

BRA駆動開発の実例：扁桃体を例に

宮本 竜也

目次

■ 背景

- 扁桃体
- 恐怖条件付けと消去

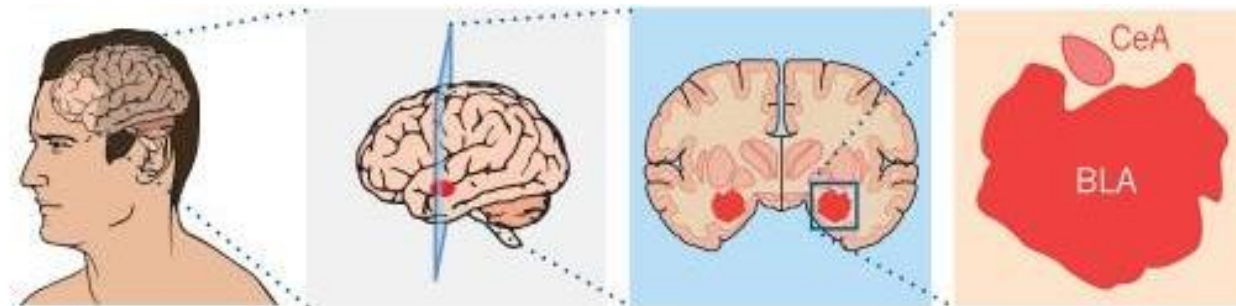
■ SCID法を用いた仮説の構築

- STEP1:BIF構築:扁桃体における事実の調査
- STEP2:ROIとTLFの設定
- STEP3:TLFの機能分解によるHCD作成

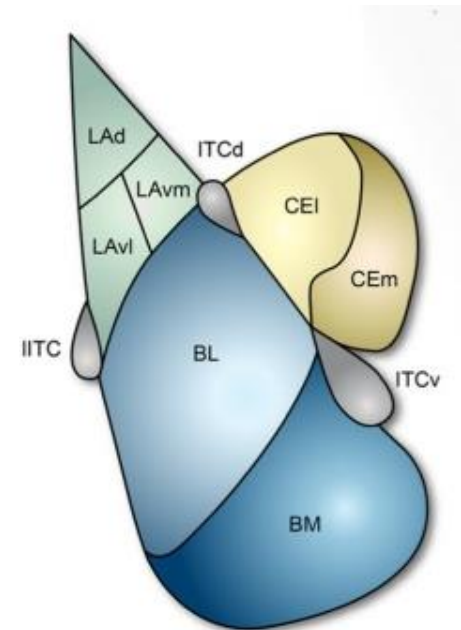
■ まとめ

扁桃体

- 側頭葉内側に位置する情動に関連する領域
- 恐怖条件付けの獲得に主要な役割を果たす
- 不均質な多数の神経核によって構成されている
 - 基底外側核(BLA)
 - 中心核(CEN)
 - インターカレレーテッド細胞核(INA)



(Patricia H. Janak, 2015)



(Lee, S., 2013)

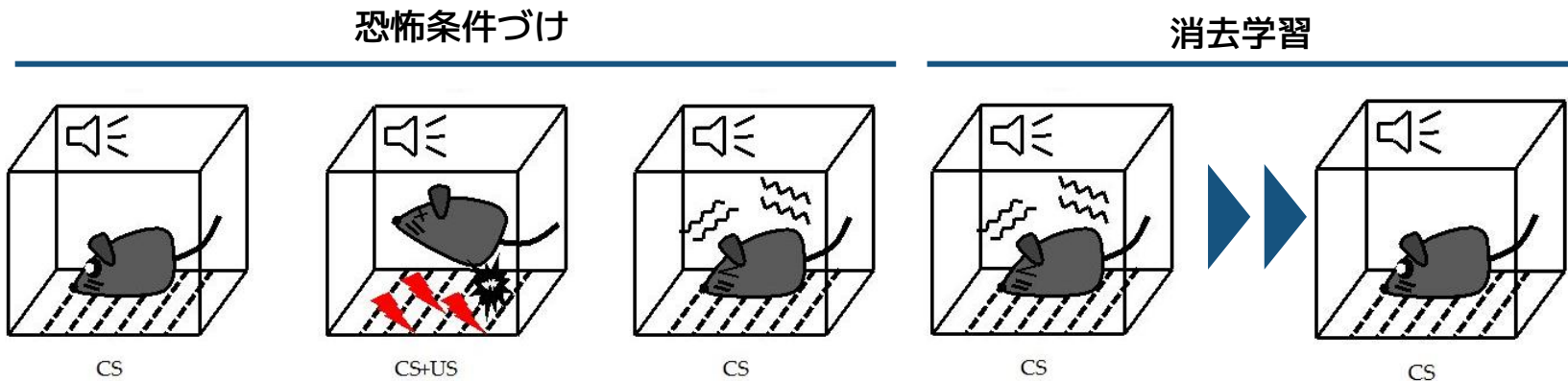
恐怖条件付けと消去学習

■ 恐怖条件付け(Conditioned Fear)

- 生物学的に重要でない中立的な刺激(条件刺激:CS)の提示と、電気刺激のような嫌悪刺激(非条件刺激:US)の提示で構成
- マウスは、CSとUSが数回同時に提示されると、CSのみで恐怖反応(Freezing、自律神経系反応、神経内分泌反応)を起こすようになる

■ 消去学習(Extinction learning)

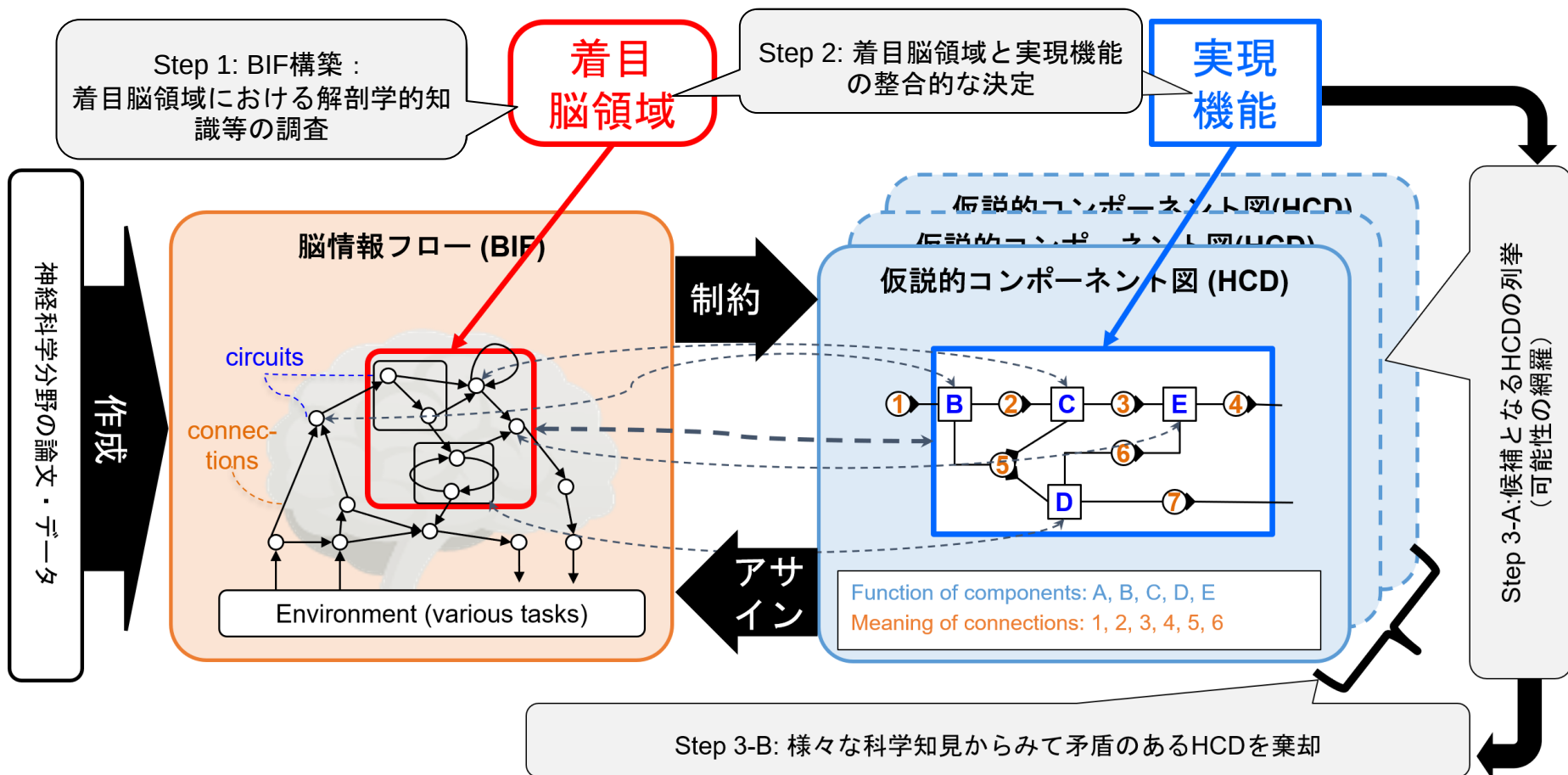
- 成立した恐怖条件付けであっても、CSのみの繰り返し刺激によって、マウスは恐怖反応を消失させる→消去(Extinction)
- 消去では、当初の条件付けを抹消するのではなく、恐怖反応を発生させなくするためのContextに依存的な新たな記憶を形成することで成立する(Bouton M. E., 2002)



