

Lecture 3

Unifying the Analysis in Control and Optimization via SDPs, Part I

Lecturer: Bin Hu, Date: 09/01/2026

An important object that has been extensively studied in the controls field is the feedback interconnection. For a dynamical system G and a mapping Δ , a feedback interconnection of G and Δ is shown in Figure 3.1 and denoted as $F_u(G, \Delta)$.

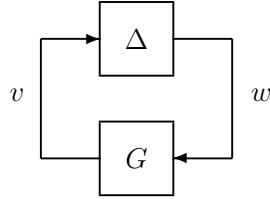


Figure 3.1. The Block-Diagram Representation for Feedback Interconnection $F_u(G, \Delta)$

The feedback interconnection states that v and w must satisfy $v = G(w)$ and $w = \Delta(v)$ simultaneously. For example, when G is an LTI system and Δ is a static nonlinearity, the feedback interconnection $F_u(G, \Delta)$ actually denotes the following recursive equations:

$$\begin{aligned}\xi_{k+1} &= A\xi_k + Bw_k \\ v_k &= C\xi_k + Dw_k \\ w_k &= \Delta(v_k)\end{aligned}\tag{3.1}$$

The first two equations in the above iterations state the fact $v = G(w)$, and the third equation enforces $w = \Delta(v)$.

Well-posedness. Clearly a basic question one should ask is whether there exists a pair of (v, w) satisfying $v = G(w)$ and $w = \Delta(v)$ simultaneously such that the feedback interconnection $F_u(G, \Delta)$ is well defined in the first place. This is the so-called well-posedness issue. Typically, one needs to prove well-posedness in a case-by-case manner.

Example: Lur'e systems. When $D = 0$, the system (3.1) is equivalent to a nonlinear autonomous system $\xi_{k+1} = A\xi_k + B\Delta(C\xi_k)$, which is the so-called Lur'e system. Therefore, the sequences $\{\xi_k\}$, $\{w_k\}$, and $\{v_k\}$ will be completely determined given an initial condition ξ_0 , and the feedback interconnection is well-posed. It is more difficult to analyze the internal stability of the nonlinear system $\xi_{k+1} = A\xi_k + B\Delta(C\xi_k)$ than the linear autonomous system $\xi_{k+1} = A\xi_k$ (which is stable if and only if A has a spectral radius smaller than 1). The nonlinear map Δ introduces some fundamental difficulty such that the spectral radius

argument cannot be applied any more. If Δ is a linear function, then the nonlinear system $\xi_{k+1} = A\xi_k + B\Delta(C\xi_k)$ becomes linear and the internal stability analysis becomes easy. However, general Δ is hard to handle. In this lecture, we will introduce the quadratic constraint approach from the controls field.

Generality of the feedback interconnection model. The above type of feedback interconnections becomes a key object for robust control study due to the fact that it can model various “perturbed” versions of linear systems. The perturbation Δ can be model uncertainty, delays, or nonlinearity. We will explain this in next section and then talk about a general robustness analysis tool called dissipation inequality.

3.1 Uncertainty Modeling in Control

In the controls field, the feedback interconnection $F_u(G, \Delta)$ is widely used to model uncertain or nonlinear systems. The idea is to separate a dynamical system into two pieces: a “nominal” part G and a perturbation Δ . The nominal part G is typically linear and easy to analyze. The perturbation Δ can be the uncertainty in the system dynamics or some troublesome element causing difficulty in the analysis. The feedback interconnection $F_u(G, \Delta)$ can be viewed as a “perturbed” version of the nominal system G . The study for such perturbed systems forms the foundation of robust control. Now let’s look at a few examples of Δ .

- Parametric uncertainty: Consider a linear system $\xi_{k+1} = A\xi_k$. We want to know whether this system is stable or not. In practice, we will not know A exactly. Typically we have $A = \bar{A} + A_\delta$ where $\bar{A}$ is some measured version of A and A_δ captures the uncertainty in the system dynamics. We do not know what A_δ is exactly equal to, but we do know that A_δ is a constant matrix whose input-output gain is bounded above by some small number. Therefore, the system dynamics becomes $\xi_{k+1} = (\bar{A} + A_\delta)\xi_k$, and can be rewritten as a special case of $F_u(G, \Delta)$ where Δ maps v to w as $w_k = (A_\delta)v_k$, and G is defined as

$$\begin{aligned}\xi_{k+1} &= \bar{A}\xi_k + w_k \\ v_k &= \xi_k\end{aligned}$$

Although we do not know what A_δ is equal to, it is still possible that we can use the bound on A_δ to establish the stability of such a feedback interconnection.

- Time-varying parameters: In the above example, we can further allow A_δ to change over time, i.e. $w_k = (A_\delta^{(k)})v_k$. We can absorb the time-varying element into Δ and treat it as a perturbation.
- Time delay: Consider a control system $\xi_{k+1} = A\xi_k + Bu_k$ where the state feedback controller is affected by a delay, i.e. $u_k = K\xi_{k-\tau_k}$. Ideally, the control input should be determined based on the current state information. However, there may be a time

delay in the system and eventually u_k is calculated based on a past state measurement $x_{k-\tau_k}$. Here τ_k is the delay at step k . We can choose G as an LTI system governed by $\xi_{k+1} = A\xi_k + BKw_k$ and $v_k = \xi_k$. Then the control system can be modeled as $F_u(G, \Delta)$ where Δ is a delay operator mapping v to w as $w_k = v_{k-\tau_k}$. Notice G and Δ should be thought as operators that map real sequences to real sequences.

