



UPPSALA
UNIVERSITET



UNIVERSITY OF
CAMBRIDGE

Steering in Theory Space: Representation Engineering of a Lagrangian Transformer

Yong Sheng Koay¹, Alessandro Favero², Sven Krippendorff²

[Ongoing-work:260X.XXXXXX]

¹Department of Physics and Astronomy, Uppsala University, Sweden

²Department of Applied Mathematics and Theoretical Physics, University of Cambridge, UK





UPPSALA
UNIVERSITET



UNIVERSITY OF
CAMBRIDGE

Steering in Theory Space: Representation Engineering of a Lagrangian Transformer

Yong Sheng Koay¹, Alessandro Favero², Sven Krippendorff²

¹Department of Physics and Astronomy, Uppsala University, Sweden

²Department of Applied Mathematics and Theoretical Physics, University of Cambridge, UK

[Ongoing-work:260X.XXXXX]

Mentioned work in this talk:

Generating particle physics Lagrangians with transformers

YSK, Rikard Enberg¹, Stefano Moretti^{1,2} and Eliel Camargo-Molina^{1,3†}

¹ Department of Physics and Astronomy, Uppsala University, Sweden

² School of Physics and Astronomy, University of Southampton, UK

³ Department of Game Design, Uppsala University, Sweden

[arxiv: 2501.09729]



How does language works?

How does language works?

v o a g u p

Alphabets

How does language works?

you

Alphabets
Words

How does language works?

give you up

Alphabets

Words

Grammar

How does language works?

Never gonna give you up, ...

Alphabets

Words

Grammar

Meaning

How does language works?

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets

Words

Grammar

Meaning

How does language works?

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets
Words
Grammar
Meaning

How does equations work?



How does language works?

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets ~ Symbols

Words

Grammar

Meaning

How does equations work?

sin x ³ + - !

How does language works?

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets ~ Symbols

Words ~ Terms

Grammar

Meaning

How does equations work?

$$\frac{x^3}{3!}$$

How does language works?

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets ~ Symbols

Words ~ Terms

Grammar ~ Math Structures

Meaning

How does equations work?

$$x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots$$

How does language works?

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets ~ Symbols

Words ~ Terms

Grammar ~ Math Structures

Meaning ~ Real World

How does equations work?

$$\sin(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots$$

How does language works?

Never gonna give you up, ...

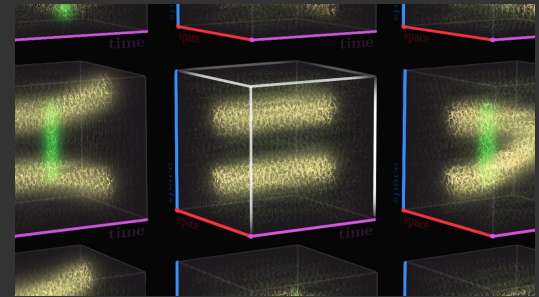
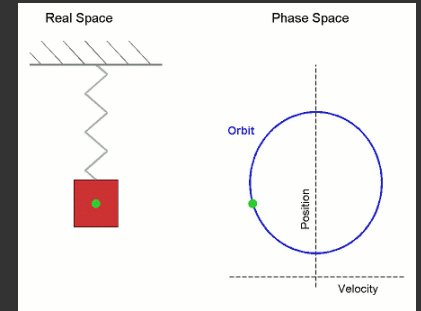


Rick Astley - Never Gonna Give You Up

Alphabets ~ Symbols
Words ~ Terms
Grammar ~ Math Structures
Meaning ~ Real World

How does equations work?

$$\sin(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots$$



language

Never gonna give you up, ...



Rick Astley - Never Gonna Give You Up

Alphabets ~ Symbols

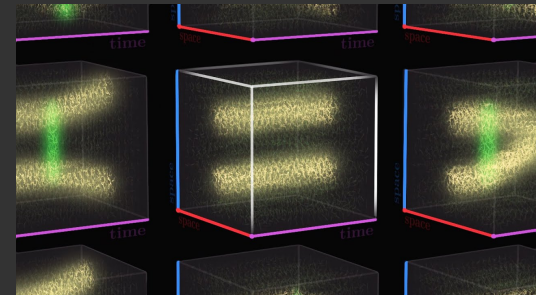
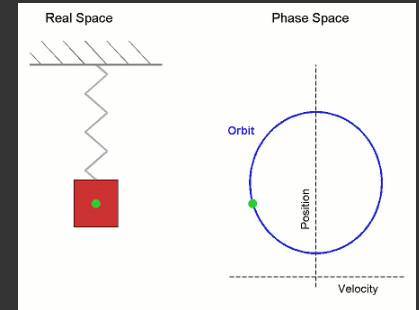
Words ~ Terms

Grammar ~ Math Structures

Meaning ~ Real World

equations

$$\sin(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots$$



language

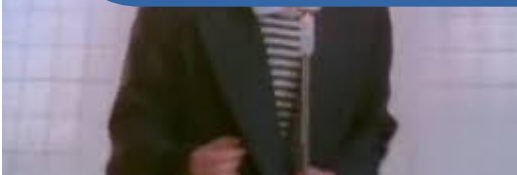
Never gonna give you up

Transformer+Language



Crazy stuff since then!

Meaning



Rick Astley - Never Gonna Give You Up



equations

$$\sin(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots$$

Transformer+Equations?



Sym



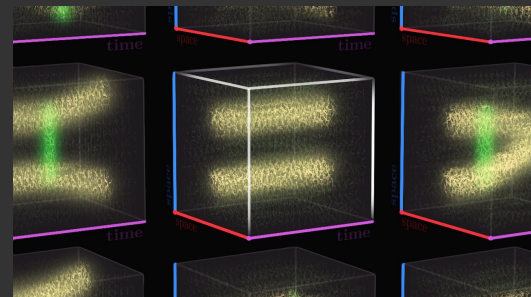
Term



Math

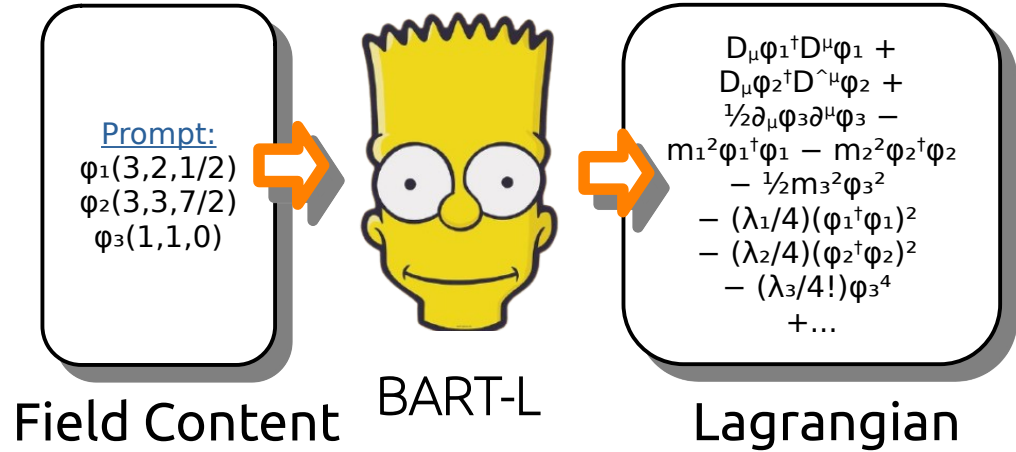


Real World



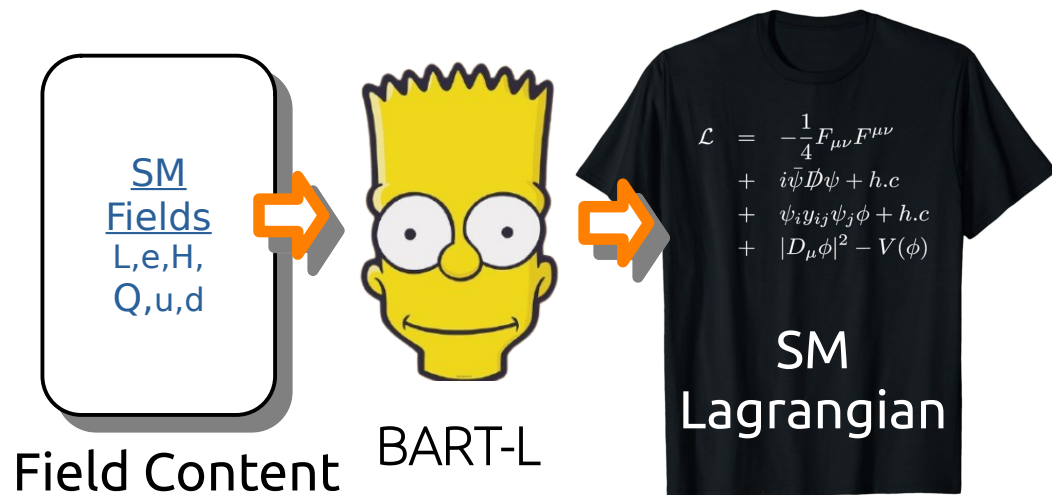
Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729



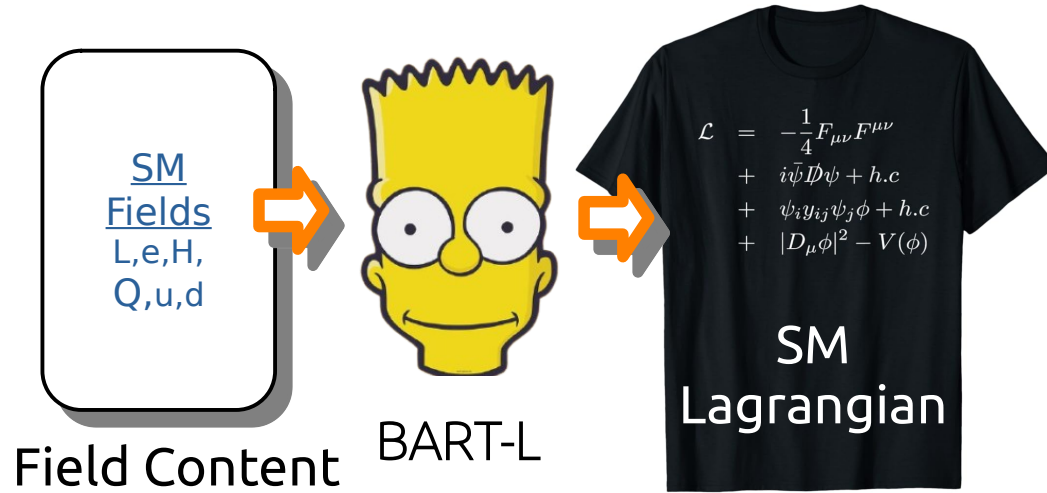
Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729



Lagrangian Transformer, BART-L

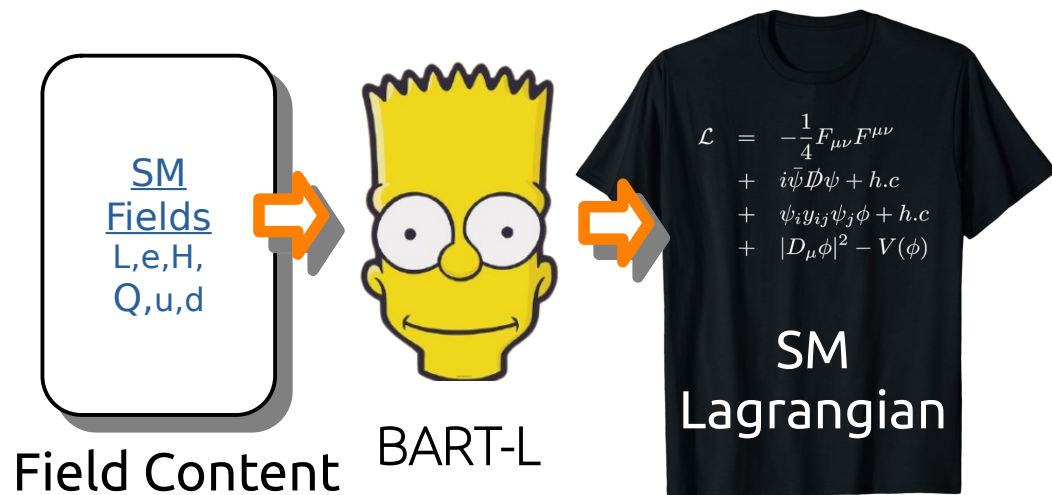
YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729



- Encoder-Decoder [357M parameters]
- Trained from Scratch, Doesn't speak English
- SU(3)xSU(2)xU(1) group, up to 6 matter fields
- 90~97% success rate
- Multitask by nature → THEP foundation model

Lagrangian Transformer, BART-L

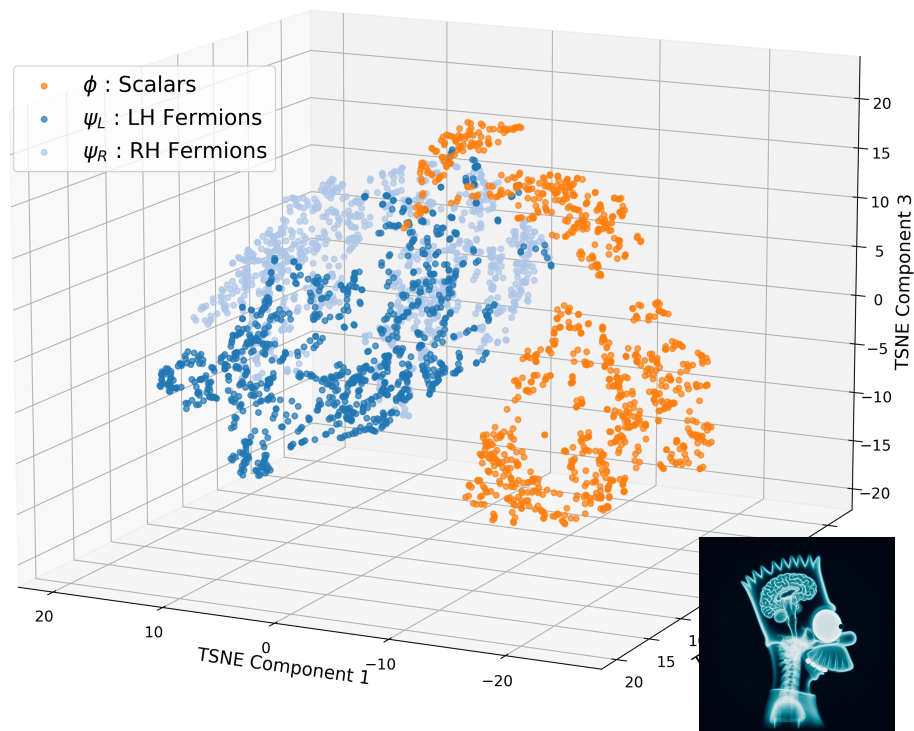
YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729



- Encoder-Decoder [357M parameters]
- Trained from Scratch, Doesn't speak English
- $SU(3) \times SU(2) \times U(1)$ group, up to 6 matter fields
- 90~97% success rate
- Multitask by nature → THEP foundation model

Embedding Analysis

(latent/activation space)



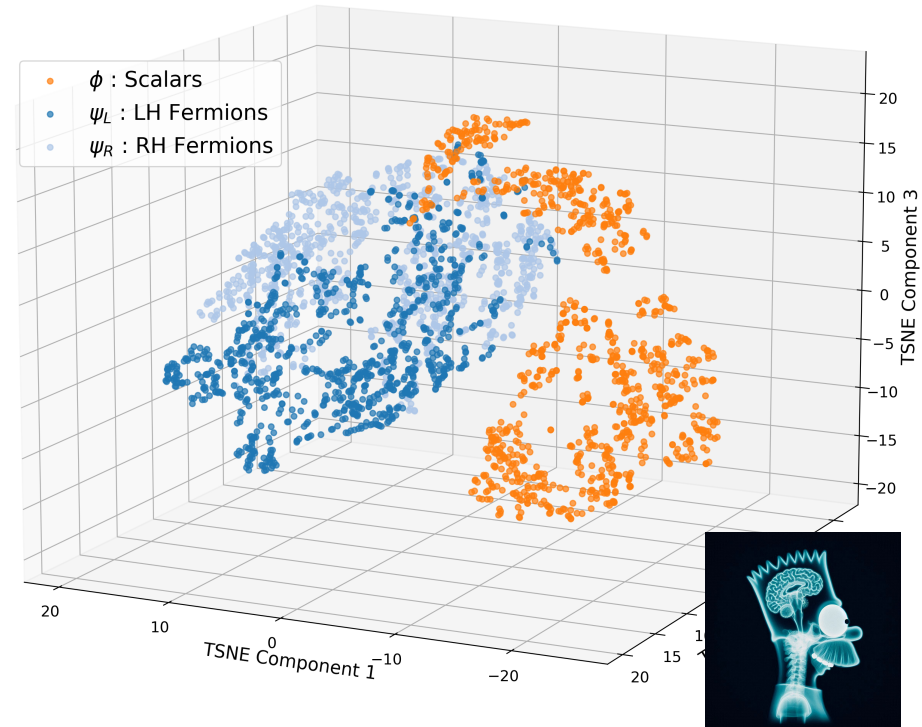
Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

Latent space $\sim$ Learned Theory Space
Each point = one model

Embedding Analysis

(latent/activation space)



Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

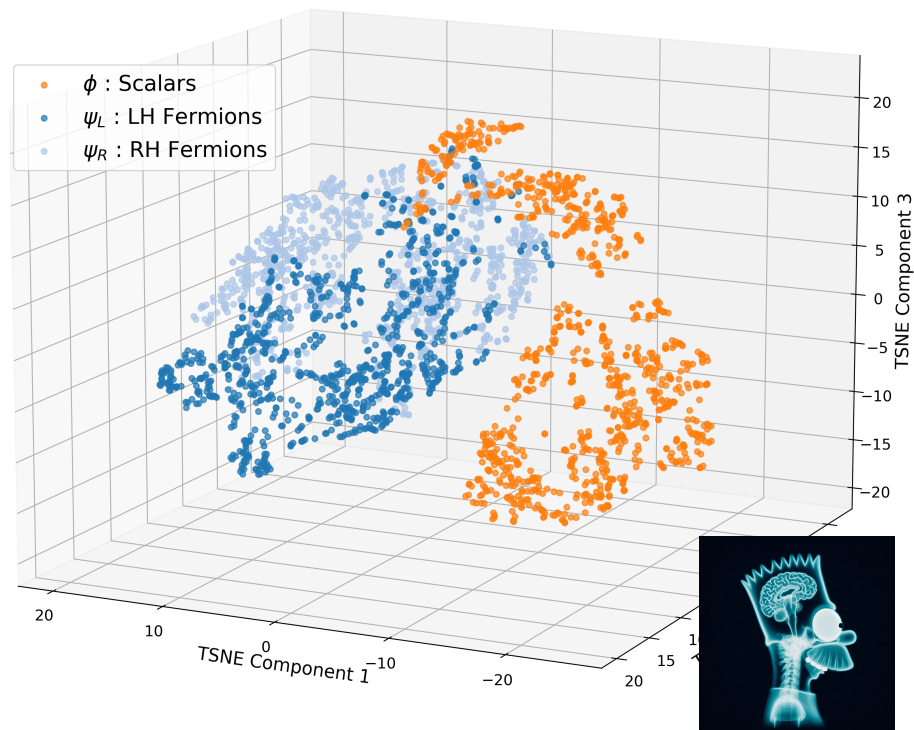
Latent space $\sim$ Learned Theory Space
Each point = one model

Interpretability

Is this a good learned theory space?

Embedding Analysis

(latent/activation space)



Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

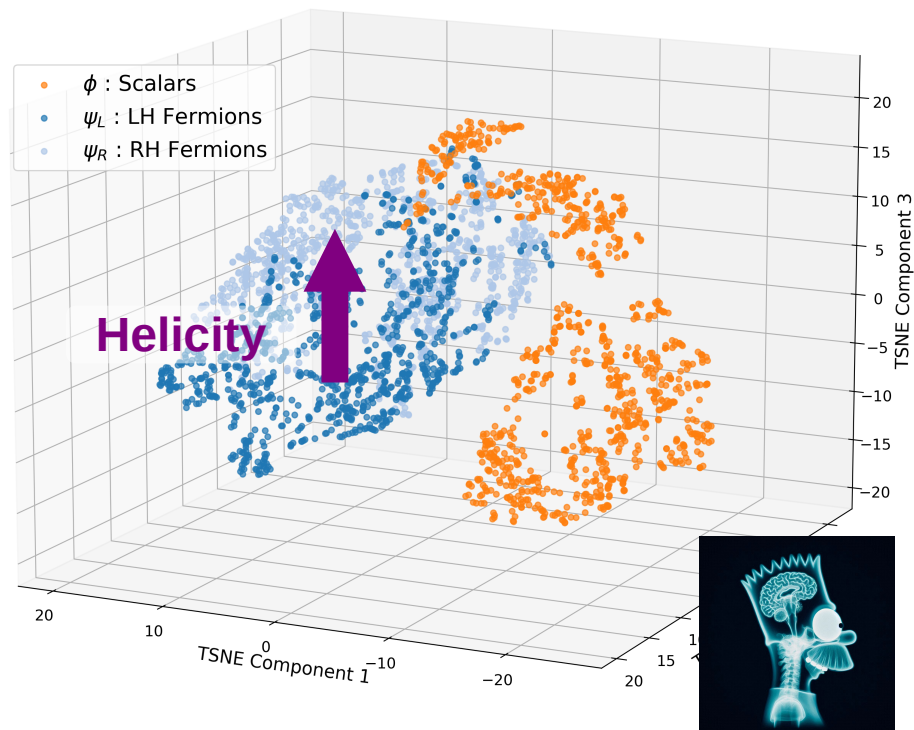
Latent space $\sim$ Learned Theory Space
Each point = one model

Interpretability

Is this a good learned theory space?
Does it has physics relation in it?

Embedding Analysis

(latent/activation space)



Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

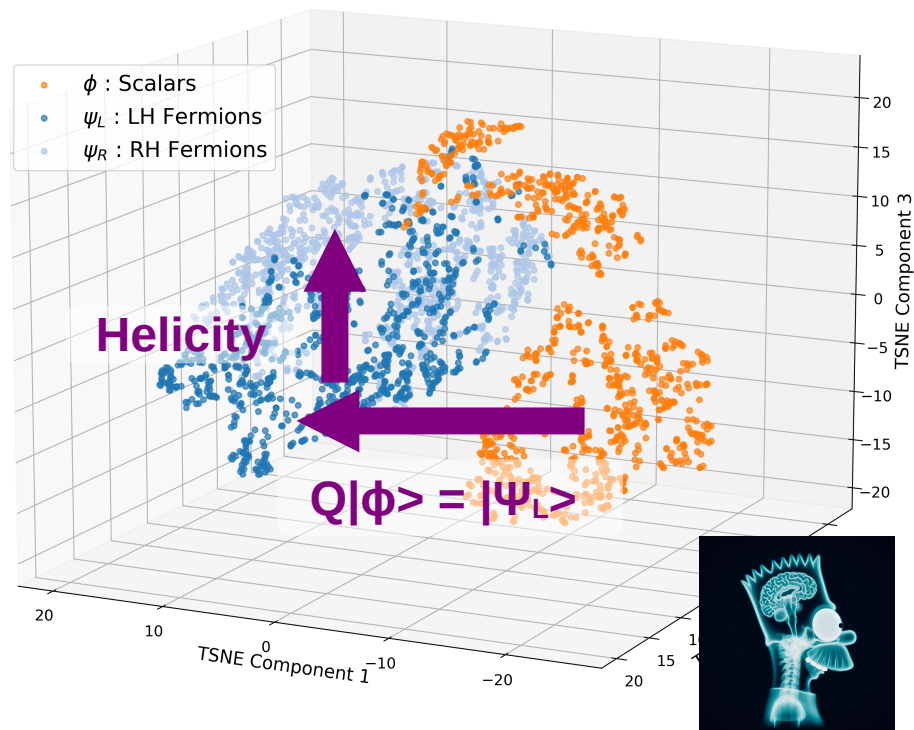
Latent space $\sim$ Learned Theory Space
Each point = one model

Interpretability

Is this a good learned theory space?
Does it has physics relation in it?

Embedding Analysis

(latent/activation space)



Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

Latent space $\sim$ Learned Theory Space
Each point = one model

Interpretability

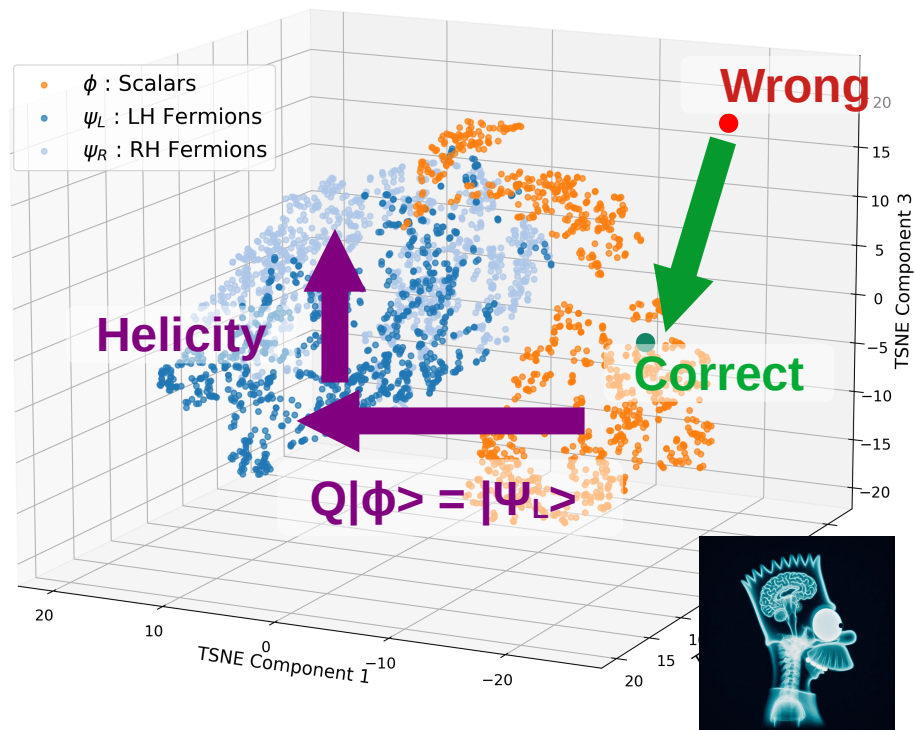
Is this a good learned theory space?
Does it has physics relation in it?

Navigation

Can we move around this space?
From wrong theories to correct one?

Embedding Analysis

(latent/activation space)



Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

Latent space $\sim$ Learned Theory Space
Each point = one model

Interpretability

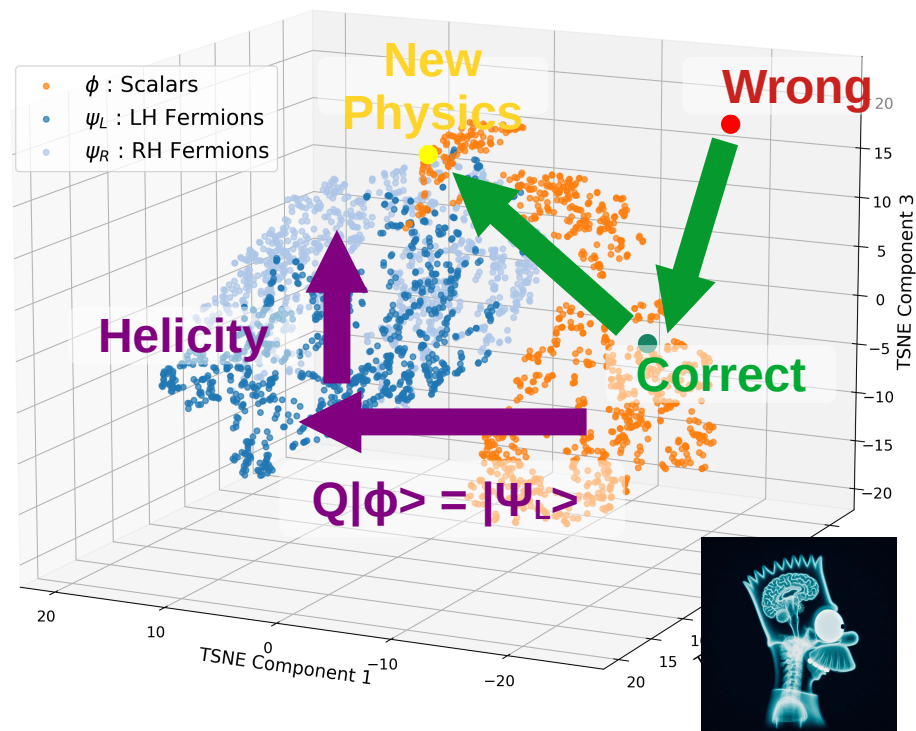
Is this a good learned theory space?
Does it has physics relation in it?

Navigation

Can we move around this space?
From wrong theories to correct one?
From correct theories to better ones?

Embedding Analysis

(latent/activation space)



Lagrangian Transformer, BART-L

YSK, Rikard Enberg, Stefano Moretti, Eliel Camargo-Molina. (2026), SciPost Phys. 21, 024. arXiv:2501.09729

Latent space ~ Learned Theory Space
Each point = one model

Embedding Analysis

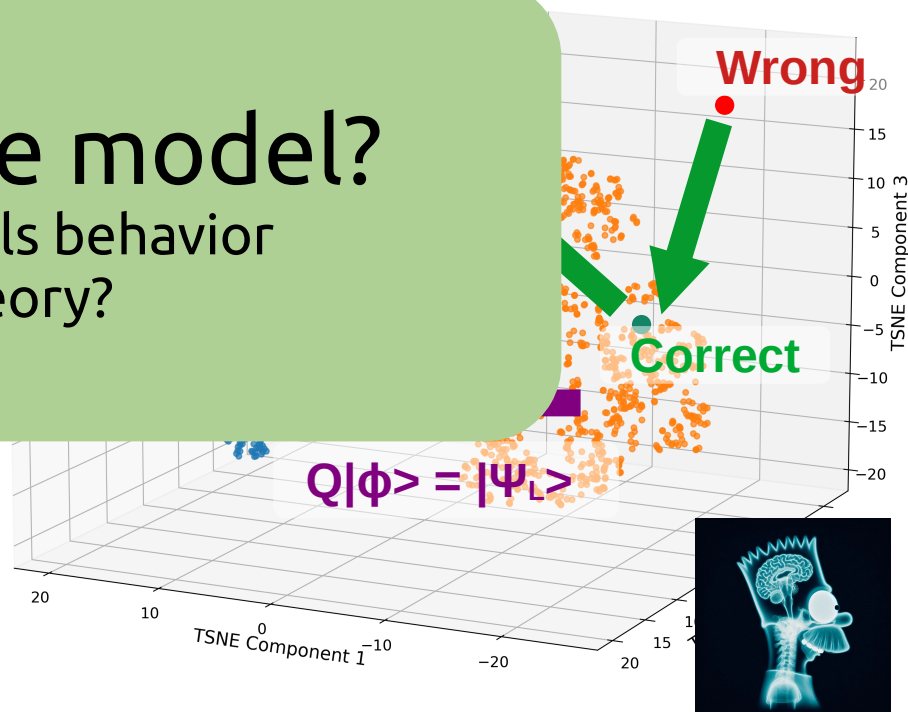
(latent/activation space)

Is this
Doe

Can we 'steer' the model?

Can we change the models behavior
as it writes the theory?

Can we move around this space?
From wrong theories to correct one?
From correct theories to better ones?



Contrastive Activation Addition

Rimsky et al 2312.06681

Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector

Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector

Input:

ψ_1, ψ_2

Output:

$$\begin{aligned} & i\psi_1^\dagger \bar{\sigma}^\mu D_\mu \psi_1 \\ & + i\psi_2^\dagger \bar{\sigma}^\mu D_\mu \psi_2 + \\ & \frac{1}{2}(m_1\psi_1\psi_1 + \\ & m_2\psi_2\psi_2 + \text{h.c.}) \end{aligned}$$



Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector

Input:

ψ_1, ψ_2

Output:

$$i\psi_1^\dagger \bar{\sigma}^\mu D_\mu \psi_1 + i\psi_2^\dagger \bar{\sigma}^\mu D_\mu \psi_2 + \frac{1}{2}(m_1\psi_1\psi_1 + m_2\psi_2\psi_2 + \text{h.c.})$$



Input:

ϕ_1, ϕ_2

Output:

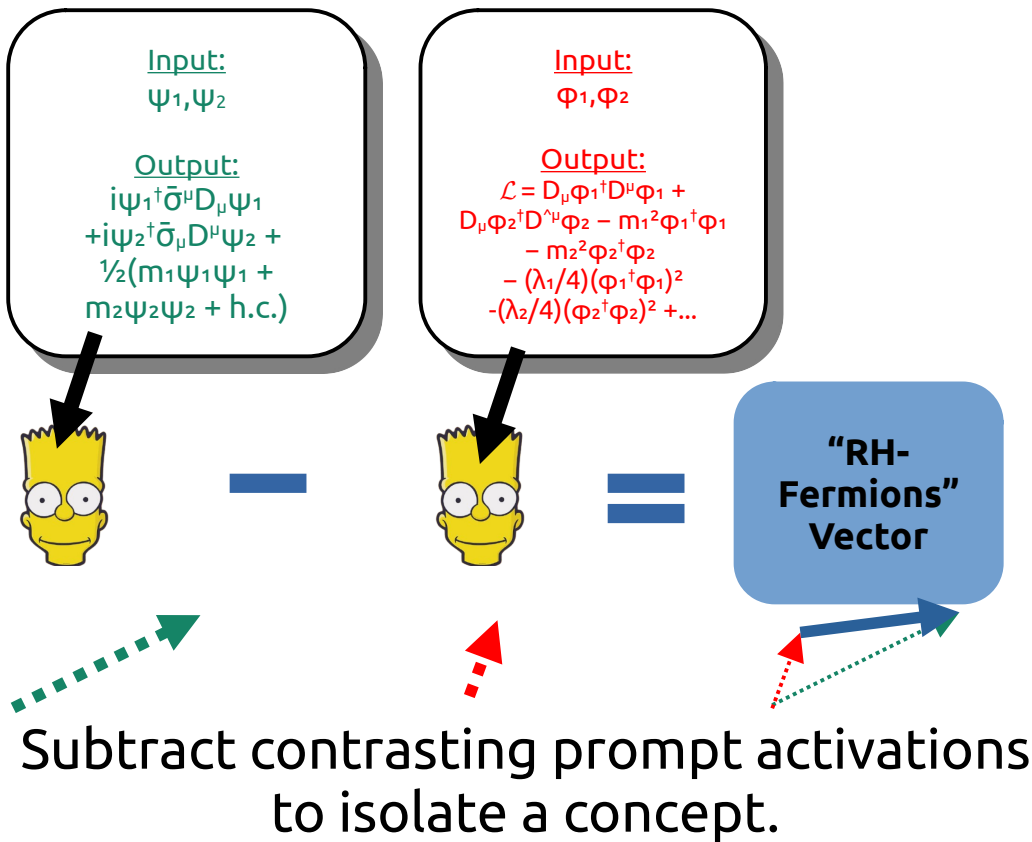
$$\begin{aligned} \mathcal{L} = & D_\mu \phi_1^\dagger D^\mu \phi_1 + D_\mu \phi_2^\dagger D^\mu \phi_2 - m_1^2 \phi_1^\dagger \phi_1 \\ & - m_2^2 \phi_2^\dagger \phi_2 - (\lambda_1/4)(\phi_1^\dagger \phi_1)^2 \\ & - (\lambda_2/4)(\phi_2^\dagger \phi_2)^2 + \dots \end{aligned}$$



Contrastive Activation Addition

Rimsky et al 2312.06681

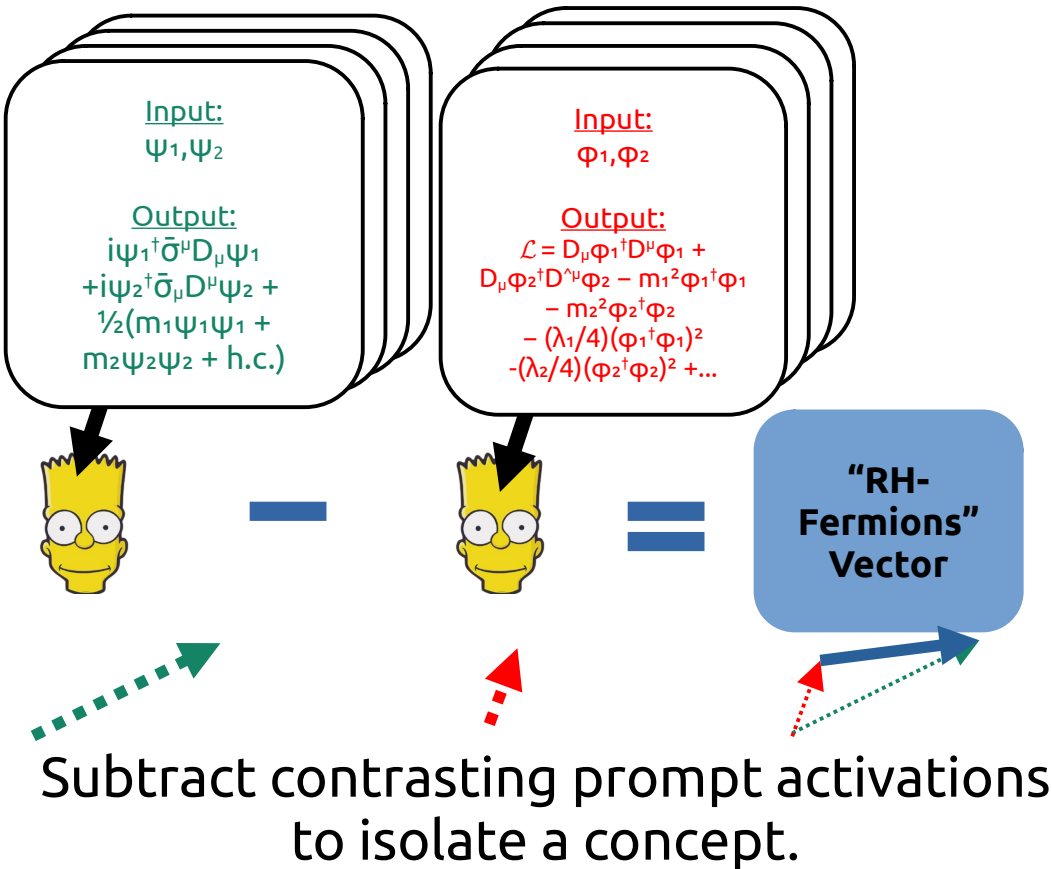
1. Getting steering vector



Contrastive Activation Addition

Rimsky et al 2312.06681

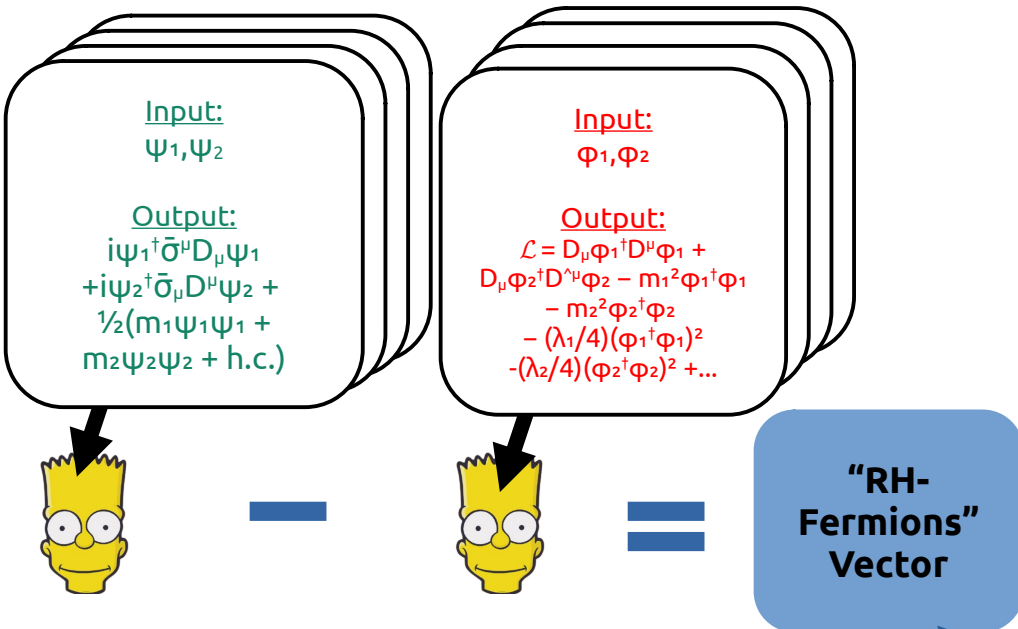
1. Getting steering vector



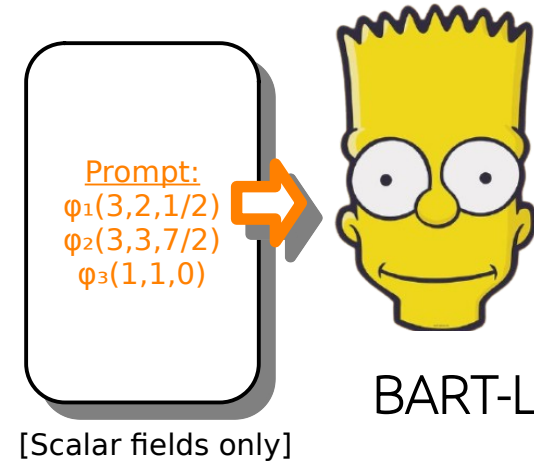
Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector



2. Steer!

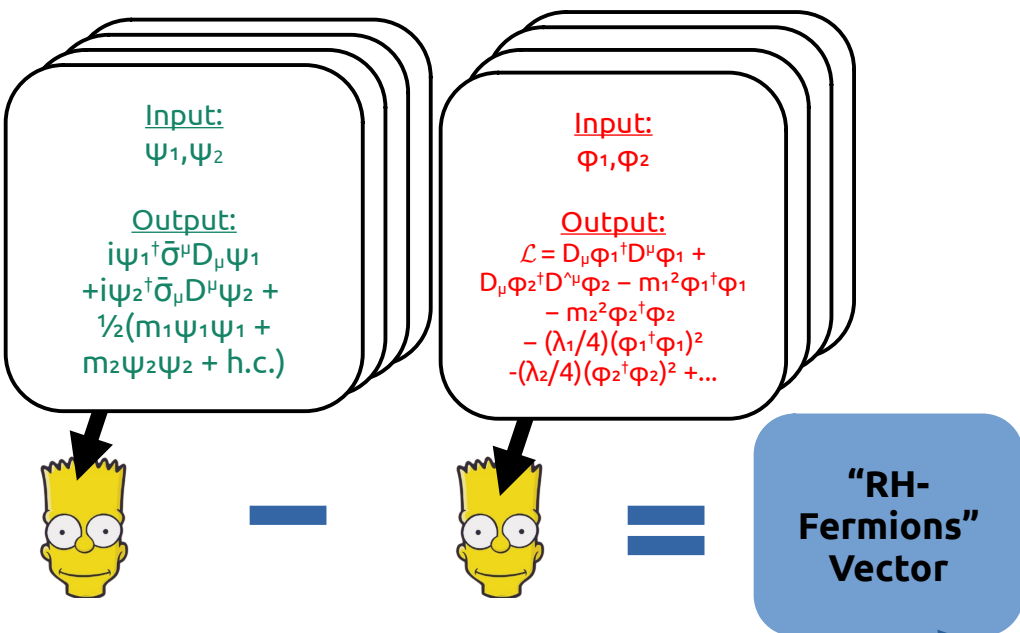


Subtract contrasting prompt activations to isolate a concept.