目的

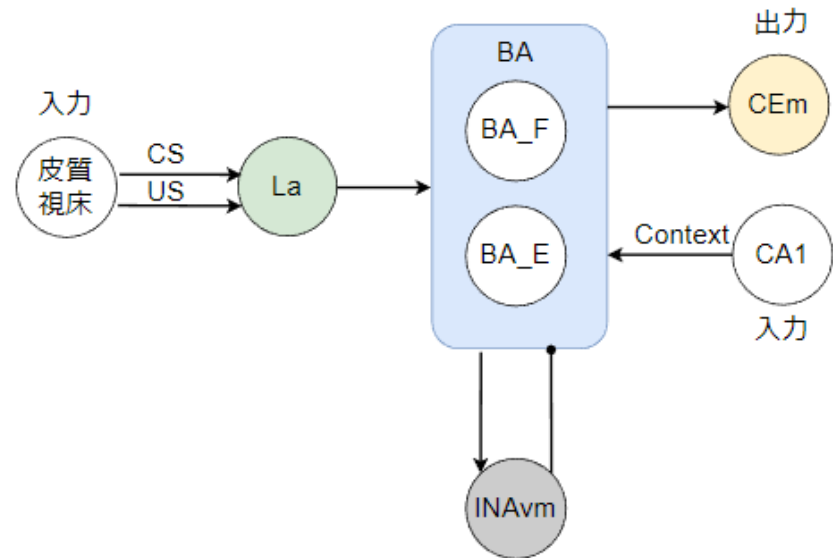
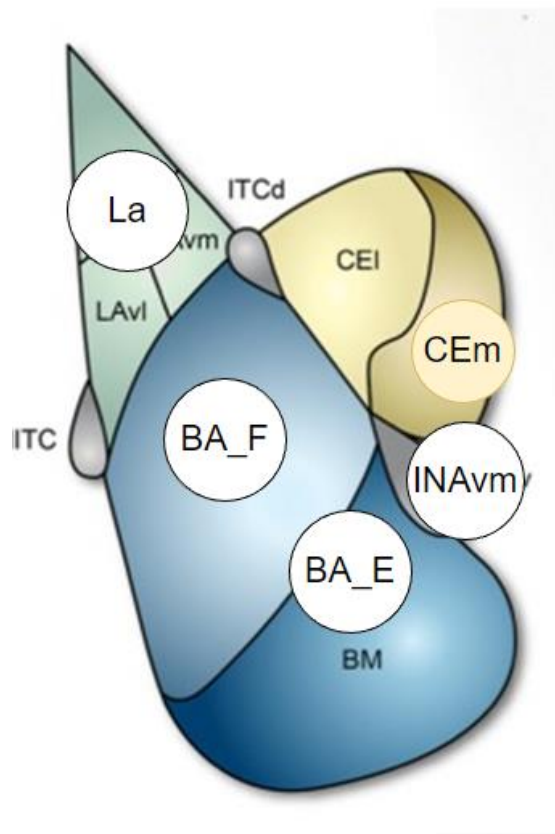
SCID法に基づいて、扁桃体で恐怖条件付けと消去学習
を実現するHCDを構築する

SCID法



STEP1: BIF(Brain Information Flow)構築

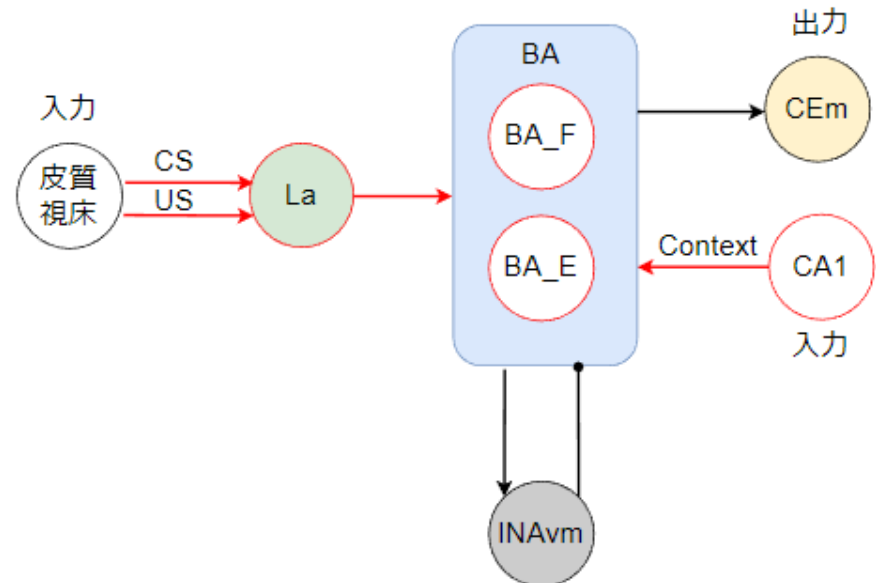
- 解剖学的、生理学的知見によって構築されたBIF
- 恐怖条件付けと関連が深い扁桃体の領域は、BLAに含まれる外側核(La)、基底核(BA)、中心核(CEN)、インターカレテッド細胞核(INA)



(Lee, S., 2013 から改変)

STEP1: BIF(Brain Information Flow)構築

- Laには皮質及び視床の感覚情報(CSおよびUS)が入力し、ヘブ則による条件付け成立を担うニューロンが存在する (Duvarci, S., & Pare, D., 2014)
- LaはBAに投射を行っていることが解剖学的に確認されている (Duvarci, S., & Pare, D., 2014)
- BAの主細胞には学習によって異なる振る舞いをする主細胞がある
 - 恐怖反応を起こすときに発火する**恐怖細胞(Fear Cell:BA_F)** (Amano, T. et al.,2011)
 - 消去学習後に活動する**消去細胞(Extinction Cell:BA_E)** (Amano, T. et al.,2011)
 - BA 内部の解剖学的接続は分かっていない
- BAは海馬CA1 細胞の投射を受けている (Pitkänen, A. et al, 2000)
 - 消去学習および回復に重要な Context に関する情報は、海馬によってもたらされると仮定