- **Dynamical uncertainty:** Sometimes even the order of the model may not be correct. For example, one may use a rigid body model for control purposes when there are flexible modes in the true dynamics. In this case, Δ is a dynamical system satisfying some norm bound. Specially, w_k is not completely determined by v_k . The computation of w_k may require the past information of the sequence $\{v_k\}$. For example, Δ can sometimes be a LTI system itself:

$$\begin{aligned}\zeta_{k+1} &= A_\Delta \zeta_k + B_\Delta v_k \\ w_k &= C_\Delta \zeta_k + D_\Delta v_k\end{aligned}$$

In this case, we do not know the matrices $(A_\Delta, B_\Delta, C_\Delta, D_\Delta)$. To make things worse, we do not even know the dimension of ζ_k . We only know that the norm of Δ is bounded, i.e. we can establish a bound in the form of $\sum_{k=0}^{\infty} \|w_k\|^2 \leq \gamma^2 \sum_{k=0}^{\infty} \|v_k\|^2$ for $\zeta_0 = 0$.

- **Actuator saturation and other nonlinearity:** Sometimes a few parts of a control system can not be modeled by linear approximations and the nonlinearity has to be taken into accounts for the stability analysis. It is possible to separate the nonlinearity from the linear dynamics and absorb it into Δ . One such example is the actuator saturation. Specifically, suppose v_k is a scalar. The saturation function maps v_k to w_k as $w_k = v_k$ for $|v_k| \leq v_{\max}$ and $w_k = v_{\max}$ for $|v_k| \geq v_{\max}$. Other examples include periodically changing nonlinear functions such as $\cos$ and $\sin$.

To summarize, the perturbation Δ can model uncertain dynamics, time delay, and nonlinearity in the control system. All these perturbation operators have been extensively studied in the controls literature. Many linear matrix inequalities (LMIs) have been formulated to test the stability of feedback systems involving such perturbations.

For example, if one knows Δ is a bounded operator and $\|\Delta(v_k)\| \leq \delta \|v_k\|$ for any v_k , then one can use the following LMI condition to test the internal stability of $F_u(G, \Delta)$.

Theorem 3.1. Suppose Δ is a bounded operator and $\|\Delta(v_k)\| \leq \delta \|v_k\|$ for any v_k . If there exists a positive definite matrix P and a positive rate $0 < \rho < 1$ such that

$$\begin{bmatrix} A^\top P A - \rho^2 P & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} + \begin{bmatrix} C & D \\ 0 & I \end{bmatrix}^\top \begin{bmatrix} \delta^2 I & 0 \\ 0 & -I \end{bmatrix} \begin{bmatrix} C & D \\ 0 & I \end{bmatrix} \leq 0 \quad (3.2)$$

then for any x_0 , the feedback interconnection (3.1) satisfies $\|x_k\| \leq c\rho^k \|x_0\|$ where c is some constant.

Proof: Based on the condition (3.2), we have

$$\begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top \left(\begin{bmatrix} A^\top P A - \rho^2 P & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} + \begin{bmatrix} C & D \\ 0 & I \end{bmatrix}^\top \begin{bmatrix} \delta^2 I & 0 \\ 0 & -I \end{bmatrix} \begin{bmatrix} C & D \\ 0 & I \end{bmatrix} \right) \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} \leq 0 \quad (3.3)$$

Since $\xi_{k+1} = A\xi_k + Bw_k$, we have

$$\begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top \begin{bmatrix} A^\top P A - \rho^2 P & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} = \xi_{k+1}^\top P \xi_{k+1} - \rho^2 \xi_k^\top P \xi_k$$

We also have

$$\begin{aligned} -\|w_k\|^2 + \delta^2 \|v_k\|^2 &= -\|w_k\|^2 + \delta^2 (C\xi_k + Dw_k)^\top (C\xi_k + Dw_k) \\ &= \begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top \begin{bmatrix} C & D \\ 0 & I \end{bmatrix}^\top \begin{bmatrix} \delta^2 I & 0 \\ 0 & -I \end{bmatrix} \begin{bmatrix} C & D \\ 0 & I \end{bmatrix} \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} \end{aligned}$$

Consequently, (3.3) just leads to

$$\xi_{k+1}^\top P \xi_{k+1} - \rho^2 \xi_k^\top P \xi_k + \delta^2 \|v_k\|^2 - \|w_k\|^2 \leq 0$$

Since $\|w_k\| = \|\Delta(v_k)\| \leq \delta \|v_k\|$, we know $\delta^2 \|v_k\|^2 - \|w_k\|^2 \geq 0$, and the above inequality leads to $\xi_{k+1}^\top P \xi_{k+1} - \rho^2 \xi_k^\top P \xi_k \leq 0$. Since P is positive definite, we can immediately obtain the desired conclusion. ■

When (A, B, C, D) and ρ^2 are given, the condition (3.2) is linear in P and can be numerically solved via semidefinite programs. The key idea in the above analysis is to replace the nonlinearity Δ with a bound $\|w_k\|^2 = \|\Delta(v_k)\|^2 \leq \delta^2 \|v_k\|^2$ and then combine this bound with the linear state-space model of G to formulate an LMI condition.

Extensions. One can extend the above analysis to handle much more general systems. One can generalize the analysis for the cases where G is time-varying or even stochastic. We will discuss this in future lectures.