Contrastive Activation Addition

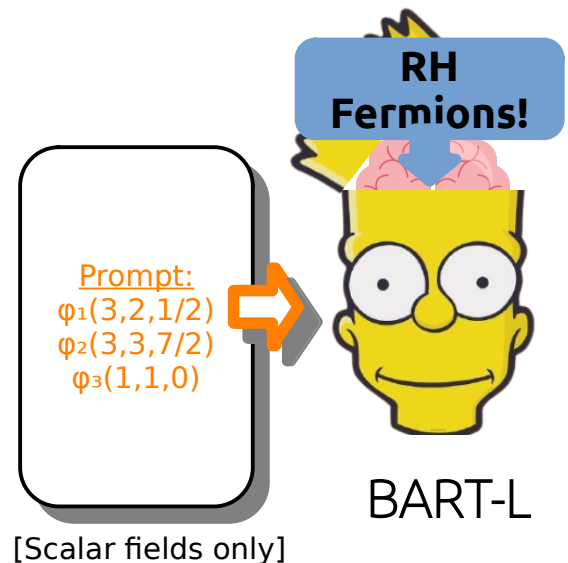
Rimsky et al 2312.06681

1. Getting steering vector



Subtract contrasting prompt activations to isolate a concept.

2. Steer!

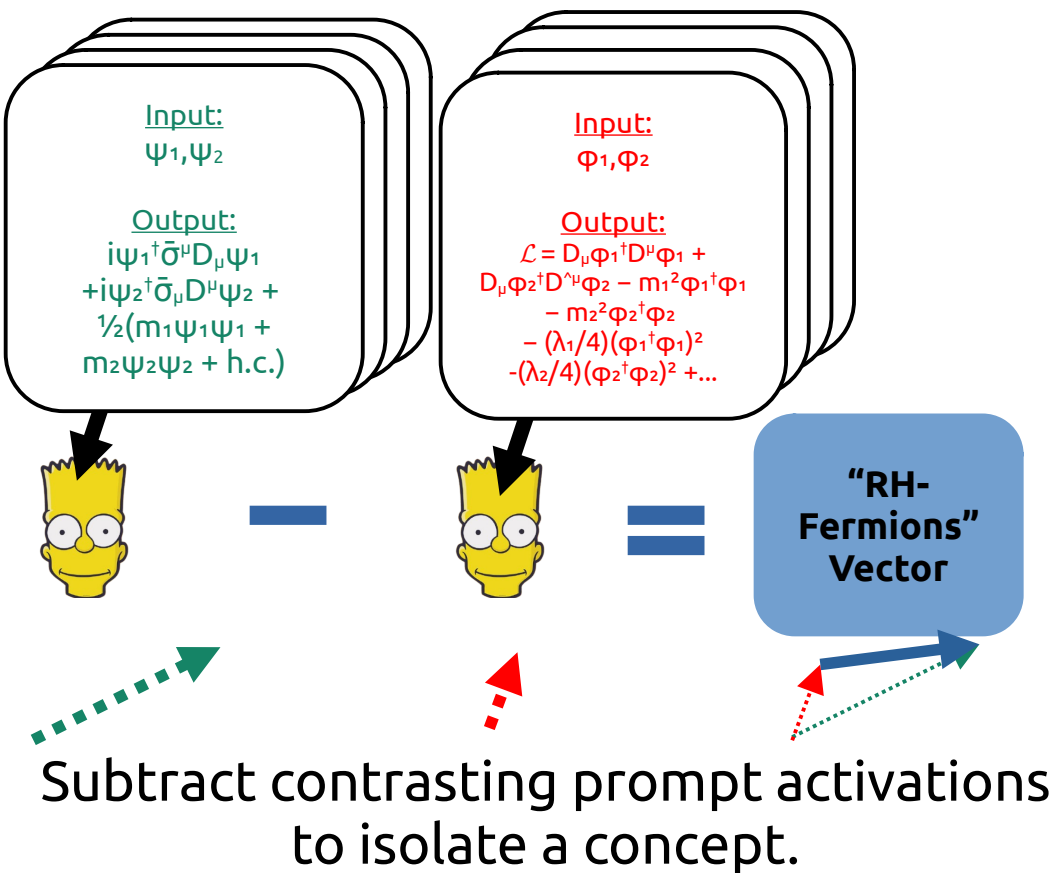


Inject the steering vector into model activations during generation.

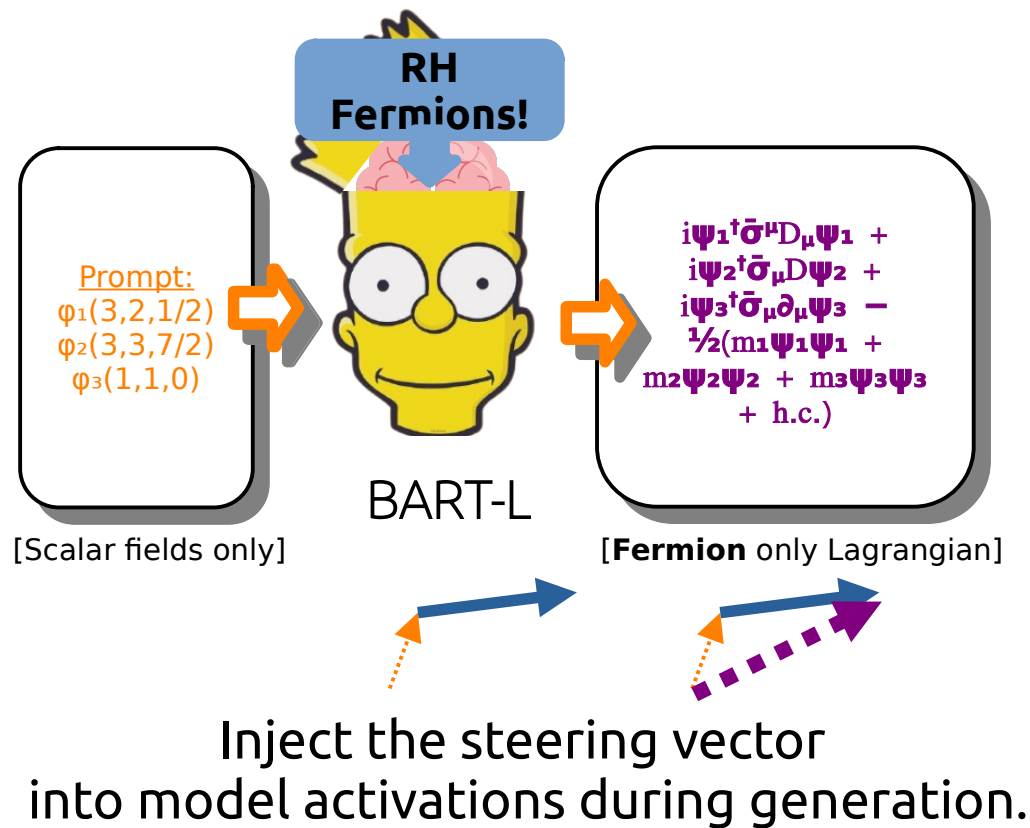
Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector



2. Steer!

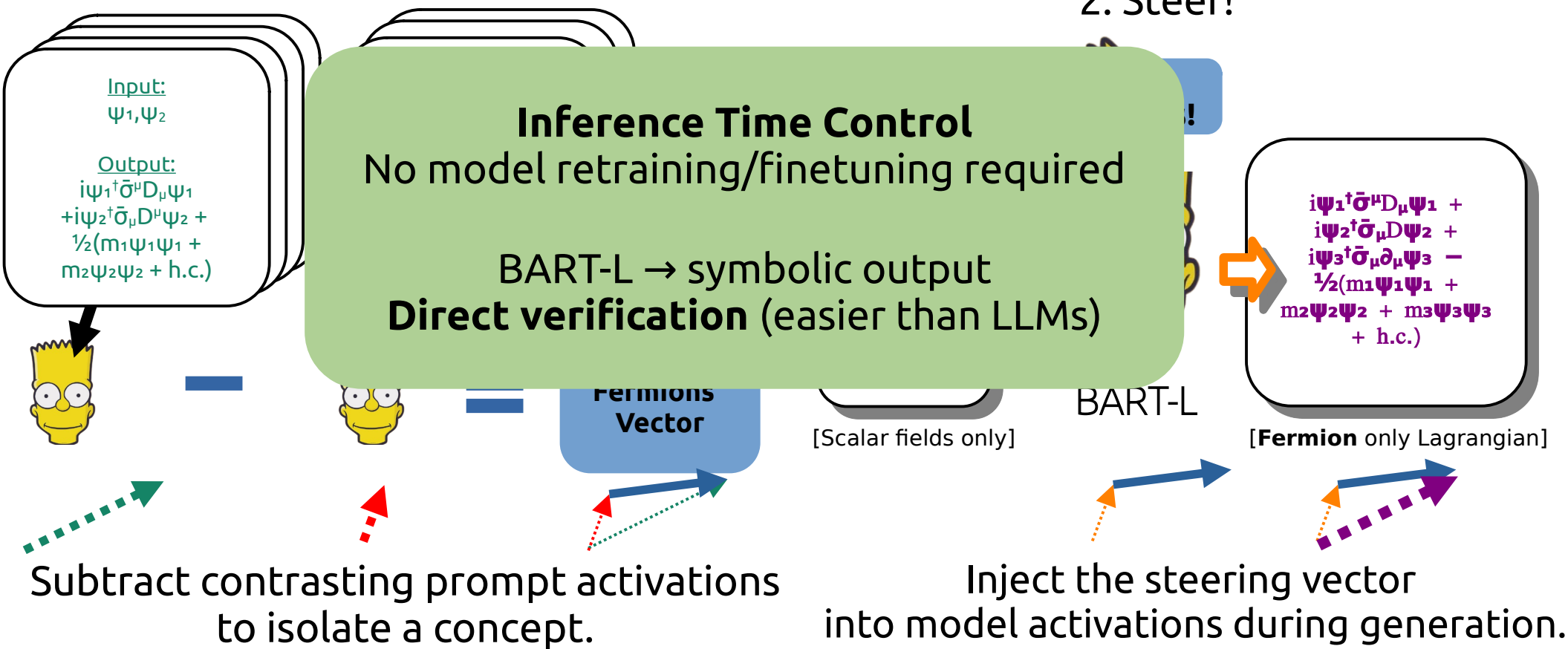


Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector

2. Steer!

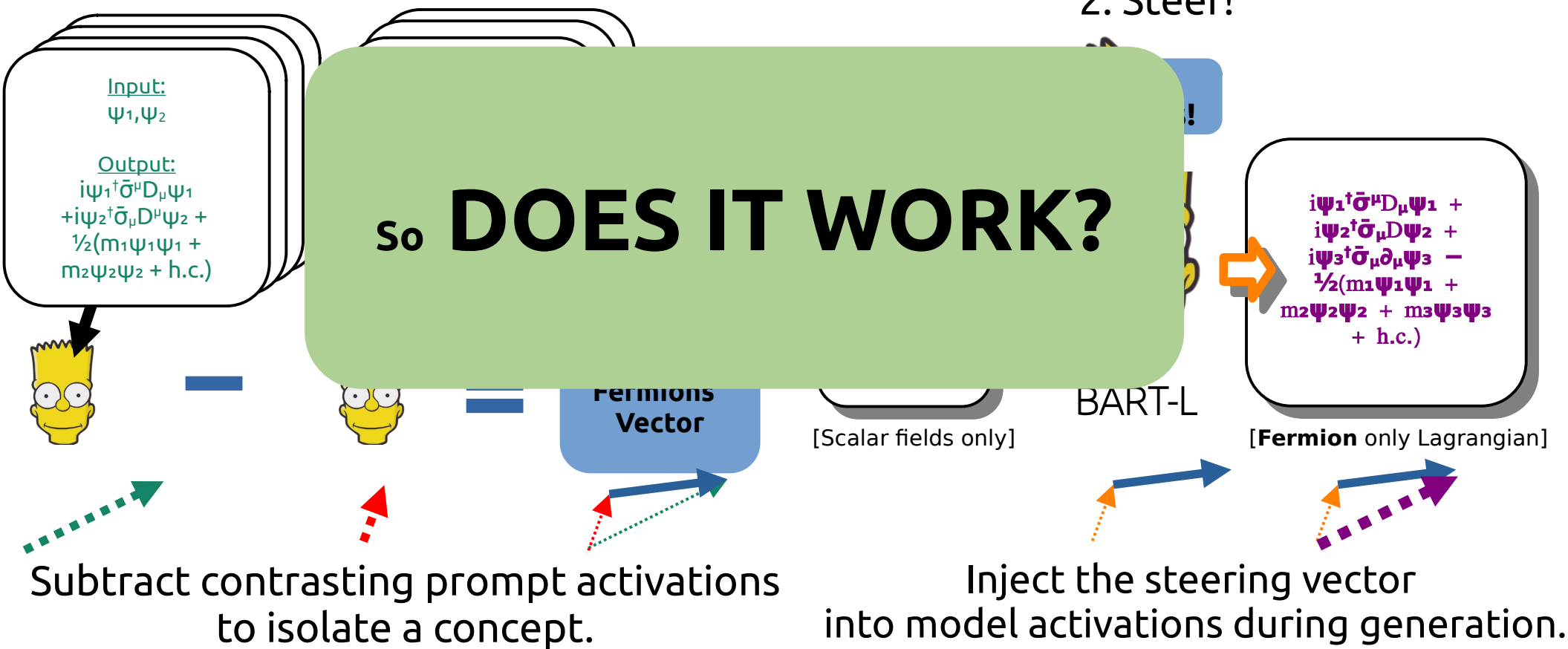


Contrastive Activation Addition

Rimsky et al 2312.06681

1. Getting steering vector

2. Steer!

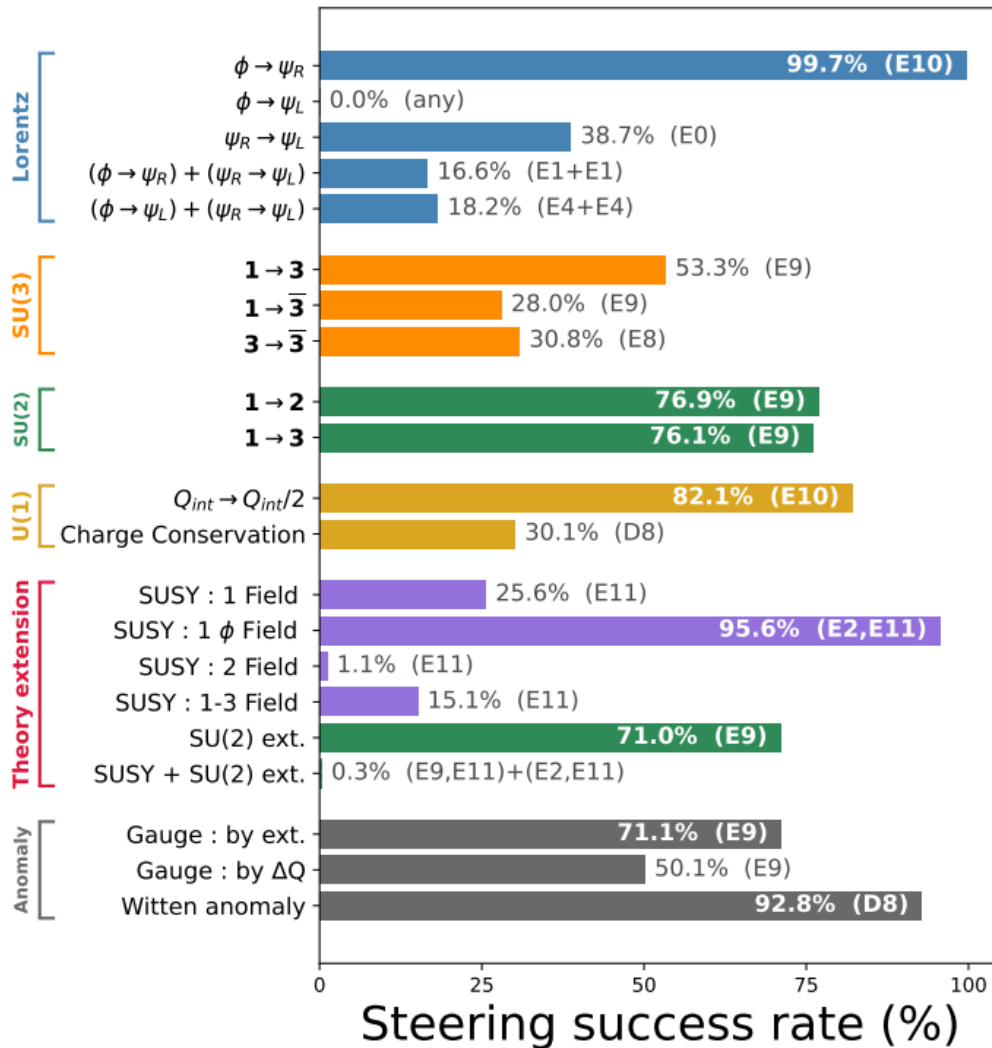


MOSTLY YES

But when it cant,
theres always a hint of why!

Interpretability

Navigation



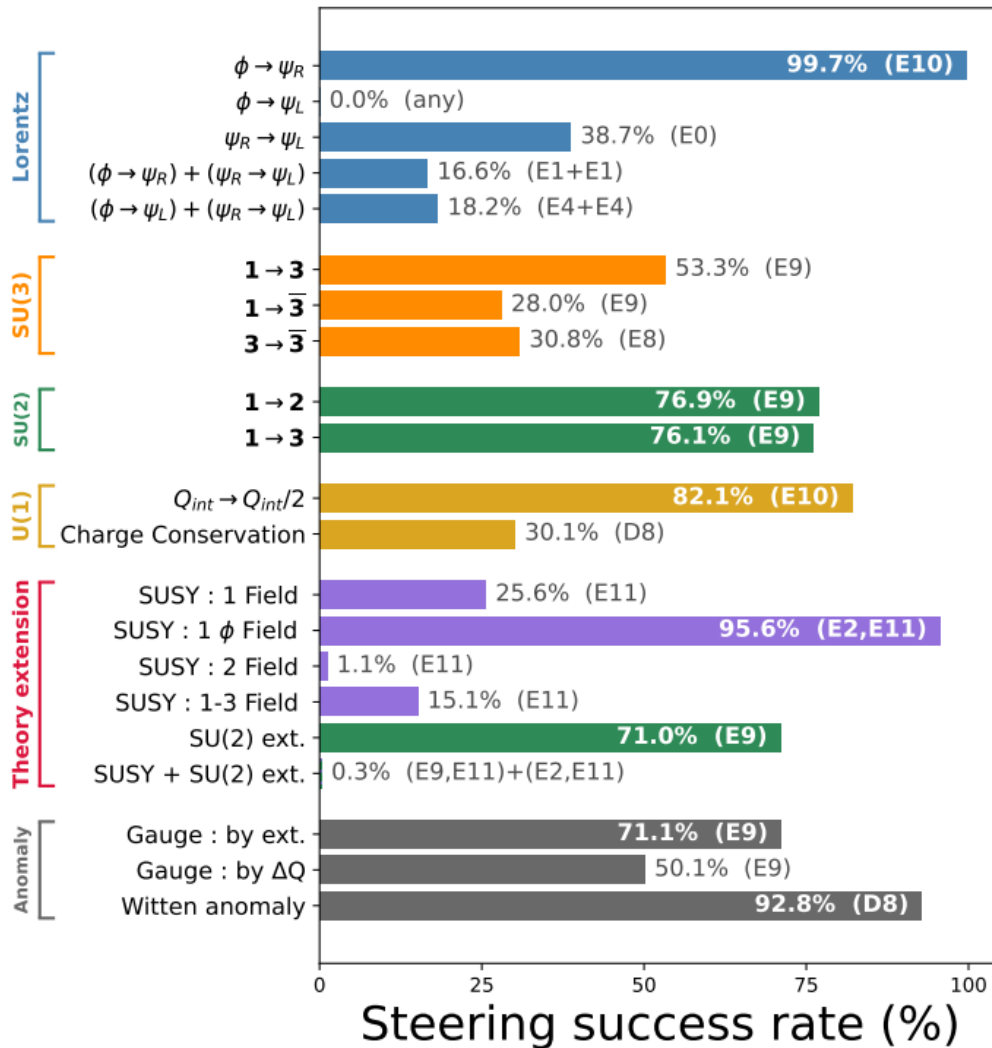
MOSTLY YES

But when it cant,
theres always a hint of why!

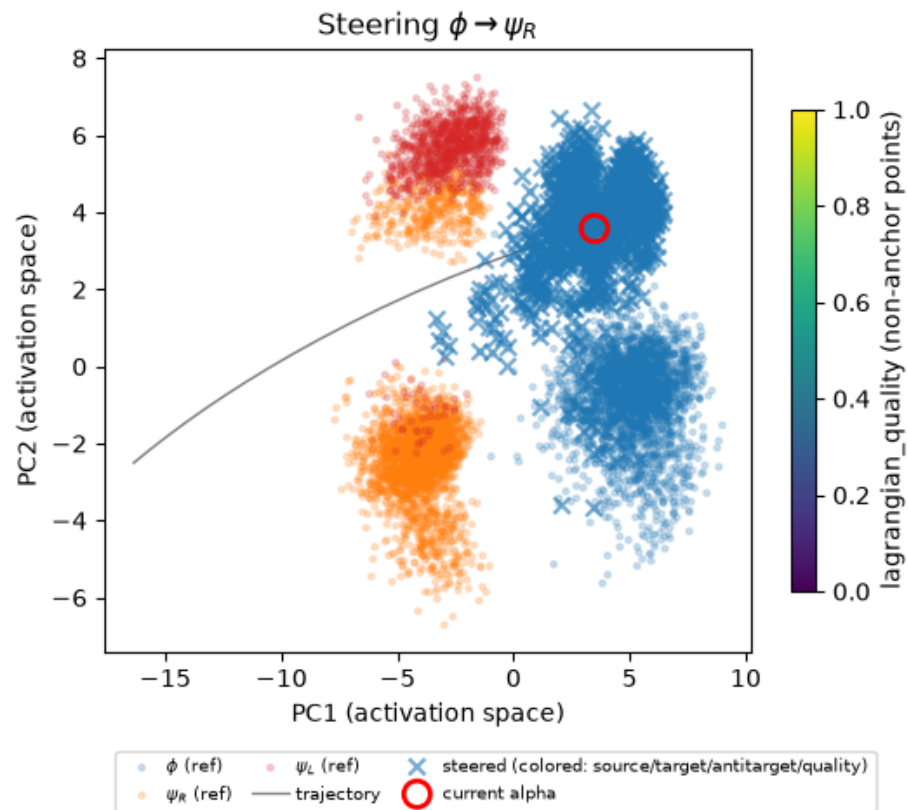
Success rate :
Exact Target Match
(Cumulative)
(except for anomaly section)

Interpretability

Navigation



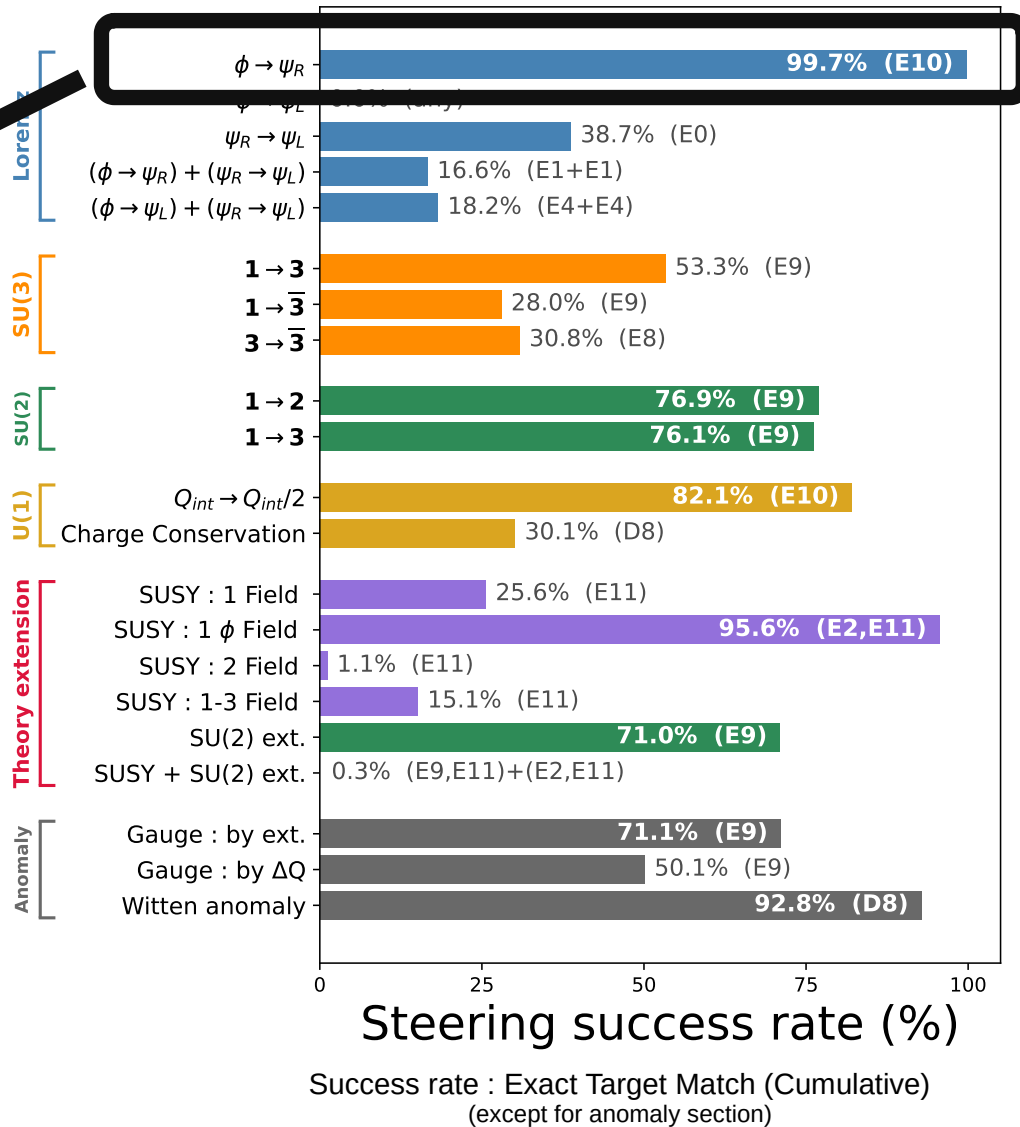
injected @ enc L10 -> viewed @ enc L11 (final layer)



alpha = 0.00

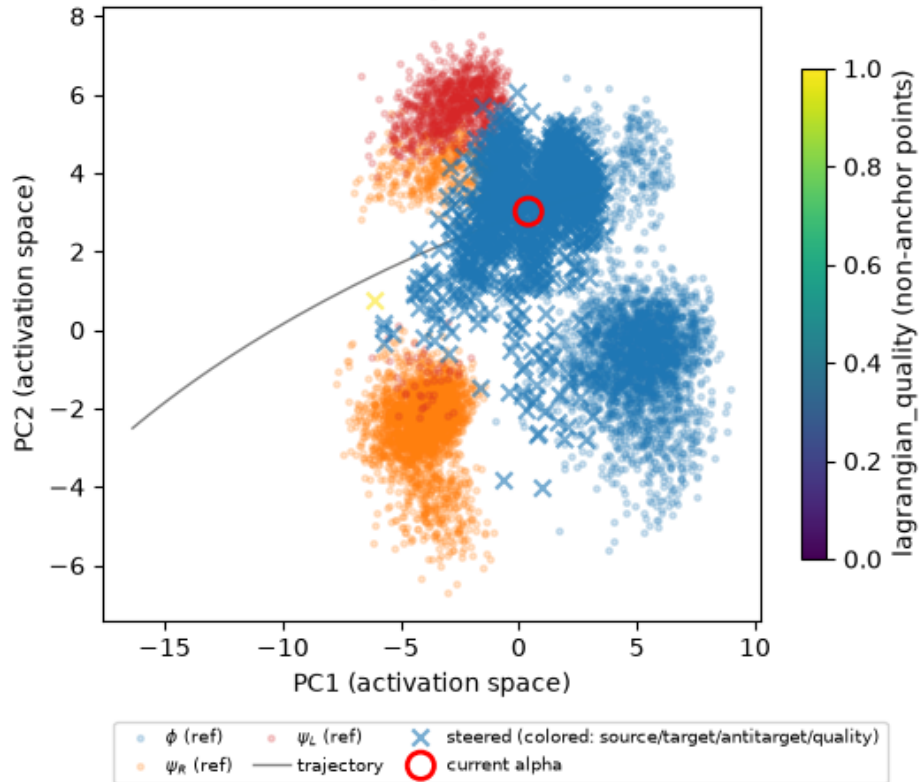
success = 0%

cumulative = 0%



injected @ enc L10 -> viewed @ enc L11 (final layer)

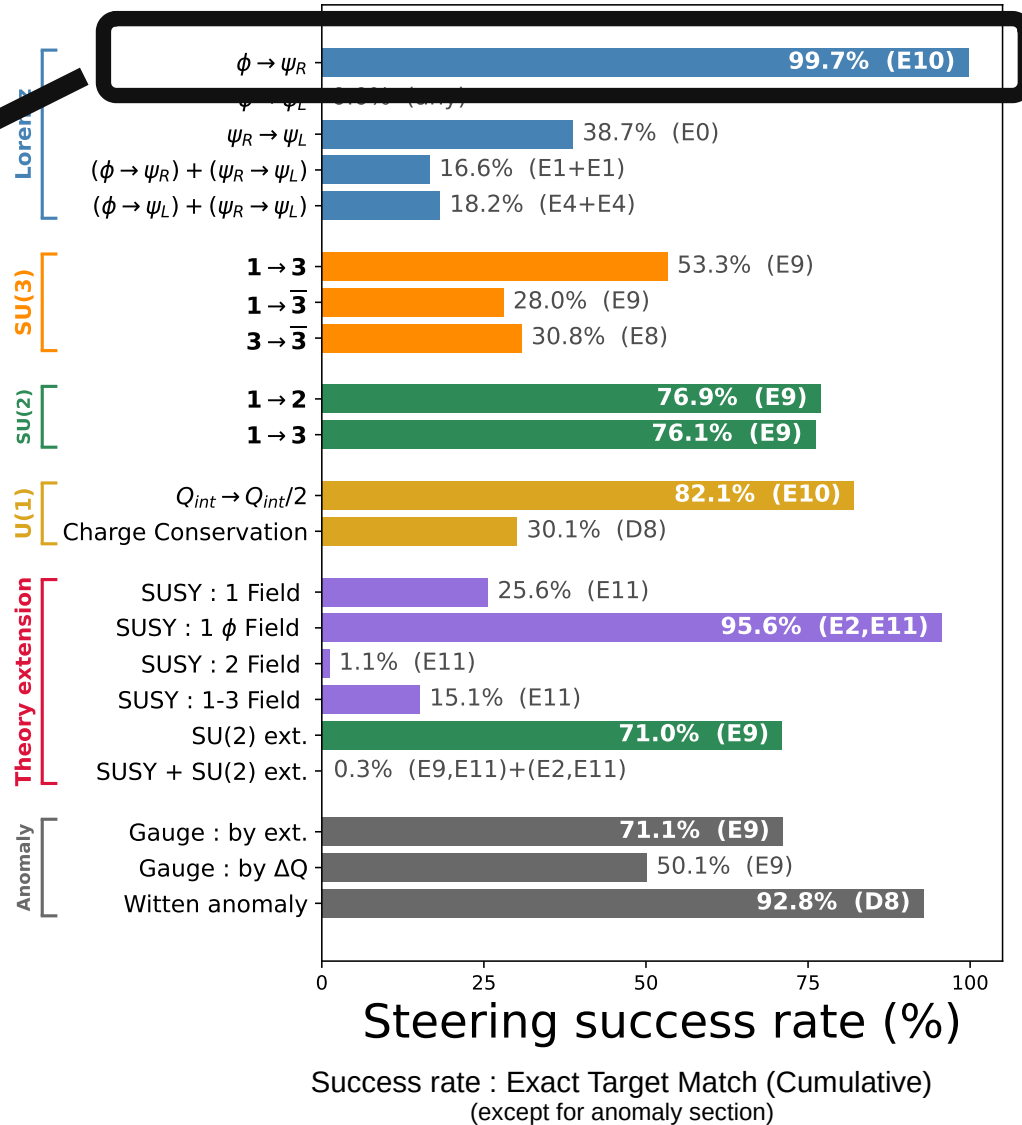
Steering $\phi \rightarrow \psi_R$



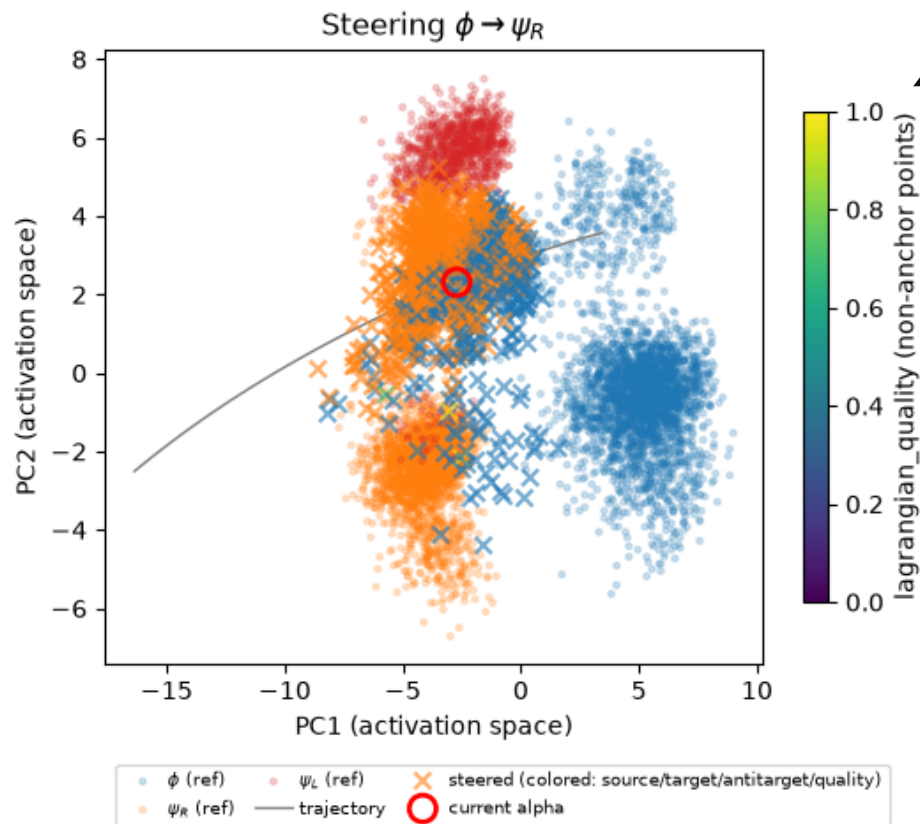
alpha = 0.20

success = 0%

cumulative = 0%



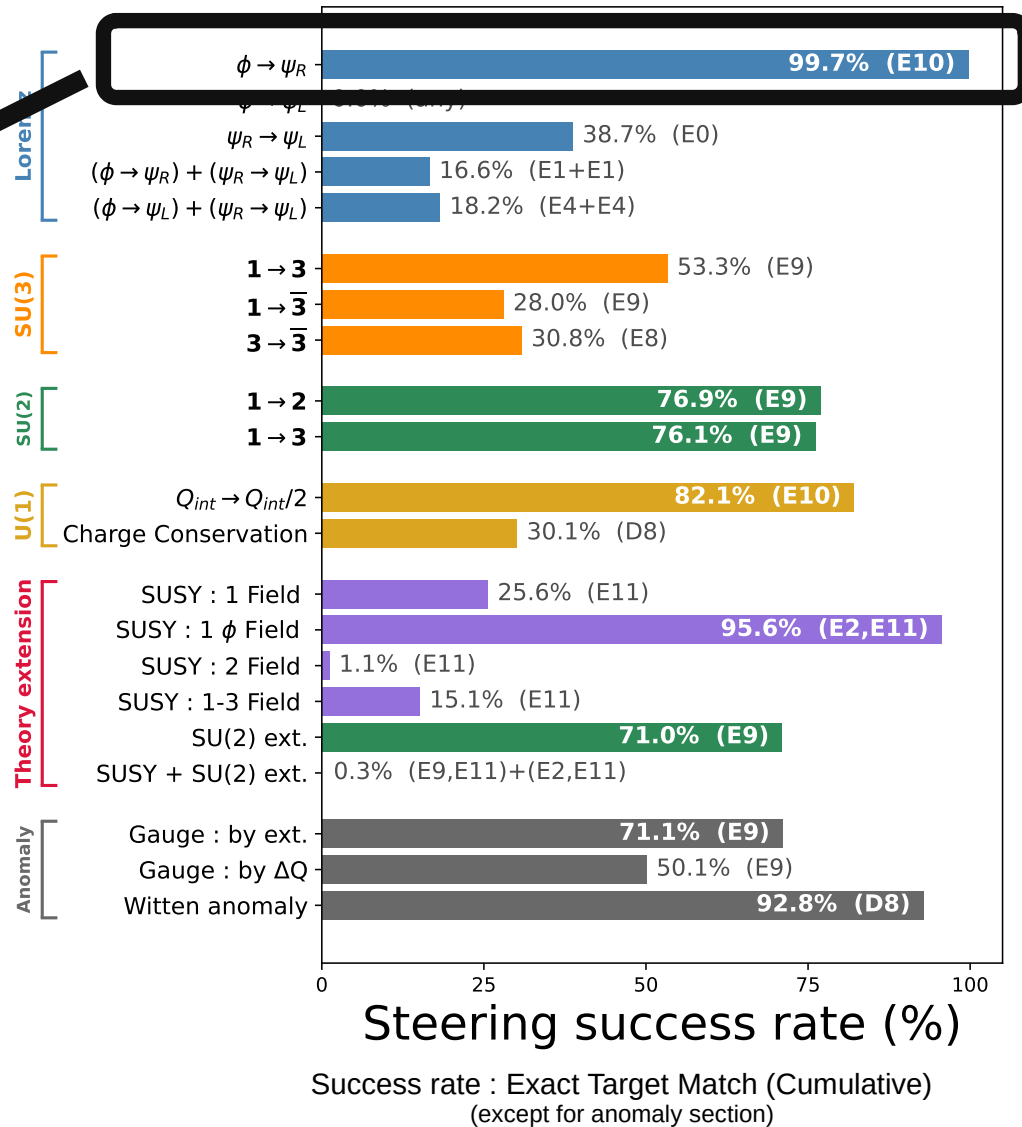
injected @ enc L10 -> viewed @ enc L11 (final layer)