STEP1: BIF(Brain Information Flow)構築

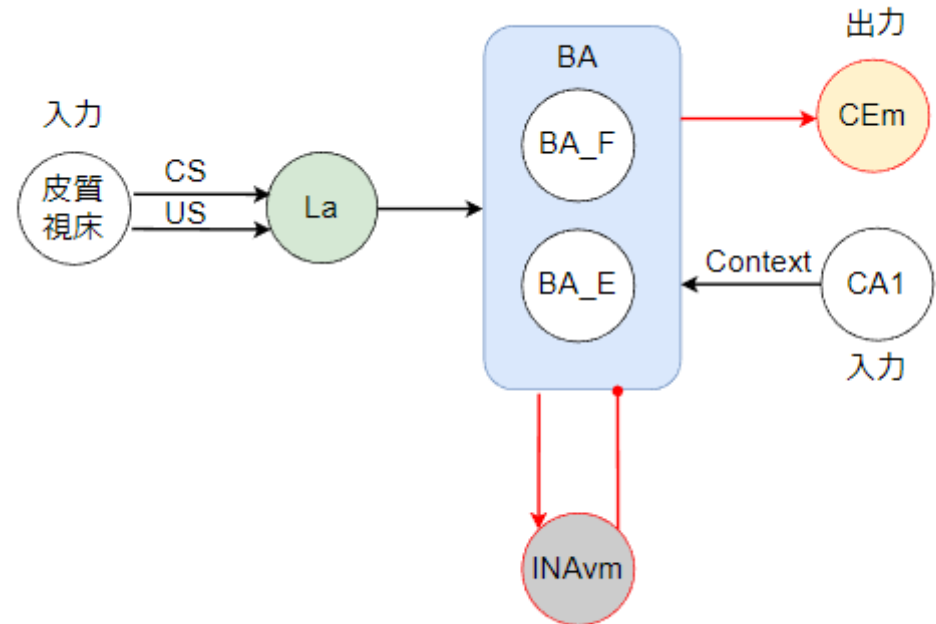
■ CEN は外側CEN(CEI)および内側CEN(CEm)に分けられる

- CEmは恐怖条件付けにおける主な出力ニューロンであり、恐怖反応を引き起こす領域に直接投射する、恐怖反応を実行するためのエフェクターである(Duvarci, S., & Pare, D., 2014)

■ INA には外側INA(INAI)と背内側INA(INAdm)と腹内側INA(INAvm)が含まれる

- INAvmは低恐怖状態で活性化される(Hagihara, K.M., et al, 2021)

■ BAからINAvmへの興奮性の投射、INAvmからBAへの抑制性の投射が解剖学的に確認されている(Hagihara, K.M., et al, 2021,) (Duvarci, S., & Pare, D., 2014)



STEP2: TLF(Top level Function)の設定(TLF1とTLF2)

■ 今回扱うTLFとして、2種類のTLFを設定した

- 恐怖条件付けのみを対象としたTLF1
- 恐怖条件付けと消去学習を対象としたTLF2

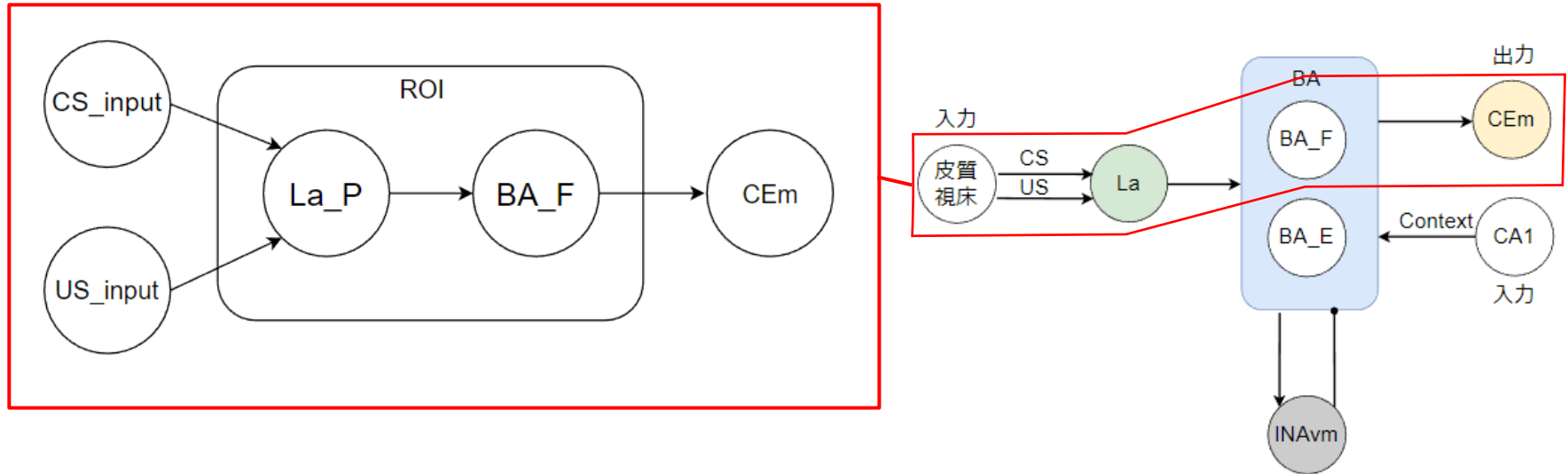
TLF1 個体の生存にとって価値のあるUSを用いて中立のCSに価値付けを行い、恐怖反応を引き起こす機能

TLF2 TLF1に加え、無条件刺激がない状態でCSが繰り返し与えられることで、恐怖反応をContext依存的に抑制する機能

▶ TLF1とTLF2について、SCID法を行ってHCDを作成してみる

STEP2: ROI(Region of Interest)の設定 (TLF1)