3.2 Optimization Methods as Feedback Systems

In recent years, it has been recognized that many first-order optimization methods for large-scale problems are just special cases of feedback systems. In this section, we will look at a few examples including the gradient descent method, the Heavy-ball method, and Nesterov's accelerated method.

To minimize a function $f(x)$, the gradient method iterates as $x_{k+1} = x_k - \alpha \nabla f(x_k)$. The Heavy-ball method iterates as

$$x_{k+1} = x_k - \alpha \nabla f(x_k) + \beta(x_k - x_{k-1}) \quad (3.4)$$

The extra term $\beta(x_k - x_{k-1})$ is the so-called “momentum term.” One needs to choose the stepsize α and the momentum β , and also initialize the method at x_0 and x_{-1} . Then based on this iteration, one can compute $x_1, x_2, \dots$

Nesterov’s accelerated method has a similar iterative form:

$$\begin{aligned} y_k &= x_k + \beta(x_k - x_{k-1}) \\ x_{k+1} &= y_k - \alpha \nabla f(y_k) \end{aligned}$$

We can simply rewrite Nesterov’s method as

$$x_{k+1} = x_k - \alpha \nabla f((1 + \beta)x_k - \beta x_{k-1}) + \beta(x_k - x_{k-1}) \quad (3.5)$$

This looks very similar to the Heavy-ball method. The difference is that Nesterov’s accelerated method uses a gradient evaluated at $(1 + \beta)x_k - \beta x_{k-1}$ while the Heavy-ball method uses a gradient evaluated at x_k . The Heavy-ball method and Nesterov’s method only use the first-order derivative (gradient) and do not require evaluating the second-order derivative (Hessian). Hence they belong to “first-order optimization methods.”

All the above methods can be modeled as feedback interconnection $F_u(G, \Delta)$ where G is an LTI system with $D = 0$ and Δ is just the gradient ∇f . In this case, $F_u(G, \Delta)$ becomes the following feedback model

$$\begin{aligned} \xi_{k+1} &= A\xi_k + Bw_k \\ v_k &= C\xi_k \\ w_k &= \nabla f(v_k) \end{aligned} \quad (3.6)$$

where A , B , and C are matrices with compatible dimensions. In this general model, we can choose (A, B, C) accordingly to recover various first-order methods.

1. For gradient method, we choose $A = I$, $B = -\alpha I$, $C = I$, and $\xi_k = x_k$. Then $v_k = C\xi_k = x_k$, and $w_k = \nabla f(v_k) = \nabla f(x_k)$. The iteration $\xi_{k+1} = A\xi_k + Bw_k$ just becomes $x_{k+1} = Ax_k + Bw_k = x_k - \alpha \nabla f(x_k)$, which is exactly the gradient method.
2. For the Heavy-ball method, we choose $A = \begin{bmatrix} (1 + \beta)I & -\beta I \\ I & 0 \end{bmatrix}$, $B = \begin{bmatrix} -\alpha I \\ 0 \end{bmatrix}$, $C = [I \ 0]$, and $\xi_k = \begin{bmatrix} x_k \\ x_{k-1} \end{bmatrix}$. Then $v_k = C\xi_k = [I \ 0] \begin{bmatrix} x_k \\ x_{k-1} \end{bmatrix} = x_k$, and $w_k = \nabla f(v_k) = \nabla f(x_k)$. The iteration $\xi_{k+1} = A\xi_k + Bw_k$ becomes

$$\begin{bmatrix} x_{k+1} \\ x_k \end{bmatrix} = \begin{bmatrix} (1 + \beta)I & -\beta I \\ I & 0 \end{bmatrix} \begin{bmatrix} x_k \\ x_{k-1} \end{bmatrix} + \begin{bmatrix} -\alpha I \\ 0 \end{bmatrix} \nabla f(x_k) = \begin{bmatrix} (1 + \beta)x_k - \beta x_{k-1} - \alpha \nabla f(x_k) \\ x_k \end{bmatrix}$$

which is exactly the iteration for the Heavy-ball method.

3. For Nesterov's accelerated method, we choose $A = \begin{bmatrix} (1+\beta)I & -\beta I \\ I & 0 \end{bmatrix}$, $B = \begin{bmatrix} -\alpha I \\ 0 \end{bmatrix}$, $C = \begin{bmatrix} (1+\beta)I & -\beta I \end{bmatrix}$, and $\xi_k = \begin{bmatrix} x_k \\ x_{k-1} \end{bmatrix}$. Then $v_k = C\xi_k = \begin{bmatrix} (1+\beta)I & -\beta I \end{bmatrix} \begin{bmatrix} x_k \\ x_{k-1} \end{bmatrix} = (1+\beta)x_k - \beta x_{k-1}$, and $w_k = \nabla f(v_k) = \nabla f((1+\beta)x_k - \beta x_{k-1})$. The iteration $\xi_{k+1} = A\xi_k + Bw_k$ becomes

$$\begin{aligned} \begin{bmatrix} x_{k+1} \\ x_k \end{bmatrix} &= \begin{bmatrix} (1+\beta)I & -\beta I \\ I & 0 \end{bmatrix} \begin{bmatrix} x_k \\ x_{k-1} \end{bmatrix} + \begin{bmatrix} -\alpha I \\ 0 \end{bmatrix} \nabla f(v_k) \\ &= \begin{bmatrix} (1+\beta)x_k - \beta x_{k-1} - \alpha \nabla f((1+\beta)x_k - \beta x_{k-1}) \\ x_k \end{bmatrix} \end{aligned}$$

which is exactly the iteration (3.5) for Nesterov's accelerated method.