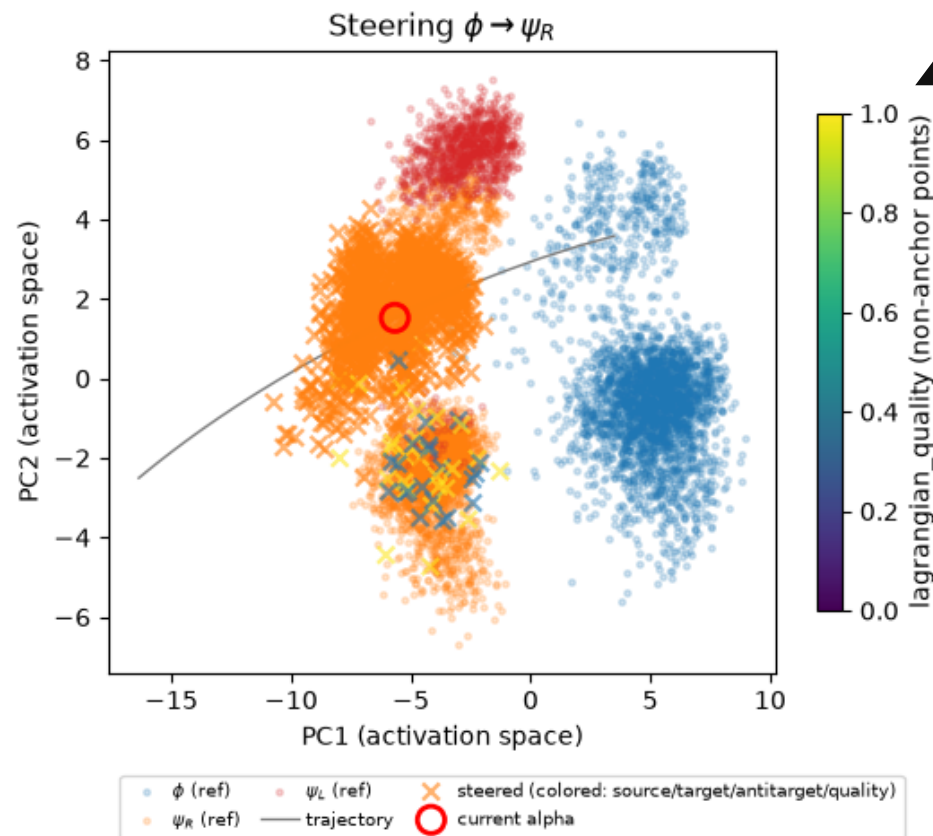
alpha = 0.40

success = 63%

cumulative = 63%



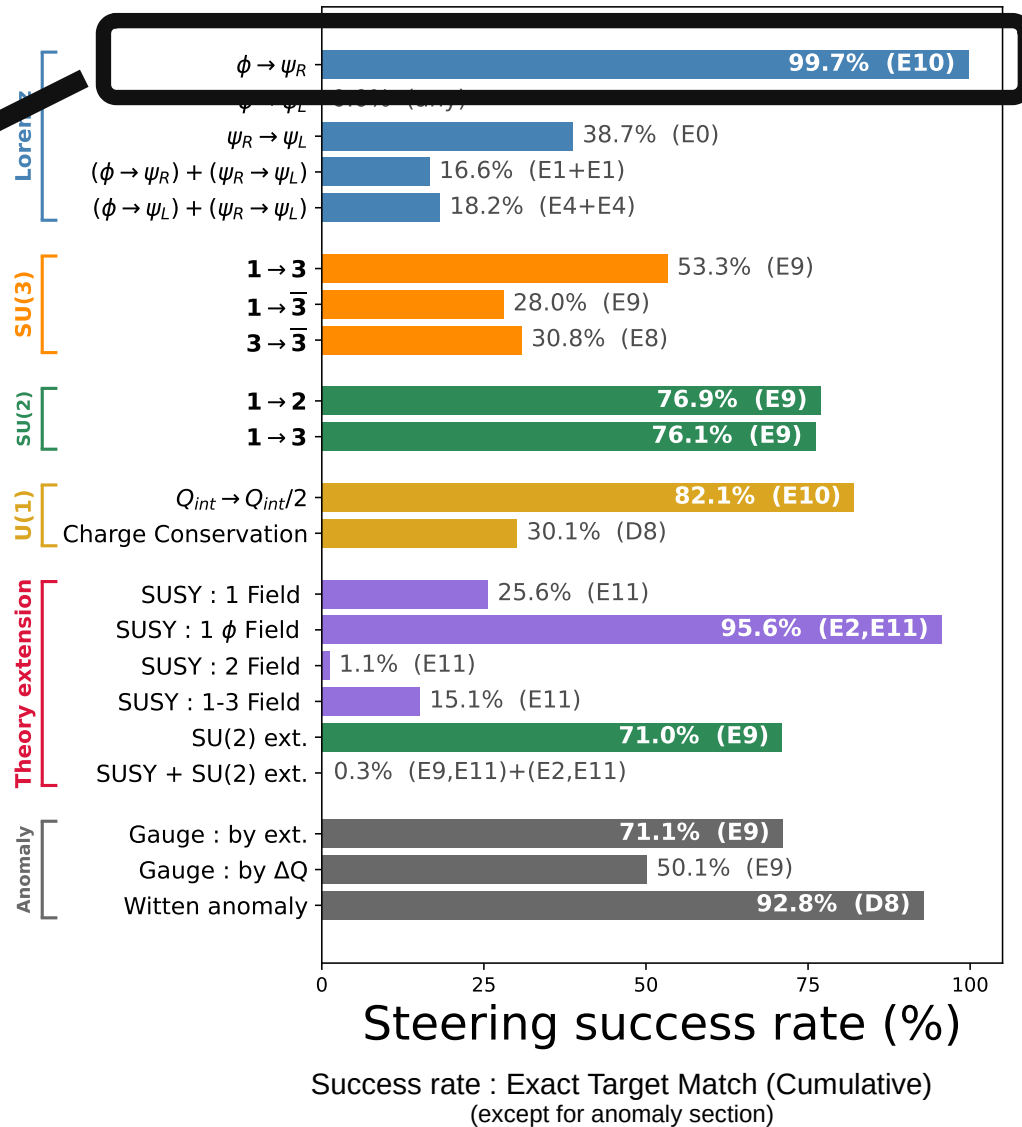
injected @ enc L10 -> viewed @ enc L11 (final layer)



alpha = 0.60

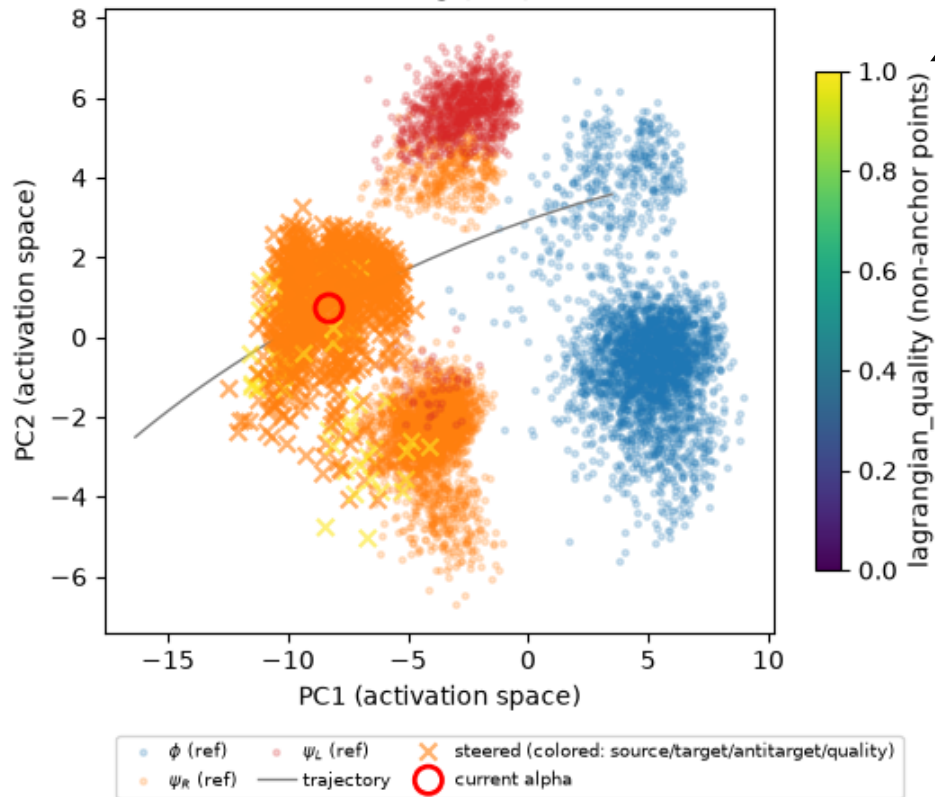
success = 95%

cumulative = 95%



injected @ enc L10 -> viewed @ enc L11 (final layer)

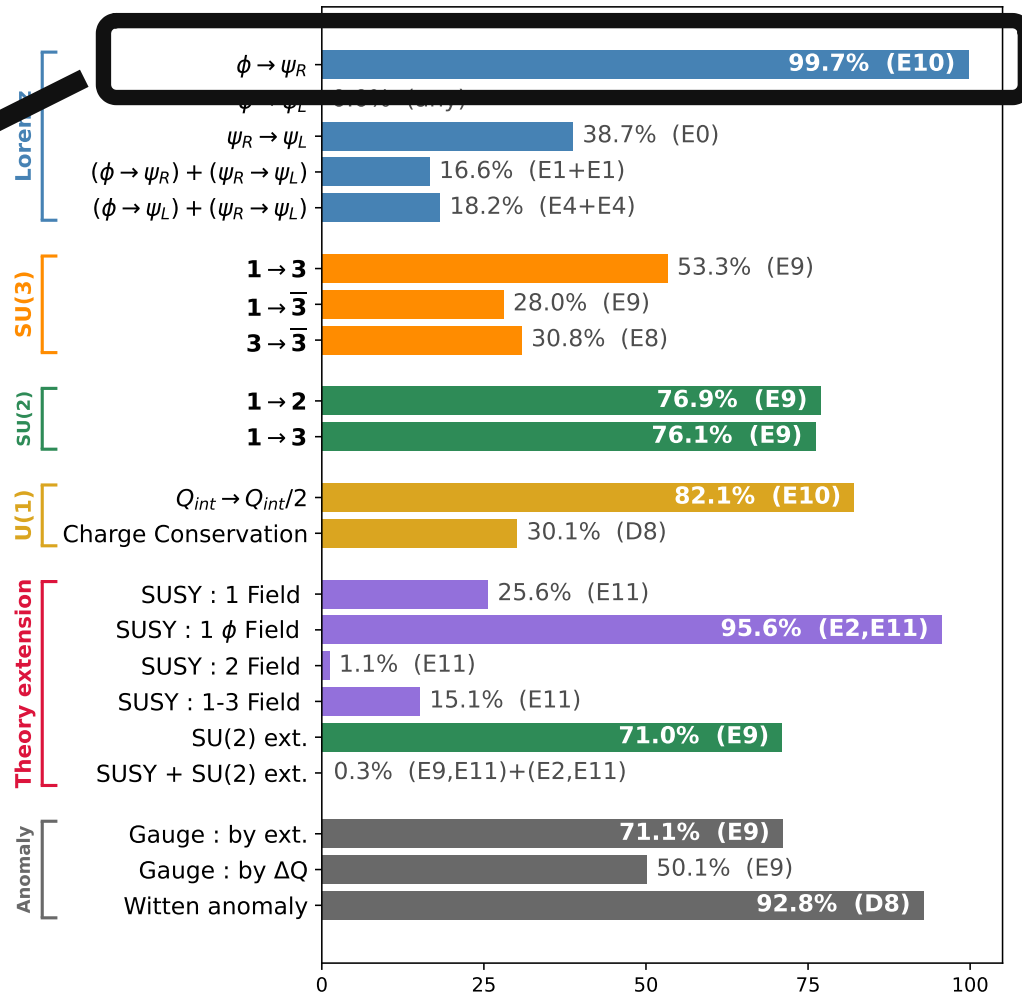
Steering $\phi \rightarrow \psi_R$



alpha = 0.80

success = 94%

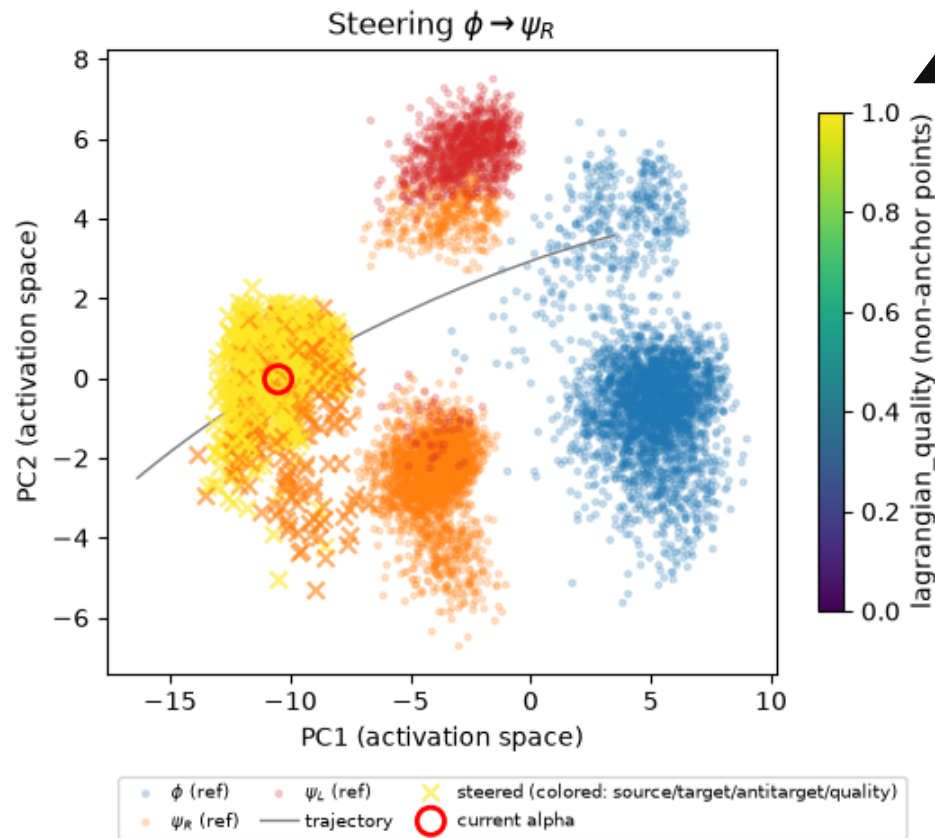
cumulative = 98%



Steering success rate (%)

Success rate : Exact Target Match (Cumulative)
(except for anomaly section)

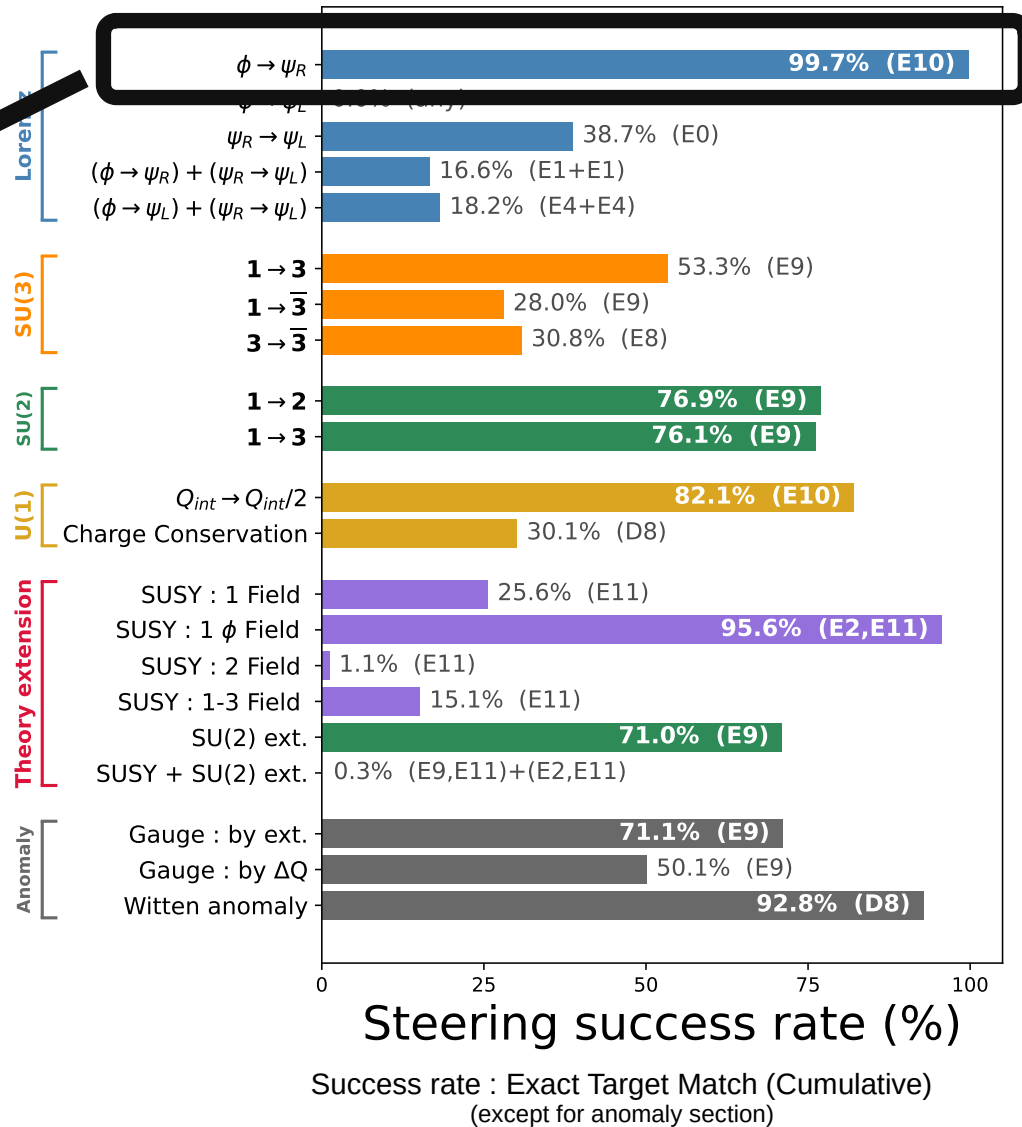
injected @ enc L10 -> viewed @ enc L11 (final layer)



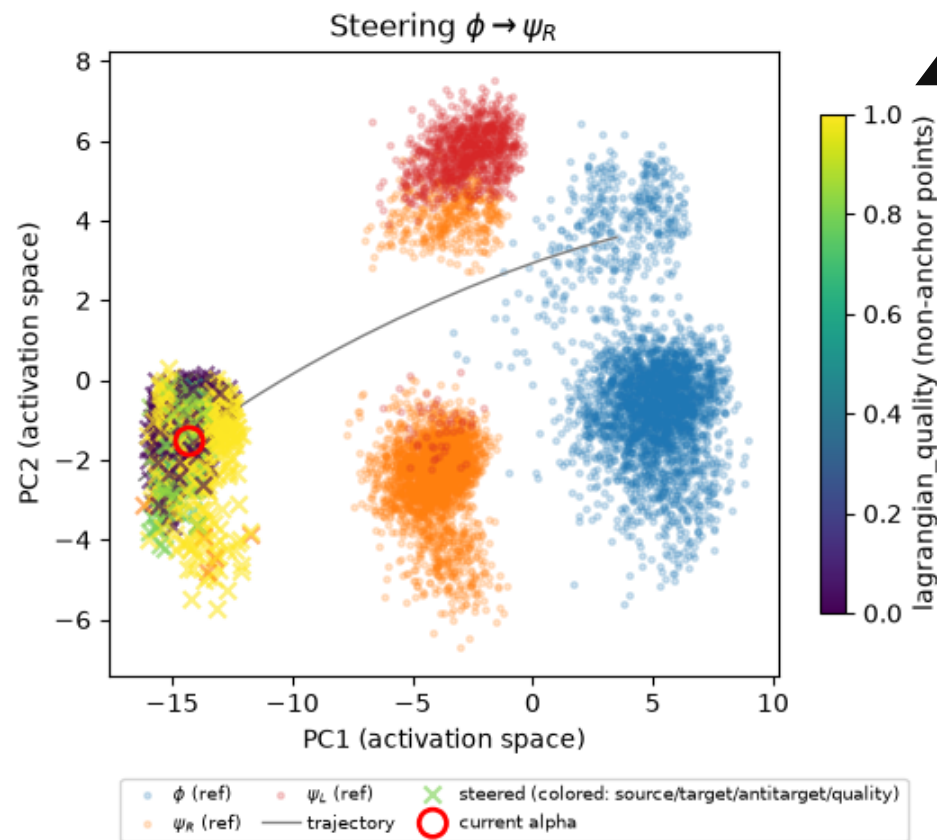
alpha = 1.00

success = 25%

cumulative = 100%



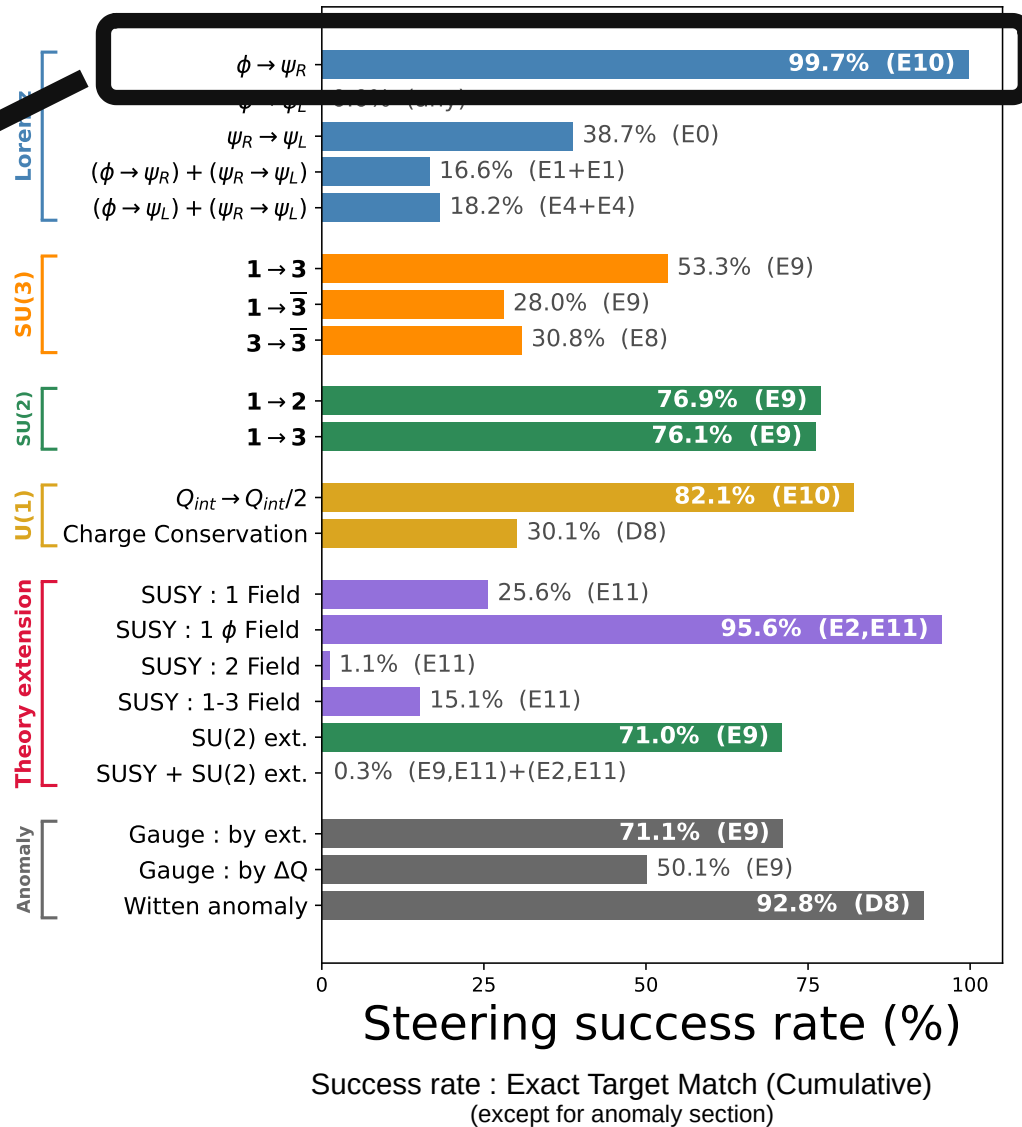
injected @ enc L10 -> viewed @ enc L11 (final layer)



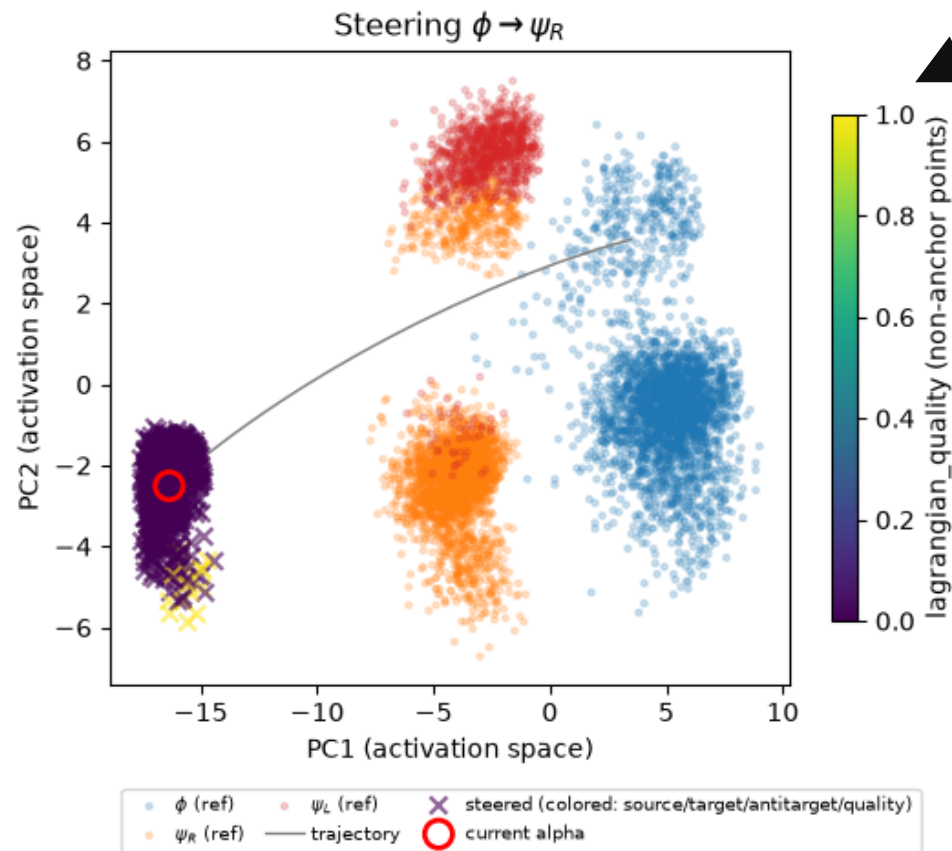
alpha = 1.50

success = 0%

cumulative = 100%



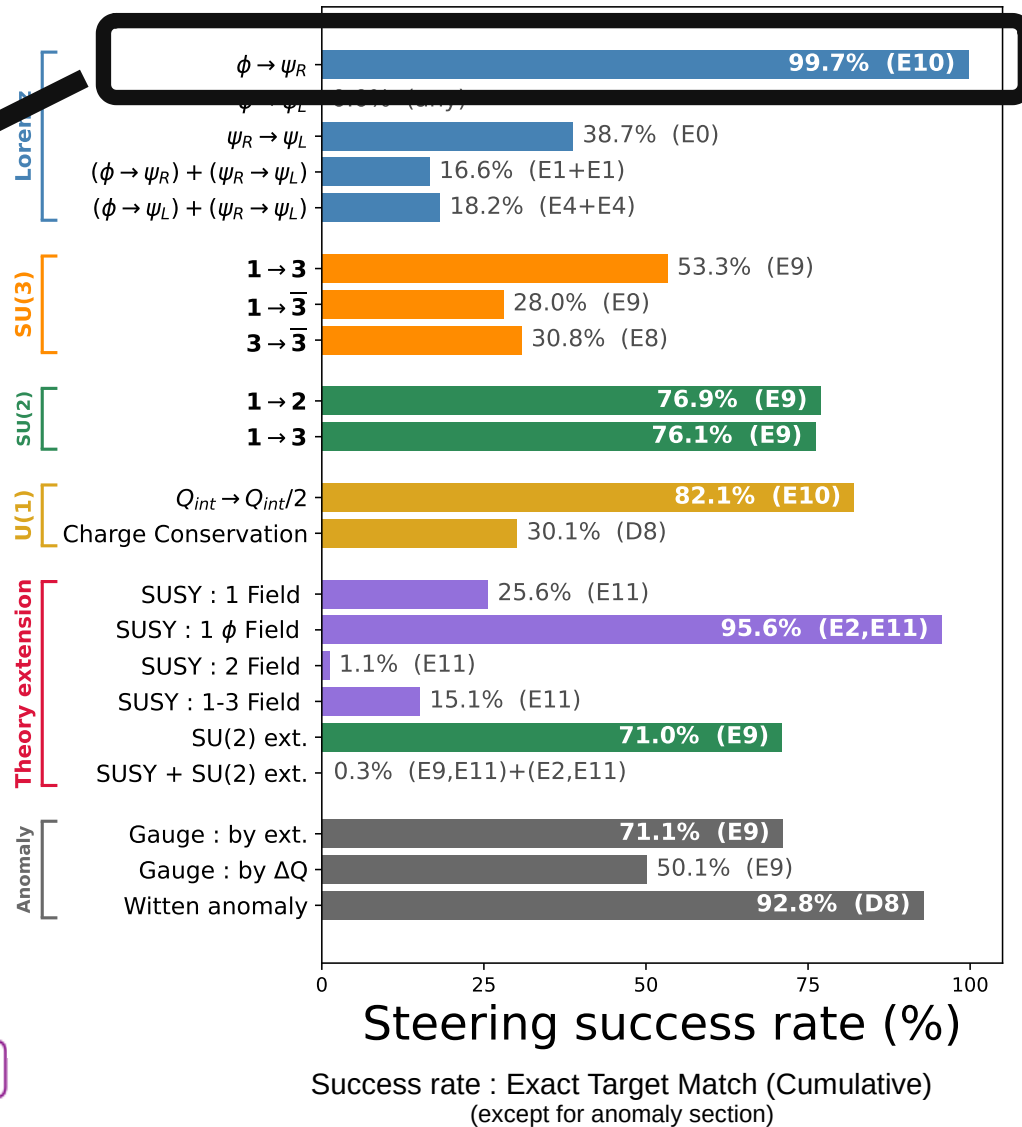
injected @ enc L10 -> viewed @ enc L11 (final layer)



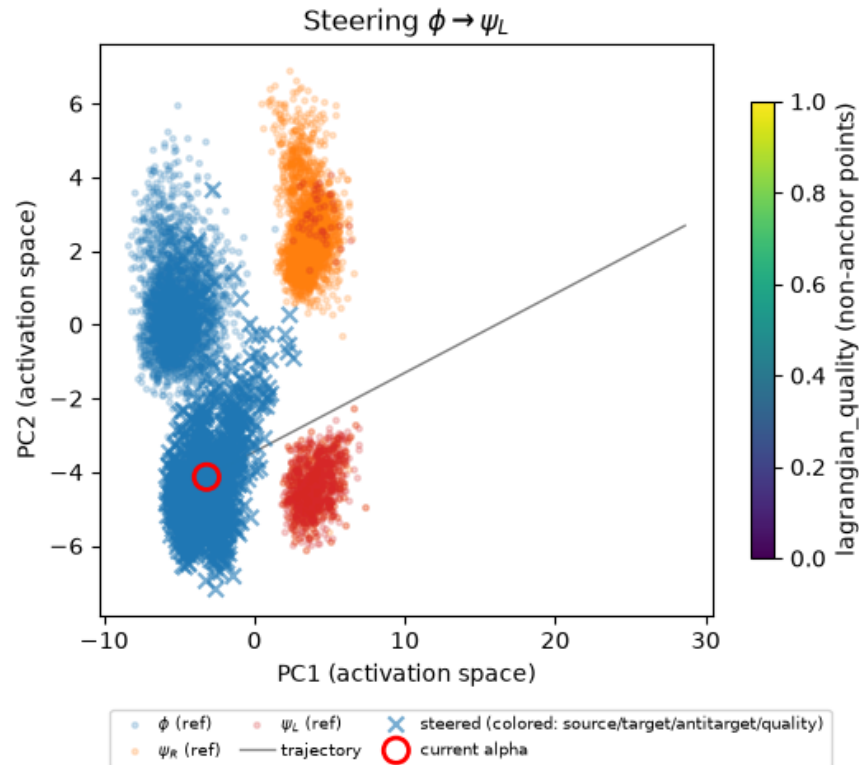
alpha = 2.00

success = 0%

cumulative = 100%



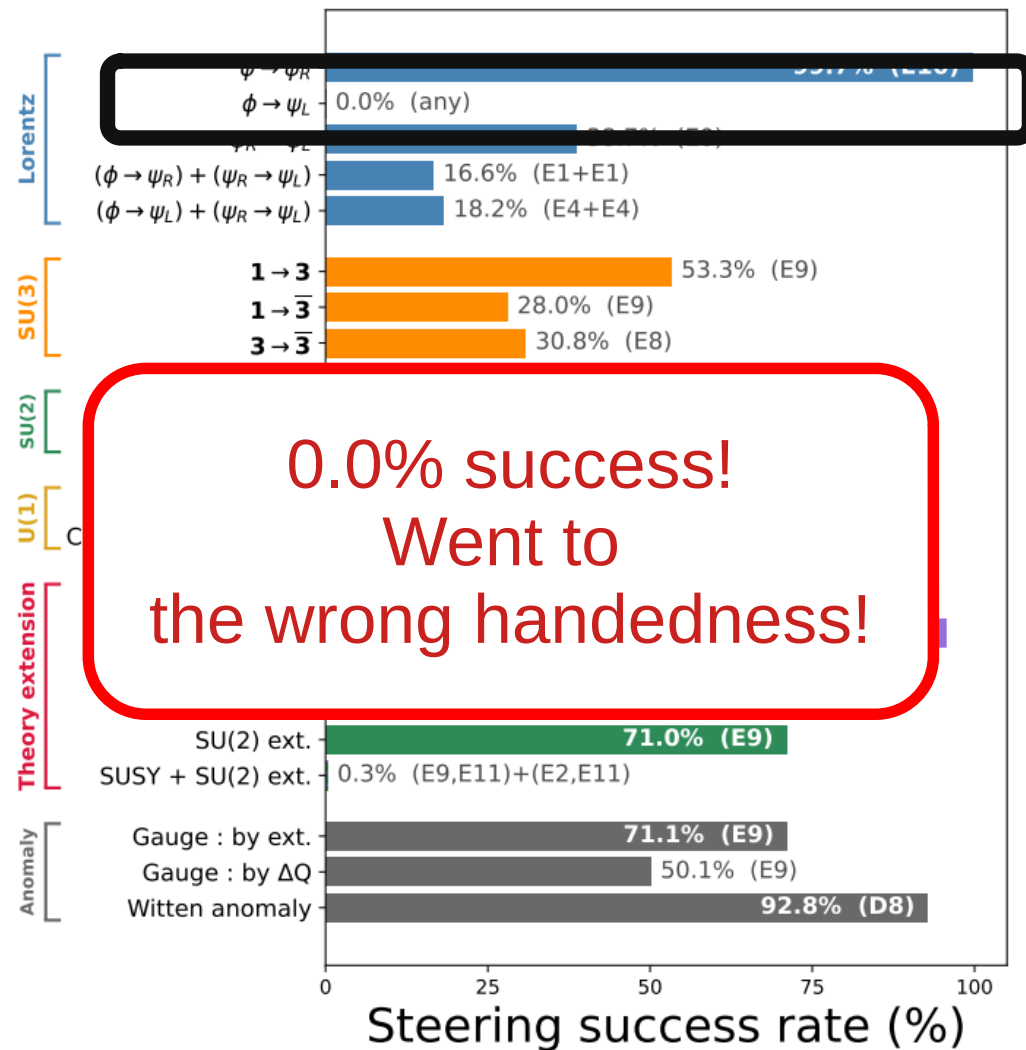
injected @ enc L11 -> viewed @ enc L11 (final layer)



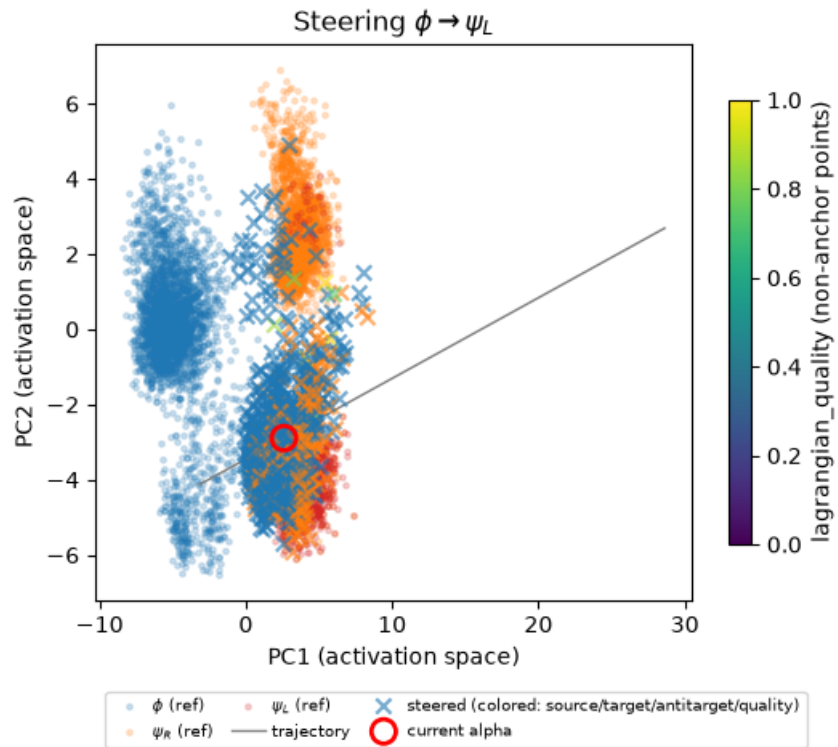
alpha = 0.00

success = 0%

cumulative = 0%



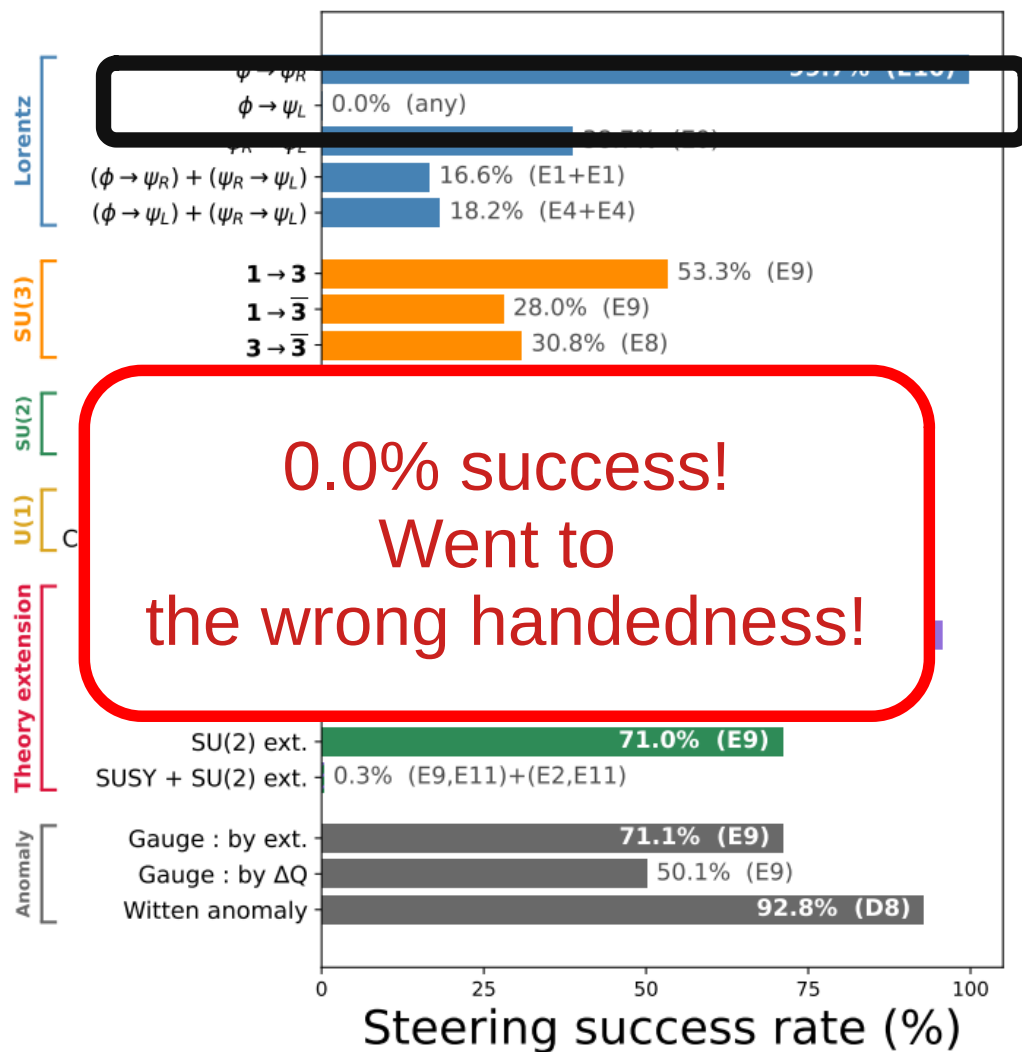
injected @ enc L11 -> viewed @ enc L11 (final layer)



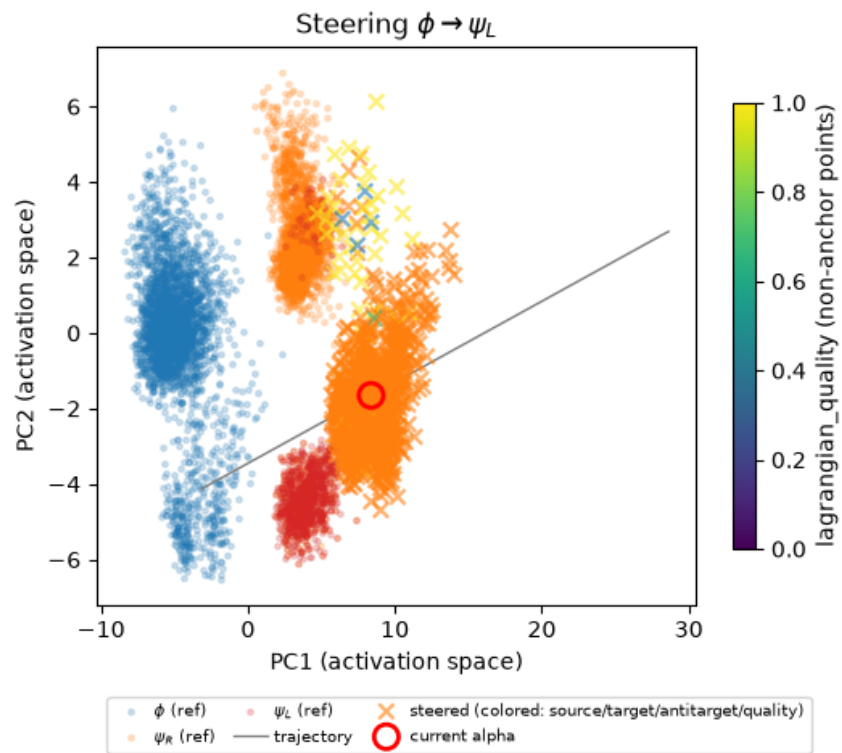
alpha = 0.20

success = 0%

cumulative = 0%



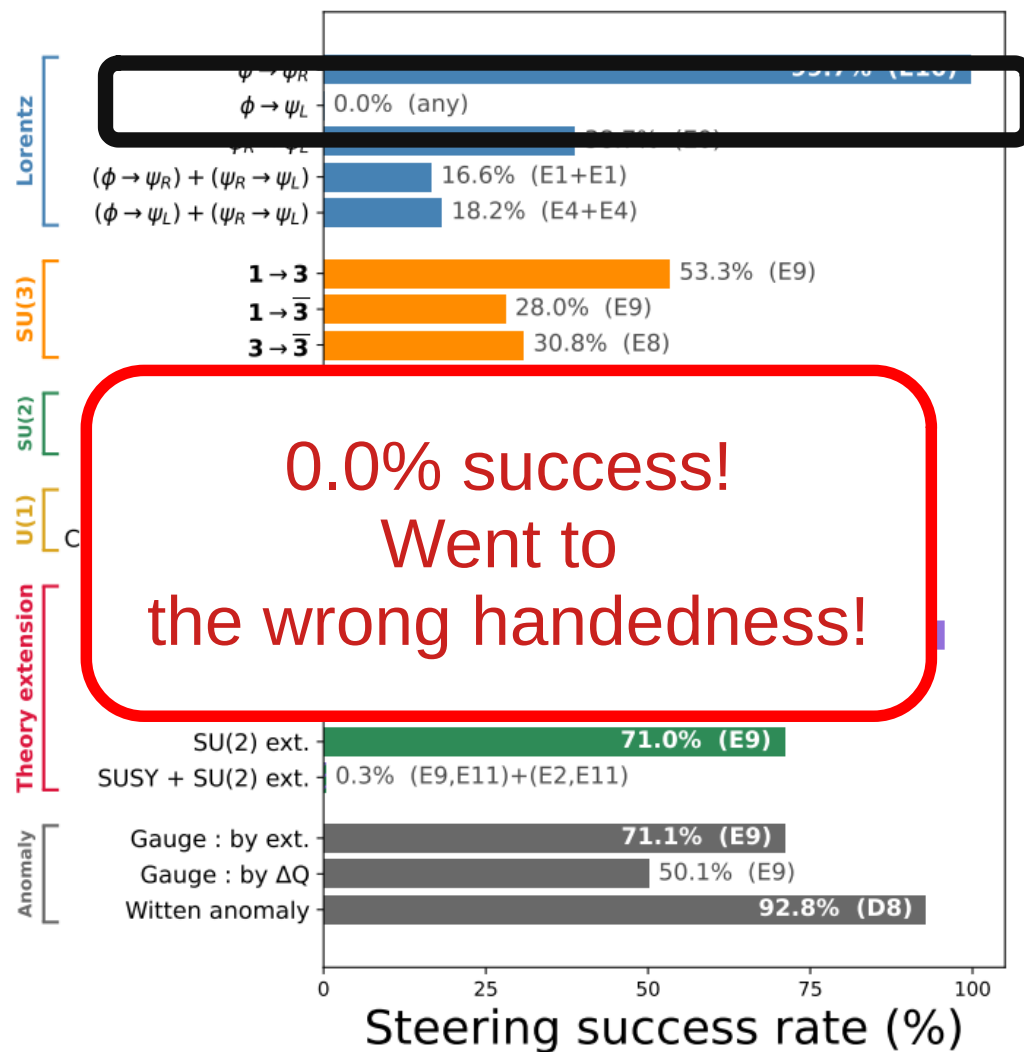
injected @ enc L11 -> viewed @ enc L11 (final layer)