- 恐怖条件付けを行うTLF1 を実現するためのROI は、La とBAの恐怖細胞によって構成されると推測される

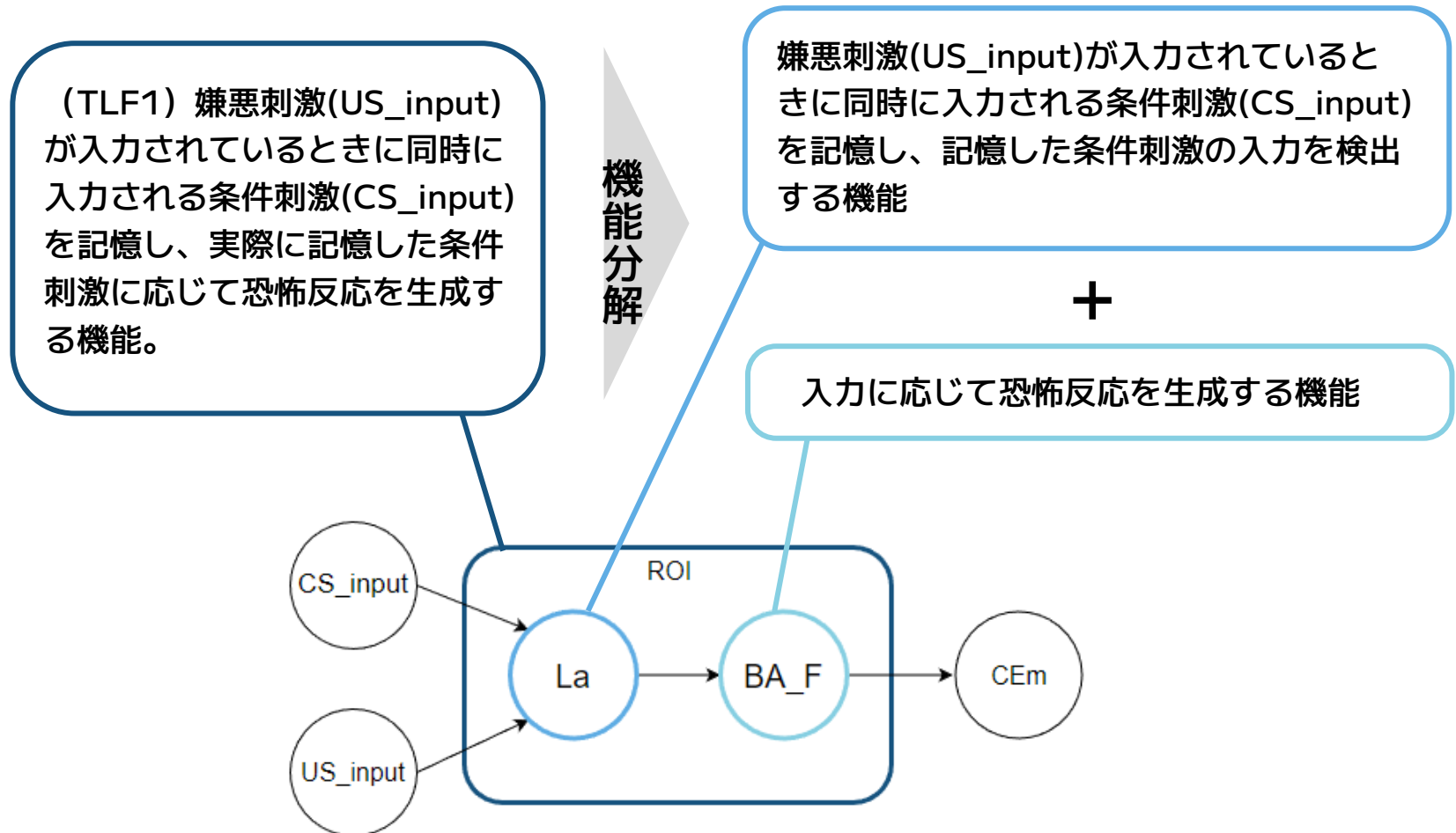


(TLF1) 個体の生存にとって価値のあるUSを用いて中立のCSに価値付けを行い、恐怖反応を引き起こす機能

- このTLFを、ROIによって定義される構造制約に従って、機能を分解する。

STEP3: 機能分解(TLF1)

- 2つのcircuitがROIで定義されているので、2つの機能に分解するのが妥当



STEP3: ProcessとOutput Semanticsとは？

■ Processとは

- コンポーネント内で行われる計算処理
- 一つ以上の入力 “I(Circuit ID)”を用いて、出力 “O(Circuit ID)”を得るための手続き（もしくはアルゴリズム）を文として記述する。
- ソフトウェア開発者が実装するコードを規定する情報

■ Output Semantics(出力の意味)とは

- コンポーネントの出力に対する外部からみた解釈
- 空間的側面 [推奨]
 - ▶ 座標系（例：網膜座標系、身体部位中心座標系、外界中心座標系）
 - ▶ 表現の形式（例：ベクトル表現、マップ表現）
 - ▶ 有効な次元数
- 時間的側面 [任意]
 - ▶ 神経活動の時間的性質（tonic/burst/phasicなど）
 - ▶ モデル化に関わる時定数（Processカラムに記することも可能）

■ Laでの例

Function

嫌悪刺激が入力されているときに同時に入力される条件刺激を記憶し、記憶した条件刺激の入力を検出する機能



Process

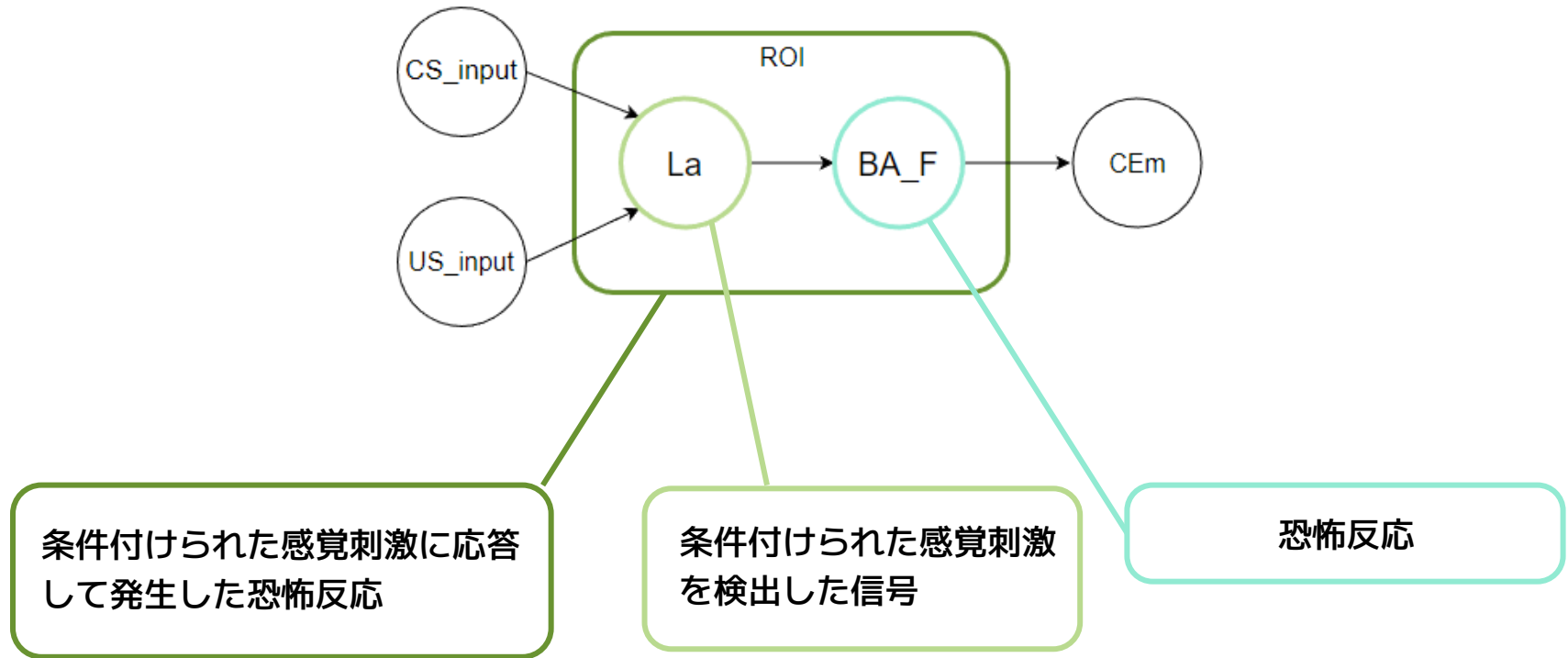
O(La)は、I(US_input)が入力されているときに同時に入力される I(CS_input)を記憶し、記憶した I(CS_input)の入力を検出することで生成される