We can see that the only difference between Nesterov's accelerated method and the Heavy-ball method is the choice of C . The different choices of C lead to completely different performance guarantees for these two methods when applied to smooth strongly-convex objective functions.

3.3 Robustness Analysis via SDPs

The impacts of the perturbation Δ on the performance of the closed-loop system $F_u(G, \Delta)$ can be assessed by various robustness analysis tools in the controls literature. One such analysis routine is provided by dissipation inequalities and quadratic constraints.

Let us first look at $\xi_{k+1} = A\xi_k + Bw_k$. Dissipation inequality just describes how the input w_k changes the energy of the state ξ_k .

Definition 1. The system $\xi_{k+1} = A\xi_k + Bw_k$ is dissipative with respect to the supply rate $S(\xi, w)$ if there exists $V : \mathbb{R}^{n_\xi} \mapsto \mathbb{R}^+$ such that

$$V(\xi_{k+1}) - V(\xi_k) \leq S(\xi_k, w_k) \quad (3.7)$$

for all k . The function V is called a storage function, which quantifies the energy stored in the state ξ . The supply rate S is a function that quantifies the energy supplied to the state ξ_k by the input w_k . In addition, (3.7) is called the dissipation inequality.

The dissipation inequality (3.7) states that the internal energy increase is equal to the sum of the supplied energy and the energy dissipation. Since there will always be some energy dissipating from the system, hence the internal energy increase (which is exactly $V(\xi_{k+1}) - V(\xi_k)$) is always bounded above by the energy supplied to the system (which is exactly $S(\xi_k, w_k)$).

One important variant of the original dissipation inequality is the so-called exponential dissipation inequality:

$$V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k) \quad (3.8)$$

where $0 < \rho^2 < 1$. The dissipation inequality (3.8) just states that at least a $(1 - \rho^2)$ fraction of the internal energy will dissipate at every step, and hence the internal energy at step $k + 1$ is bounded above by the sum of the remaining energy $\rho^2 V(\xi_k)$ and the supply energy S .

3.3.1 How to use dissipation inequality?

Suppose we can construct the dissipation inequality (3.8). What are we going to do about it? The answer is that the dissipation inequality (3.8) can be used to prove stability or convergence rate bounds for $F_u(G, \Delta)$. To make things concrete, let's focus on (3.6) which is a general model for optimization methods.

Notice by definition $V_k \geq 0$ (the internal energy should be non-negative). Typically V is chosen to be a distance metric between ξ_k and the equilibrium point ξ^* . For example, for gradient method, V is chosen as $V = \|x - x^*\|^2$. When applied to analyze optimization methods, the dissipation inequality is typically used to prove two types of bounds.

1. If one already knows $S \leq 0$, then the dissipation inequality (3.8) states $V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k) \leq 0$. This gives a bound $V(\xi_{k+1}) \leq \rho^2 V(\xi_k)$. This proves a linear convergence rate ρ when V is used as a distance metric. We will present such an example by analyzing the gradient method.
2. If one already knows $S \leq \rho^2(f(x_k) - f(x^*)) - (f(x_{k+1}) - f(x^*))$, then the dissipation inequality (3.8) states $V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k) \leq \rho^2(f(x_k) - f(x^*)) - (f(x_{k+1}) - f(x^*))$. This gives a bound $V(\xi_{k+1}) + f(x_{k+1}) - f(x^*) \leq \rho^2(V(\xi_k) + f(x_k) - f(x^*))$. This proves a linear convergence rate ρ when $V(\xi_k) + f(x_k) - f(x^*)$ is used as a distance metric. There is going to be one lecture devoting to cover such an argument for the convergence rate analysis of Nesterov's accelerated method.

3.3.2 How to choose supply rate?

The supply rate S typically takes a form of a quadratic function:

$$S(\xi, w) = \begin{bmatrix} \xi - \xi^* \\ w \end{bmatrix}^T X \begin{bmatrix} \xi - \xi^* \\ w \end{bmatrix} \quad (3.9)$$

where X is some given matrix. The key issue is how to choose X .

Recall that the feedback dynamics $F_u(G, \Delta)$ consists of two parts: $v = G(w)$ and $w = \Delta(v)$. If we want to choose X to guarantee the supply rate S satisfying some inequality, e.g. $S \leq 0$, we need to use the property of Δ .

For example, consider the gradient method. Here Δ is just ∇f . If f is L -smooth and m -strongly convex¹, we know the following inequality holds for any $w_k = \nabla f(C\xi_k)$

$$\begin{bmatrix} C\xi_k - C\xi^* \\ w_k \end{bmatrix} \begin{bmatrix} -2mLI & (m+L)I \\ (m+L)I & -2I \end{bmatrix} \begin{bmatrix} C\xi_k - C\xi^* \\ w_k \end{bmatrix} \geq 0. \quad (3.10)$$

We can simply choose $X = \begin{bmatrix} C^\top & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} 2mLI & -(m+L)I \\ -(m+L)I & 2I \end{bmatrix} \begin{bmatrix} C & 0 \\ 0 & I \end{bmatrix}$ and then the supply rate (3.9) satisfies $S \leq 0$ due to the fact $w_k = \nabla f(C\xi_k)$.