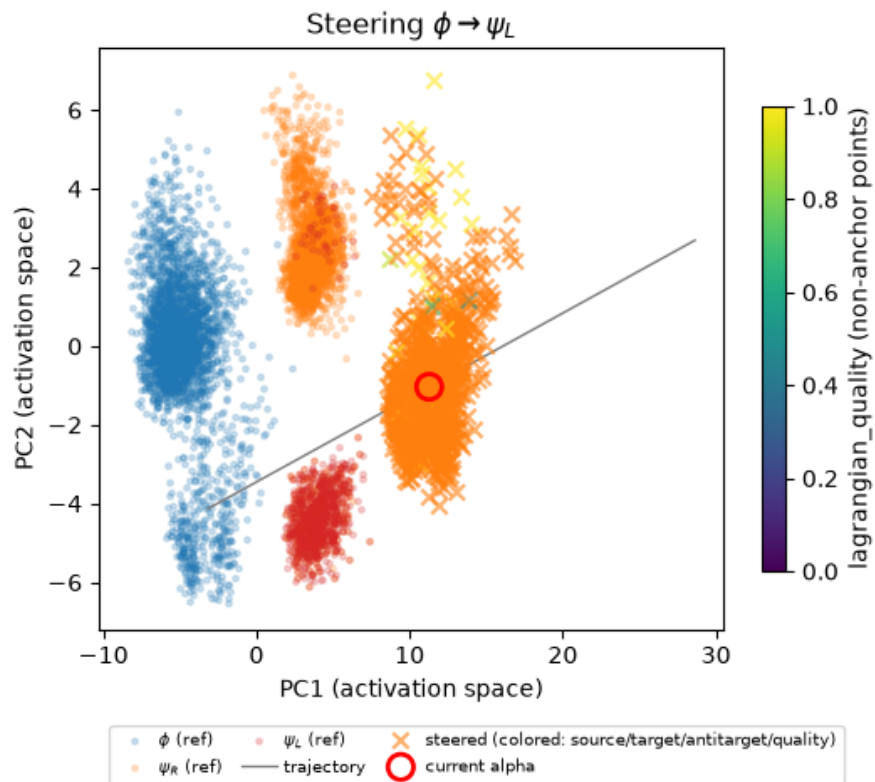
alpha = 0.40

success = 0%

cumulative = 0%



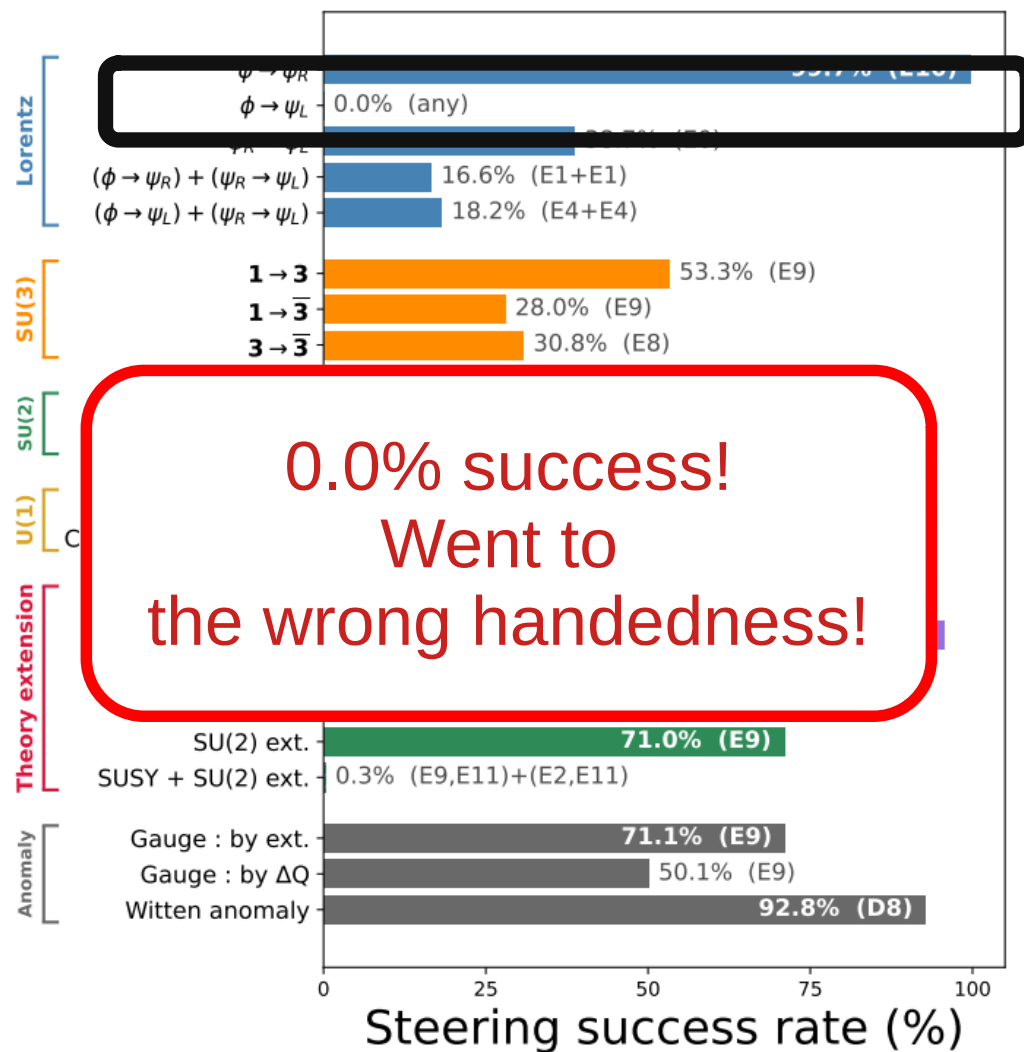
injected @ enc L11 -> viewed @ enc L11 (final layer)



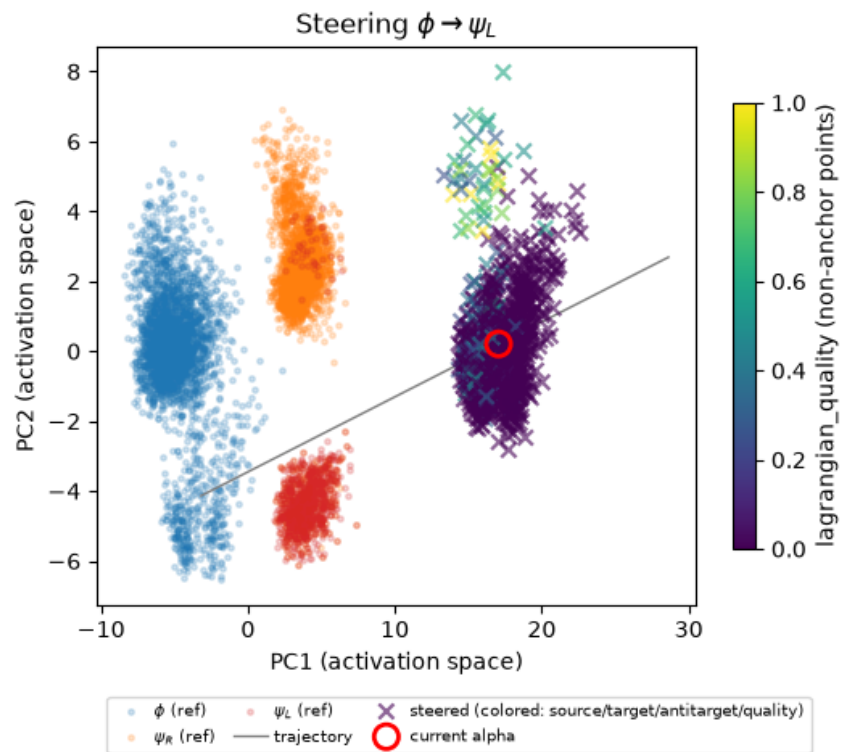
alpha = 0.50

success = 0%

cumulative = 0%



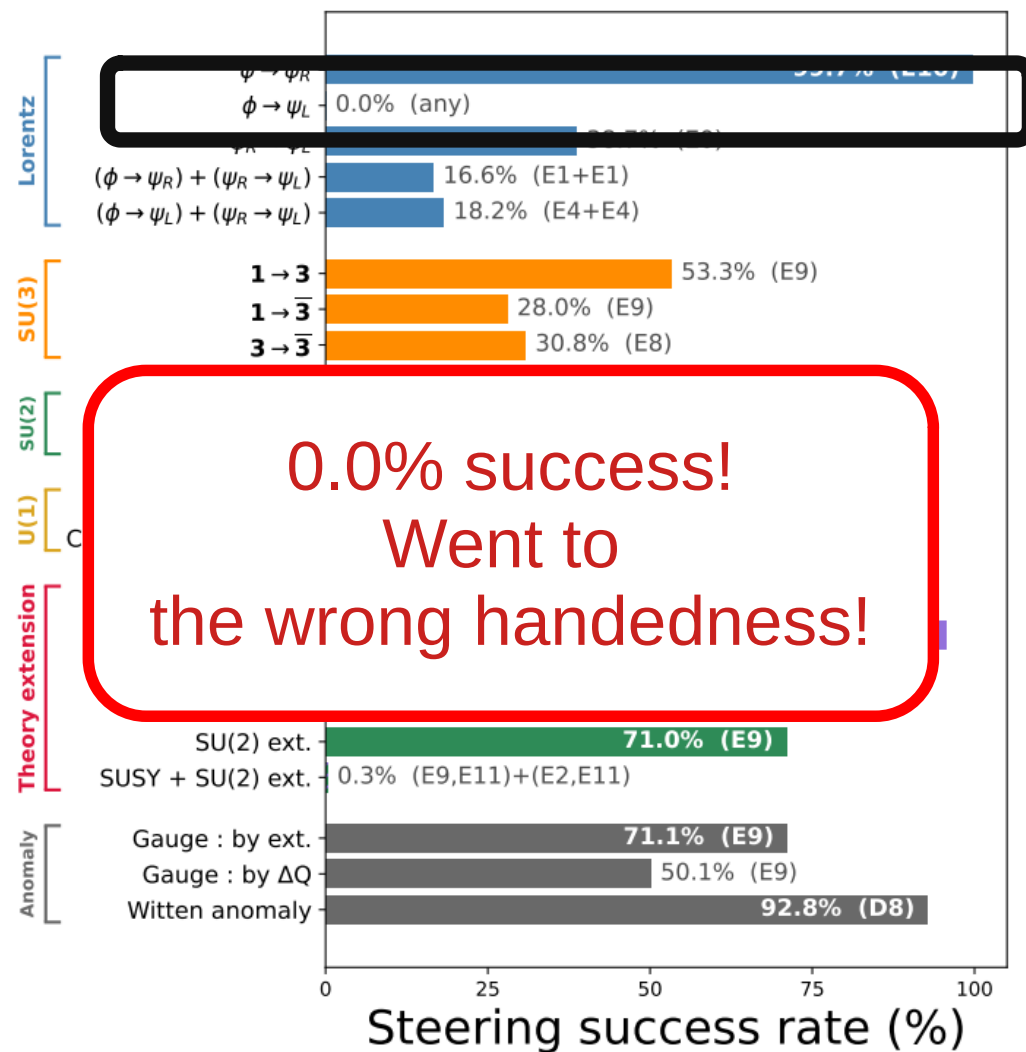
injected @ enc L11 -> viewed @ enc L11 (final layer)



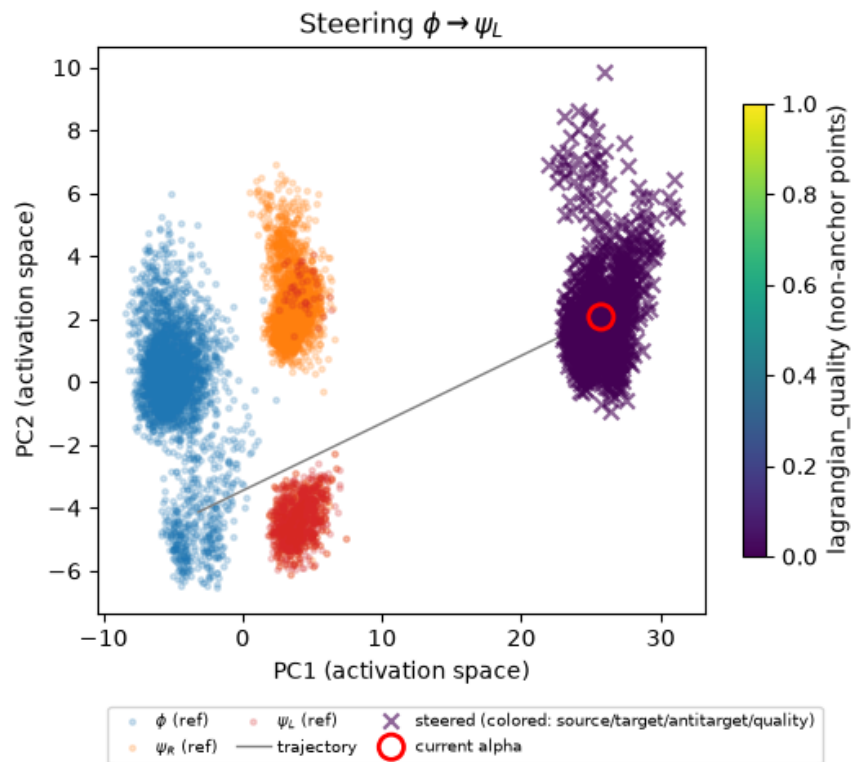
alpha = 0.70

success = 0%

cumulative = 0%



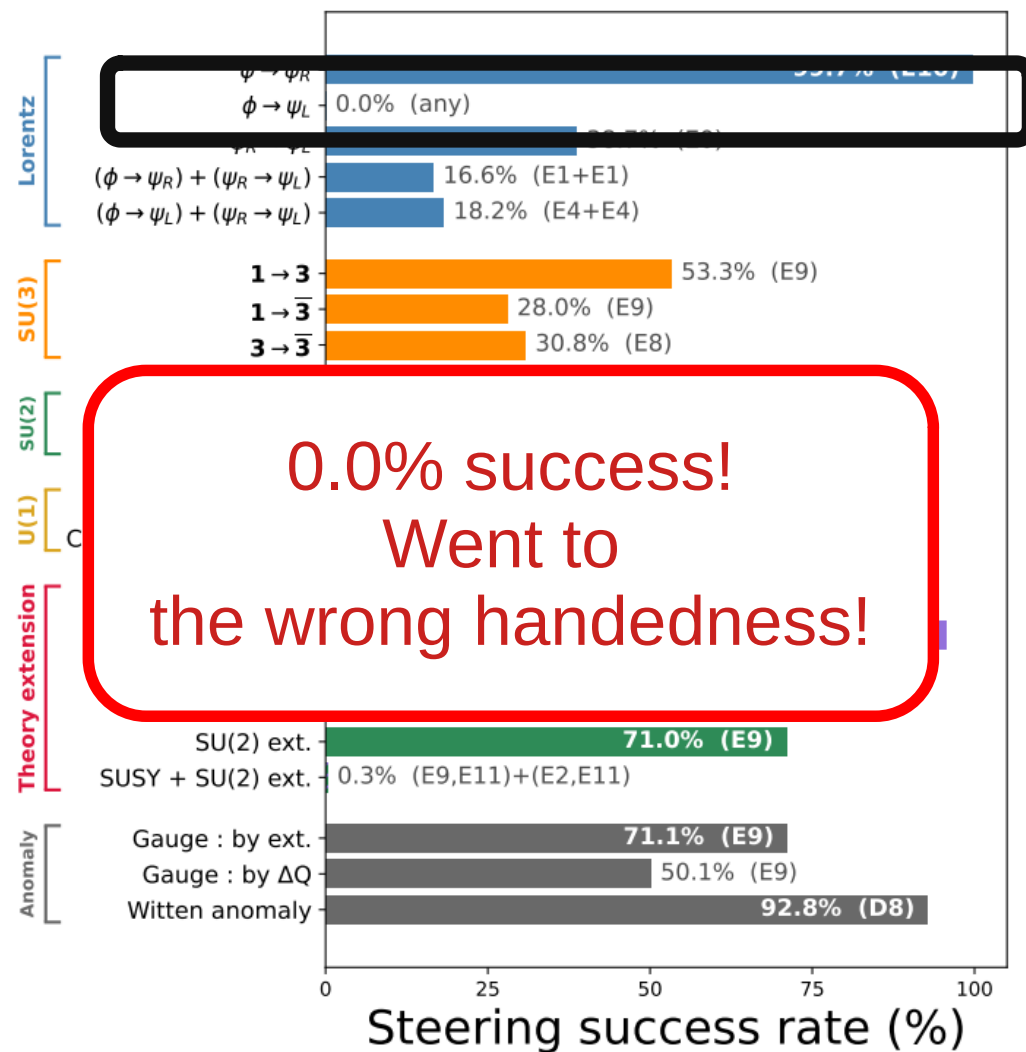
injected @ enc L11 -> viewed @ enc L11 (final layer)

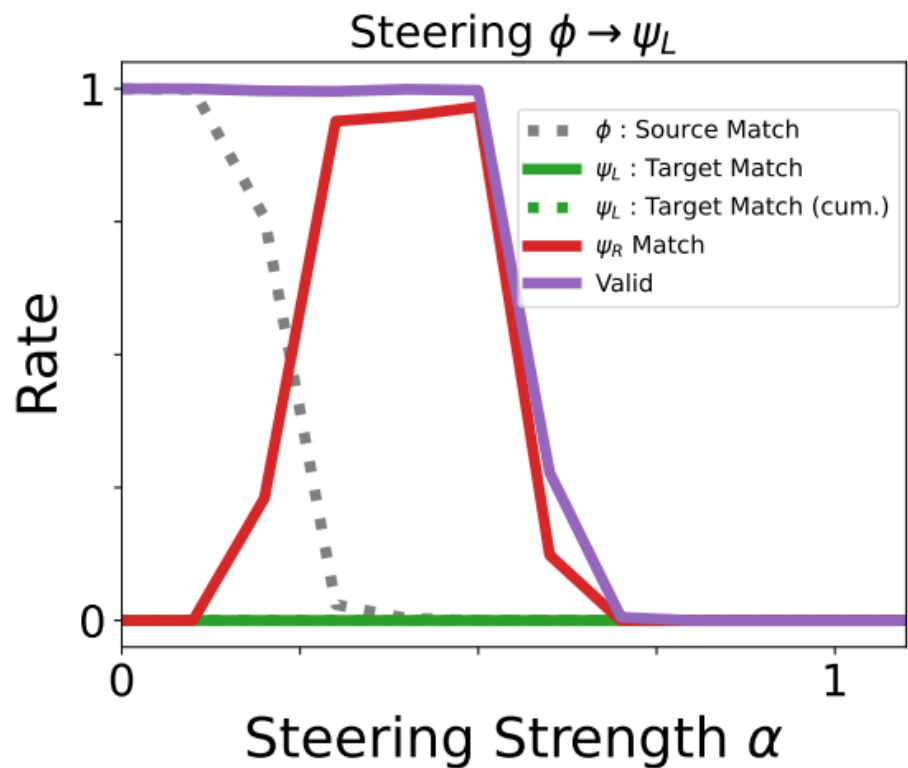


alpha = 1.00

success = 0%

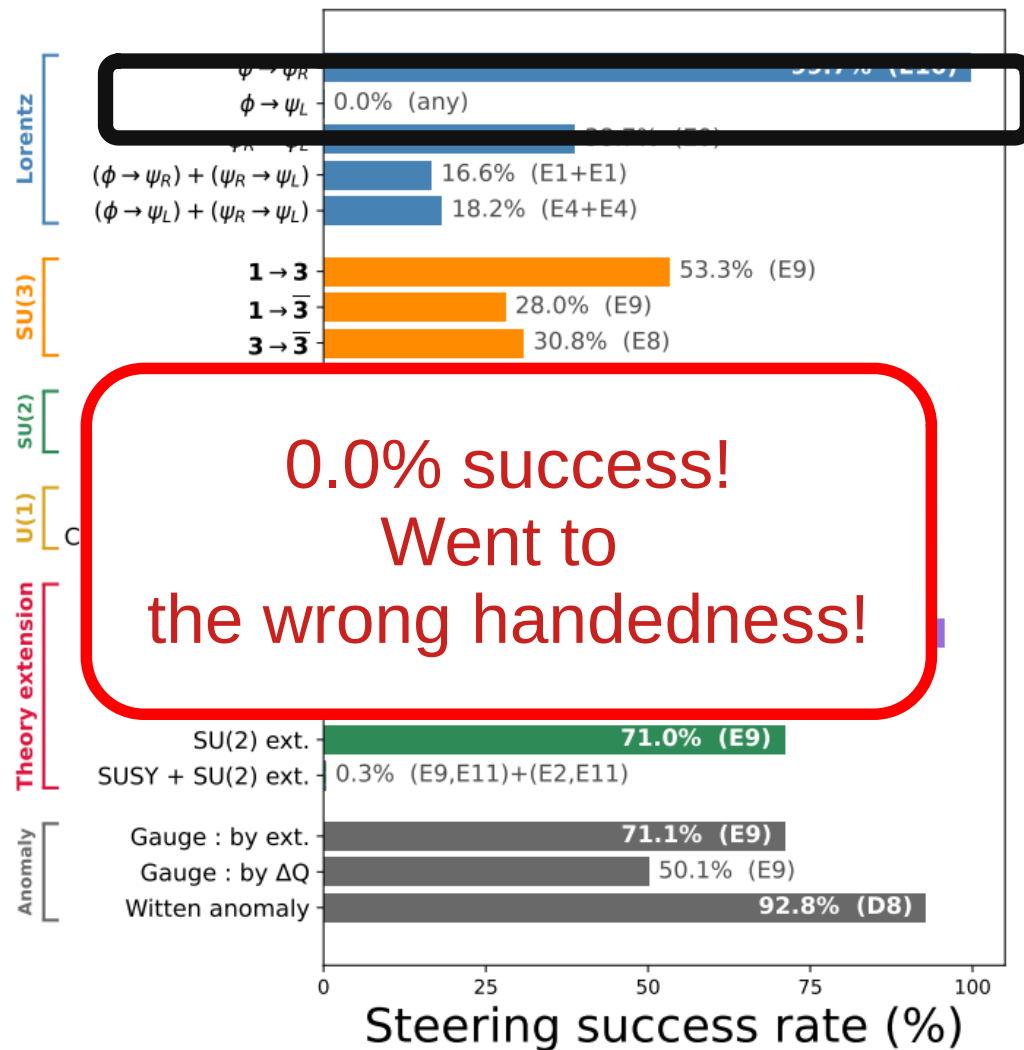
cumulative = 0%

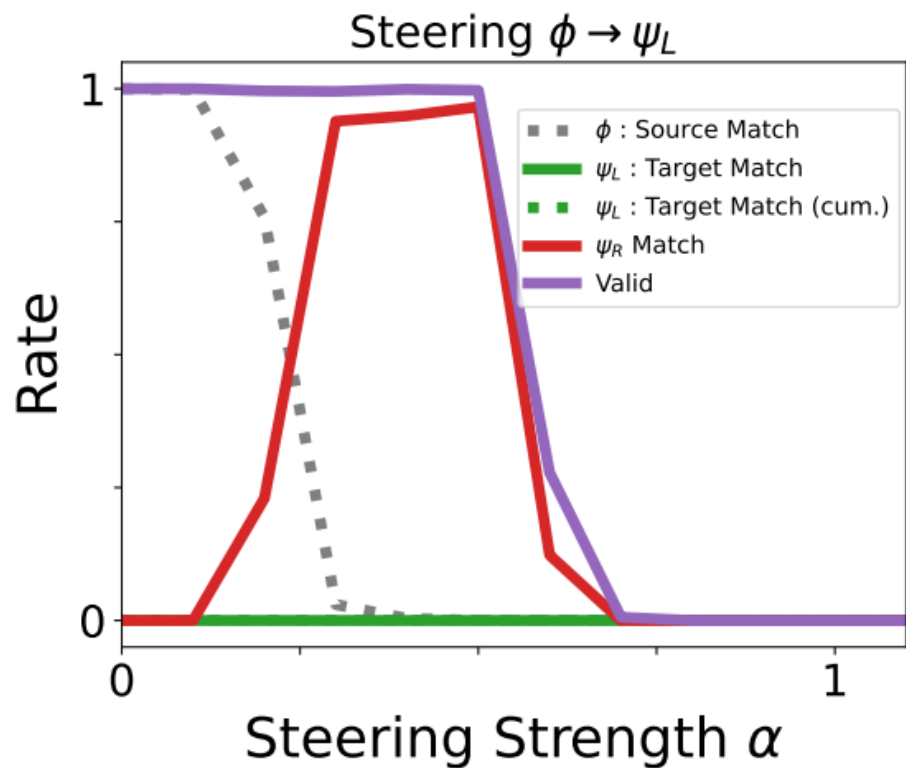




Behavior

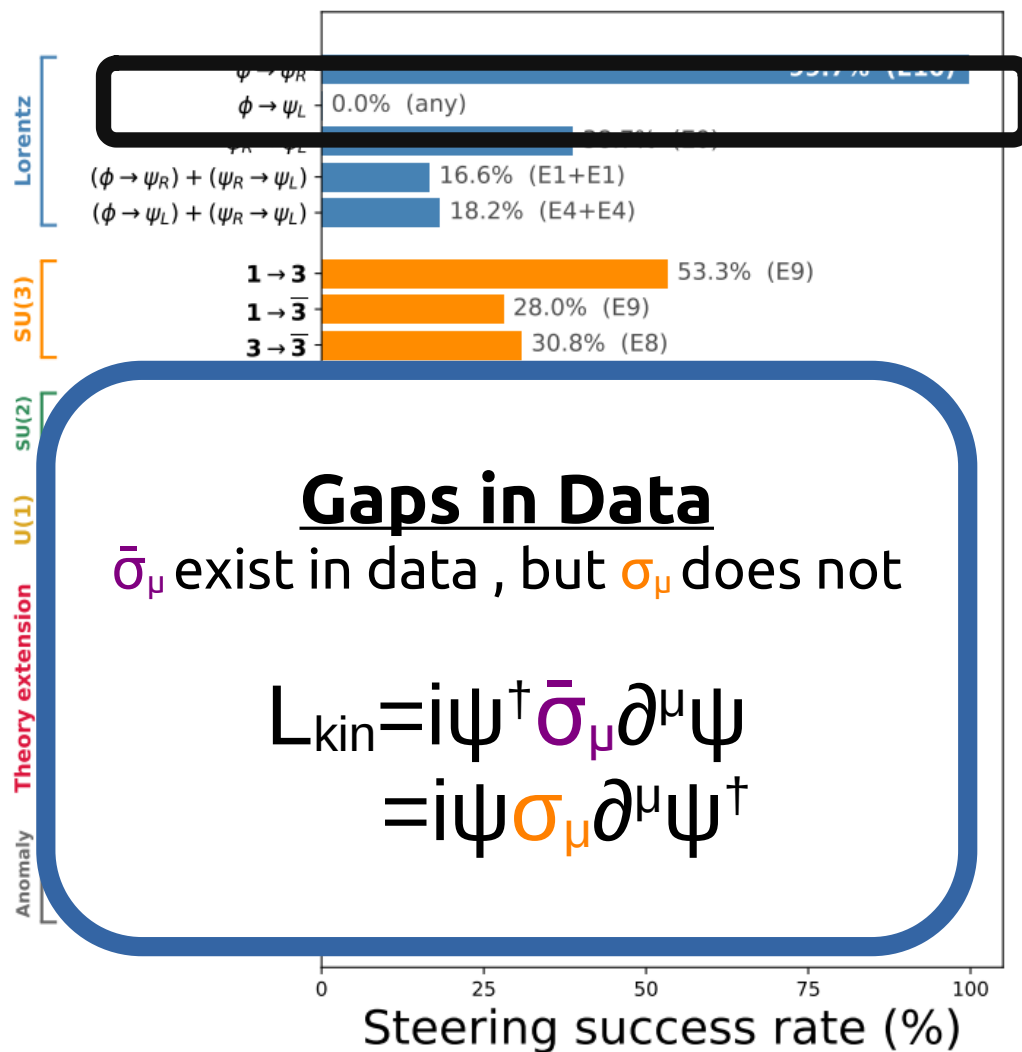
It went to ψ_R , instead of ψ_L !





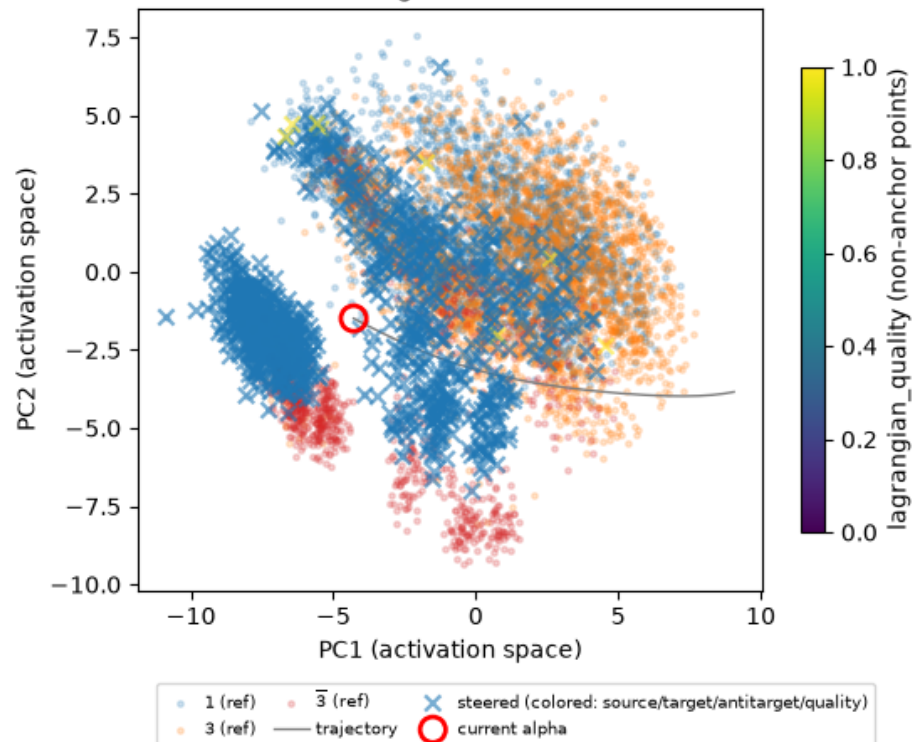
Behavior

It went to ψ_R , instead of ψ_L !



injected @ enc L9 -> viewed @ enc L11 (final layer)

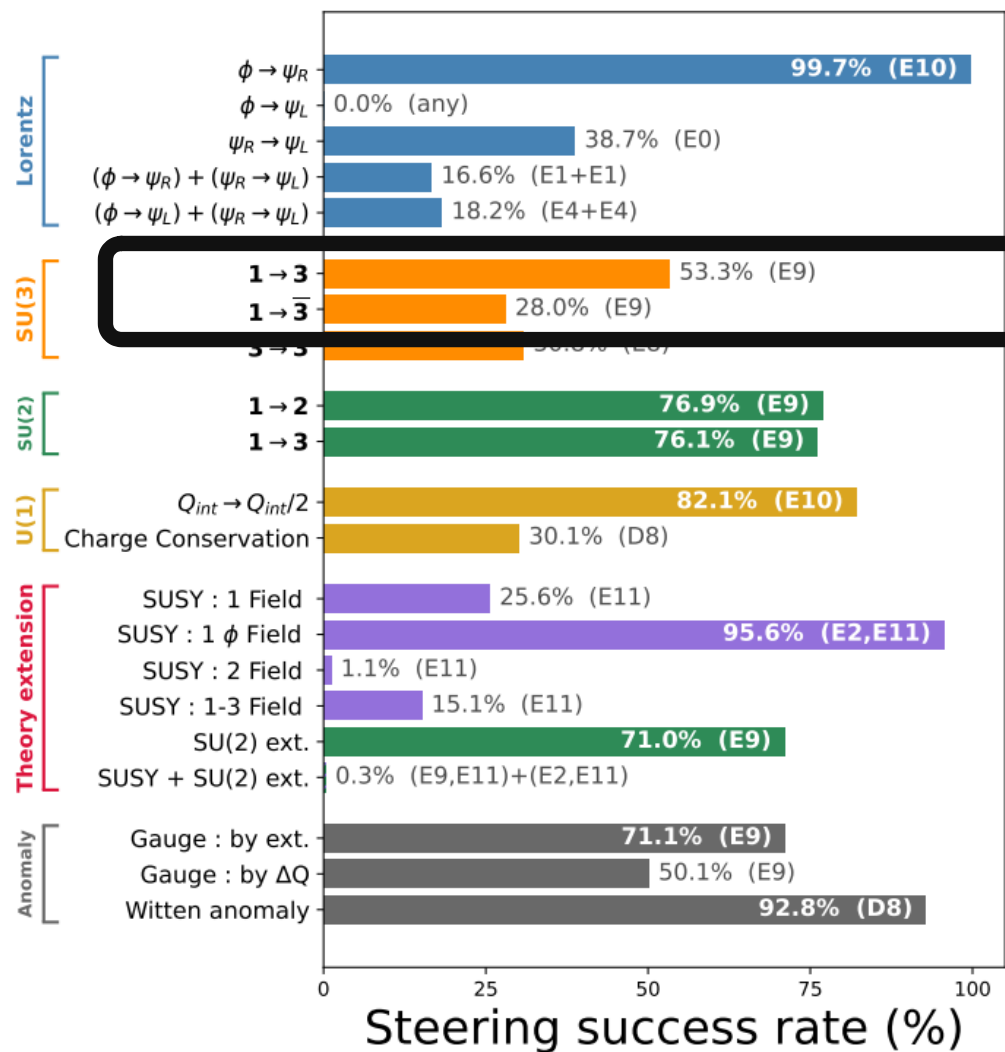
Steering $1 \rightarrow 3$ (SU(3))



alpha = 0.00

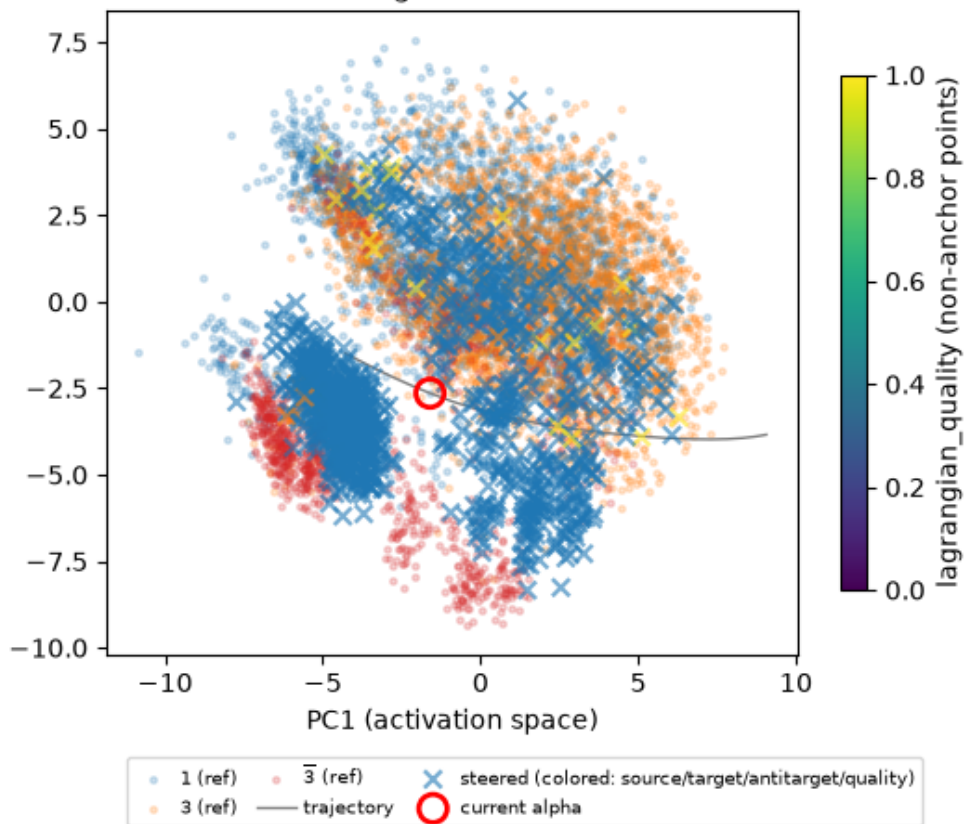
success = 0%

cumulative = 0%



injected @ enc L9 -> viewed @ enc L11 (final layer)

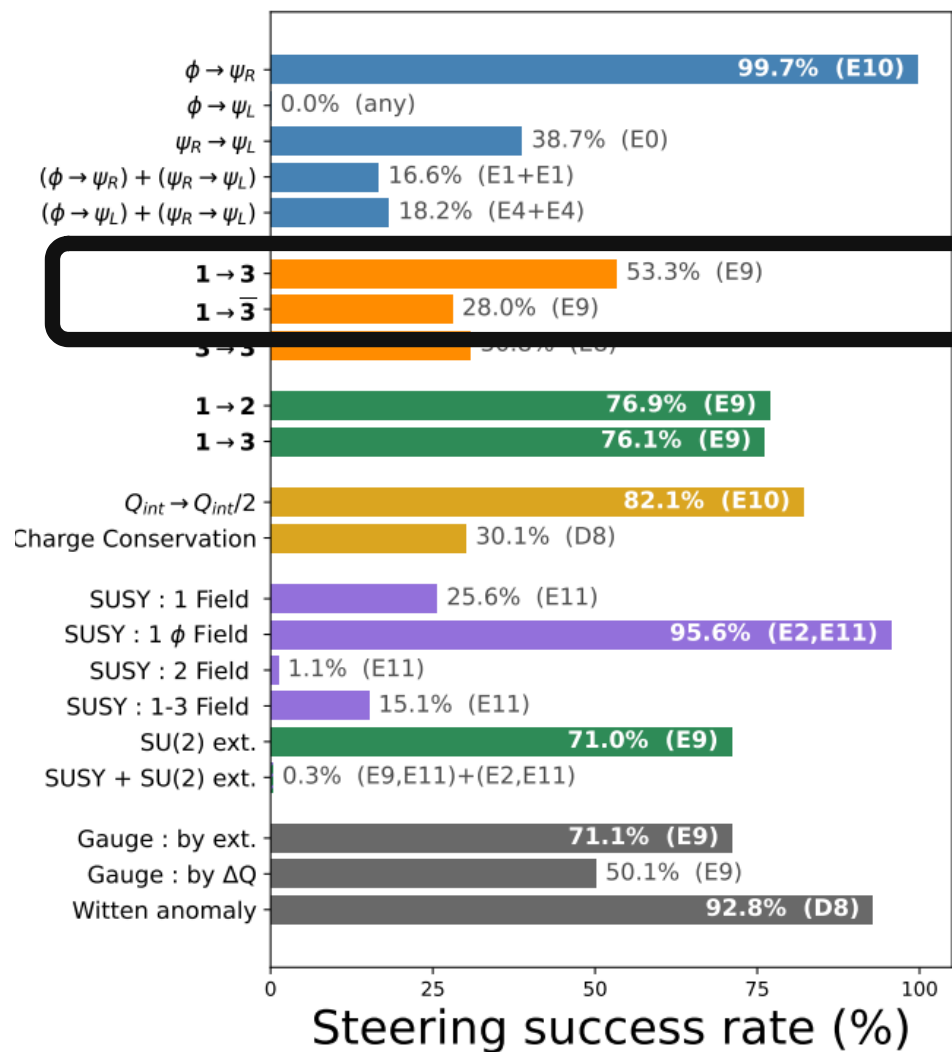
Steering $1 \rightarrow 3$ (SU(3))



alpha = 0.60

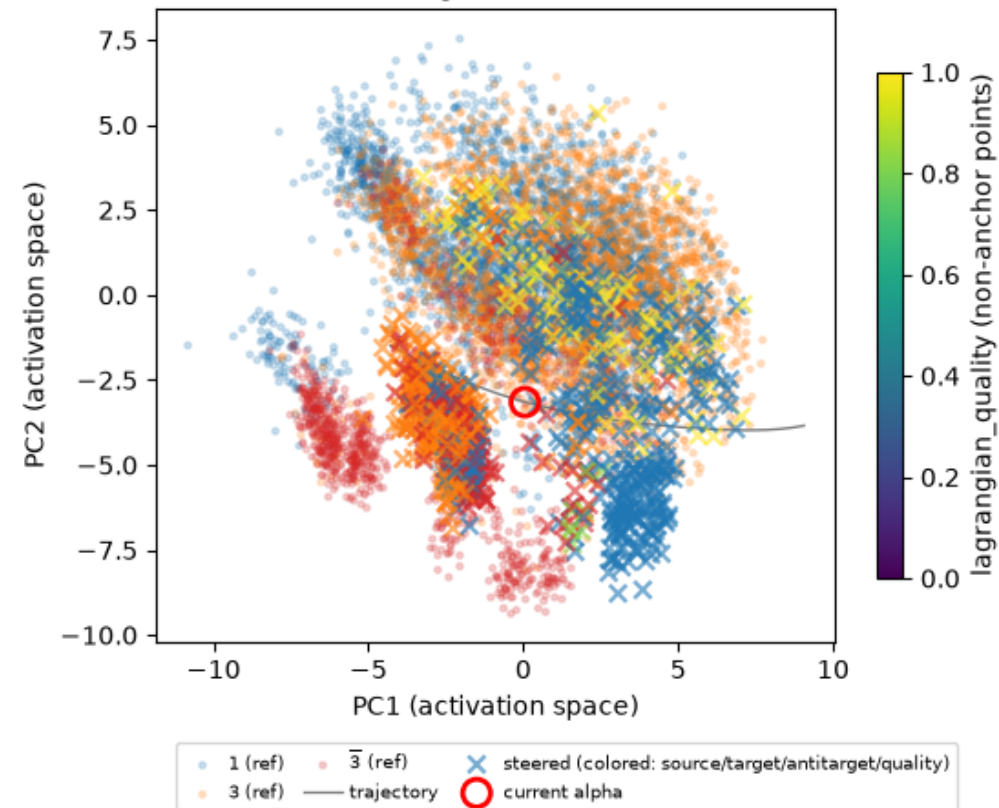
success = 1%

cumulative = 1%



injected @ enc L9 -> viewed @ enc L11 (final layer)