Output Semantics

条件付けられた感覚刺激を検出した信号

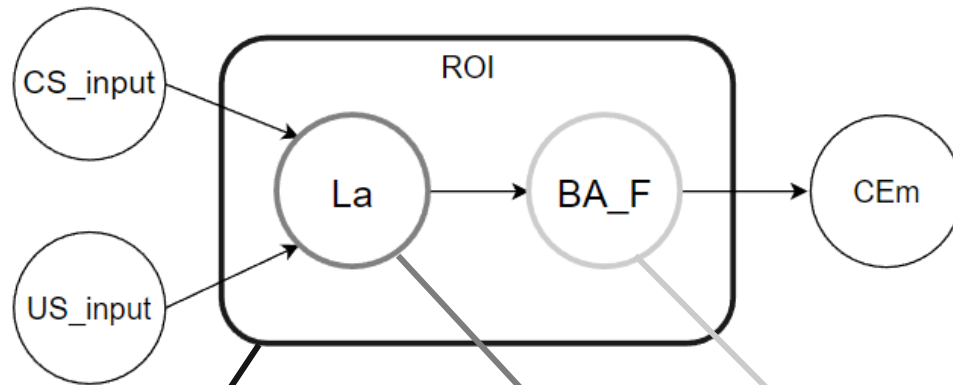
STEP3: 出力の意味(Output Semantics)(TLF1)

- output semanticsは、機能の分解によって、各Circuitにおいては以下のように決定できる



STEP3: Process(TLF1)

- output semanticsは、機能の分解によって、各Circuitにおいては以下のように決定できる



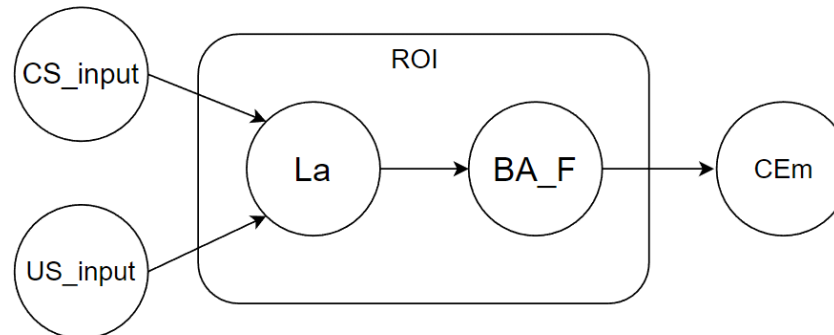
O(TLF)は、 $I(US_input)$ が入力されているときに同時に入力される $I(CS_input)$ を記憶し、実際に記憶した $I(CS_input)$ が入力されると生成される。

O(La)は $I(US_input)$ が入力されているときに同時に入力される $I(CS_input)$ を記憶し、記憶した $I(CS_input)$ が入力されると生成される

O(BA_F)は $I(La)$ が入力されると生成される

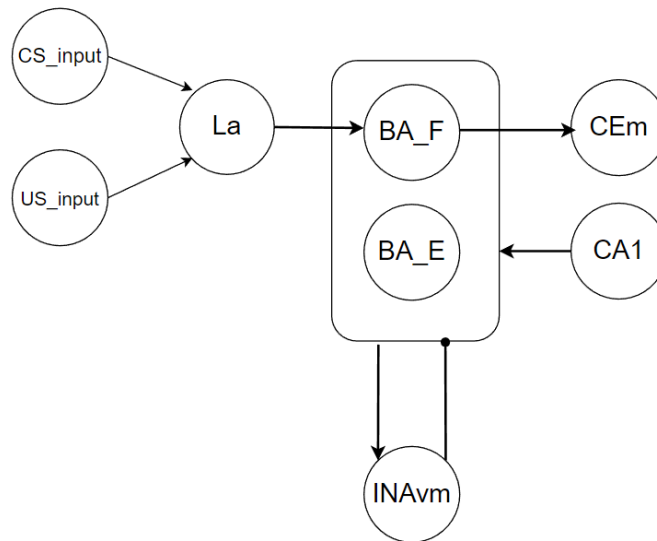
TLF1のHCDにおけるFunction Tree Viewer

Region Category	Depth = 0	Depth = 1	Depth = 2	Depth = 3	Depth = 4
Input	US_input (US_input) [U] -- Input Circuits -- - -- Output Circuits --	(Undefined interface process)	嫌悪刺激	Confirmation Summary: - Output Semantics is - Process is - Comment :	-
Input	CS_input (CS_input) [U] -- Input Circuits -- - -- Output Circuits -- La;	(Undefined interface process)	中立的な感覚刺激	Confirmation Summary: - Output Semantics is - Process is - Comment :	-
Input	CA1 (CA1) -- Input Circuits -- - -- Output Circuits -- BA_E;	(Undefined interface process)	文脈情報	Confirmation Summary: - Output Semantics is - Process is - Comment :	-
Output	CEm (CEm) -- Input Circuits -- BA_F -- Output Circuits --	(Undefined interface process)		Confirmation Summary: - Output Semantics is - Process is - Comment :	-
ROI	TLF1 (TLF1) -- Input Circuits -- - -- Output Circuits -- ?	[TLF] 嫌悪刺激(US_input)が入力されているときに同時に入力される条件刺激(CS_input)を記憶し、記憶した条件刺激のみによって恐怖反応を発生することができる能力を獲得し、実際に条件刺激に応じて恐怖反応を生成する機能。	O(BA_F): 条件付けられた感覚刺激に反応して発生した恐怖反応	Confirmation Summary: - Output Semantics is - Process is - Comment :	-
inROI	====>	La (La) -- Input Circuits -- CS_input -- Output Circuits -- BA_F;	O(La)はI(US_input)が入力されているときに同時に入力されるI(CS_input)を記憶し、記憶したI(CS_input)が入力されると生成される	条件付けられた感覚刺激に反応した情報	Confirmation Summary: - Output Semantics is - Process is - Comment :
inROI	====>	BA_F (BA_F) [U] -- Input Circuits -- La INAv -- Output Circuits -- CEm;	O(BA_F)はI(La)が入力されると生成される	恐怖反応	Confirmation Summary: - Output Semantics is - Process is - Comment :



STEP2: ROI(Region of Interest)の設定 (TLF1)

- 恐怖条件付けと消去を行うTLF2 を実現するためのROIを、La とBAの恐怖細胞 (BA_F)に加え、BAの消去細胞(BA_E)とINAdmとして設定する



(TLF2) TLF1に加え、無条件刺激がない状態でCSが繰り返し与えられることで、恐怖反応をContext依存的に抑制する機能

STEP3:機能分解(TLF2)

- TLF2ではTLF1に加えて消去の機能が追加されているので、TLF1の機能とそれ以外で分解してみる

嫌悪刺激(US_input)が入力されているときに同時に入力される条件刺激(CS_input)を記憶し、実際に記憶した条件刺激に応じて恐怖反応を生成でき、ある文脈情報(CA1)が入力されているときに、嫌悪刺激(US_input)を伴わない記憶済みの条件刺激(CS_input)が繰り返し入力すると、その文脈情報(CA1)を記憶して、この文脈情報(CA1)が入力している間は恐怖反応が発生しなくなる機能

機能分解

(TLF1) 嫌悪刺激(US_input)が入力されているときに同時に入力される条件刺激(CS_input)を記憶し、記憶した条件刺激のみによって恐怖反応を発生することができる能力を獲得し、実際に条件刺激に応じて恐怖反応を生成する機能。