Many control papers focus on developing X for various types of Δ . We will come back to this in next section.

3.3.3 How to construct the dissipation inequality?

Now suppose we have already constructed the supply rate (3.9) with desired properties. How can we construct the dissipation inequality (3.8) for such a supply rate? We can use the following approach.

Theorem 2. Suppose $\xi_{k+1} = A\xi_k + Bw_k$ and $\xi^* = A\xi^*$. Consider a quadratic supply rate (3.9). If there exists a positive semidefinite matrix $P \in \mathbb{R}^{n_\xi \times n_\xi}$ s.t.

$$\begin{bmatrix} A^\top PA - \rho^2 P & A^\top PB \\ B^\top PA & B^\top PB \end{bmatrix} - X \leq 0 \quad (3.11)$$

then we have $V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k)$ with $V(\xi) = (\xi - \xi^*)^\top P(\xi - \xi^*)$.

Proof: Based on (3.11), we directly have

$$\begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix}^\top \left(\begin{bmatrix} A^\top PA - \rho^2 P & A^\top PB \\ B^\top PA & B^\top PB \end{bmatrix} - X \right) \begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix} \leq 0$$

Notice we have $V(\xi_{k+1}) = \begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix}^\top \begin{bmatrix} A^\top PA & A^\top PB \\ B^\top PA & B^\top PB \end{bmatrix} \begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix}$. This immediately leads to the desired conclusion. ■

¹A differentiable function $f : \mathbb{R}^p \rightarrow \mathbb{R}$ is L -smooth if for all $x, y \in \mathbb{R}^p$, one has $\|\nabla f(x) - \nabla f(y)\| \leq L\|x - y\|$. In addition, f is m -strongly convex (for some $m > 0$) if for all $x, y \in \mathbb{R}^p$, one has $f(x) \geq f(y) + \nabla f(y)^\top (x - y) + \frac{m}{2}\|x - y\|^2$. A point $x^* \in \mathbb{R}^n$ is a global min of f if $f(x^*) \leq f(x)$ for all x . When f is m -strongly convex, x^* is unique and satisfies $\nabla f(x^*) = 0$. When f is L -smooth and m -strongly convex, we have $(\nabla f(x) - \nabla f(y))^\top (x - y) \geq \frac{mL}{m+L}\|x - y\|^2 + \frac{1}{m+L}\|\nabla f(x) - \nabla f(y)\|^2$ for all $x, y \in \mathbb{R}^n$. This is equivalent to

$$\begin{bmatrix} x - y \\ \nabla f(x) - \nabla f(y) \end{bmatrix}^\top \begin{bmatrix} -2mLI & (m+L)I \\ (m+L)I & -2I \end{bmatrix} \begin{bmatrix} x - y \\ \nabla f(x) - \nabla f(y) \end{bmatrix} \geq 0.$$

Example: Analysis of the gradient method. Now we apply the above theorem to analyze the gradient method. For the gradient method, we have $A = I$, $B = -\alpha I$, and $C = I$. As discussed in the last section, we can choose the following X to guarantee $S \leq 0$:

$$X = \begin{bmatrix} C^\top & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} 2mLI & -(m+L)I \\ -(m+L)I & 2I \end{bmatrix} \begin{bmatrix} C & 0 \\ 0 & I \end{bmatrix} = \begin{bmatrix} 2mLI & -(m+L)I \\ -(m+L)I & 2I \end{bmatrix}$$

Now it is straightforward to verify that the condition (3.11) leads to the following condition

$$\begin{bmatrix} 1 - \rho^2 & -\alpha \\ -\alpha & \alpha^2 \end{bmatrix} + \lambda \begin{bmatrix} -2mL & m+L \\ m+L & -2 \end{bmatrix} \leq 0 \quad (3.12)$$

if we choose $P = \frac{1}{\lambda}I$. Now we can apply this condition to obtain the convergence rate ρ for the gradient method with various stepsize choices.

- Case 1: If we choose $\alpha = \frac{1}{L}$, $\rho = 1 - \frac{m}{L}$, and $\lambda = \frac{1}{L^2}$, we have

$$\begin{bmatrix} 1 - \rho^2 & -\alpha \\ -\alpha & \alpha^2 \end{bmatrix} + \lambda \begin{bmatrix} -2mL & m+L \\ m+L & -2 \end{bmatrix} = \begin{bmatrix} -\frac{m^2}{L^2} & \frac{m}{L^2} \\ \frac{m}{L^2} & -\frac{1}{L^2} \end{bmatrix} = \frac{1}{L^2} \begin{bmatrix} -m^2 & m \\ m & -1 \end{bmatrix} \quad (3.13)$$

The right side is clearly negative semidefinite due to the fact that $\begin{bmatrix} a \\ b \end{bmatrix}^\top \begin{bmatrix} -m^2 & m \\ m & -1 \end{bmatrix} \begin{bmatrix} a \\ b \end{bmatrix} = -(ma - b)^2 \leq 0$. Therefore, the gradient method with $\alpha = \frac{1}{L}$ converges as

$$\|x_k - x^*\| \leq \left(1 - \frac{m}{L}\right)^k \|x_0 - x^*\| \quad (3.14)$$

- Case 2: If we choose $\alpha = \frac{2}{m+L}$, $\rho = \frac{L-m}{L+m}$, and $\lambda = \frac{2}{(m+L)^2}$, we have

$$\begin{bmatrix} 1 - \rho^2 & -\alpha \\ -\alpha & \alpha^2 \end{bmatrix} + \lambda \begin{bmatrix} -2mL & m+L \\ m+L & -2 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} \quad (3.15)$$

