs can be fused over time; it is not part of the paper’s central static-grid contribution. The per-voxel Bayesian log-odds accumulator is

$$L_t = \alpha \cdot L_{t-1} + (1 - \alpha) \cdot \log \frac{\hat{p}_t}{1 - \hat{p}_t}, \quad L_t \in [-5, 5], \quad (4)$$

where $\hat{p}_t$ is the BP-corrected probability at frame t , $\alpha = 0.7$ is the EMA coefficient, and the clamp at ± 5 prevents lock-in. Changed voxels undergo soft decay ($\times 0.3$) rather than hard reset.

Table 5: Temporal accumulation on a 20-frame morphing scene with 8% surface tampering per frame. No neural decoder; uses spectral prior + surface BP only. From `reproduce_results.py` Section 9. This deterministic benchmark is identical across seeds.

Decoding strategy	Mean IoU	After 5 frames
Single-frame surface BP	0.847	0.847
+ Bayesian accumulation	0.958	0.975

References

- [1] R. G. Gallager. Low-density parity-check codes. *IRE Trans. Inf. Theory*, 8(1):21–28, 1962.
- [2] R. M. Tanner. A recursive approach to low complexity codes. *IEEE Trans. Inf. Theory*, 27(5):533–547, 1981.

- [3] F. R. Kschischang, B. J. Frey, and H.-A. Loeliger. Factor graphs and the sum-product algorithm. *IEEE Trans. Inf. Theory*, 47(2):498–519, 2001.
- [4] E. J. Candès, J. Romberg, and T. Tao. Robust uncertainty principles: exact signal reconstruction from highly incomplete frequency information. *IEEE Trans. Inf. Theory*, 52(2):489–509, 2006.
- [5] E. Nachmani, Y. Be’ery, and D. Burshtein. Learning to decode linear codes using deep learning. In *Allerton Conference*, pp. 341–346, 2016.
- [6] E. Nachmani et al. Deep learning methods for improved decoding of linear codes. *IEEE J. Sel. Topics Signal Processing*, 12(1):119–131, 2018.
- [7] J. Ballé, V. Laparra, and E. P. Simoncelli. End-to-end optimized image compression. In *ICLR*, 2017.
- [8] D. Minnen, J. Ballé, and G. D. Toderici. Joint autoregressive and hierarchical priors for learned image compression. In *NeurIPS*, pp. 10771–10780, 2018.
- [9] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *CVPR*, pp. 770–778, 2016.
- [10] C. B. Choy, D. Xu, J. Gwak, K. Chen, and S. Savarese. 3D-R2N2: A unified approach for single and multi-view 3D object reconstruction. In *ECCV*, pp. 628–644, 2016.
- [11] Google. Draco: 3D data compression library. <https://github.com/google/draco>, 2017.
- [12] A. D. Wyner and J. Ziv. The rate-distortion function for source coding with side information at the decoder. *IEEE Trans. Inf. Theory*, 22(1):1–10, 1976.
- [13] D. Slepian and J. K. Wolf. Noiseless coding of correlated information sources. *IEEE Trans. Inf. Theory*, 19(4):471–480, 1973.