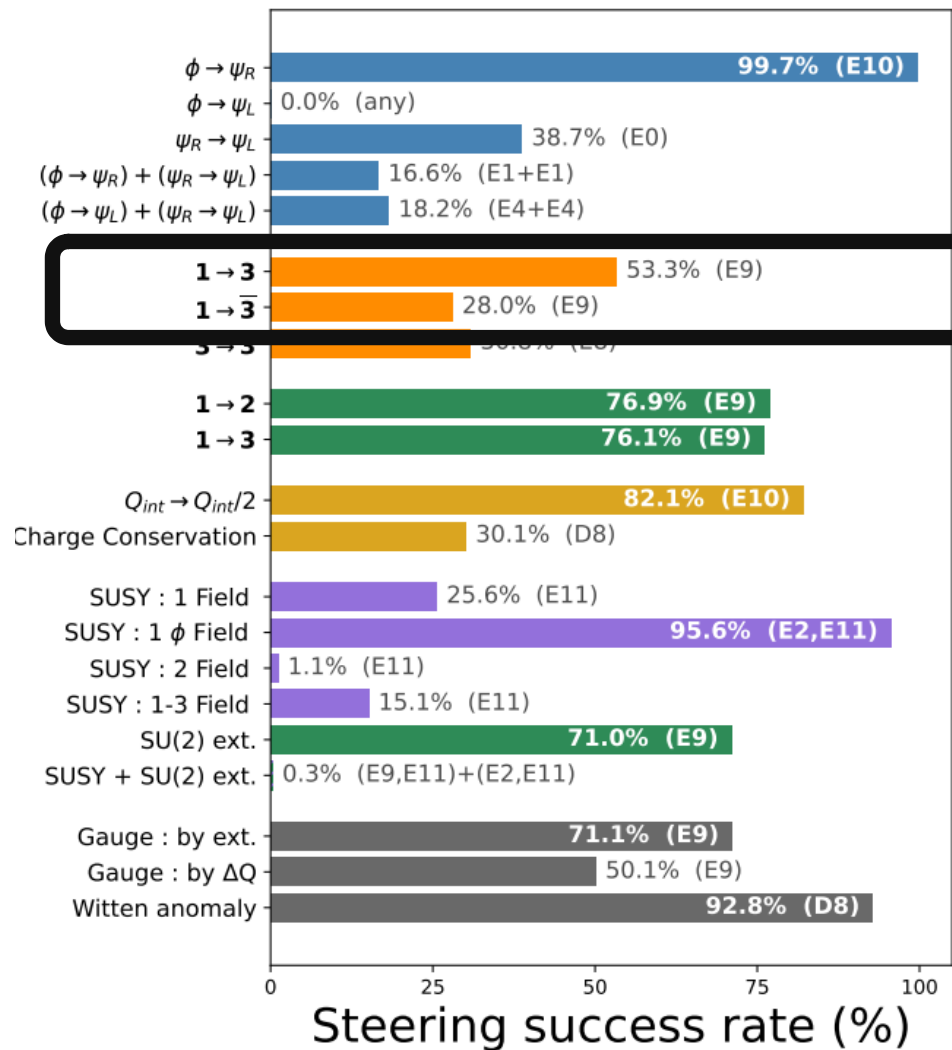
Steering $1 \rightarrow 3$ (SU(3))



alpha = 0.80

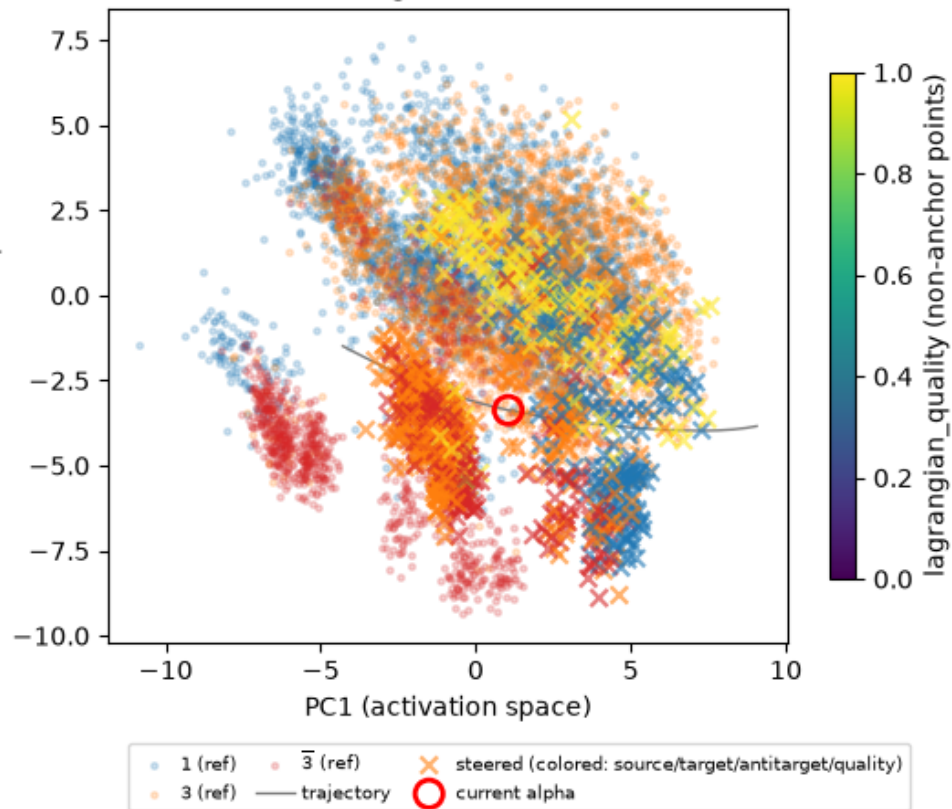
success = 27%

cumulative = 27%



injected @ enc L9 -> viewed @ enc L11 (final layer)

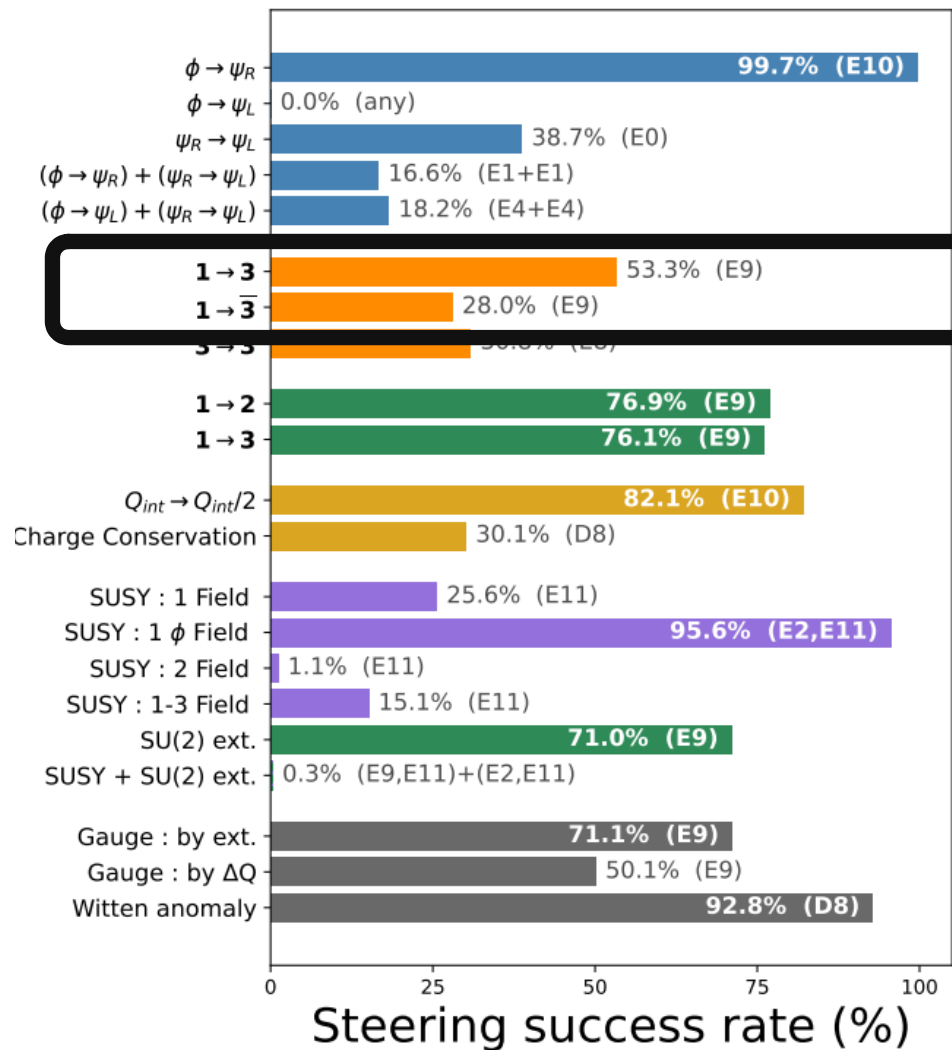
Steering $1 \rightarrow 3$ (SU(3))



alpha = 0.90

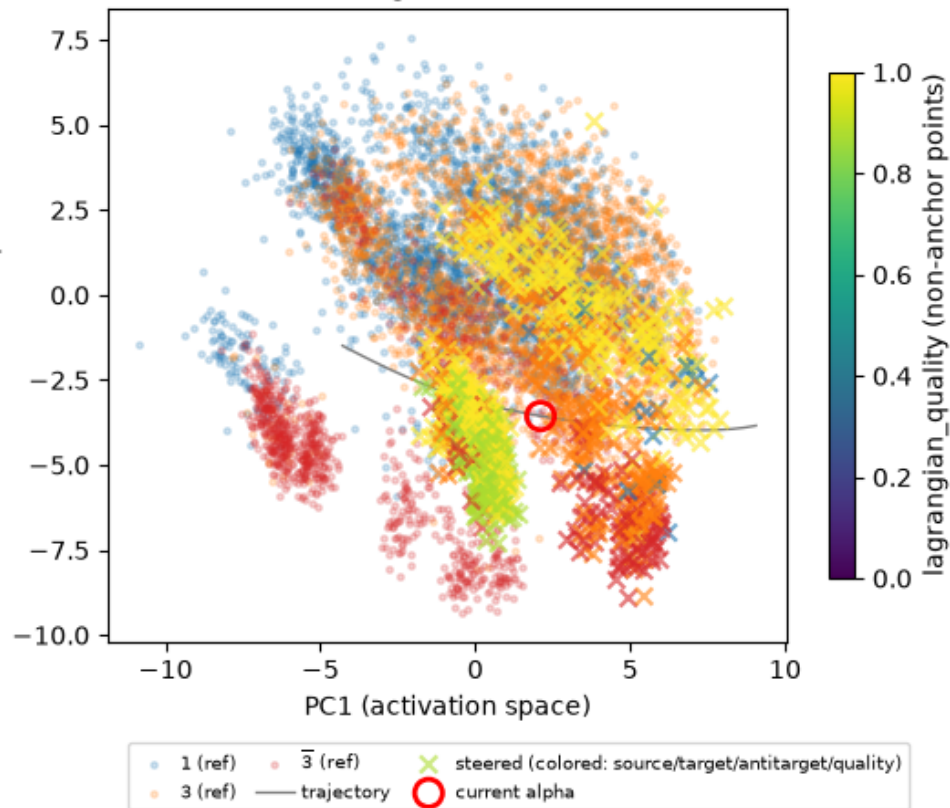
success = 36%

cumulative = 40%



injected @ enc L9 -> viewed @ enc L11 (final layer)

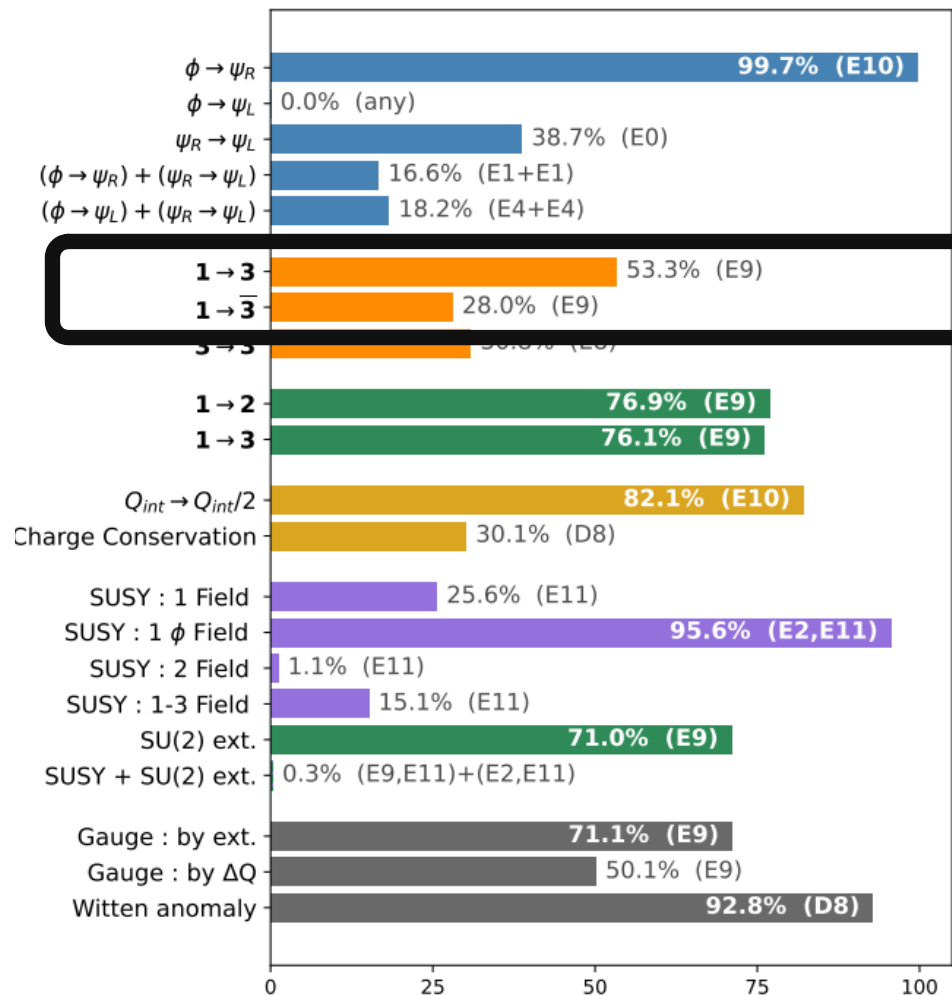
Steering $1 \rightarrow 3$ (SU(3))



alpha = 1.00

success = 21%

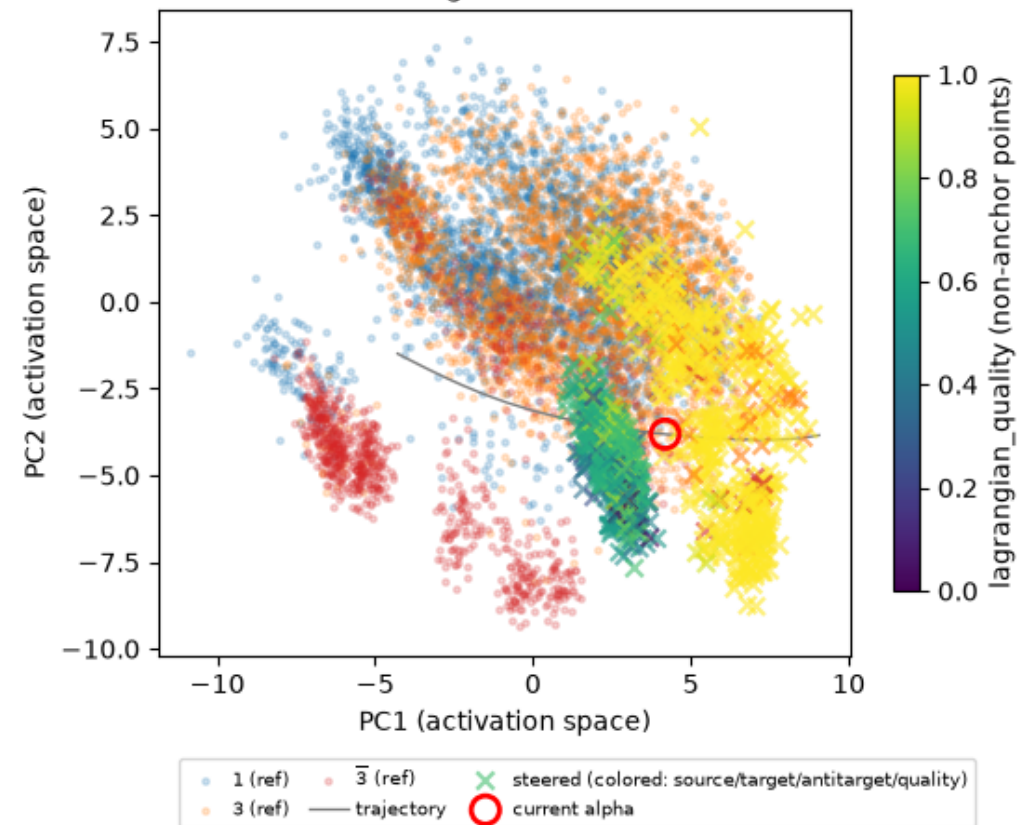
cumulative = 50%



Steering success rate (%)

injected @ enc L9 -> viewed @ enc L11 (final layer)

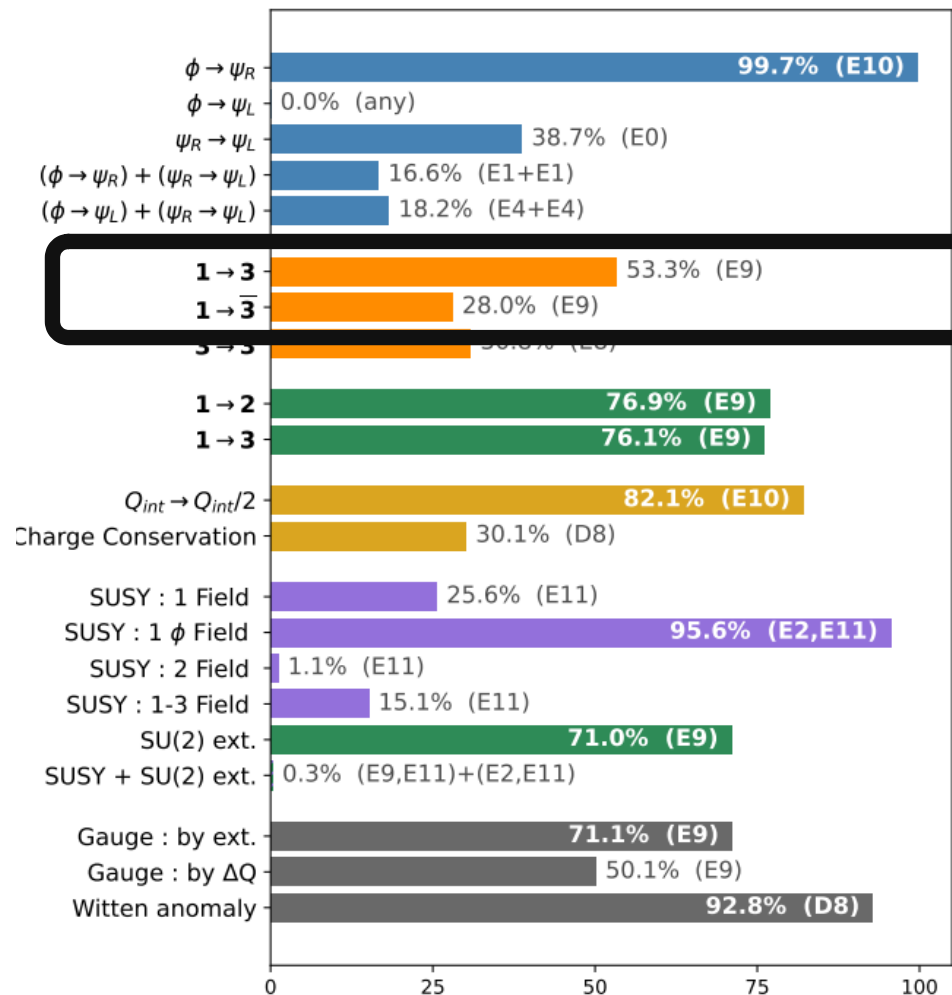
Steering $1 \rightarrow 3$ (SU(3))



alpha = 1.20

success = 4%

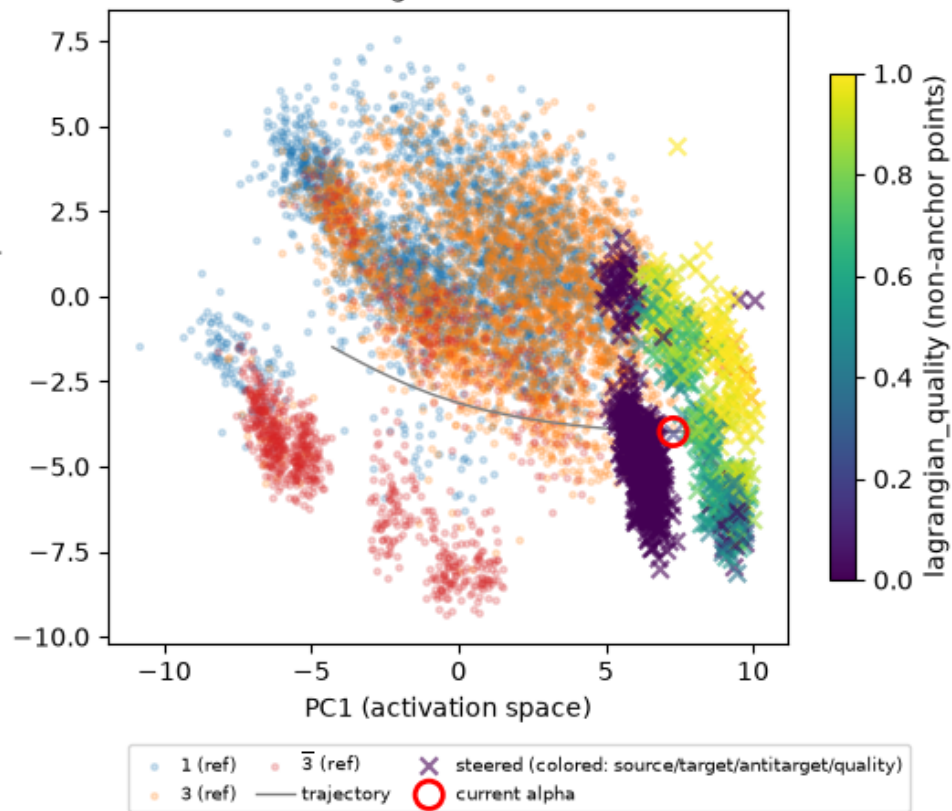
cumulative = 53%



Steering success rate (%)

injected @ enc L9 -> viewed @ enc L11 (final layer)

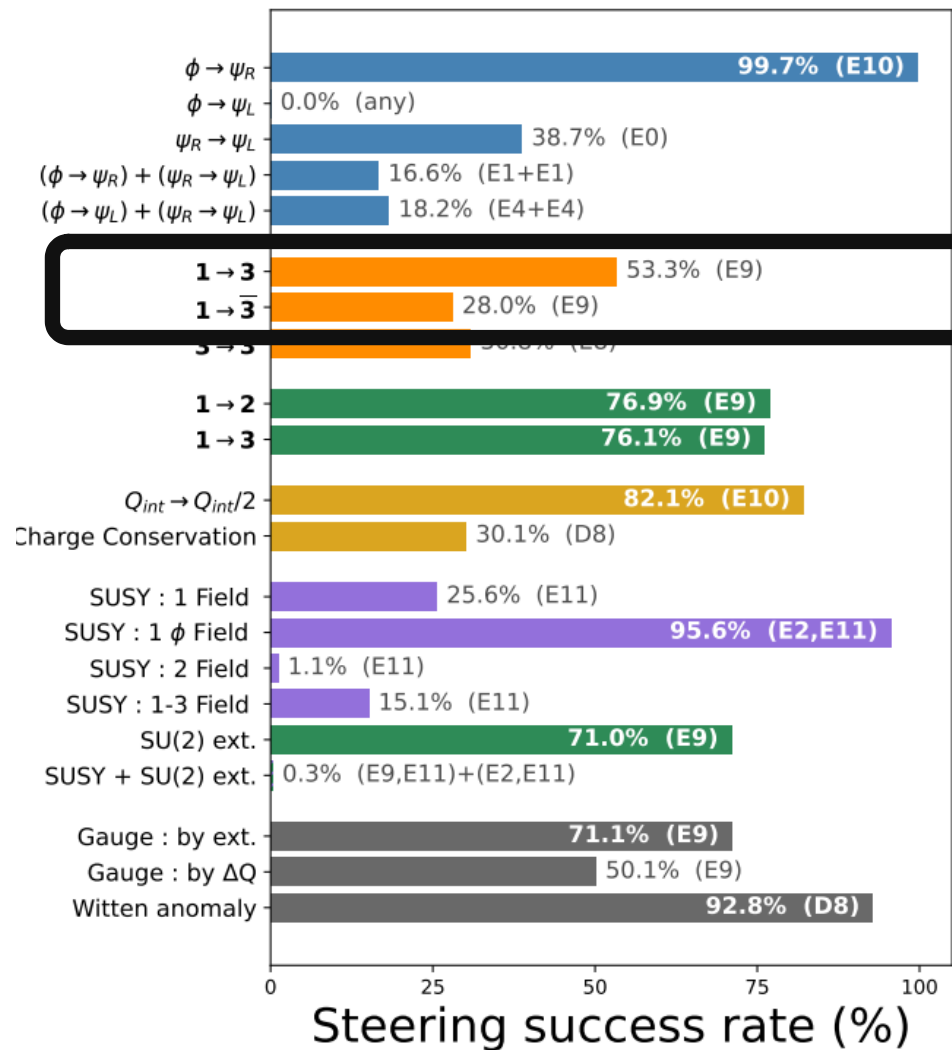
Steering $1 \rightarrow 3$ (SU(3))



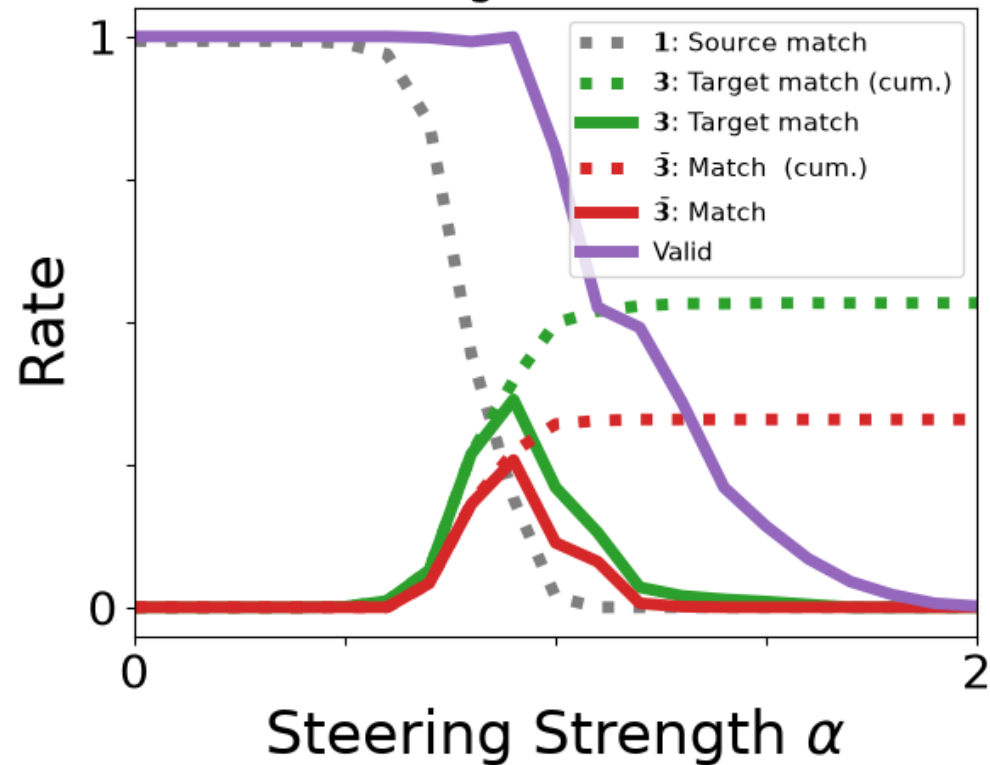
alpha = 1.60

success = 0%

cumulative = 53%

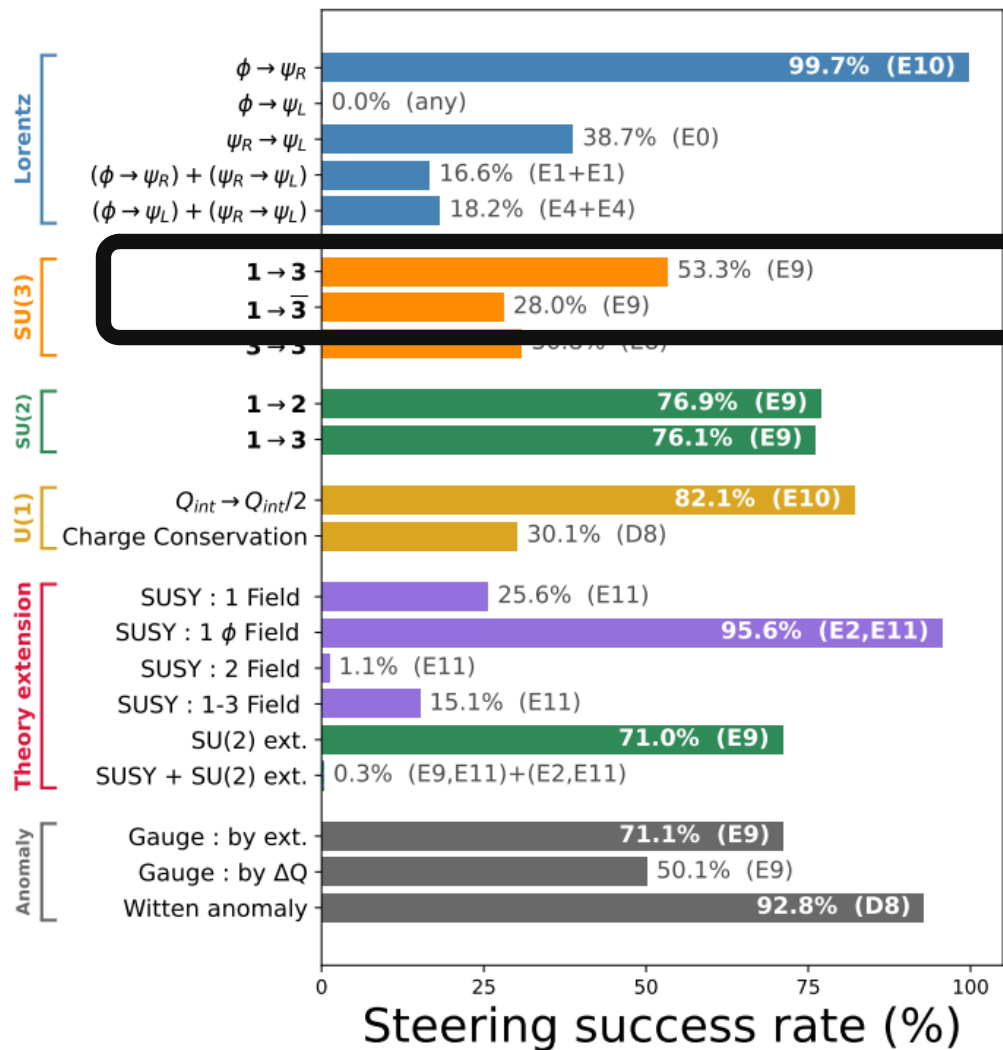


Steering SU(3) : $\mathbf{1} \rightarrow \mathbf{3}$

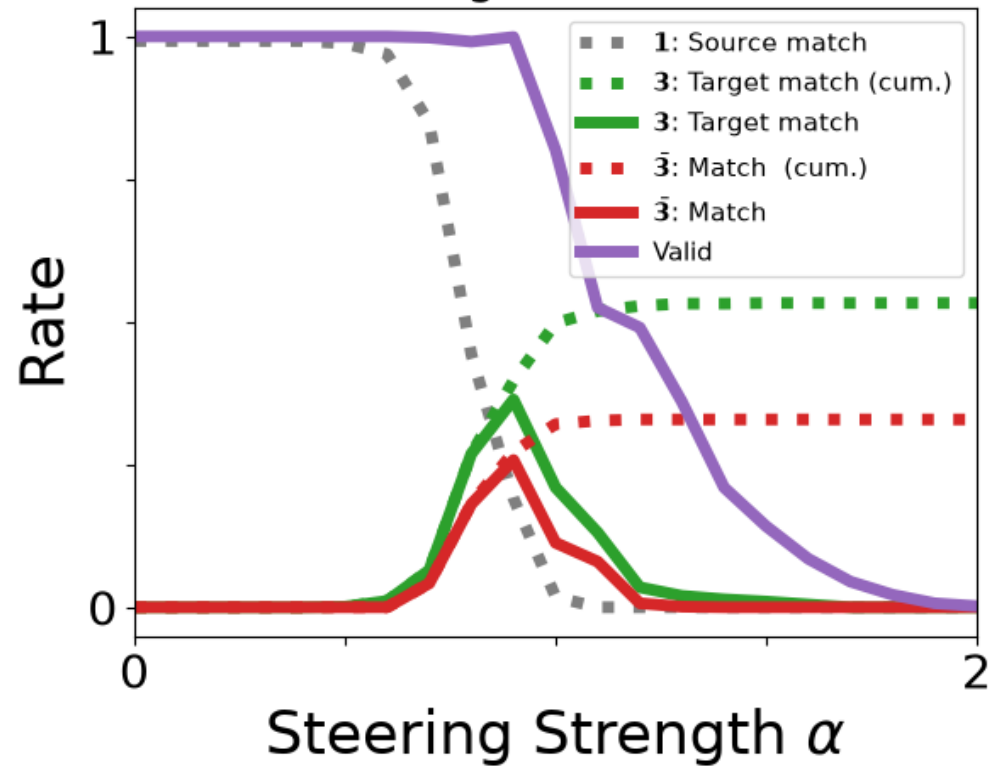


Behavior

It goes to 3,
but when it doesn't, it goes to $\bar{3}$

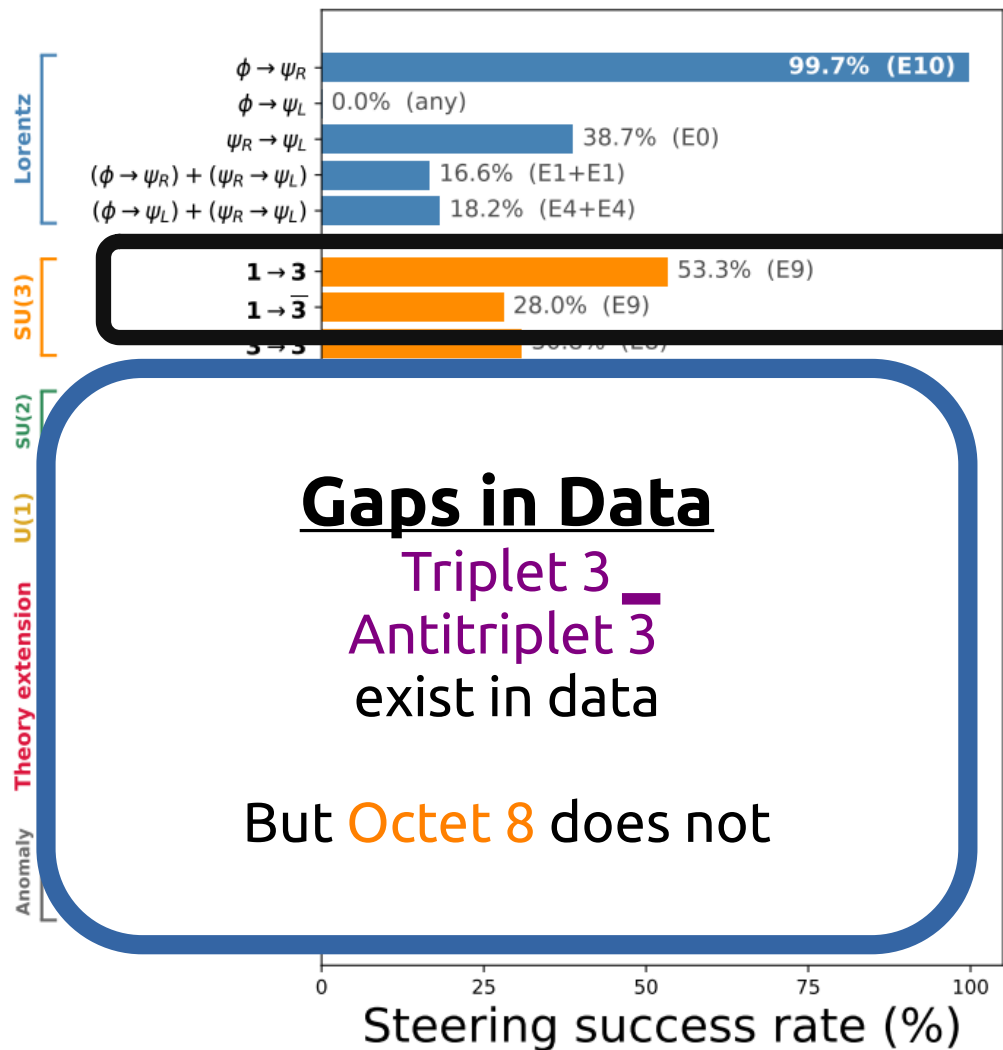


Steering SU(3) : $\mathbf{1} \rightarrow \mathbf{3}$

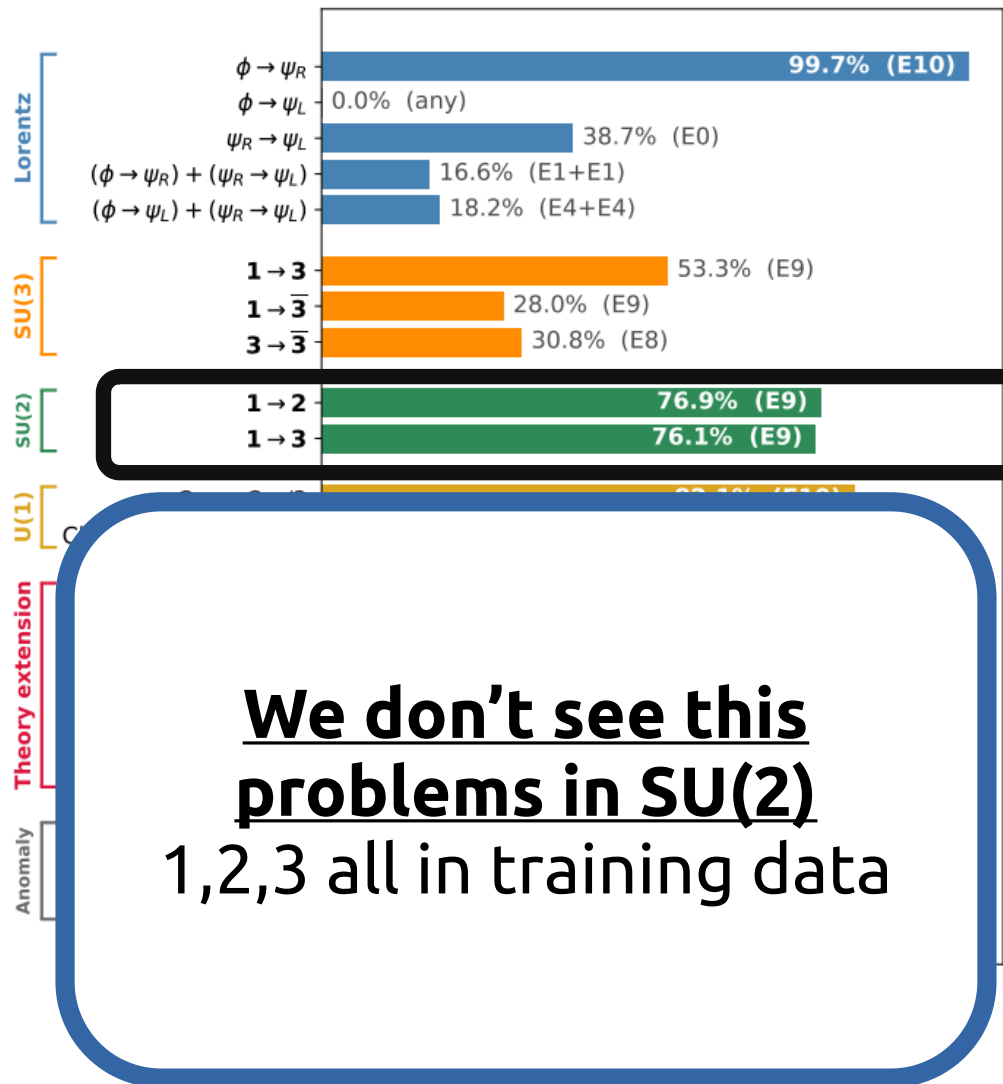
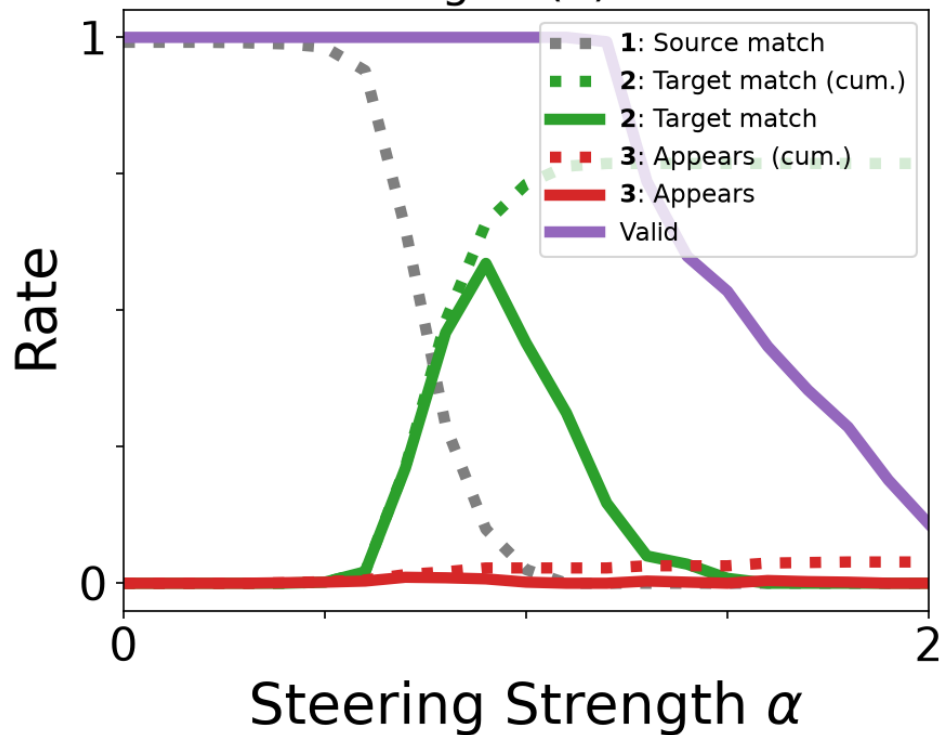