The zero matrix is clearly negative semidefinite. Therefore, the gradient method with $\alpha = \frac{2}{m+L}$ converges as

$$\|x_k - x^*\| \leq \left(\frac{L-m}{L+m}\right)^k \|x_0 - x^*\| \quad (3.16)$$

Notice $L \geq m > 0$ and hence $1 - \frac{m}{L} \geq \frac{L-m}{L+m}$. This means the gradient method with $\alpha = \frac{2}{m+L}$ converges slightly faster than the case with $\alpha = \frac{1}{L}$. However, m is typically unknown in practice. The step choice of $\alpha = \frac{1}{L}$ is also more robust (we will discuss this in later sections). The most popular choice for α is still $\frac{1}{L}$.

The key message in the above example is that to apply the dissipation inequality for linear convergence rate analysis, one typically follows two steps:

1. Choose a proper quadratic supply rate function S satisfying certain desired properties, e.q. $S(\xi_k, w_k) \leq 0$.
2. Find a positive semidefinite matrix P satisfying (3.11) and obtain a quadratic storage function V which is then used to construct the dissipation inequality.

More importantly, many other optimization methods in the form of (3.1) with well-chosen (A, B, C) can be analyzed using similar ideas. We will discuss this in next lectures.

3.3.4 Graphical Interpretation and Conservatism Reduction

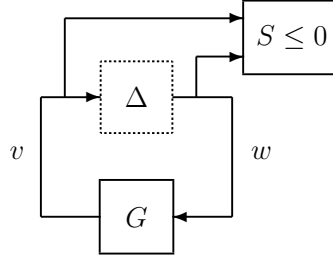


Figure 3.2. Removing Δ by Enforcing the Supply Rate Condition $S \leq 0$

When analyzing $F_u(G, \Delta)$, we aim to draw conclusions on the pair (v, w) in the set $\{(v, w) : v = G(w), w = \Delta(v)\}$. If for any $w = \Delta(v)$, we have $S \leq 0$, then we have

$$\{(v, w) : v = G(w), w = \Delta(v)\} \subset \{(v, w) : v = G(w), S \leq 0\} \quad (3.17)$$

If we can prove ξ_k converges at a certain linear rate for any pair (v, w) in the set $\{(v, w) : v = G(w), S \leq 0\}$, then we guarantee that ξ_k converges at the same linear rate for any pair (v, w) satisfying $v = G(w)$ and $w = \Delta(v)$ simultaneously. Hence we can completely remove the troublesome element Δ from our analysis by enforcing the condition $S \leq 0$. A graphical interpretation for this idea is shown in Figure 3.2. We still have $v = G(w)$. But we remove Δ by enforcing the inequality $S \leq 0$.

Obviously, we are looking at a bigger set $\{(v, w) : v = G(w), S \leq 0\}$. How conservative such a relaxation is depends on whether the worst-case trajectories in these two sets are closed or not. One way to reduce the conservatism in the analysis is to use multiple supply rate functions. Suppose for any $w = \Delta(v)$, we have $S_j \leq 0$ for all $j = 1, 2, \dots, J$. Obviously, if (v, w) satisfies $S_j \leq 0$ for all j , then they also satisfy $S_1 \leq 0$. Hence the set $\{(v, w) : v = G(w), S_j \leq 0 \forall j\}$ is contained in the set $\{(v, w) : v = G(w), S_1 \leq 0\}$. We have

$$\{(v, w) : v = G(w), w = \Delta(v)\} \subset \{(v, w) : v = G(w), S_j \leq 0 \forall j\} \subset \{(v, w) : v = G(w), S_1 \leq 0\}$$

Analyzing the trajectories in $\{(v, w) : v = G(w), S_j \leq 0 \forall j\}$ can potentially lead to less conservative results. Then the key question is how to construct a dissipation inequality when multiple supply rate conditions are available. We can apply the following theorem.

Theorem 3. Suppose $\xi_{k+1} = A\xi_k + Bw_k$ and $\xi^* = A\xi^*$. In addition, the following supply rate functions are given

$$S_j(\xi, w) = \begin{bmatrix} \xi - \xi^* \\ w \end{bmatrix}^T X_j \begin{bmatrix} \xi - \xi^* \\ w \end{bmatrix}, \text{ for } j = 1, 2, \dots, J$$

If there exists a positive semidefinite matrix $P \in \mathbb{R}^{n_\xi \times n_\xi}$ and non-negative scalars λ_j s.t.

$$\begin{bmatrix} A^\top P A - \rho^2 P & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} - \sum_{j=1}^J \lambda_j X_j \leq 0 \quad (3.18)$$

then we have $V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k)$ with $V = (\xi - \xi^*)^\top P (\xi - \xi^*)$ and $S = \sum_{j=1}^J \lambda_j S_j$. In addition, if $S_j \leq 0$, then $V(\xi_{k+1}) \leq \rho^2 V(\xi_k)$.