+

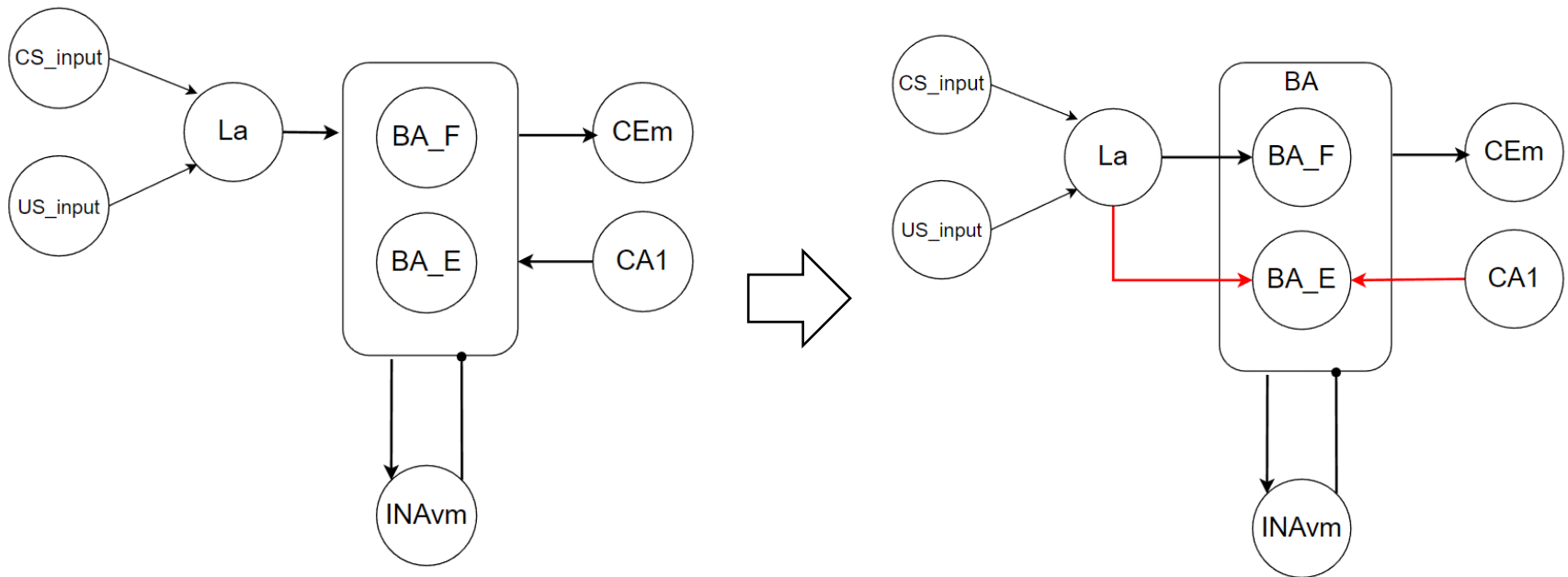
無条件刺激がない状態で学習済みCSが繰り返し与えられたときのContextを記憶し、このContextの入力によって恐怖反応を抑制する機能

STEP3:機能分解(TLF2)

無条件刺激がない状態で学習済みCSが繰り返し与えられたときのContextを記憶し、このContextの入力によって恐怖反応を抑制する機能

■ Laの情報とContextの情報が収束するCircuitが必要

- 候補：BA_E, BA_F
 - ▶ 消去が起こるときに活動するBA_Eが最も妥当そう



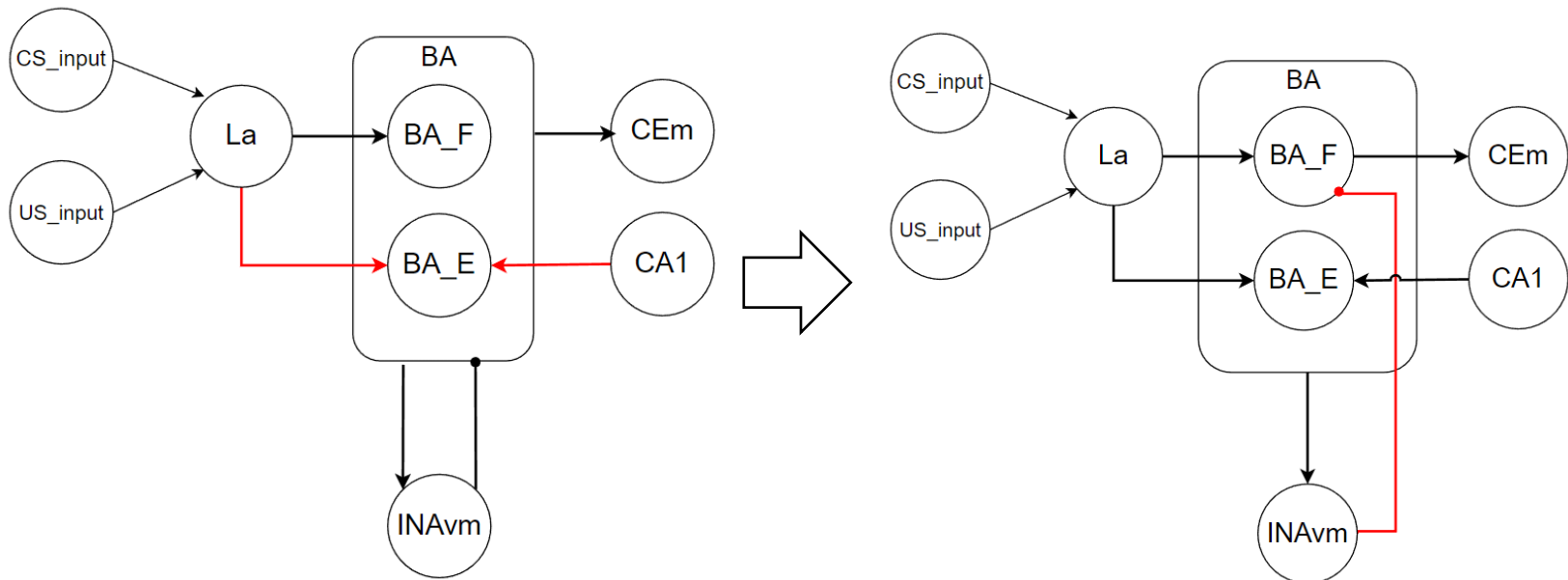
STEP3:機能分解(TLF2)

無条件刺激がない状態で学習済みCSが繰り返し与えられたときのContextを記憶し、このContextの入力によって**恐怖反応を抑制する機能**

■ 恐怖反応を抑制するCircuitが必要

● 候補：INAvmのみ

- ▶ INAvmが「恐怖反応を起こすとき活動するBA_F」に投射していれば恐怖反応を抑制できそう



STEP3:機能分解(TLF2)

- TLF2消去にかかわる機能をさらに機能分解することができた

無条件刺激がない状態でCSが繰り返し与えられることで、恐怖反応をContext依存的に抑制する機能

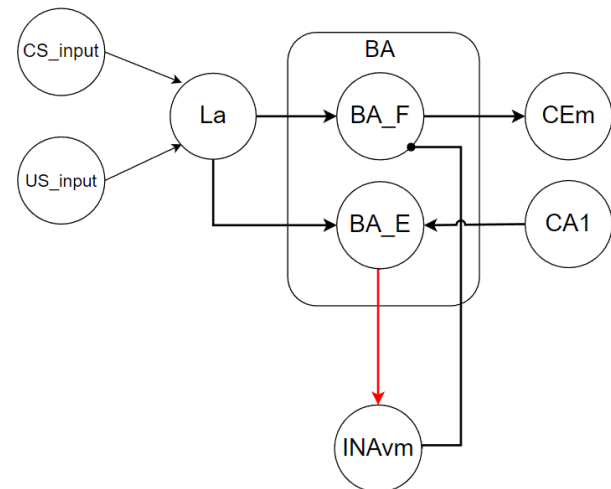
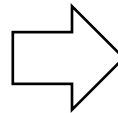
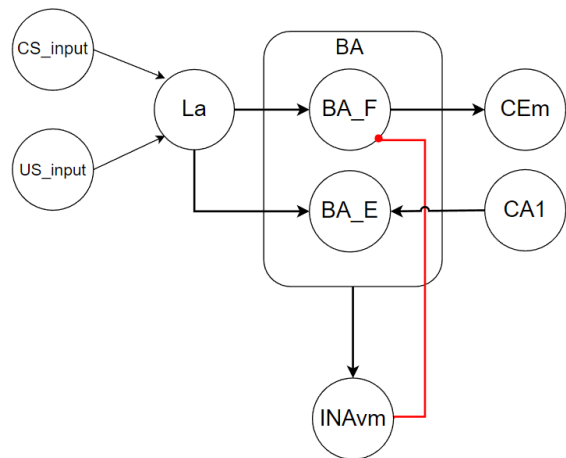
機能分解