Behavior

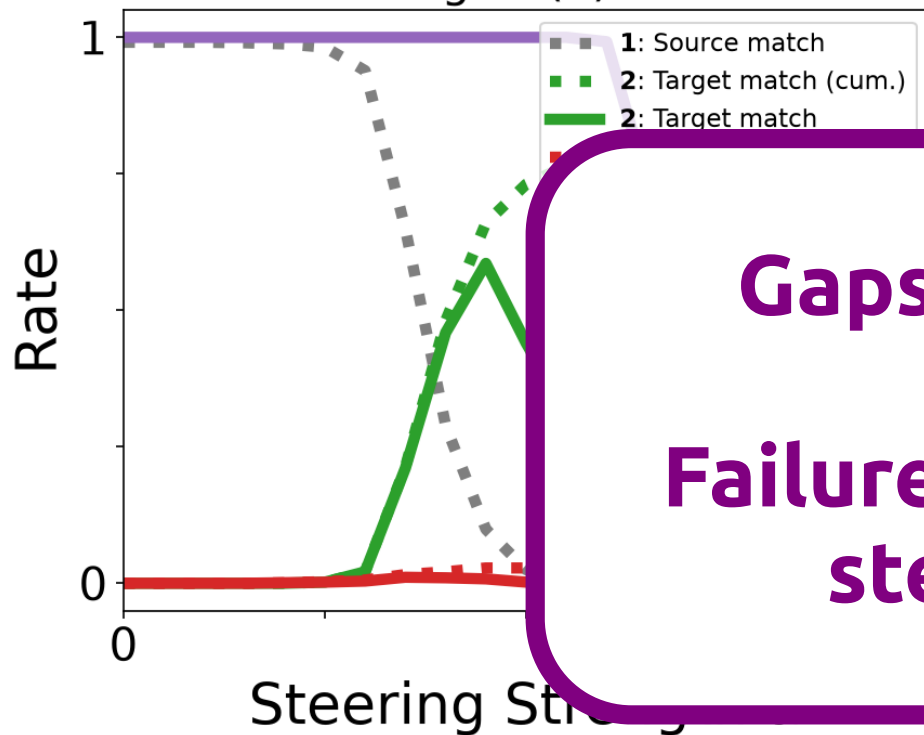
It goes to $\mathbf{3}$,
but when it doesn't, it goes to $\bar{\mathbf{3}}$



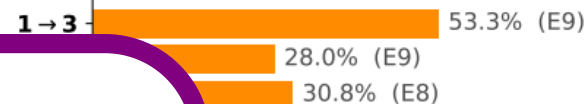
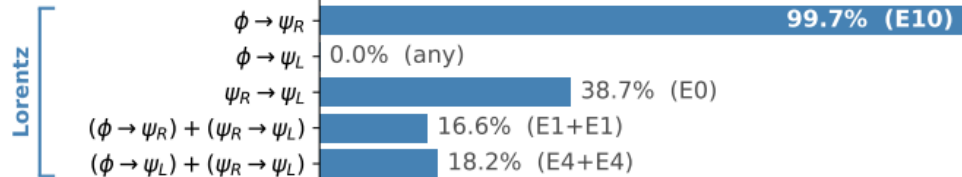
Steering SU(2) : $1 \rightarrow 2$



Steering SU(2) : $\mathbf{1} \rightarrow \mathbf{2}$



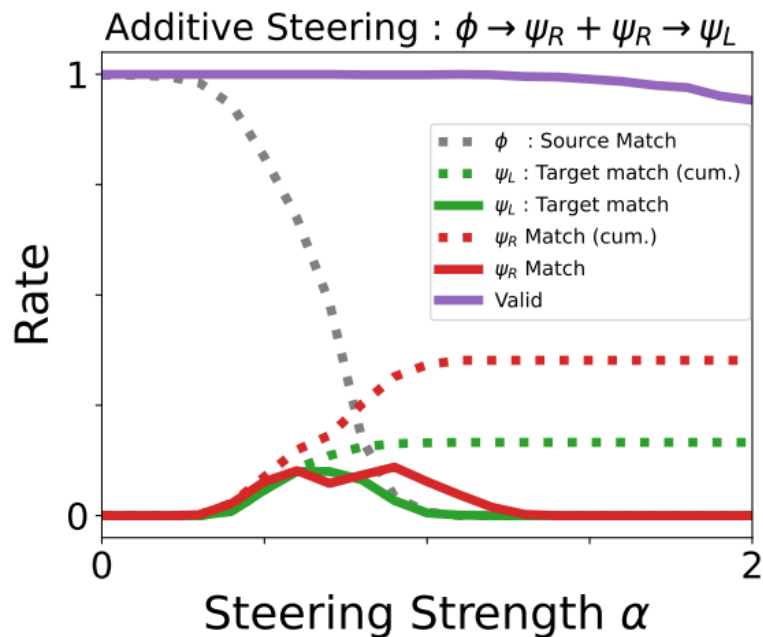
Gaps in data
~
Failure Modes in steering



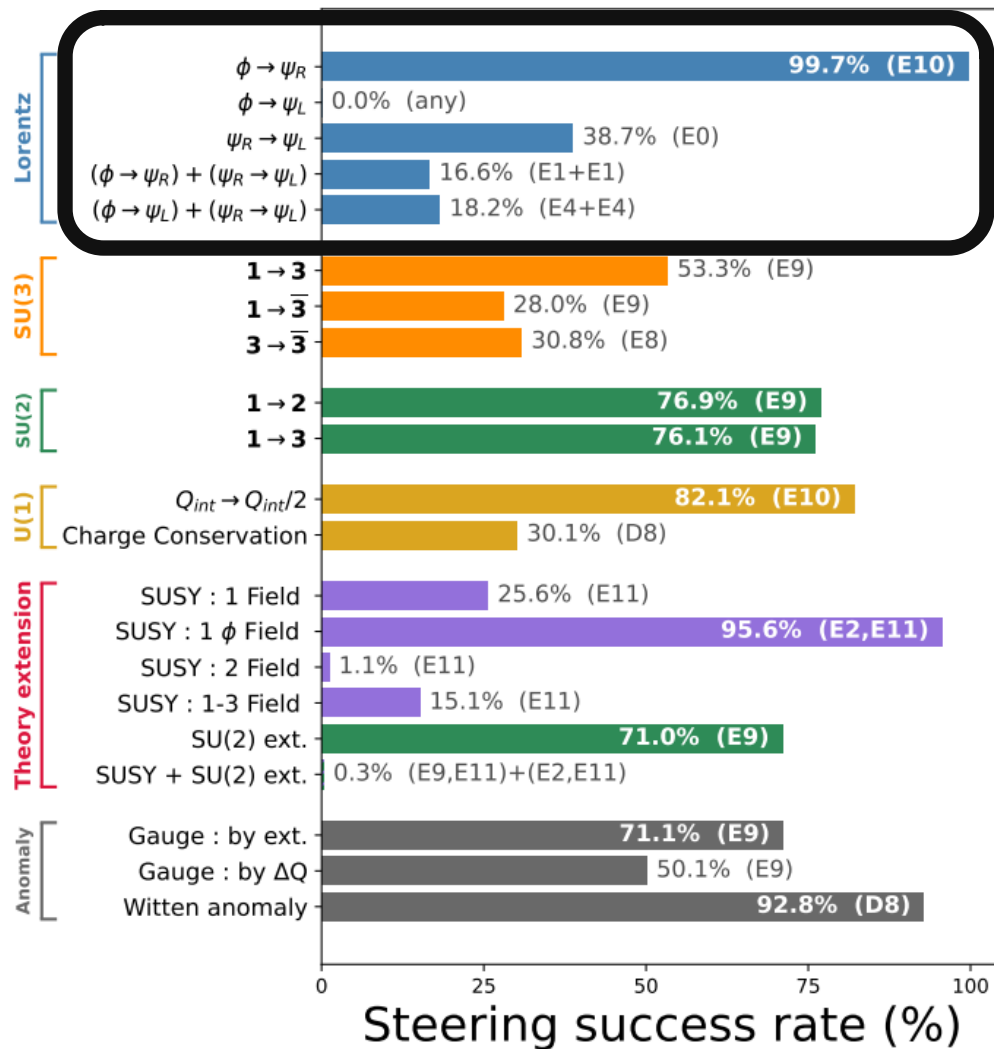
Don't see this
in SU(2)

1, 2, 3 all in training data

Steering can be additive:



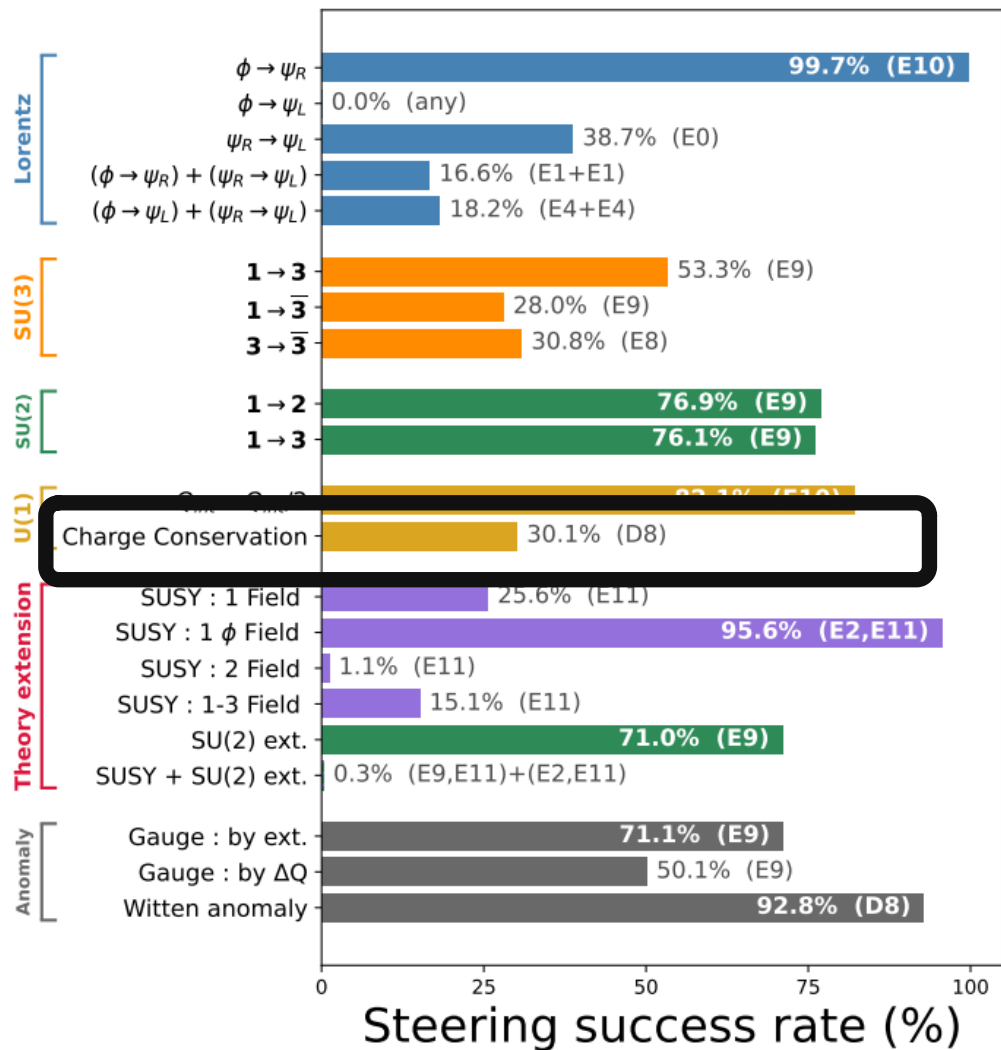
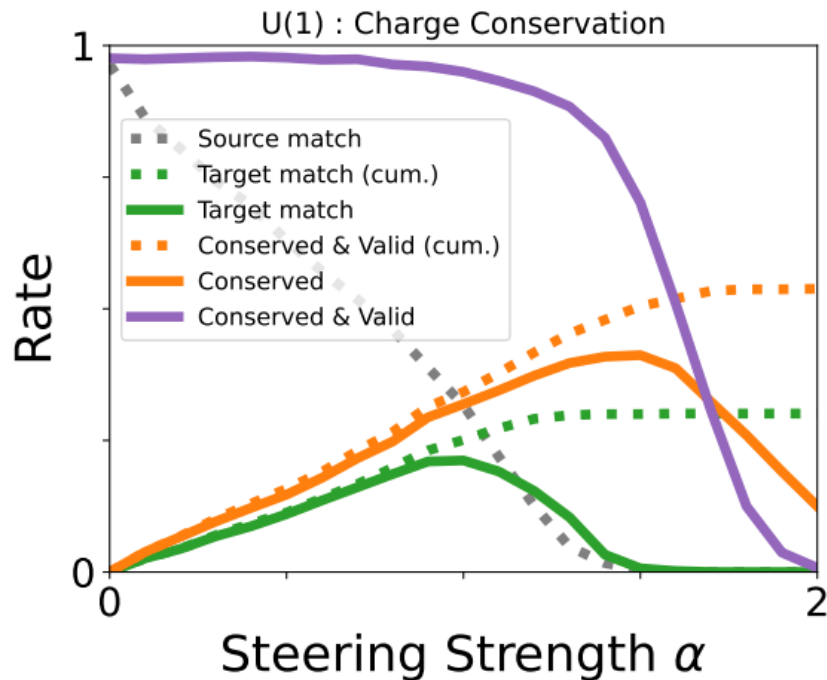
ψ_L was inaccessible by a single direction.
Additive steering reaches it,
but still leaks to ψ_R



Post Training Correction or Enforcing symmetries

$$\mathcal{L}_{\text{broken}} \rightarrow \mathcal{L}_{\text{conserving}}$$

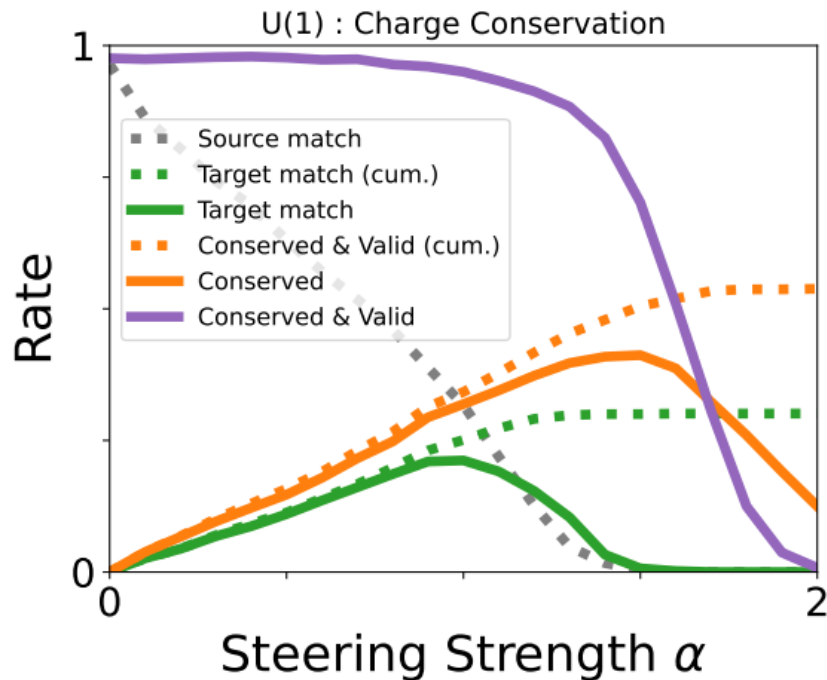
$$\Sigma Q \neq 0 \rightarrow \Sigma Q = 0$$



Post Training Correction or Enforcing symmetries

$$\mathcal{L}_{\text{broken}} \rightarrow \mathcal{L}_{\text{conserving}}$$

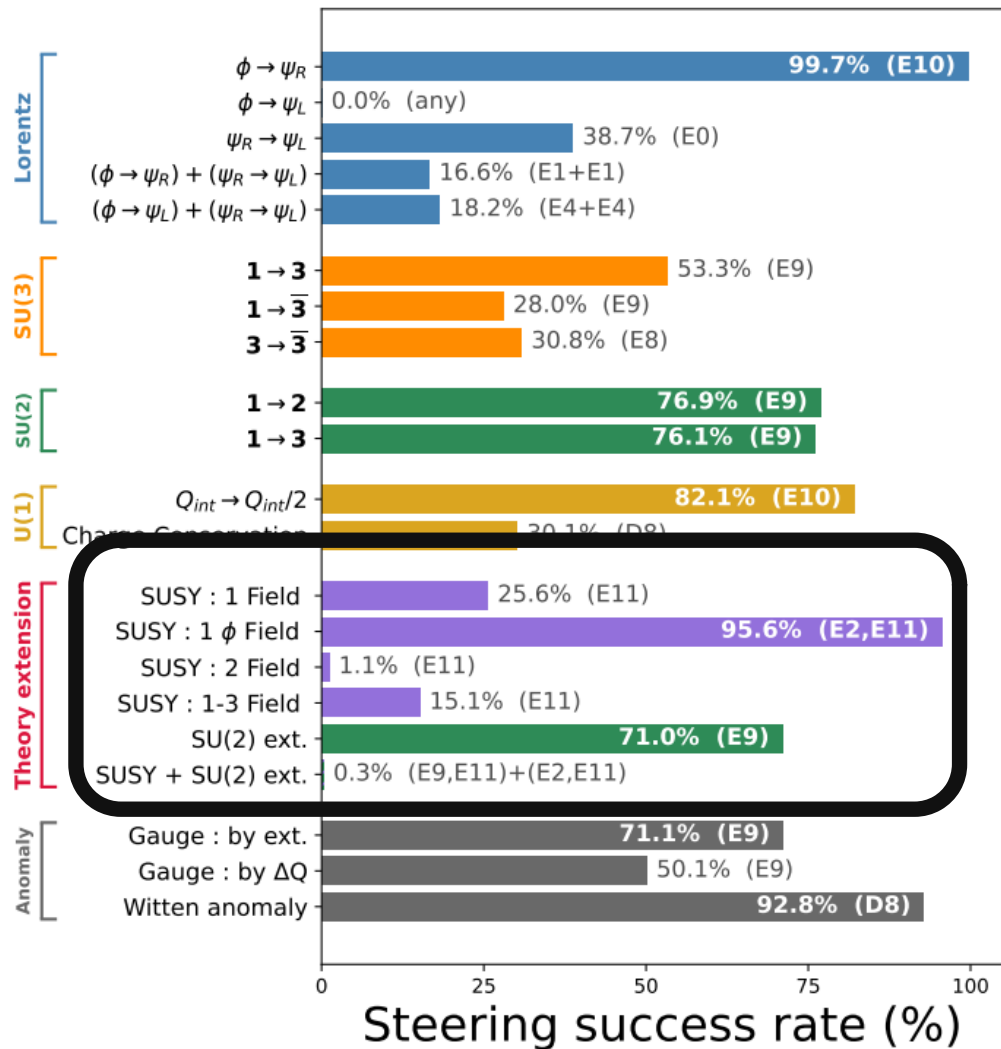
$$\Sigma Q \neq 0 \rightarrow \Sigma Q = 0$$



Changing Theory Domains Model Extensions:

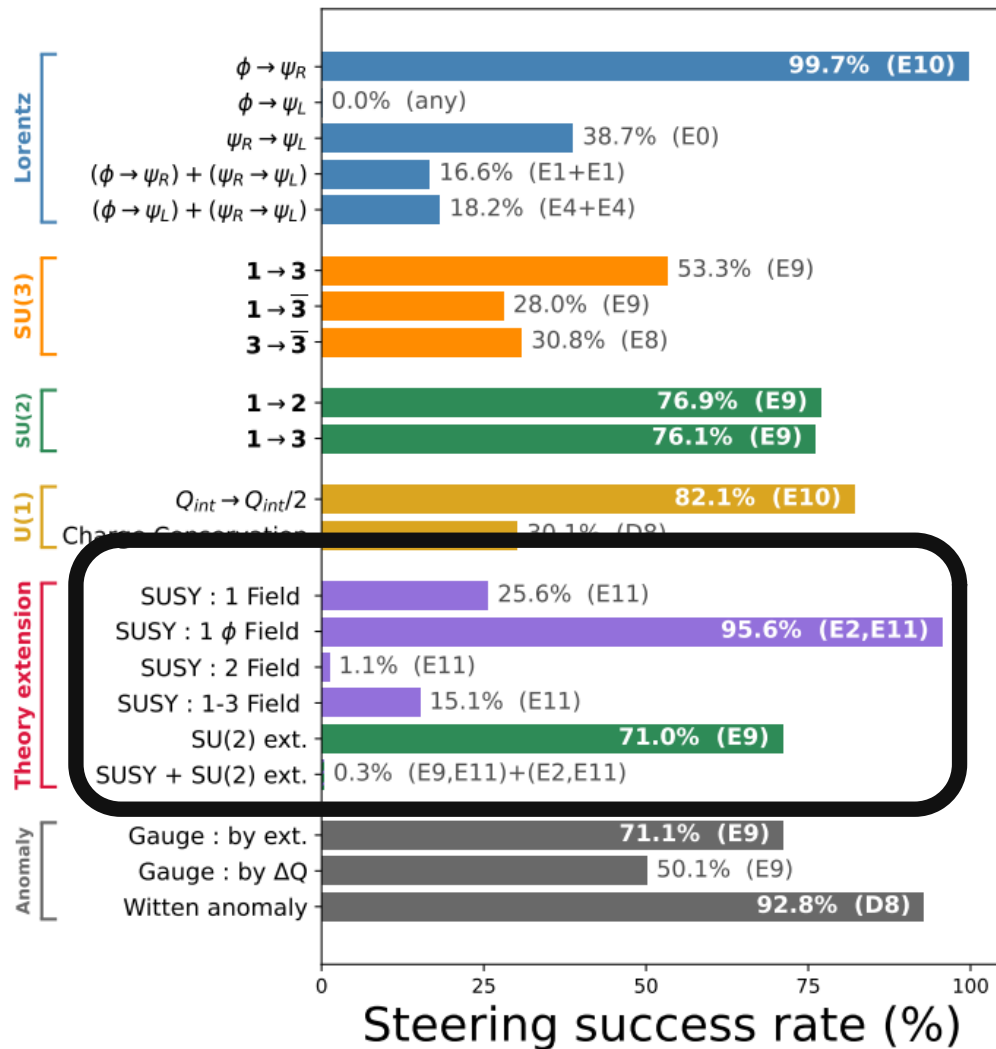
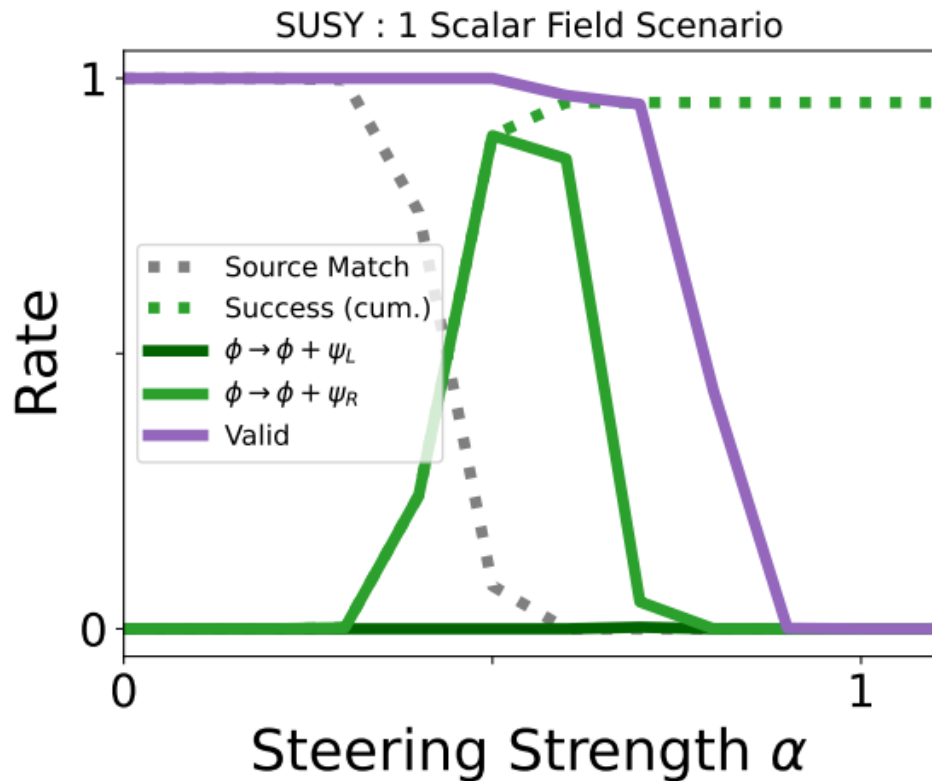
SUSY-like extension :
 $\mathcal{L}(F) \rightarrow \mathcal{L}(F, \tilde{F})$

SUSY-like : extending to superpartner of
 input fields. Did not include
 Gaugino Sector.



Changing Theory Domains

Model Extensions:



Changing Theory Domains Model Extensions:

SUSY-like extension :

$$\mathcal{L}(F) \rightarrow \mathcal{L}(F, \tilde{F})$$

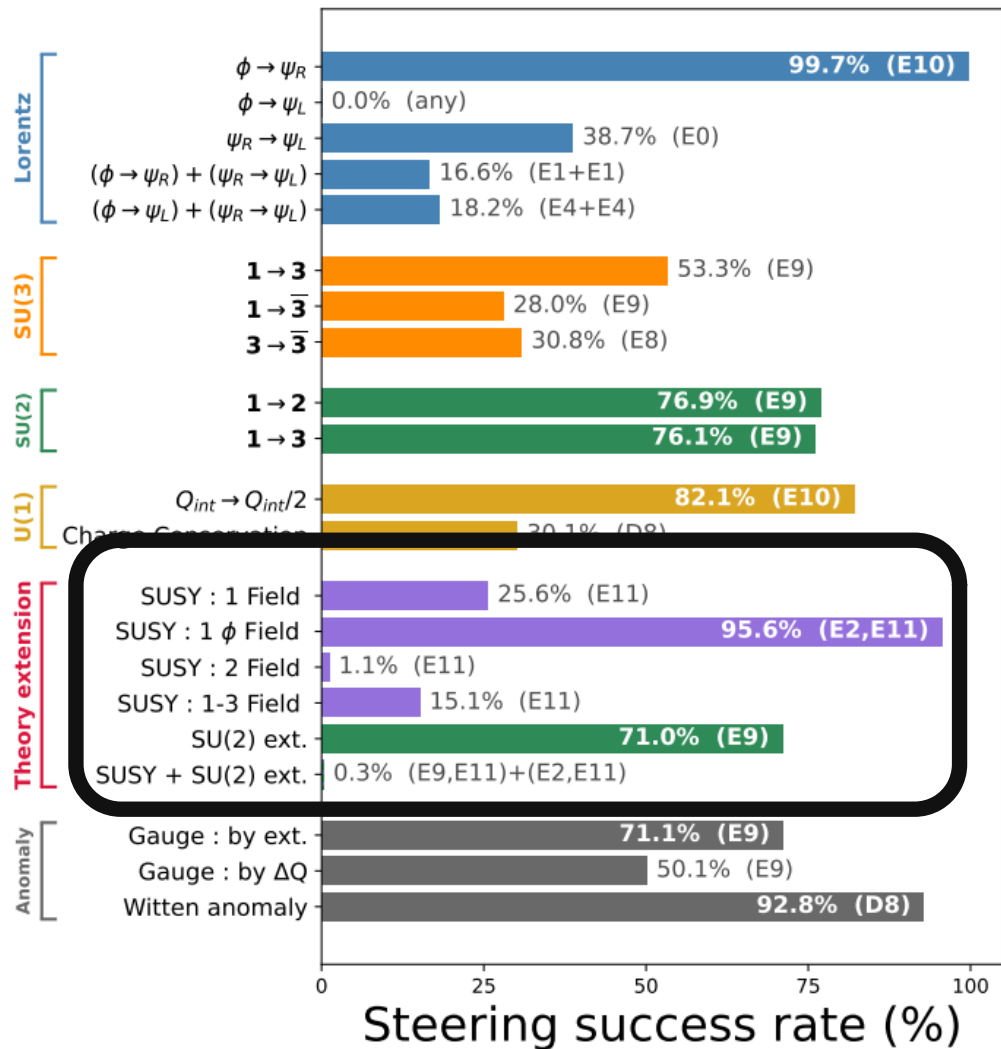
Extending Gauge Group:

$$\mathcal{L}_G \rightarrow \mathcal{L}_{G \times SU(2)}$$

Both

$$\mathcal{L}_G(F) \rightarrow \mathcal{L}_{G \times SU(2)}(F, \tilde{F})$$

SUSY-like : extending to superpartner of input fields. Did not include Gaugino Sector.



Locating Theory Solutions

$$\mathcal{L}_{\text{anom}} \rightarrow \mathcal{L}_{\text{anom.free}}$$

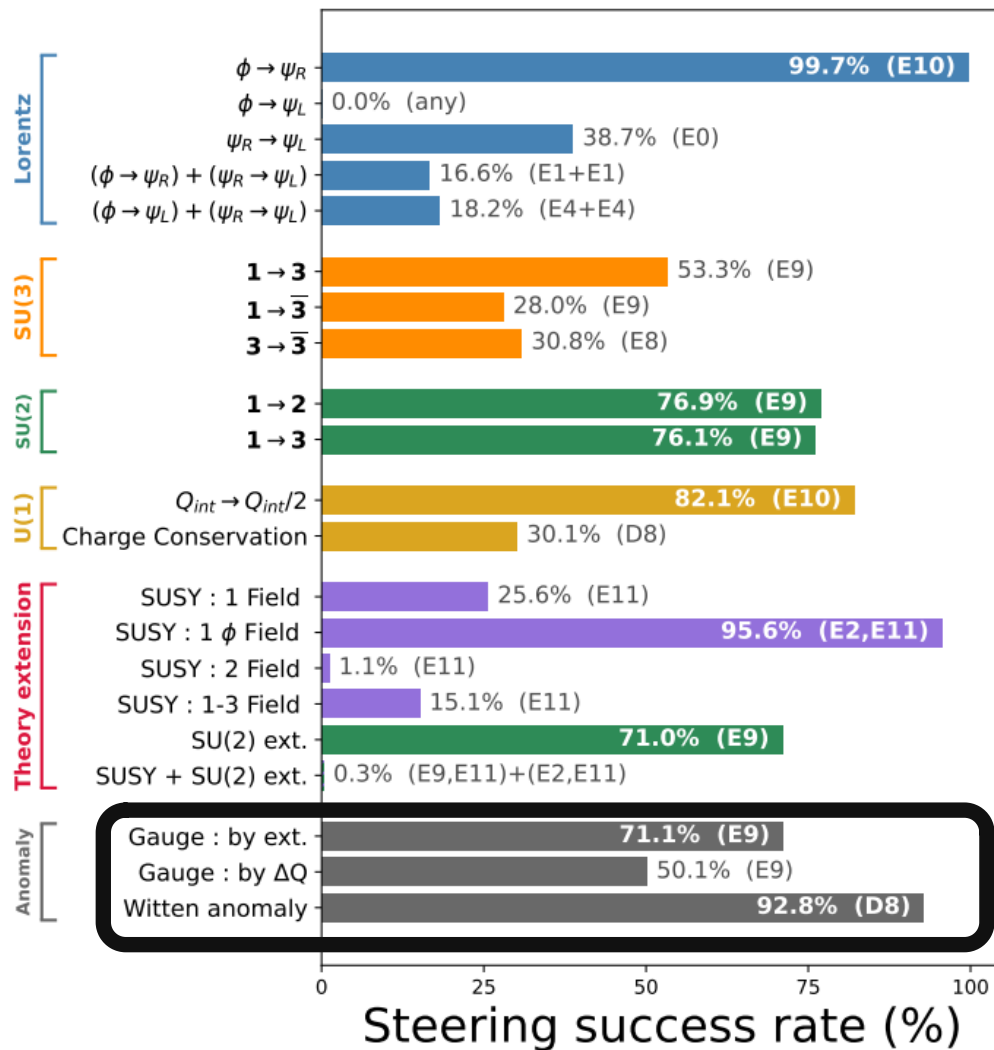
Gauge Anomalies

Anomaly Cancellation Conditions*:

$$A[L] = 0$$

Witten Anomaly

Number of SU2 doublet fermion
must be even



Locating Theory Solutions

$$\mathcal{L}_{\text{anom}} \rightarrow \mathcal{L}_{\text{anom.free}}$$

Gauge Anomalies

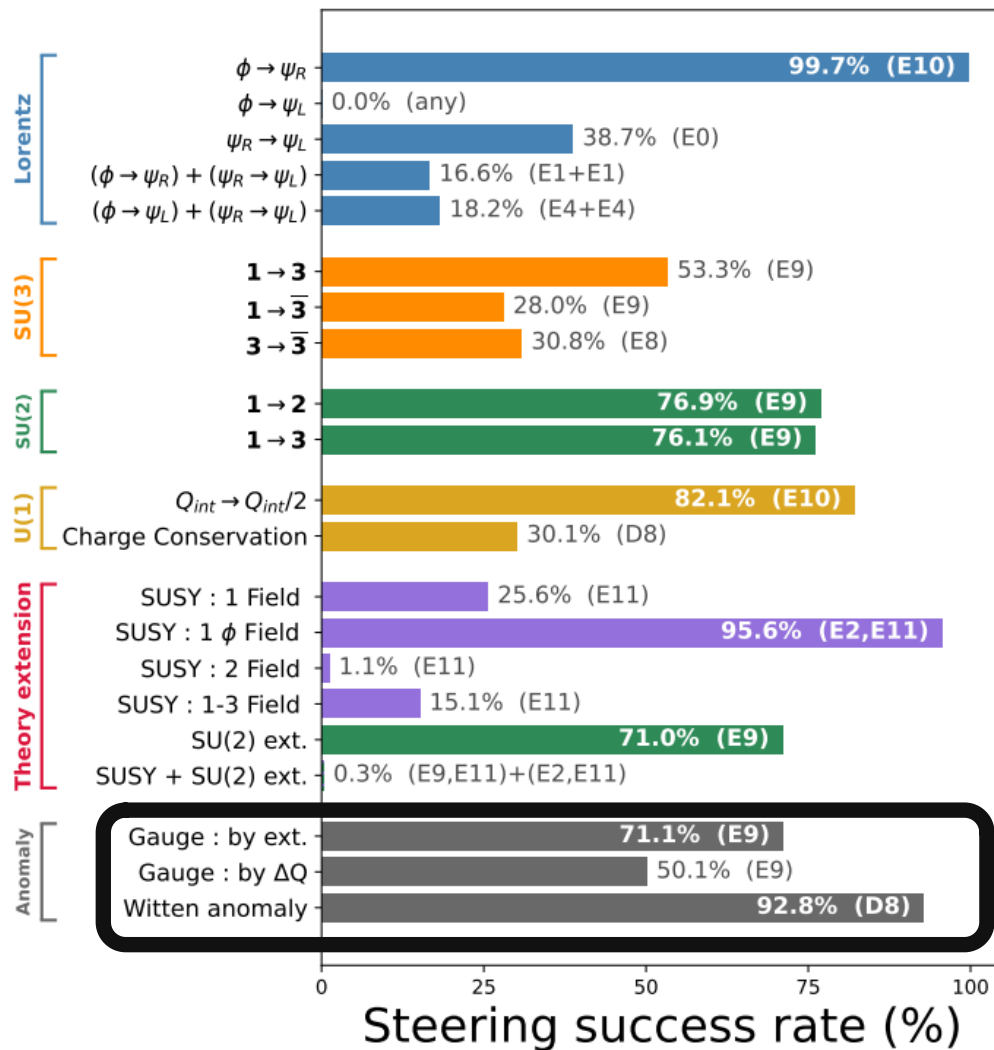
Anomaly Cancellation Conditions*:

$$A[L] = 0$$

Witten Anomaly

Number of SU2 doublet fermion
must be even

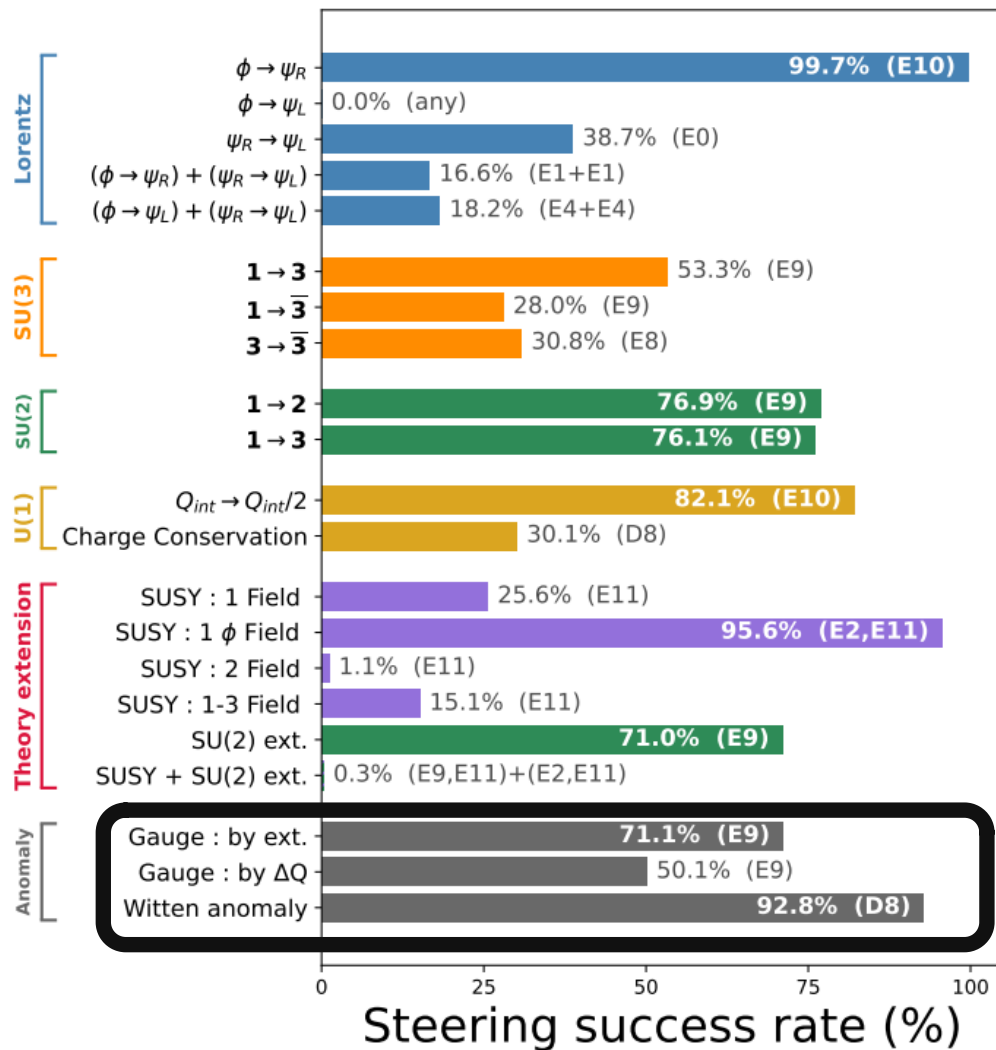
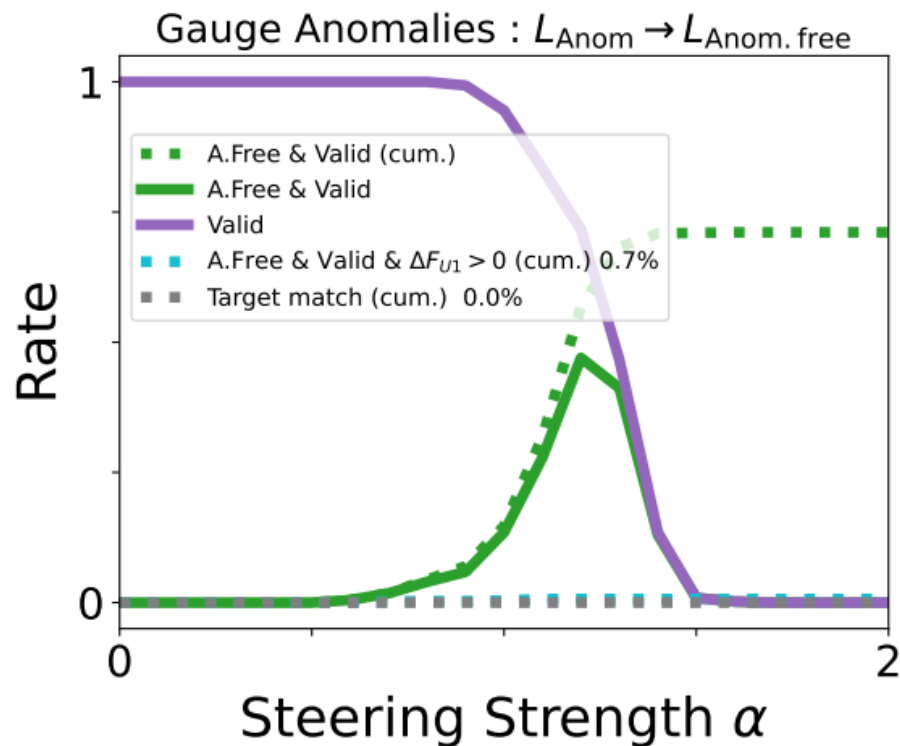
Success:
Anomaly Free and Valid lagrangians
(More than one way to solve anomalies)