Proof: Again, we do the same calculations as before.

$$\begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix}^\top \left(\begin{bmatrix} A^\top P A - \rho^2 P & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} - \sum_{j=1}^J \lambda_j X_j \right) \begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix} \leq 0$$

Notice we have $V(\xi_{k+1}) = \begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix}^\top \begin{bmatrix} A^\top P A & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} \begin{bmatrix} \xi_k - \xi^* \\ w_k \end{bmatrix}$. This immediately leads to the desired conclusion. ■

For fixed (A, B, X_j, ρ) , the condition (3.18) is linear in P and λ_j . Hence it is still an LMI which can be solved efficiently with SDP solvers.

Why is (3.18) less conservative than (3.11)? If only one supply rate condition is used (let's say we just use X_1), the resultant LMI condition is just (3.11) with $X = X_1$. In this case, if (3.11) is feasible, then (3.18) is also feasible with $\lambda_1 = 1$, and $\lambda_j = 0$ ($j \neq 1$) (we just choose the same P). The reverse direction is not true. If (3.18) is feasible, (3.11) with $X = X_1$ may not be feasible. Introducing multiple supply rate conditions helps in many situations. In addition, implementing (3.18) is as easy as implementing (3.11). Therefore, it is almost free to include extra supply rate conditions if we only care about obtaining numerical rate certifications. Of course, adding more decision variables could cause trouble for analytical rate proofs. Therefore, a more practical way of doing things is to first use numerical implementation to figure out a minimum number of relevant supply rate conditions and then start analytical proofs with those supply rate conditions.

3.4 Supply Rate and Quadratic Constraints

Now we give more discussions on how to construct supply rate conditions. Many supply rate conditions have already been documented in the controls literature. We will look at a

few basic ones including small gain, passivity, and sector bound. These conditions are also called “quadratic constraints”. For simplicity, we will first talk about the pointwise versions of these conditions. Then we will briefly discuss the “integral” versions of these conditions which are the so-called integral quadratic constraints (IQCs).

3.4.1 Pointwise Quadratic Constraints

Consider a perturbation operator Δ that maps v to w in a static manner, i.e. w_k is completely determined by v_k . The pointwise quadratic constraint just enforces the following inequality on the input/output pair of Δ :

$$\begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix}^\top M \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix} \leq 0, \quad (3.19)$$

where M is a symmetric matrix, and (w^*, v^*) are typically determined by the fixed points of the feedback interconnection $F_u(G, \Delta)$. The terminology “pointwise” just means that we require the above inequality to hold for all k . Clearly, many supply rate conditions that we have used so far are in the form of such pointwise quadratic constraints. Suppose $v_k - v^* = C(\xi_k - \xi^*)$. Then the quadratic constraint (3.19) just gives the following supply rate condition

$$\begin{bmatrix} \xi_k - \xi^* \\ w_k - w^* \end{bmatrix}^\top \left(\begin{bmatrix} C & 0 \\ 0 & I \end{bmatrix}^\top M \begin{bmatrix} C & 0 \\ 0 & I \end{bmatrix} \right) \begin{bmatrix} \xi_k - \xi^* \\ w_k - w^* \end{bmatrix} \leq 0.$$

For now, we just focus on how to obtain the quadratic constraint (3.19).

3.4.2 Small Gain

Suppose Δ is bounded in the sense that we have $\|w_k - w^*\| \leq L\|v_k - v^*\|$. The parameter L can be viewed as the input-output gain of the operator Δ . The small gain bound $\|w_k - w^*\| \leq L\|v_k - v^*\|$ is equivalent to the quadratic inequality $\|w_k - w^*\|^2 - L^2\|v_k - v^*\|^2 \leq 0$ which can be rewritten as the following quadratic constraint:

$$\begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix}^\top \begin{bmatrix} -L^2 I & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix} \leq 0. \quad (3.20)$$

This is the most commonly-used quadratic constraint. Now let’s see a few examples.

- Uncertainty in a multiplicative form: Let Δ map v to w as $w_k = \delta_k v_k$ where δ_k is a matrix changing with k . If we know the Frobenius norm of δ_k is bounded above by L for all k , then we have the small gain bound (3.20) for $(v^*, w^*) = (0, 0)$.
- Gradients of L -smooth functions: Let Δ map v to w as $w_k = \nabla f(v_k)$ where f is L -smooth. Then we have the small gain bound (3.20) holds for any reference point (v^*, w^*) satisfying $w^* = \nabla f(v^*)$.

3.4.3 Passivity

In its simplest form, passivity can be used to describe a function that is in the first and third quadrants. For illustrative purposes, consider a scalar case. Suppose $w_k = \phi(v_k)$ where the function $\phi : \mathbb{R} \rightarrow \mathbb{R}$ satisfies $\phi(0) = 0$ and is in the first and third quadrants. Clearly v_k and w_k are both scalars in this case. If $v_k \geq 0$, we have $w_k \geq 0$. If $v_k \leq 0$, we have $w_k \leq 0$. Hence we always have $w_k^\top v_k \geq 0$. This is the basic form of passivity.

A slightly more general form of passivity gives the constraint $(w_k - w^*)^\top (v_k - v^*) \geq 0$ when (v_k, w_k) are vectors and (potentially non-zero) reference points (v^*, w^*) are used. The passivity condition can be rewritten as the following quadratic constraint (verify it!):

$$\begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix}^\top \begin{bmatrix} 0 & -I \\ -I & 0 \end{bmatrix} \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix} \leq 0. \quad (3.21)$$

Example: Gradients of Convex Functions. Let Δ map v to w as $w_k = \nabla f(v_k)$ where f is a convex function. By definitions, the following inequalities hold for any (v_k, v^*) :

$$\begin{aligned} f(v_k) - f(v^*) &\geq \nabla f(v^*)^\top (v_k - v^*) \\ f(v^*) - f(v_k) &\geq \nabla f(v_k)^\top (v^* - v_k) \end{aligned}$$

Summing the above two inequalities directly leads to the passivity condition $(w_k - w^*)^\top (v_k - v^*) \geq 0$. Therefore, gradients of convex functions satisfy the passivity condition.