無条件刺激がない状態でCSが繰り返し与えられるときのContextを記憶し、そのContextを検出する機能

+

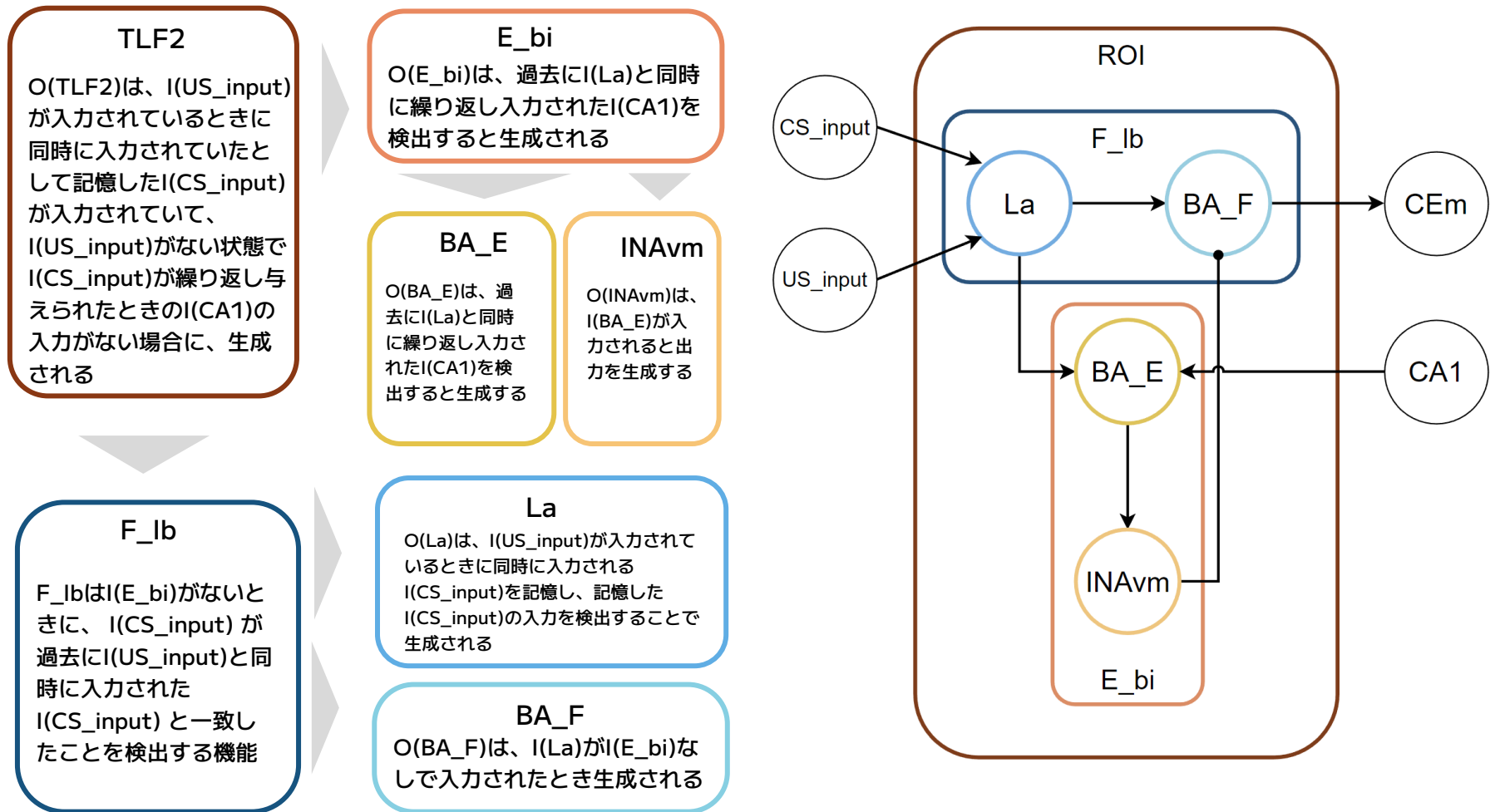
恐怖反応を抑制する機能

- 情報の流れとしてはBA_EからINAvmへ投射があれば問題なさそう



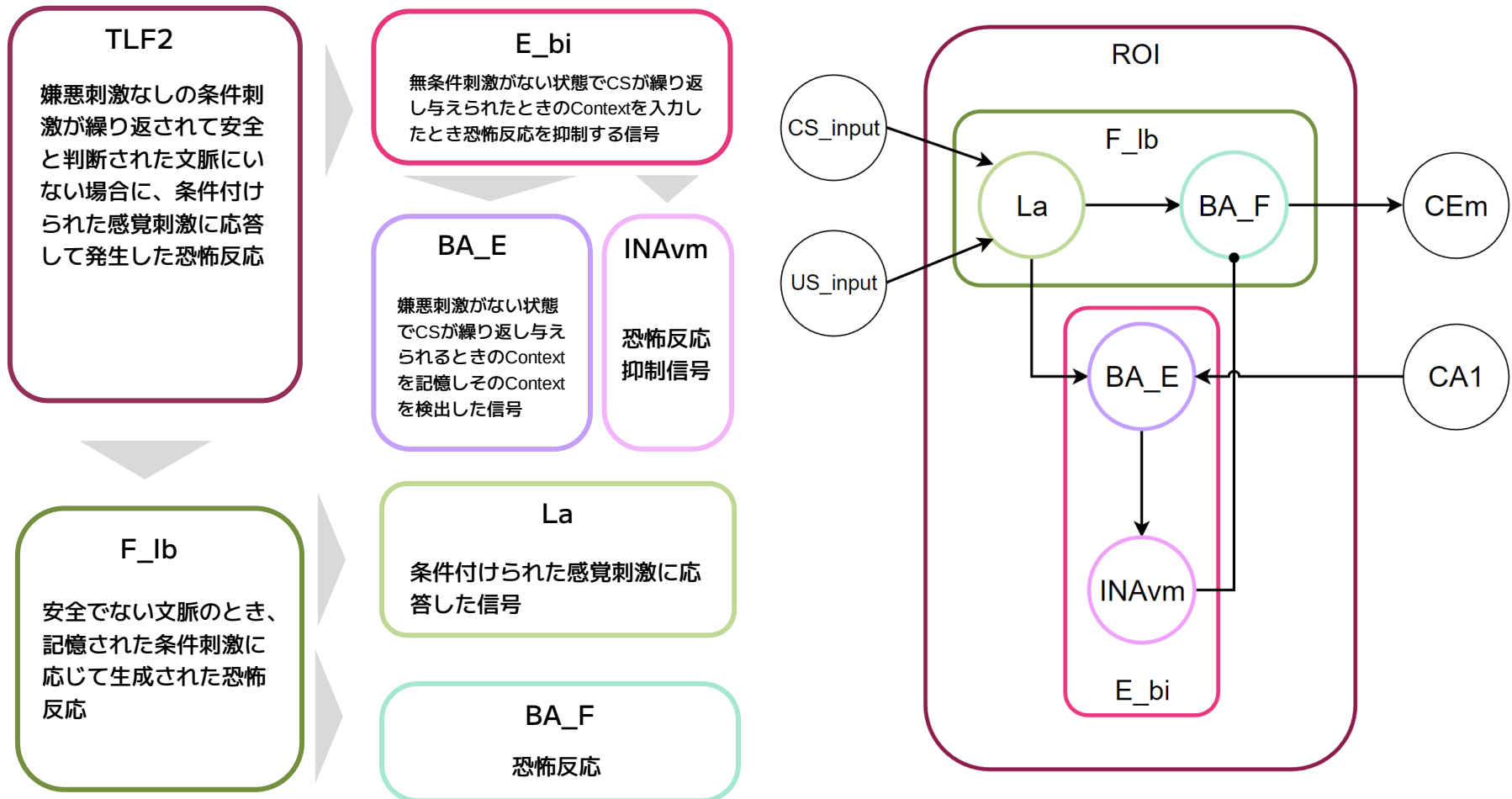
STEP3:Processの決定(TLF2)