Gauge Anomalies

Anomaly Cancellation Conditions*:

$$A[L] = 0$$

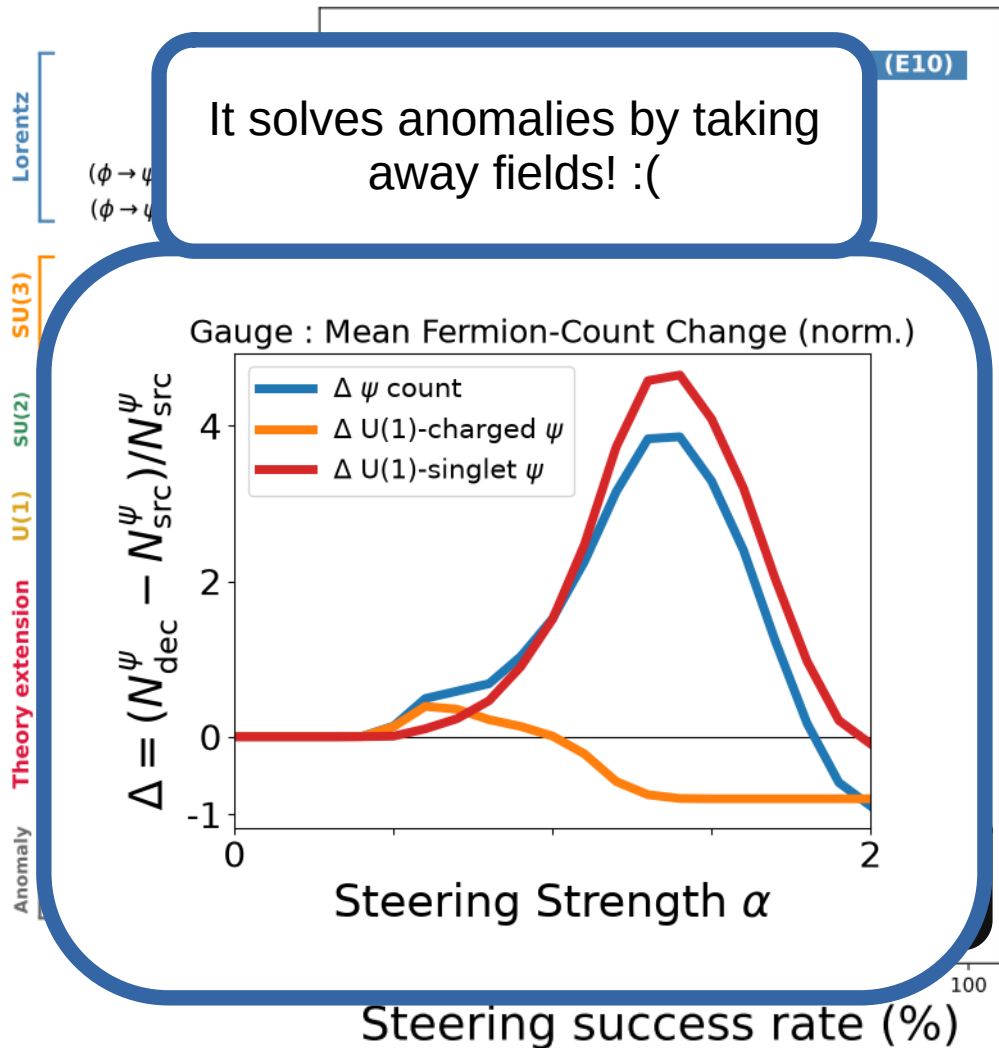
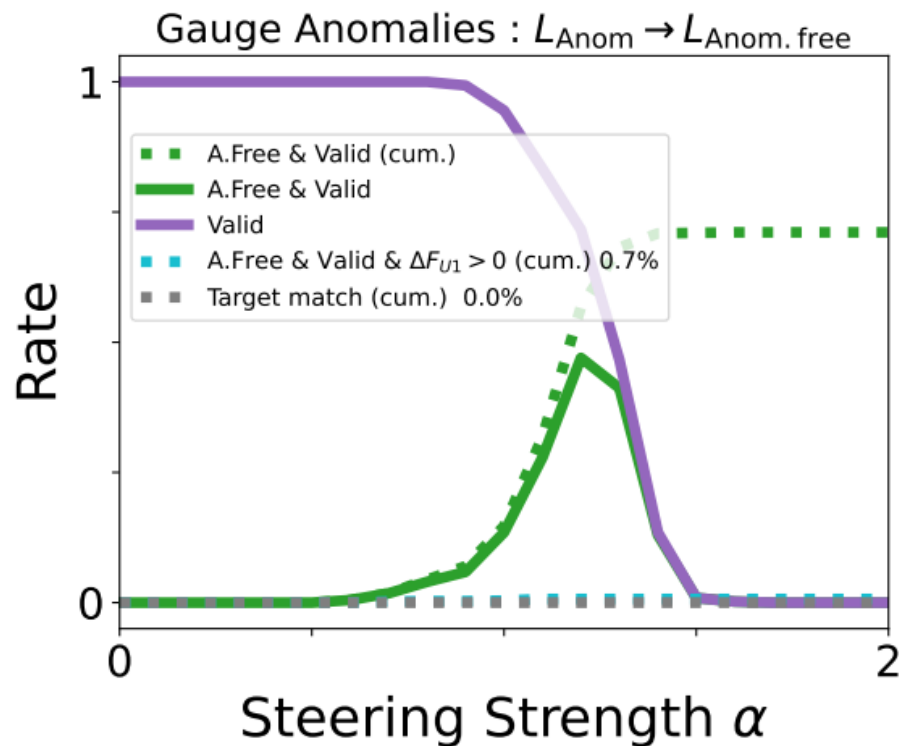


*Only working with trivial ACC solutions for now

Gauge Anomalies

Anomaly Cancellation Conditions*:

$$A[L] = 0$$

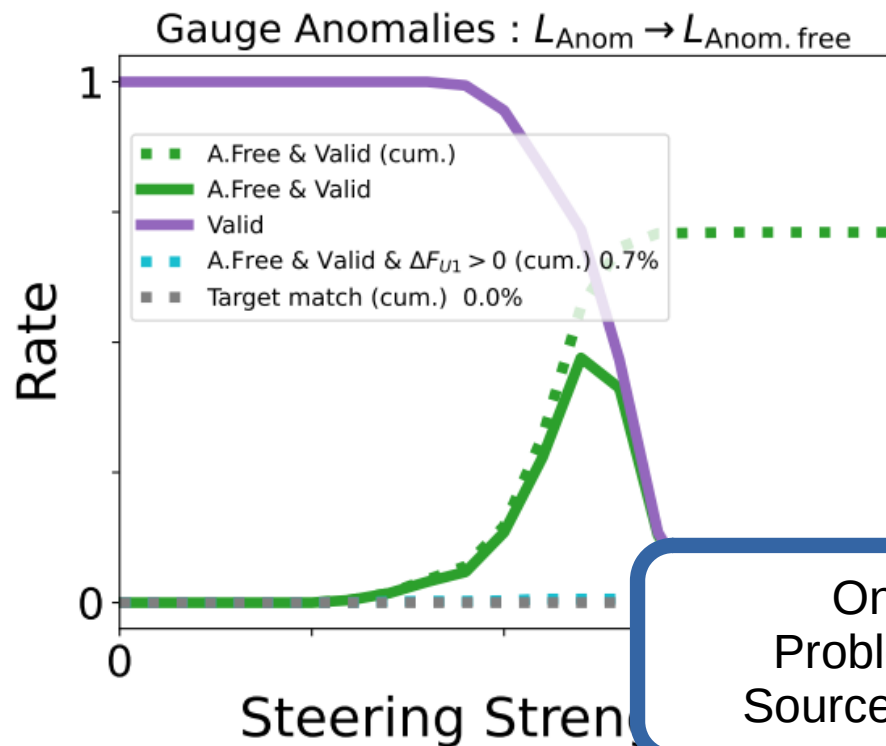


*Only working with trivial ACC solutions for now

Gauge Anomalies

Anomaly Cancellation Conditions*:

$$A[L] = 0$$



Lorentz

$(\phi \rightarrow \psi)$
 $(\phi \rightarrow \psi)$

(E10)

It solves anomalies by taking away fields! :(

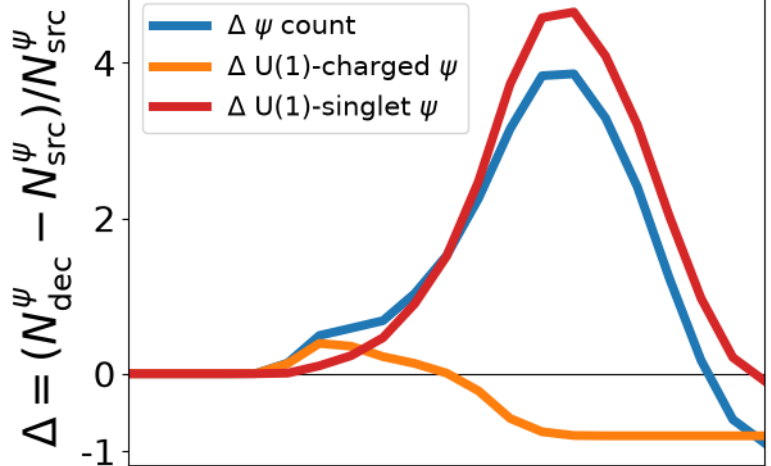
SU(3)

SU(2)

U(1)

Theory extension

Gauge : Mean Fermion-Count Change (norm.)



On going work!
Problem might be in
Source and target data!

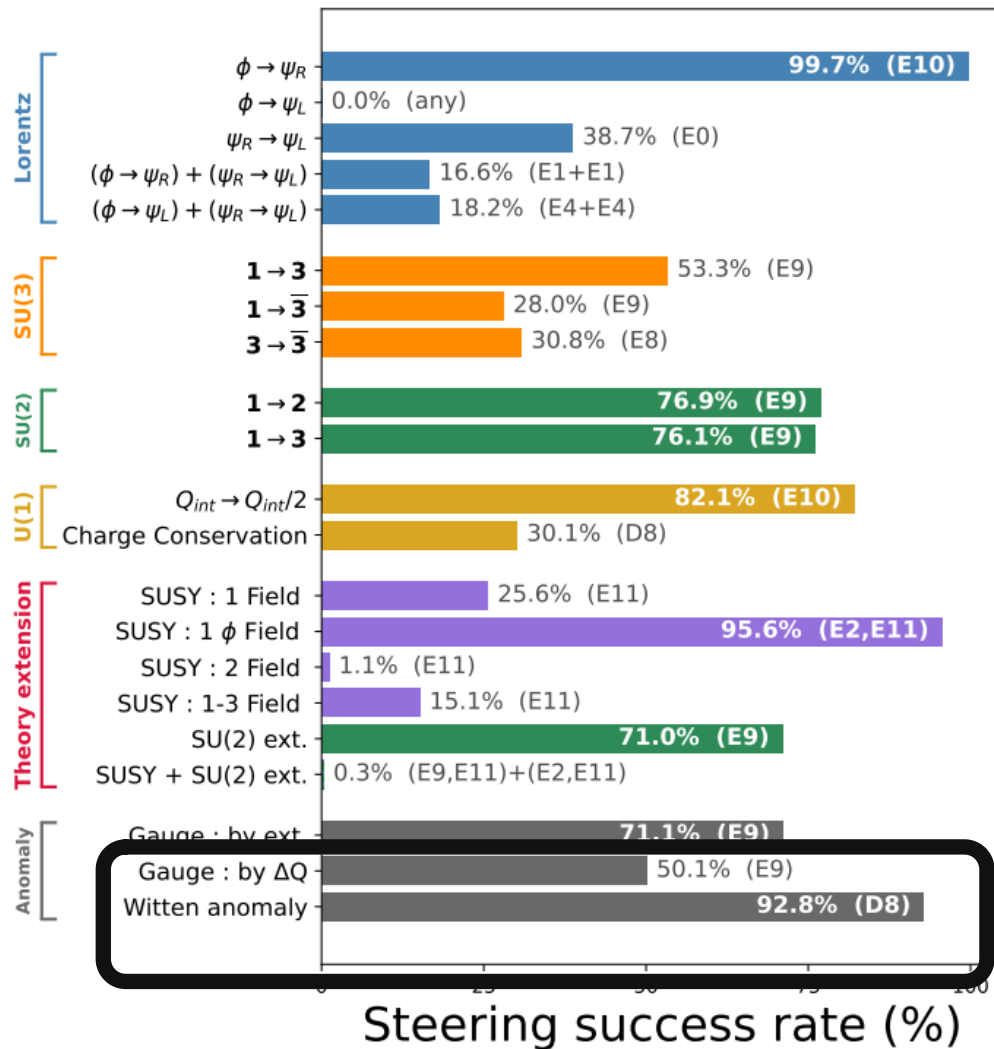
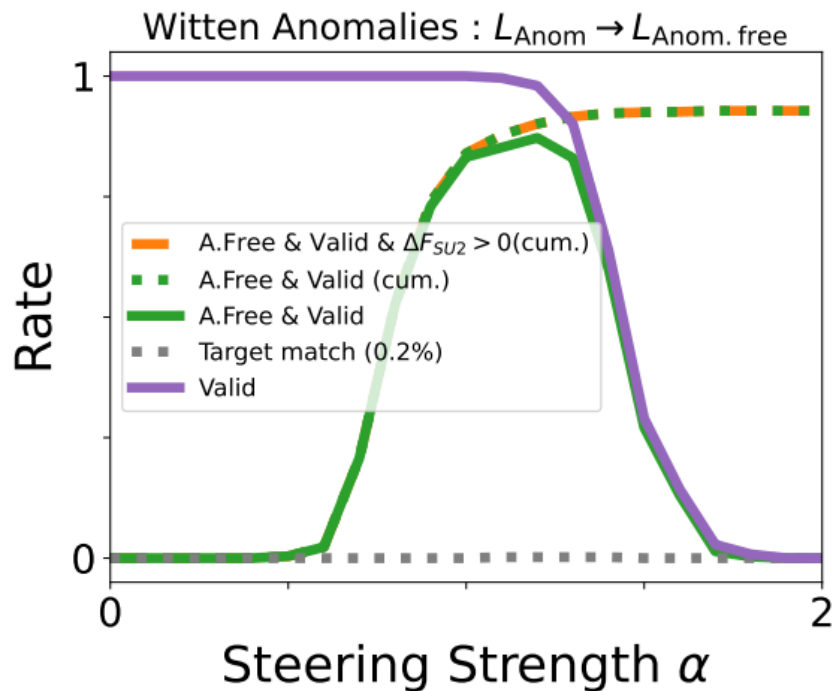
Steering Strength α

Steering success rate (%)

*Only working with trivial ACC solutions for now

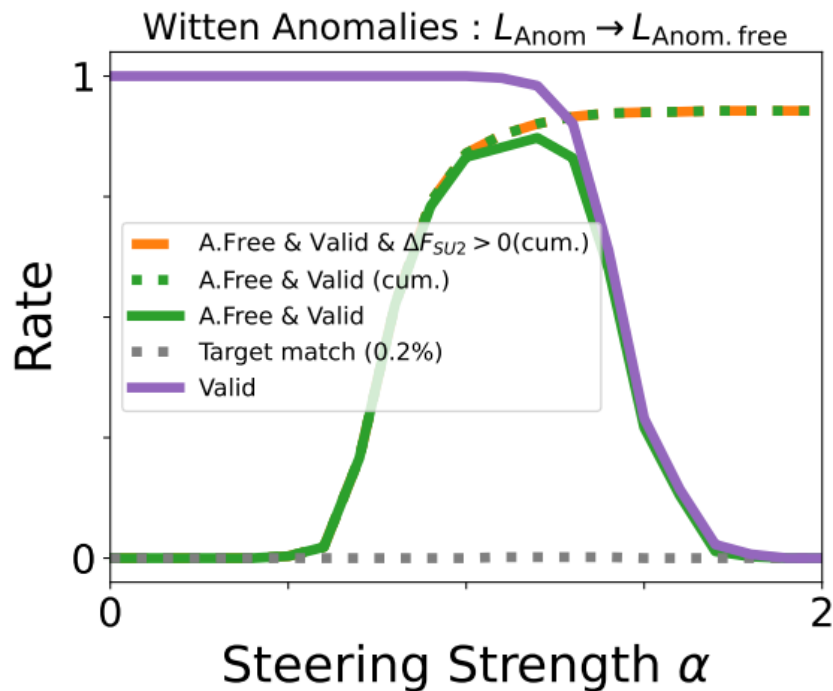
Witten Anomaly

Number of SU2 doublet fermion must be even

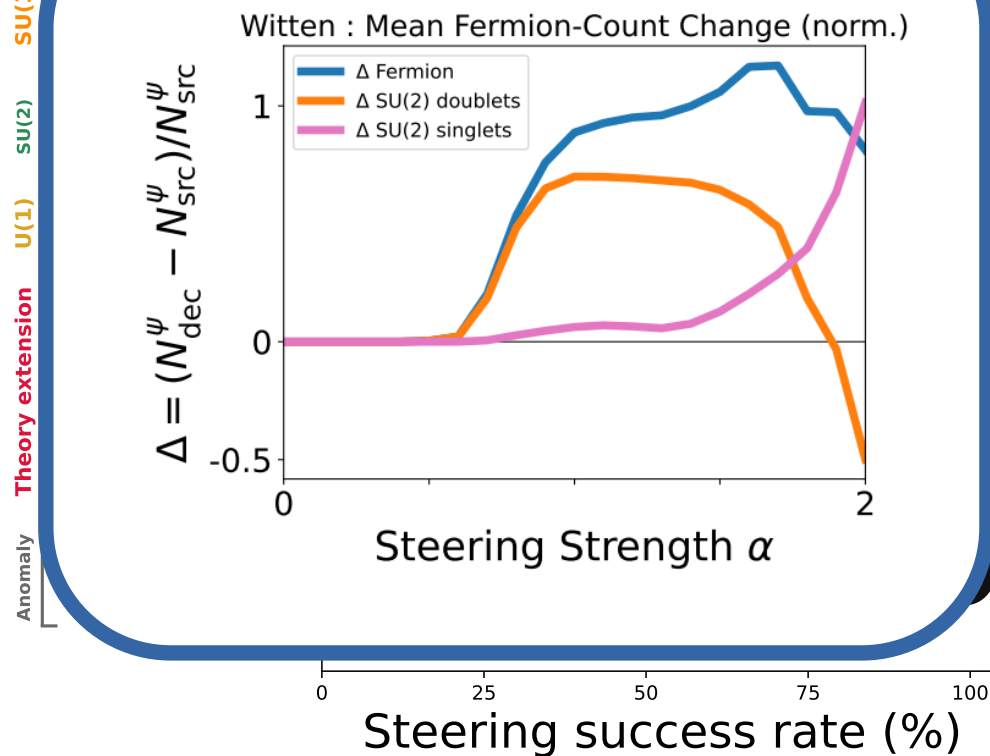


Witten Anomaly

Number of SU2 doublet fermion
must be even



This is GREAT!



Conclusion

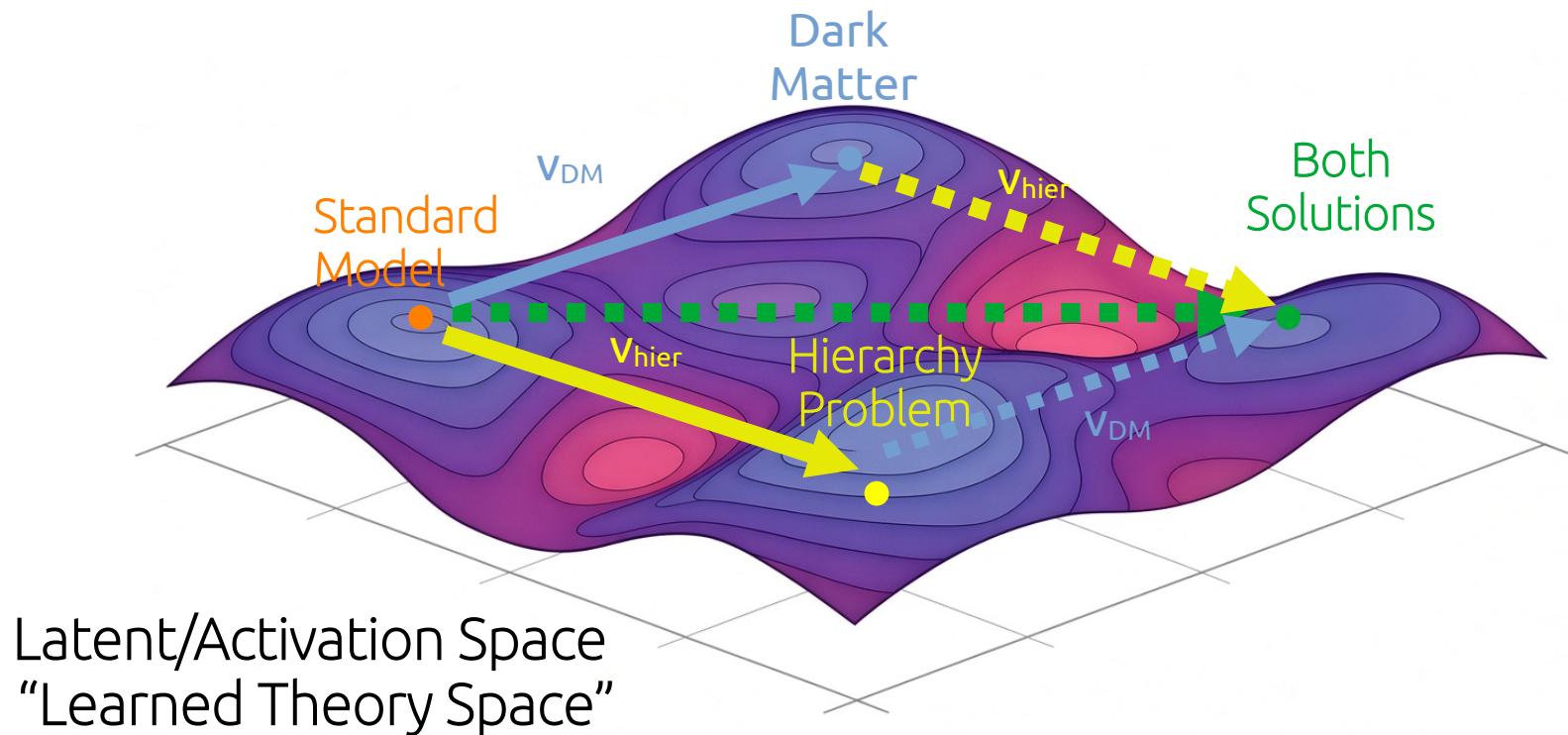
- BART-L writes Lagrangians!
- Latent Space ~ Learned Theory Space
- Steering
 - is cheap (Inference-time control)
 - Easier in Math than Language
 - works if model and data is good
 - Gap in Data ~ Gaps in Steerability
 - Can be additive
- Finally...



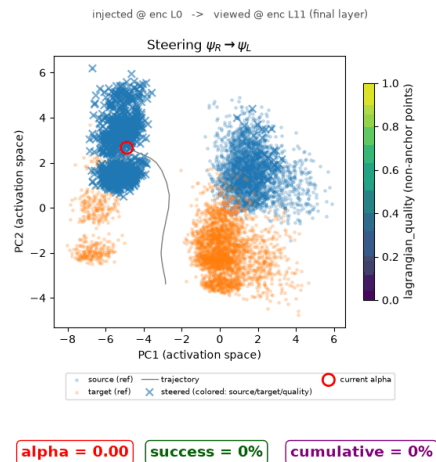
Outlook

Search For New Physics is now starting to be directional!

$$h_{\text{sm}} \rightarrow h_{\text{BSM}} = h_{\text{sm}} + V_{\text{DM}} + V_{\text{hier}} + V_{\text{v-mass}} + \dots$$



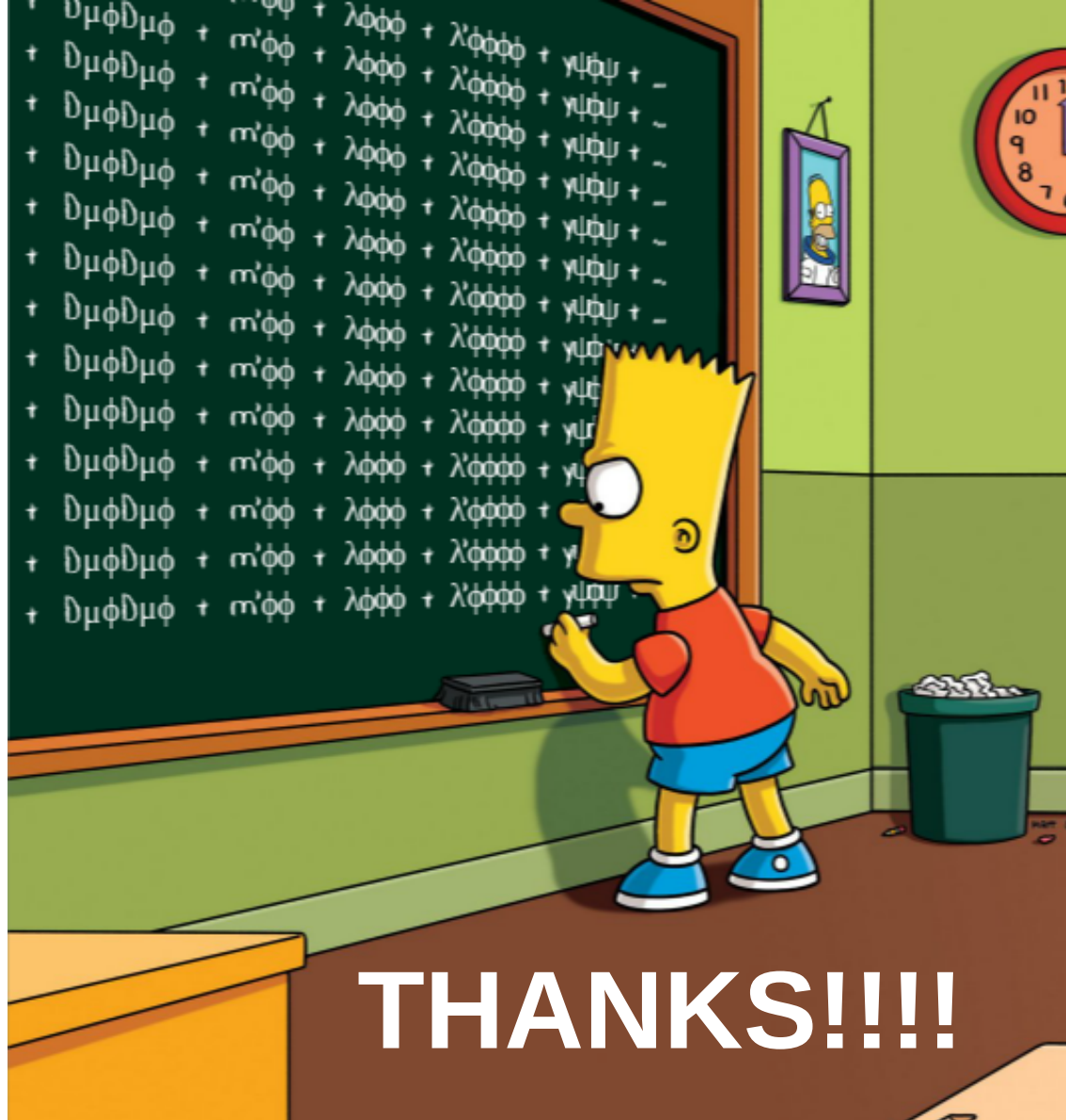
For more Steering GIFs:



Try BART-L yourself!



Generate some
Lagrangians!



BACKUP

How do we tokenize it?

Table 2: Example of object-tokenization.

Object	Tokens
$\phi(\bar{\mathbf{3}}, \mathbf{2}, -1/3)$	FIELD, SPIN, 0, SU3, -, 3, SU2, 2, U1, -, 1, /, 3, ID3
$\psi_L(\bar{\mathbf{1}}, \mathbf{3}, 0)$	FIELD, SPIN, 1, /, 2, SU2, 3, HEL, 1, /, 2, ID9
$\phi^\dagger(\mathbf{1}, \mathbf{1}, 7/5)$	FIELD, SPIN, 0, U1, 7, /, 5, DAGGER, ID7
$D_\mu^{(\text{SU}(2), \text{U}(1))}$	DERIVATIVE, SU2, U1, ID4
$D_\mu^{(\text{SU}(3), \text{SU}(2), \text{U}(1))}$	DERIVATIVE, SU3, SU2, U1, ID1
∂_μ	DERIVATIVE, ID5
$\bar{\sigma}^\mu$	SIGMA_BAR, ID2
$g^{\mu\nu}$	LORENTZ, ID6, ID4
ϵ^{ab}	SU2, ID0, ID7

akin to char-level tokenization rather than word-level tokenization

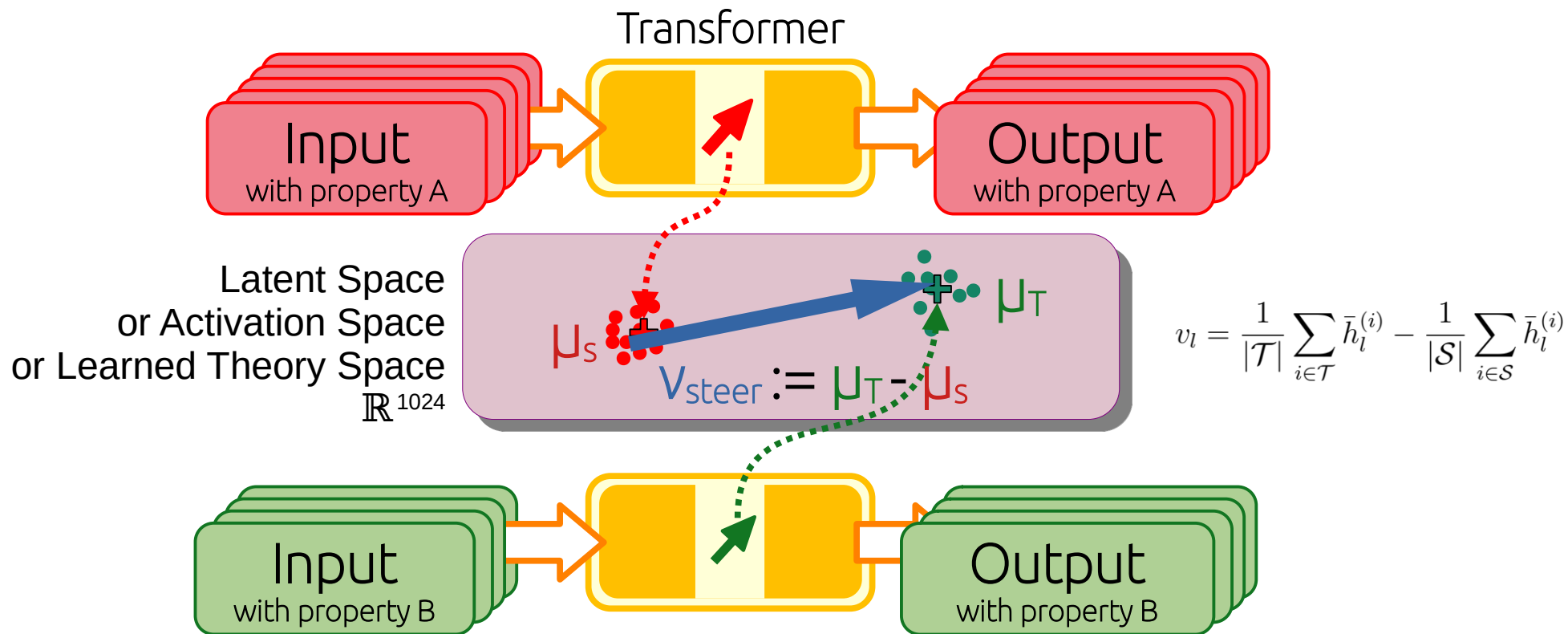
How do we tokenize it?

Example: Quartic Term

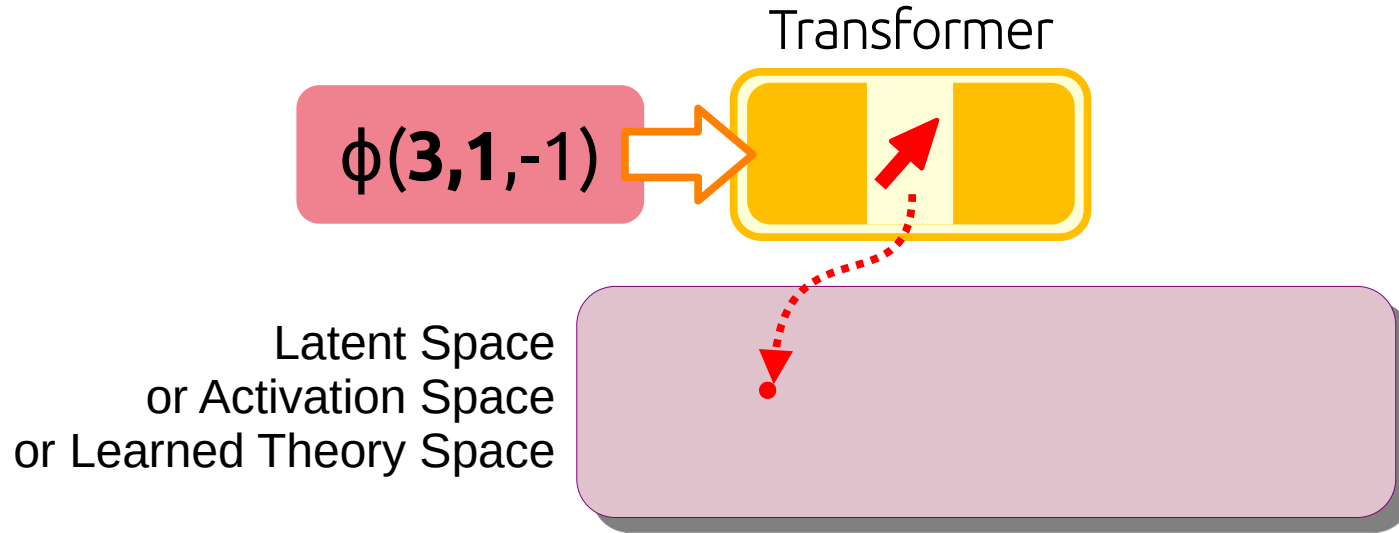
$$\text{Term} : \lambda \phi \phi^\dagger \Phi \Phi^\dagger$$

Information	Tokens
λ	
$\phi(\mathbf{3}, \mathbf{2}, 6)$	FIELD, SPIN, 0, SU3, 3, SU2, 3, U1, 6, ID5
$\phi^\dagger(\bar{\mathbf{3}}, \mathbf{2}, -6)$	FIELD, SPIN, 0, SU3, -, 3, SU2, 3, U1, -, 6, DAGGER, ID7
$\Phi(\bar{\mathbf{3}}, \mathbf{1}, -7/2)$	FIELD, SPIN, 0, SU3, -, 3, U1, -, 7, /, 2, ID8
$\Phi^\dagger(\mathbf{3}, \mathbf{1}, 7/2)$	FIELD, SPIN, 0, SU3, 3, U1, 7, /, 2, DAGGER, ID2
CONTRACTIONS	
$\epsilon^{a_1 b_2 c_2} \phi_{a_1} \phi_{b_1 b_2}^\dagger \Phi_{c_1 c_2} \Phi_{d_1}^\dagger$	SU3, ID5, ID7, ID8
$\epsilon^{b_1 c_1 d_1} \phi_{a_1} \phi_{b_1 b_2}^\dagger \Phi_{c_1 c_2} \Phi_{d_1}^\dagger$	SU3, ID7, ID8, ID2
$\epsilon^{a_1 b_1} \phi_{a_1 a_2} \phi_{b_1 b_2}^\dagger \Phi \Phi^\dagger$	SU2, ID5, ID7
$\epsilon^{a_2 b_2} \phi_{a_1 a_2} \phi_{b_1 b_2}^\dagger \Phi \Phi^\dagger$	SU2, ID5, ID7

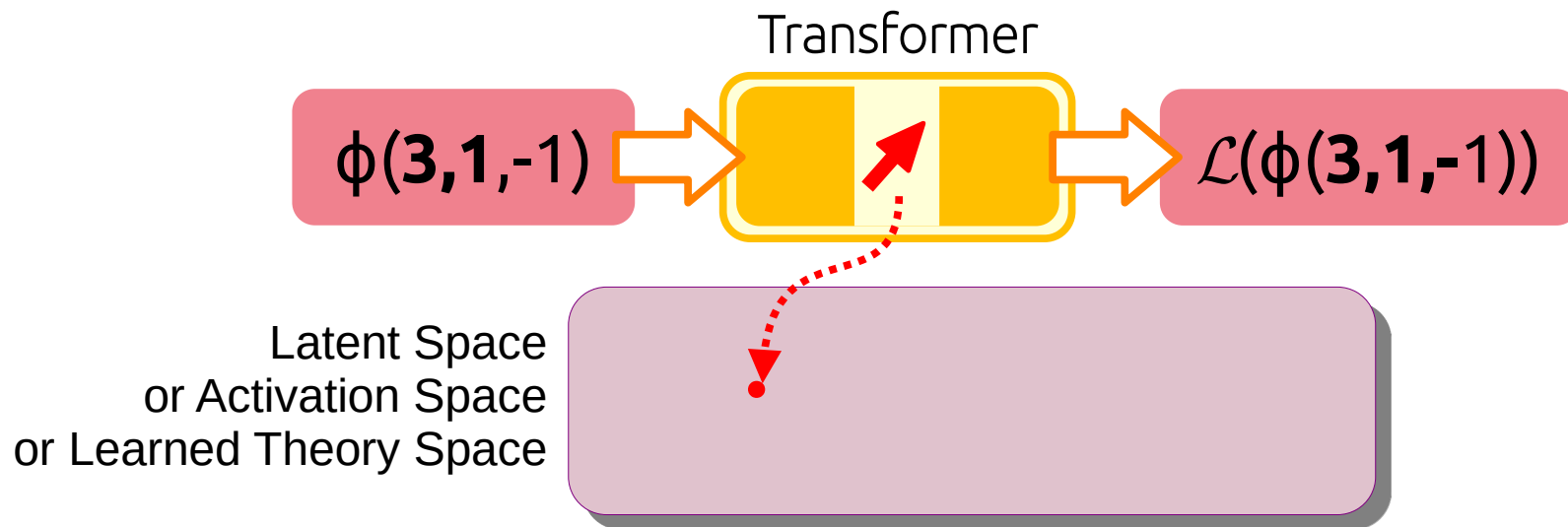
Contrastive Activation Addition (CAA)



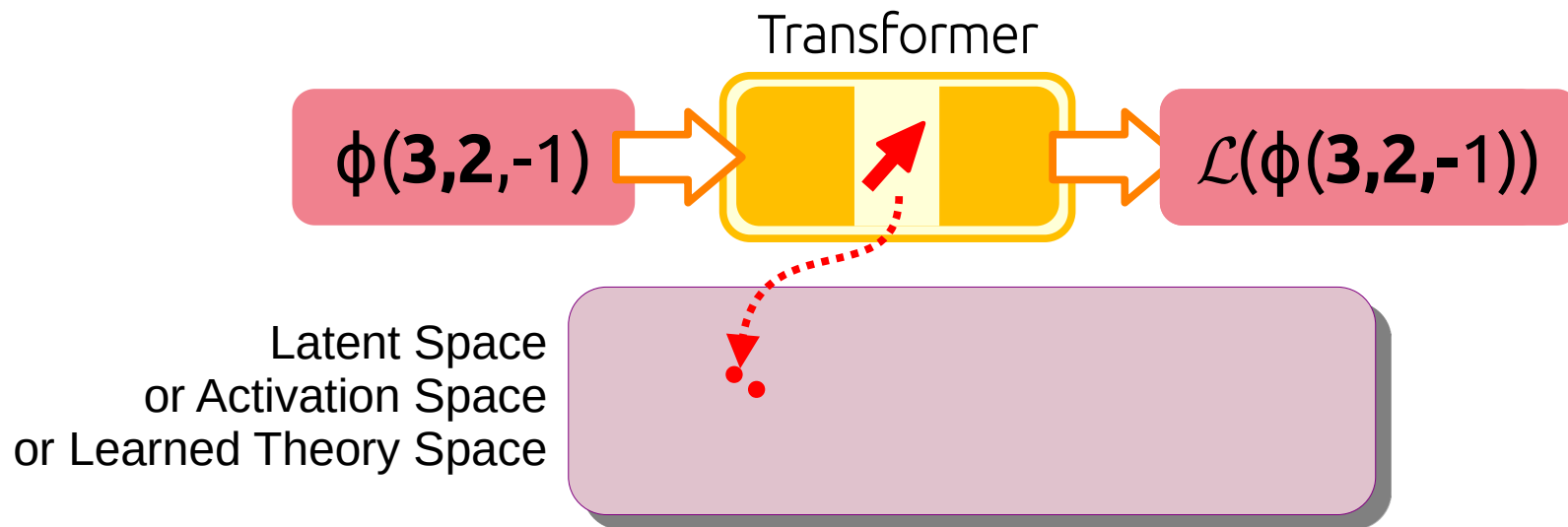
Contrastive Activation Addition (CAA)