3.4.4 Sector Bound

Originally sector bound was used to describe a function that is in a sector formed by two lines whose slopes are m and L . First we consider a scalar case. Suppose $w_k = \phi(v_k)$ where the function $\phi : \mathbb{R} \rightarrow \mathbb{R}$ satisfies $\phi(0) = 0$ and is in a sector formed by two lines whose slopes are m and L . For simplicity, we assume $L \geq m$. Clearly the sector assumption just ensures $(Lv_k - w_k)^\top (w_k - mv_k) \geq 0$. This is the basic form of the sector bound condition.

Now we can introduce the more general form of the sector bound condition that gives the constraint $(L(v_k - v^*) - (w_k - w^*))^\top (w_k - w^* - m(v_k - v^*)) \geq 0$ when (v_k, w_k) are vectors and the reference points (v^*, w^*) are allowed to be non-zero. The sector bound condition can be rewritten as the following quadratic constraint (verify it!):

$$\begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix}^\top \begin{bmatrix} 2mLI & -(m+L)I \\ -(m+L)I & 2I \end{bmatrix} \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix} \leq 0. \quad (3.22)$$

Example: Gradients of L -Smooth m -Strongly Convex Functions. Let Δ map v to w as $w_k = \nabla f(v_k)$ where f is L -smooth and m -strongly convex. Then Δ satisfies the sector bound condition (3.22). We have used this condition to prove the linear convergence rate of the gradient method in the previous lectures.

3.4.5 Integral Quadratic Constraints

In general, Δ is just an operator that maps a sequence $\{v_k\}$ to another sequence $\{w_k\}$. In controls literature, we typically confine Δ to be a causal operator in the sense that w_k is completely determined by $\{v_0, v_1, \dots, v_k\}$. Here Δ is not static anymore. There may be dynamics involved in Δ . Examples include norm-bounded LTI uncertainty and time-varying delays. For such type of Δ , the pointwise quadratic constraints no longer hold. However, the quadratic constraints may hold when we sum them. Specifically, the integral quadratic constraints (IQCs) just enforce the following inequality for any N ,

$$\sum_{k=0}^N \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix}^\top M \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix} \leq 0, \quad (3.23)$$

Originally the above type of quadratic constraints were developed in continuous-time domain where the quadratic forms are integrated over the time horizon. So it is called “integral” quadratic constraints. For discrete-time operators, we just sum things up. We require (3.23) to hold for any N . In controls literature, this type of constraints are “hard” IQCs. We will briefly talk about “soft” IQCs in some future lecture when we discuss the KYP lemma. For now, we focus on hard IQCs that are in the form of (3.23). Hard IQCs can be directly incorporated into the dissipation inequality framework. Typically hard IQCs lead to a supply rate condition $\sum_{k=0}^N S(\xi_k, w_k) \leq 0$. Suppose we have constructed a dissipation inequality $V(\xi_{k+1}) - V(\xi_k) \leq S(\xi_k, w_k)$. Now we do not have $S \leq 0$ for all k . However, we can first sum up the dissipation inequality from $k = 0$ to N to get $V(\xi_{N+1}) \leq V(\xi_0) + \sum_{k=0}^N S(\xi_k, w_k)$. Now using the new supply rate condition $\sum_{k=0}^N S(\xi_k, w_k) \leq 0$, we obtain $V(\xi_{N+1}) \leq V(\xi_0)$. Hence the internal energy is bounded. The physical interpretation is that as long as the total energy supplied to the system (which is equal to $\sum_{k=0}^N S(\xi_k, w_k)$) is non-positive, the internal energy is not going to be larger than the initial energy. We will talk about how to use IQCs for convergence rate analysis later. There is a routine for that.

IQCs are more general than pointwise quadratic constraints. Whenever we have the pointwise quadratic constraint (3.19), we immediately have an IQC in the form of (3.23) by summing the constraints from $k = 0$ to N . The reverse direction is not always true. When Δ has dynamics and memory, it is very common that we will only be able to construct IQCs.

Example: A general version of small gain bound. Consider a general causal operator Δ . A general version of the small gain bound enforces the following inequality for the input/output pair of Δ :

$$\sum_{k=0}^N \|w_k - w^*\|^2 \leq L^2 \sum_{k=0}^N \|v_k - v^*\|^2.$$

This bound is equivalent to the following IQC:

$$\sum_{k=0}^N \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix}^T \begin{bmatrix} -L^2 I & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} v_k - v^* \\ w_k - w^* \end{bmatrix} \leq 0. \quad (3.24)$$

In general, when Δ is the so-called “bounded” operator, we will always have the above small gain IQC. For example, if Δ is an unknown stable LTI system whose $\mathcal{H}_\infty$ norm is L , then we will not have a pointwise small gain bound but (3.24) still holds with $v^* = w^* = 0$. You can verify a similar fact when Δ is a time-varying delay. In many situations, even for static Δ , we can construct useful IQCs to complement the use of pointwise constraints. We will see more examples in future lectures.