■ これらの機能分解によって得られたProcess



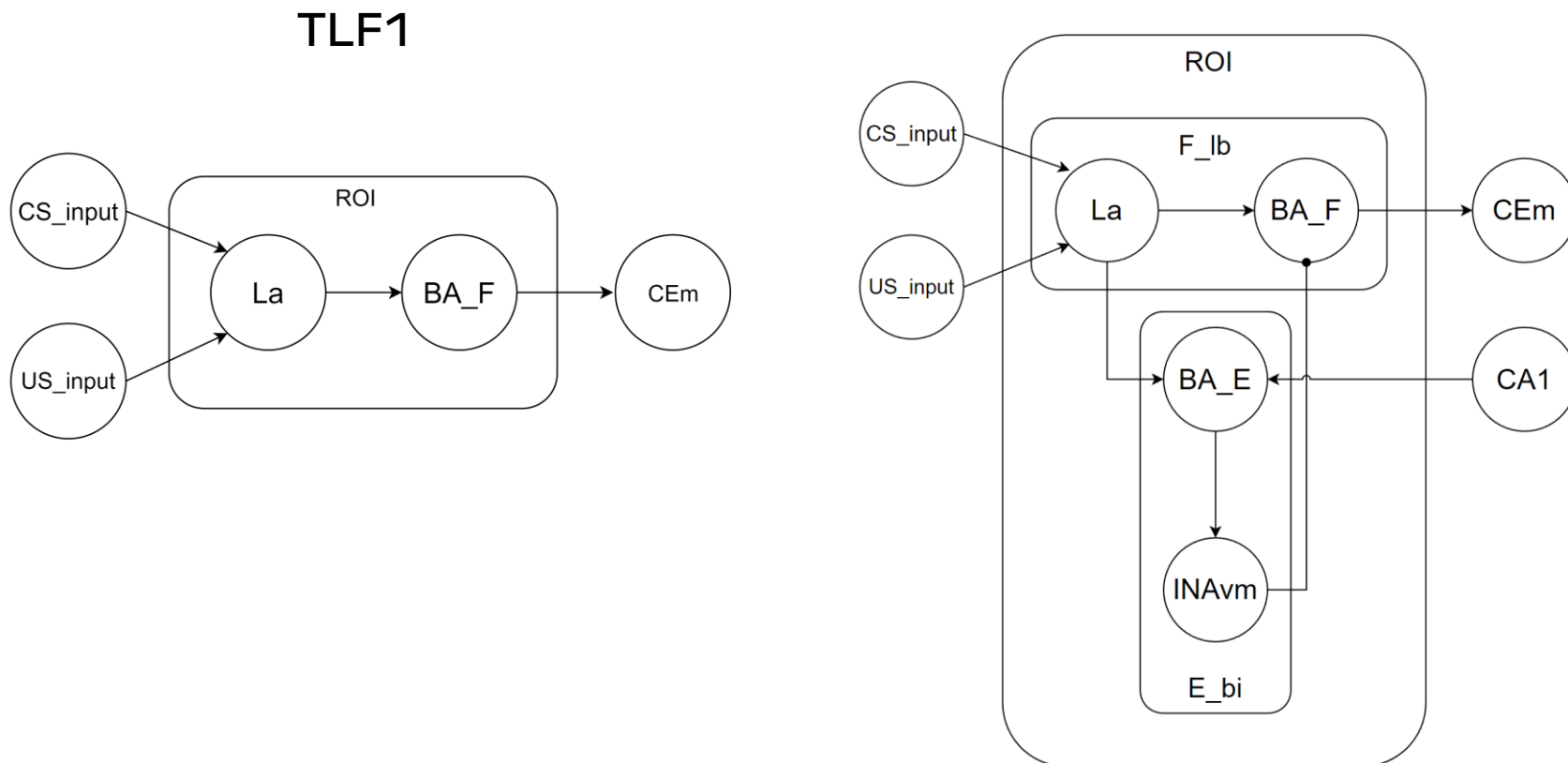
STEP3:Output semanticsの決定(TLF2)

■ これらの機能分解によってOutput Semantics



まとめ

- SCID法を用いて構造制約的にTLFを分解することで、扁桃体において未知の神経接続を推測しながら妥当なHCDを作成できた



参考文献

- Janak, P. H., & Tye, K. M. (2015). From circuits to behaviour in the amygdala. *Nature*, 517(7534), 284–292. <https://doi.org/10.1038/nature14188>
- Lee, S., Kim, S. J., Kwon, O. B., Lee, J. H., & Kim, J. H. (2013). Inhibitory networks of the amygdala for emotional memory. *Frontiers in neural circuits*, 7, 129. <https://doi.org/10.3389/fncir.2013.00129>
- Bouton M. E. (2002). Context, ambiguity, and unlearning: sources of relapse after behavioral extinction. *Biological psychiatry*, 52(10), 976–986. [https://doi.org/10.1016/s0006-3223\(02\)01546-9](https://doi.org/10.1016/s0006-3223(02)01546-9)
- Pitkänen, A., Jolkkonen, E., & Kemppainen, S. (2000). Anatomic heterogeneity of the rat amygdaloid complex. *Folia morphologica*, 59(1), 1–23.
- Shin, R. (2012). Cellular Mechanisms in the Amygdala Involved in Memory of Fear Conditioning. In (Ed.), *The Amygdala - A Discrete Multitasking Manager*. IntechOpen. <https://doi.org/10.5772/48397>
- Amano, T., Duvarci, S., Popa, D., & Paré, D. (2011). The fear circuit revisited: contributions of the basal amygdala nuclei to conditioned fear. *The Journal of neuroscience : the official journal of the Society for Neuroscience*, 31(43), 15481–15489. <https://doi.org/10.1523/JNEUROSCI.3410-11.2011>
- Duvarci, S., & Pare, D. (2014). Amygdala microcircuits controlling learned fear. *Neuron*, 82(5), 966–980. <https://doi.org/10.1016/j.neuron.2014.04.042>
- Hagihara, K.M., Bukalo, O., Zeller, M. et al. Intercalated amygdala clusters orchestrate a switch in fear state. *Nature* 594, 403–407 (2021). <https://doi.org/10.1038/s41586-021-03593-1>