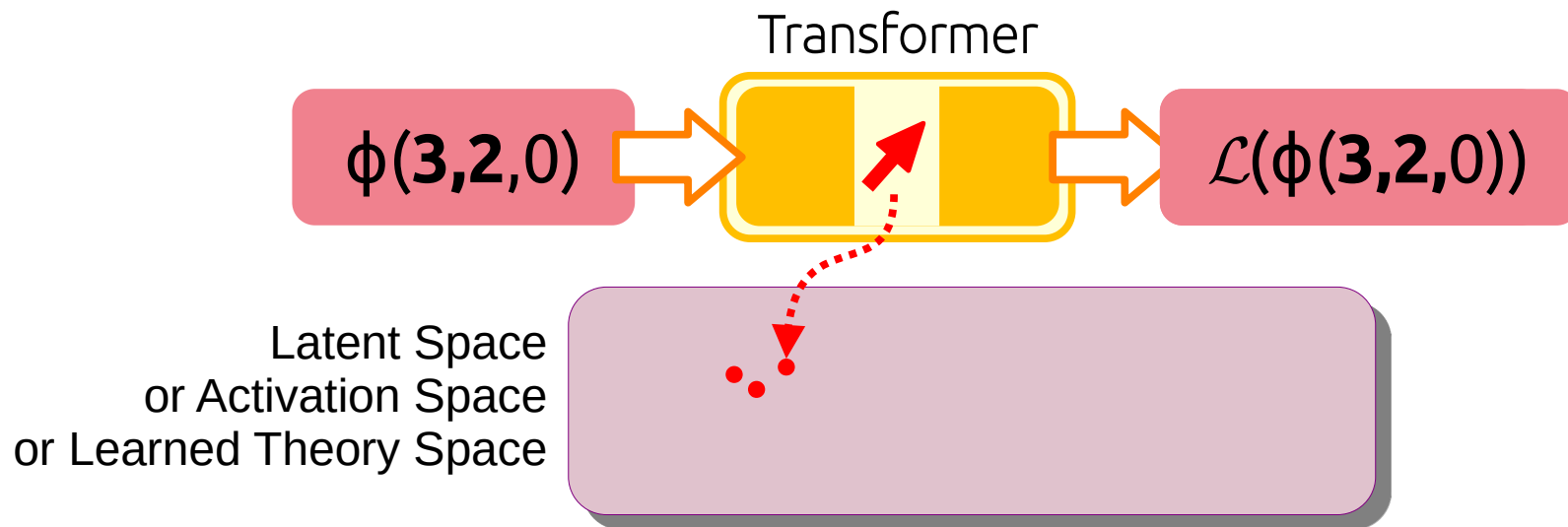
Contrastive Activation Addition (CAA)



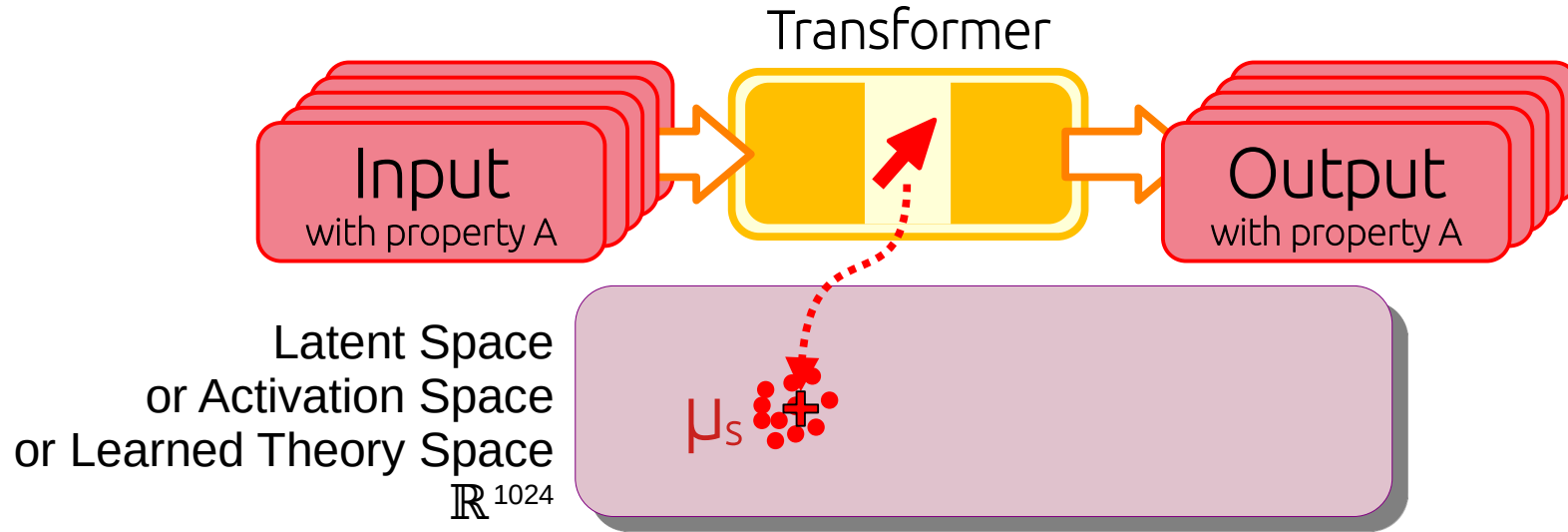
Contrastive Activation Addition (CAA)



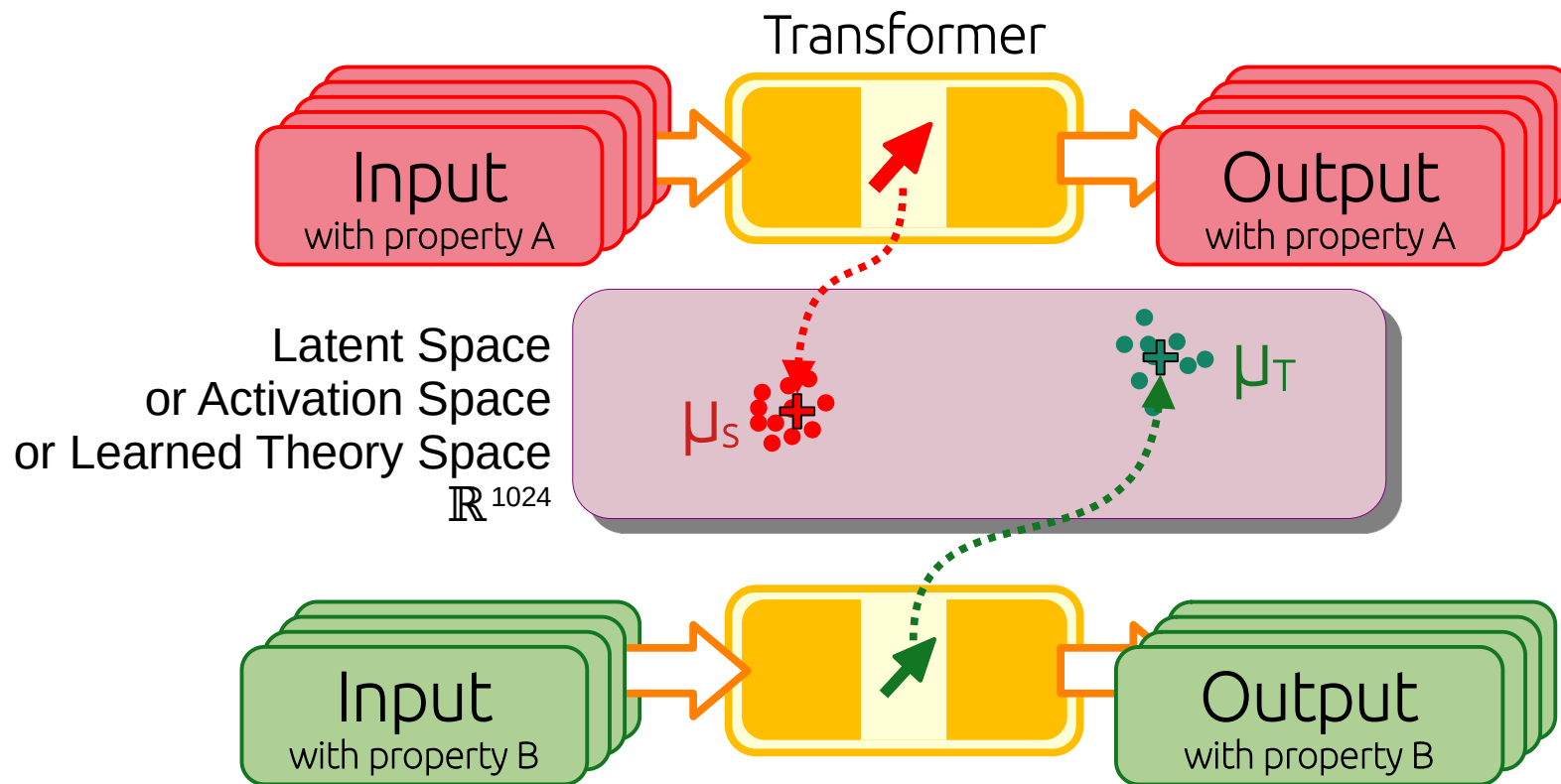
Contrastive Activation Addition (CAA)



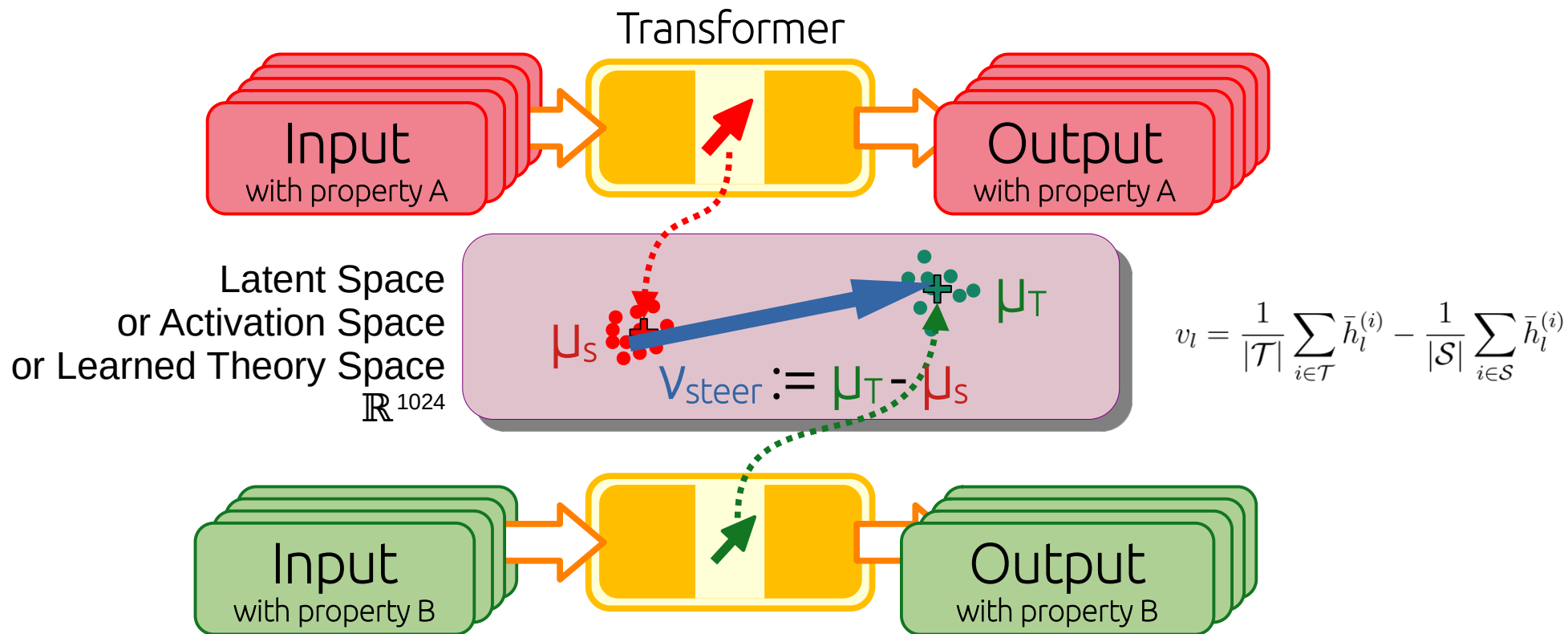
Contrastive Activation Addition (CAA)



Contrastive Activation Addition (CAA)



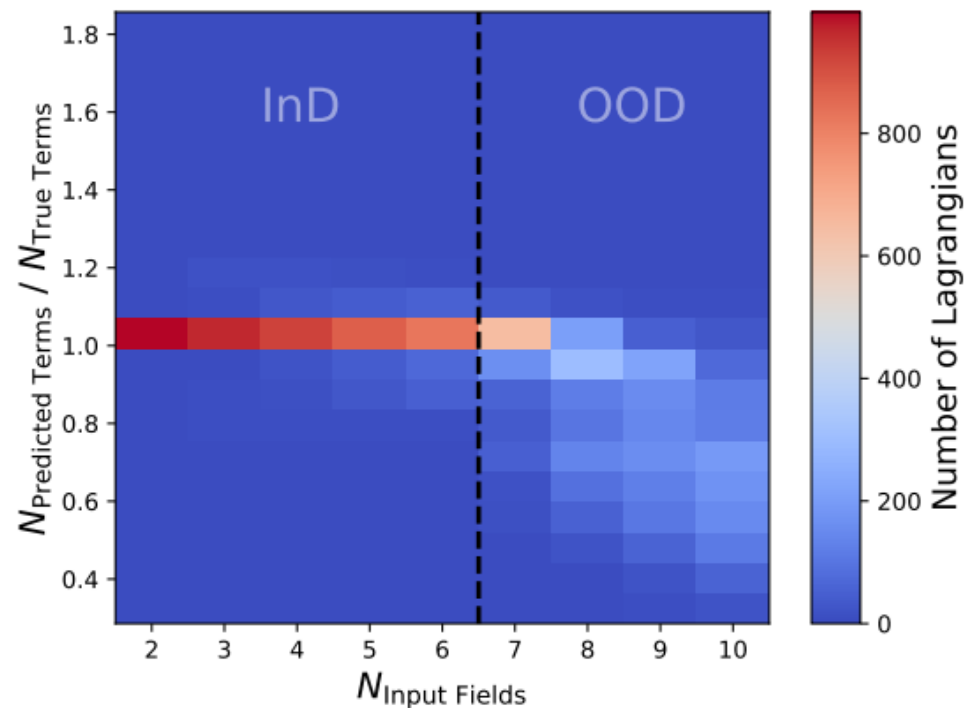
Contrastive Activation Addition (CAA)



OOD : N_{Fields} Generalization (or “Length Generalization”)

How do we check?

In Distribution (Train): ≤ 6 Fields
Out-of-Distribution (Test): 7-10 Fields



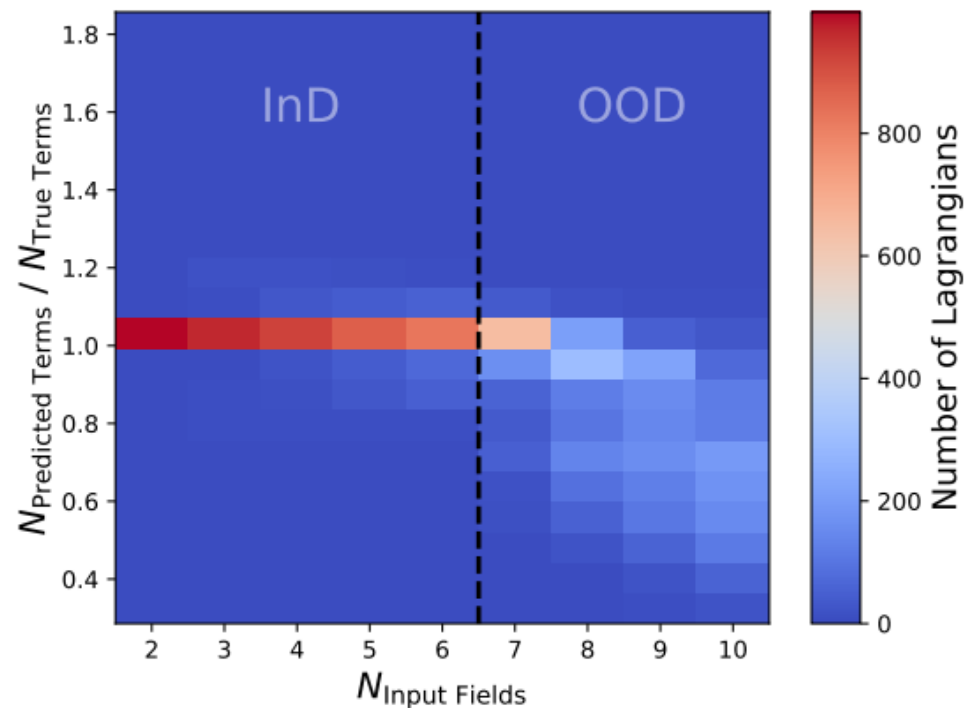
OOD : N_{Fields} Generalization (or “Length Generalization”)

How do we check?

In Distribution (Train): ≤ 6 Fields
Out-of-Distribution (Test): 7-10 Fields

At OOD:

1. Its missing terms



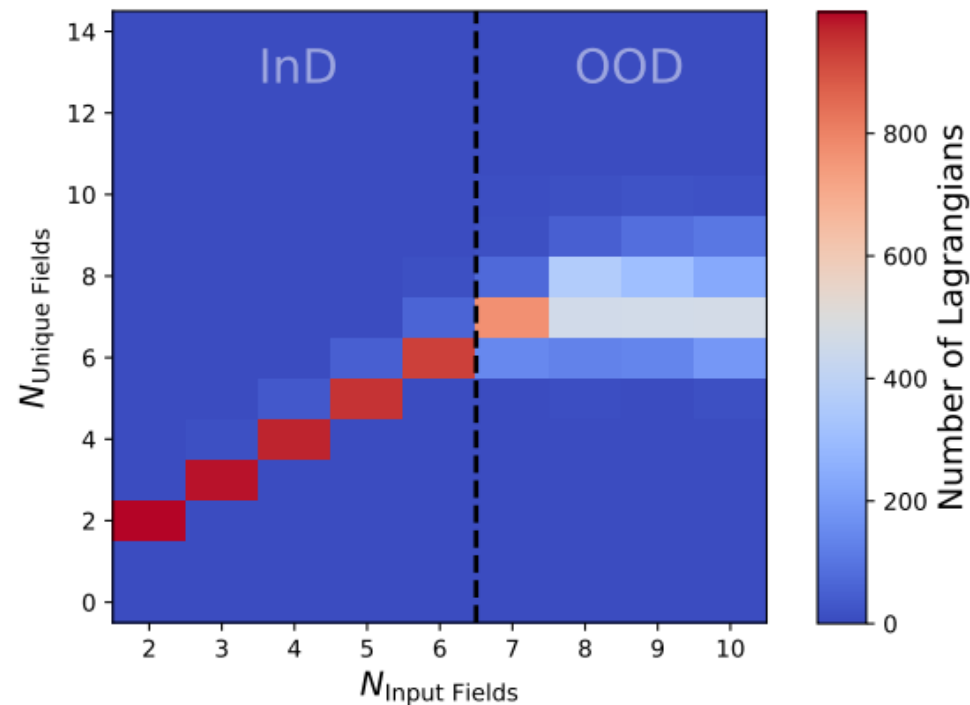
OOD : N_{Fields} Generalization (or “Length Generalization”)

How do we check?

In Distribution (Train): ≤ 6 Fields
Out-of-Distribution (Test): 7-10 Fields

At OOD:

1. Its missing terms.



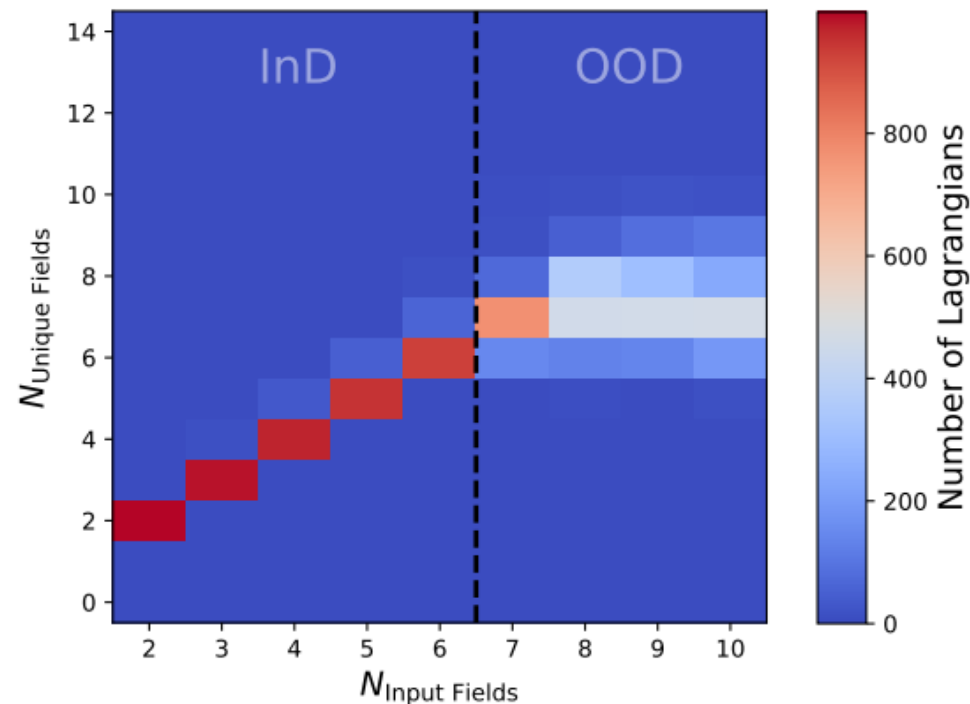
OOD : N_{Fields} Generalization (or “Length Generalization”)

How do we check?

In Distribution (Train): ≤ 6 Fields
Out-of-Distribution (Test): 7-10 Fields

At OOD:

1. Its missing terms.
2. It knows there should be more, but dont know how many



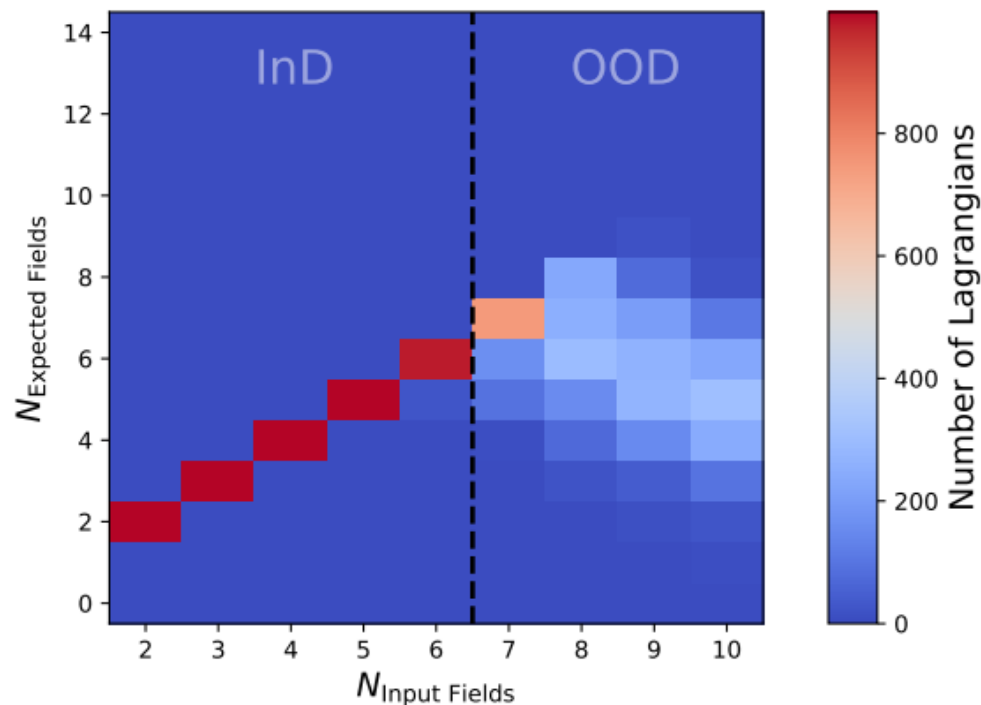
OOD : N_{Fields} Generalization (or “Length Generalization”)

How do we check?

In Distribution (Train): ≤ 6 Fields
Out-of-Distribution (Test): 7-10 Fields

At OOD:

1. Its missing terms.
2. It knows there should be more but dont know how many



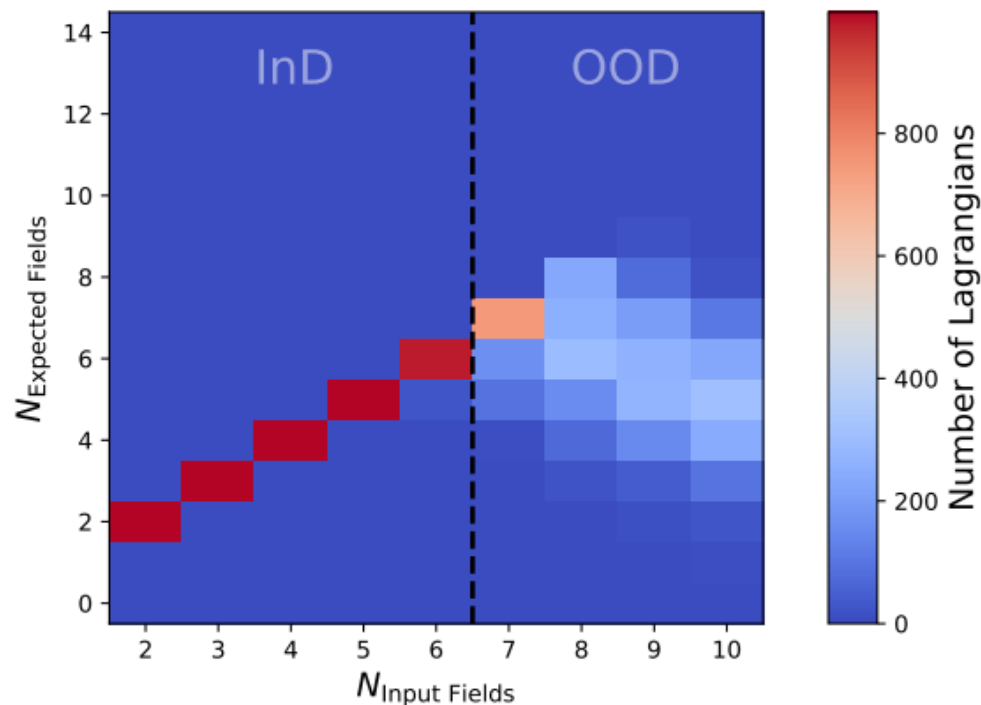
OOD : N_{Fields} Generalization (or “Length Generalization”)

How do we check?

In Distribution (Train): ≤ 6 Fields
Out-of-Distribution (Test): 7-10 Fields

At OOD:

1. Its missing terms.
2. It knows there should be more but dont know how many
3. It jumbles up the inputs when there's too much inputs



It can't count!

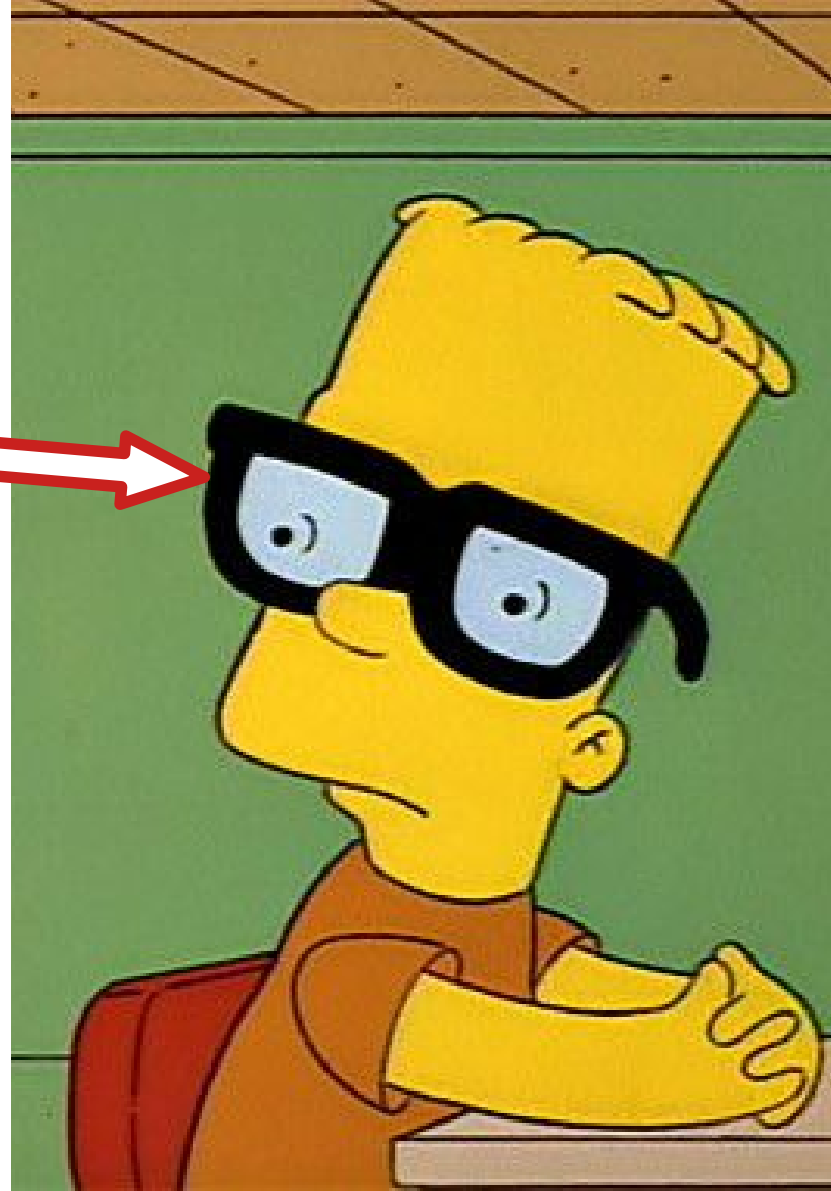
Non-Causal Encoder
with Abs. Positional Encoding
is bad for contextual counting

[2406.02585]

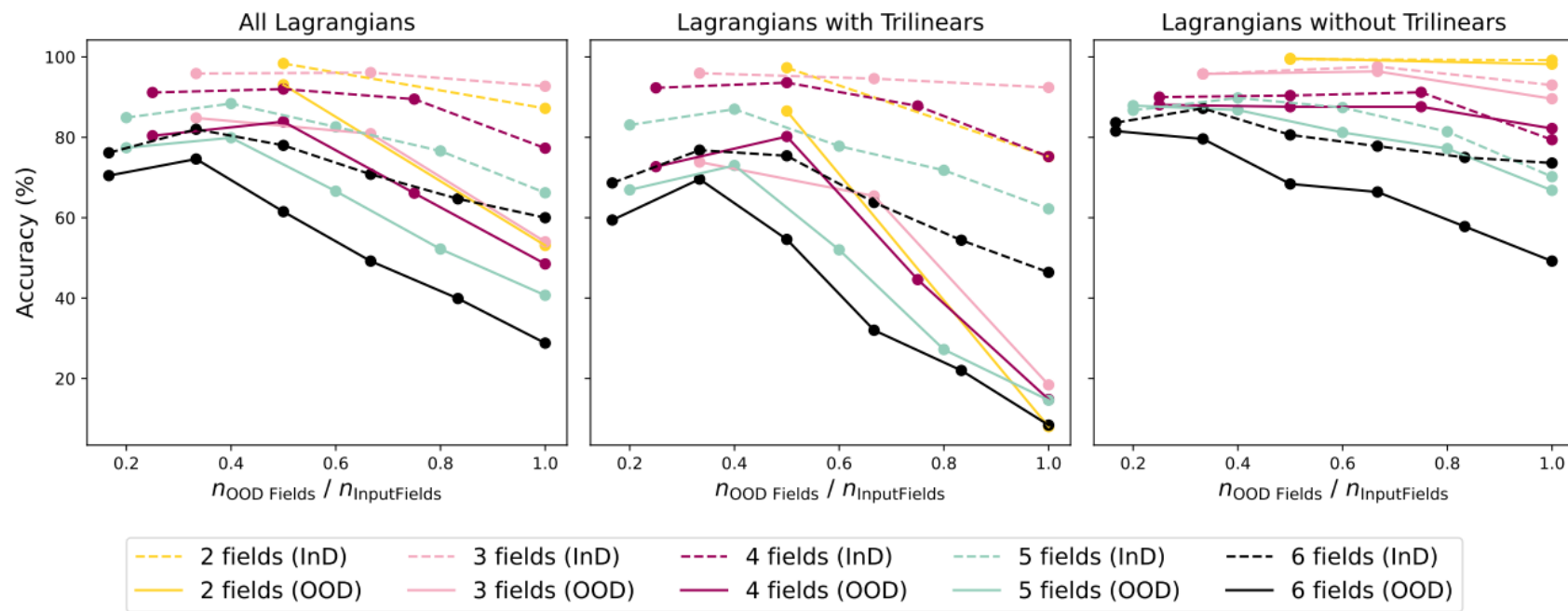


The Famous
"How many 'r' in the word 'strawberry'?"
Problem

Ask ChatGPT to count the number of 1s in a random string of 1s and 0s. No python code.
It still fails sometime.



We also did OOD U(1) Charge Generalization



InD Charges(Train): $\frac{1}{2}, 2, -\frac{7}{2}, \dots$
 OOD Charges(Test): $\frac{2}{4}, \frac{6}{2}, \frac{3}{9}, \dots$

Trilinears Terms:

$$\begin{matrix} \Phi & \Phi' & \Phi'' \\ +Q & +Q' & +Q'' = 0 \end{matrix}$$

Need arithmetic

Quartic Terms:

$$\begin{matrix} \Phi & \Phi^\dagger & \Phi' & \Phi'^\dagger \\ +Q & -Q & +Q' & -Q' = 0 \end{matrix}$$

Pattern Recognition

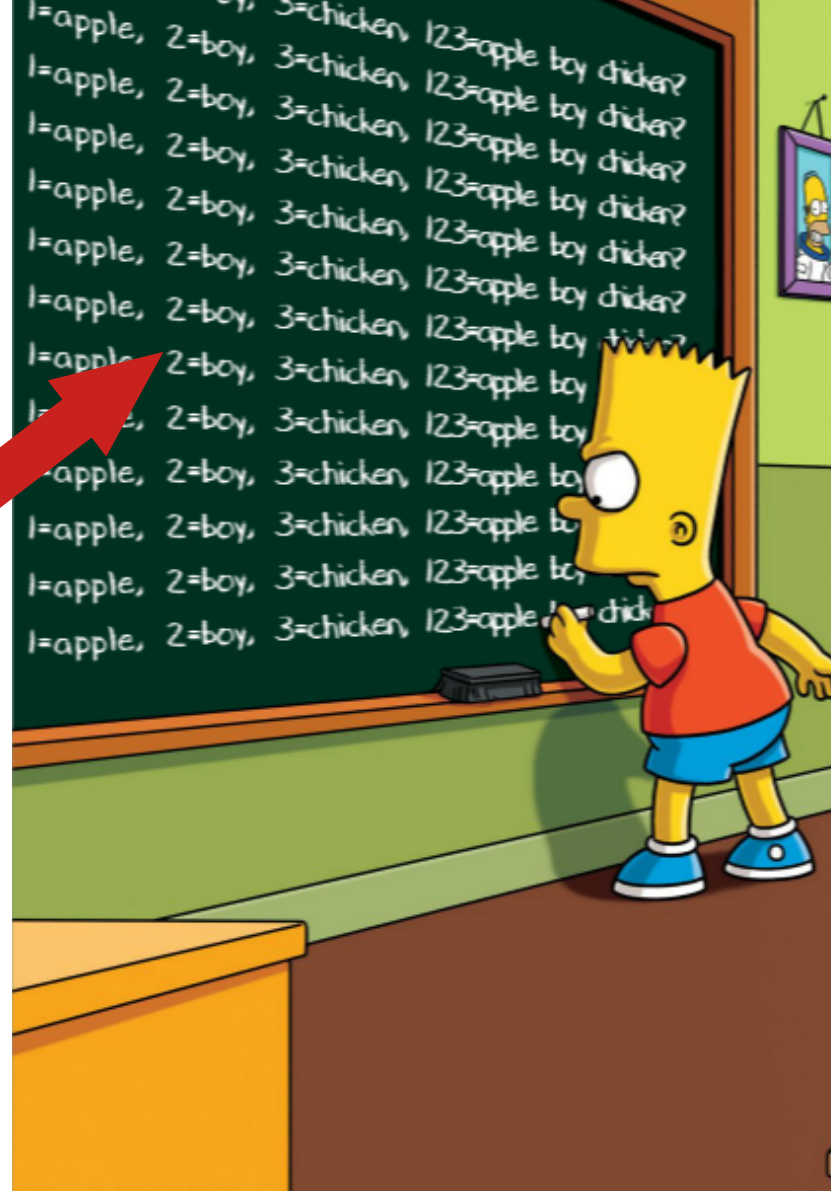
It cant do arithmetic! It doesnt understand numbers!

Text-based encoding :
It is treating numbers
like alphabets

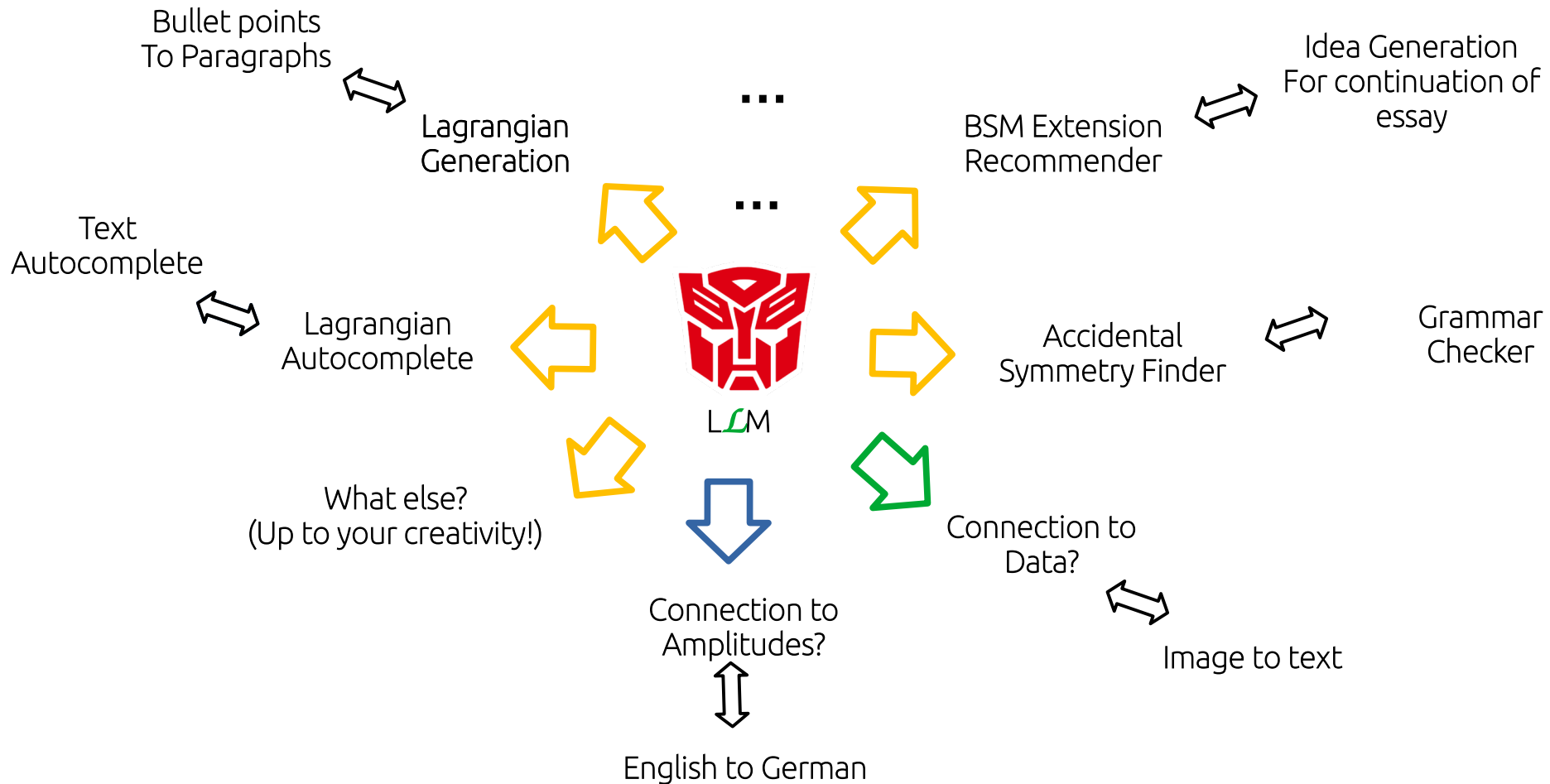
(Need continuous number tokenizing scheme)
[2310.02989]

ChatGPT cant do arithmetic precisely
without calculator or steps

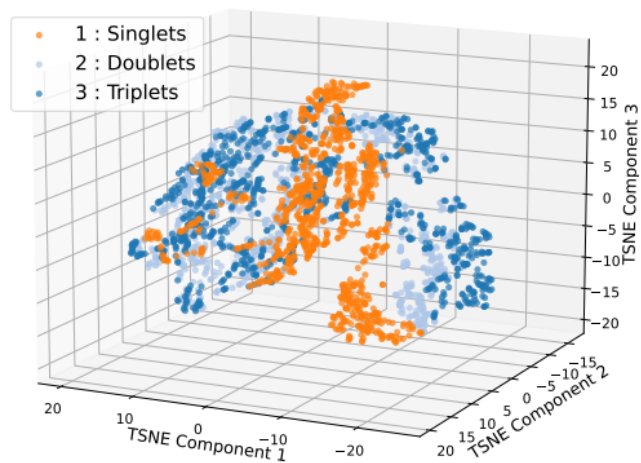
Ask ChatGPT to add two large numbers without calculator. Answer only without steps.
It fails sometimes.



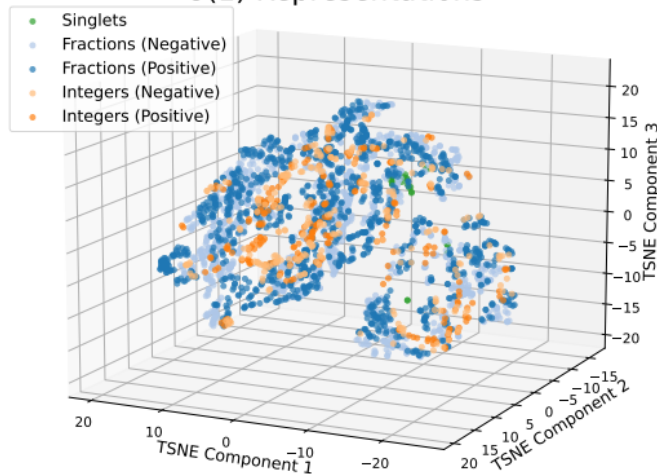
L $\mathcal{L}$ M: Large “Lagrangian” Model?



SU(2) Representations



U(1) Representations



SU(3) Representations

